

User's Guide

to

PARI / GP

(version 2.17.4)

The PARI Group

Institut de Mathématiques de Bordeaux, UMR 5251 du CNRS.
Université de Bordeaux, 351 Cours de la Libération
F-33405 TALENCE Cedex, FRANCE
e-mail: `pari@math.u-bordeaux.fr`

Home Page:

<https://pari.math.u-bordeaux.fr/>

Copyright © 2000–2024 The PARI Group

Permission is granted to make and distribute verbatim copies of this manual provided the copyright notice and this permission notice are preserved on all copies.

Permission is granted to copy and distribute modified versions, or translations, of this manual under the conditions for verbatim copying, provided also that the entire resulting derived work is distributed under the terms of a permission notice identical to this one.

PARI/GP is Copyright © 2000–2024 The PARI Group

PARI/GP is free software; you can redistribute it and/or modify it under the terms of the GNU General Public License as published by the Free Software Foundation. It is distributed in the hope that it will be useful, but WITHOUT ANY WARRANTY WHATSOEVER.

Table of Contents

Chapter 1: Overview of the PARI system	5
1.1 Introduction	5
1.2 Multiprecision kernels / Portability	6
1.3 The PARI types	7
1.4 The PARI philosophy	9
1.5 Operations and functions	11
Chapter 2: The gp Calculator	13
2.1 Introduction	13
2.2 The general gp input line	15
2.3 The PARI types	17
2.4 GP operators	29
2.5 Variables and symbolic expressions	33
2.6 Variables and Scope	36
2.7 User defined functions	39
2.8 Member functions	48
2.9 Strings and Keywords	49
2.10 Errors and error recovery	51
2.11 Interfacing GP with other languages	57
2.12 Defaults	57
2.13 Simple metacommands	58
2.14 The preferences file	62
2.15 Using readline	64
2.16 GNU Emacs and PariEmacs	66
Chapter 3: Functions and Operations Available in PARI and GP	67
3.1 Programming in GP: control statements	69
3.2 Programming in GP: other specific functions	86
3.3 Parallel programming	119
3.4 GP defaults	123
3.5 Standard monadic or dyadic operators	135
3.6 Conversions and similar elementary functions or commands	144
3.7 Combinatorics	169
3.8 Arithmetic functions	177
3.9 Polynomials and power series	245
3.10 Vectors, matrices, linear algebra and sets	273
3.11 Transcendental functions	313
3.12 Sums, products, integrals and similar functions	336
3.13 General number fields	375
3.14 Associative and central simple algebras	492
3.15 Elliptic curves	521
3.16 Hypergeometric Motives	577
3.17 L -functions	583
3.18 Modular forms	604
3.19 Modular symbols	640
3.20 Plotting functions	661
Index	671

Chapter 1:

Overview of the PARI system

1.1 Introduction.

PARI/GP is a specialized computer algebra system, primarily aimed at number theorists, but has been put to good use in many other different fields, from topology or numerical analysis to physics.

Although quite an amount of symbolic manipulation is possible, PARI does badly compared to systems like Magma, Maple, Mathematica, Maxima, or Sagemath on such tasks, e.g. multivariate polynomials, formal integration, etc. On the other hand, the three main advantages of the system are its speed, the possibility of using directly data types which are familiar to mathematicians, and its extensive algebraic number theory module (in the above-mentioned systems, Magma and Sagemath provide similar features).

Non-mathematical strong points include the possibility to program either in high-level scripting languages or with the PARI library, a mature system (development started in the mid eighties) that was used to conduct and disseminate original mathematical research, while building a large user community, linked by helpful mailing lists and a tradition of great user support from the developers. And, of course, PARI/GP is Free Software, covered by the GNU General Public License, either version 2 of the License or (at your option) any later version.

PARI is used in three different ways:

- 1) as a library `libpari`, which can be called from an upper-level language application, for instance written in ANSI C or C++;
- 2) as a sophisticated programmable calculator, named `gp`, whose language `GP` contains most of the control instructions of a standard language like C;
- 3) the compiler `gp2c` translates `GP` code to C, and loads it into the `gp` interpreter. A typical script compiled by `gp2c` runs 3 to 10 times faster. The generated C code can be edited and optimized by hand. It may also be used as a tutorial to `libpari` programming.

The present Chapter 1 gives an overview of the PARI/GP system; `gp2c` is distributed separately and comes with its own manual. Chapter 2 describes the `GP` programming language and the `gp` calculator. Chapter 3 describes all routines available in the calculator. Programming in library mode is explained in Chapters 4 and 5 in a separate booklet: *User's Guide to the PARI library* (`libpari.pdf`).

A tutorial for `gp` is provided in the standard distribution: *A tutorial for PARI/GP* (`tutorial.pdf`) and you should read this first. You can then start over and read the more boring stuff which lies ahead. You can have a quick idea of what is available by looking at the `gp` general reference card (`refcard.pdf`; other more specialized reference cards are available). In case of need, you can refer to the complete function description in Chapter 3.

How to get the latest version. Everything can be found on PARI's home page:

<https://pari.math.u-bordeaux.fr/>.

From that point you may access all sources, some binaries, version information, the complete mailing list archives, frequently asked questions and various tips. All threaded and fully searchable.

How to report bugs. Bugs are submitted online to our Bug Tracking System, available from PARI's home page, or directly from the URL

<https://pari.math.u-bordeaux.fr/Bugs/>.

Further instructions can be found on that page.

1.2 Multiprecision kernels / Portability.

The PARI multiprecision kernel comes in three non exclusive flavors. See Appendix A for how to set up these on your system; various compilers are supported, but the GNU `gcc` compiler is the definite favorite.

A first version is written entirely in ANSI C, with a C++-compatible syntax, and should be portable without trouble to any 32 or 64-bit computer having no drastic memory constraints. We do not know any example of a computer where a port was attempted and failed.

In a second version, time-critical parts of the kernel are written in inlined assembler. At present this includes

- the whole ix86 family (Intel, AMD, Cyrix) starting at the 386, up to the Xbox gaming console, including the Opteron 64 bit processor.
- three versions for the Sparc architecture: version 7, version 8 with SuperSparc processors, and version 8 with MicroSparc I or II processors. UltraSparcs use the MicroSparc II version;
- the DEC Alpha 64-bit processor;
- the Intel Itanium 64-bit processor;
- the PowerPC equipping old macintoshs (G3, G4, etc.);
- the HPPA processors (both 32 and 64 bit);
- the MIPS processors (both 32 and 64 bit);
- the RISC-V 64 bit processors.

A third version uses the GNU MP library to implement most of its multiprecision kernel. It improves significantly on the native one for large operands, say 100 decimal digits of accuracy or more. You *should* enable it if GMP is present on your system. Parts of the first version are still in use within the GMP kernel, but are scheduled to disappear.

A historical version of the PARI/GP kernel, written in 1985, was specific to 680x0 based computers, and was entirely written in MC68020 assembly language. It ran on SUN-3/xx, Sony News, NeXT cubes and on 680x0 based Macs. It is no longer part of the PARI distribution; to run PARI with a 68k assembler micro-kernel, use the GMP kernel!

Mathematical notations and conventions.

- Standard rings and fields. We denote $\mathbf{Z}$ the ring of integers, $\mathbf{Q}$ the field of rational numbers, $\mathbf{R}$ the field of real numbers and $\mathbf{C}$ the field of complex numbers (containing an element i such that $i^2 = -1$). Given a prime power q , $\mathbf{F}_q$ denotes the finite field with q elements. Given a prime number p , v_p denotes the p -adic valuation $\mathbf{Z}_p$ is ring of p -adic integers, $\mathbf{Q}_p$ the field of p -adic numbers and $\mathbf{C}_p$ the p -adic completion of the algebraic closure of $\mathbf{Q}_p$. We write $|x|_p = p^{-v_p(x)}$ for $x \in \mathbf{C}_p$.

- Intervals. We write $[a, b]$ for the closed interval $\{x \in \mathbf{R}: a \leq x \leq b\}$, $]a, b[$ for the open interval $\{x \in \mathbf{R}: a < x < b\}$ and similarly $]a, b]$ and $[a, b[$ for half-open intervals.

- Linear Algebra. Let K be some field and $m, n \geq 0$ be integers. Elements in the vector space K^n are represented as *column* vectors (of length n). Elements of the algebra $\text{Hom}_K(K^n, K^m)$ are represented as $m \times n$ matrices; due to an unfortunate historical design decision, $m \times 0$ matrices do not exist in PARI unless $m = 0$. If M is an $m \times n$ matrix, we use the notation tM to denote its transpose (an $n \times m$ matrix). The (right) *kernel* of a matrix M is the vector space $\{v \in K^n: Mv = 0\}$. Similarly, the image of M is the span of its columns.

1.3 The PARI types.

The GP language is not typed in the traditional sense; in particular, variables have no type. In library mode, the type of all PARI objects is **GEN**, a generic type. On the other hand, it is dynamically typed: each object has a specific internal type, depending on the mathematical object it represents.

The crucial word is recursiveness: most of the PARI types are recursive. For example, the basic internal type **t_COMPLEX** exists. However, the components (i.e. the real and imaginary part) of such a “complex number” can be of any type. The only sensible ones are integers (we are then in $\mathbf{Z}[i]$), rational numbers ($\mathbf{Q}[i]$), real numbers ($\mathbf{R}[i] = \mathbf{C}$), or even elements of $\mathbf{Z}/n\mathbf{Z}$ (in $(\mathbf{Z}/n\mathbf{Z})[t]/(t^2+1)$), or p -adic numbers when $p \equiv 3 \pmod{4}$ ($\mathbf{Q}_p[i]$). This feature must not be used too rashly in library mode: for example you are in principle allowed to create objects which are “complex numbers of complex numbers”. (This is not possible under **gp**.) But do not expect PARI to make sensible use of such objects: you will mainly get nonsense.

On the other hand, it *is* allowed to have components of different, but compatible, types, which can be freely mixed in basic ring operations $+$ or $\times$. For example, taking again complex numbers, the real part could be an integer, and the imaginary part a rational number. On the other hand, if the real part is a real number, the imaginary part cannot be an integer modulo n !

Let us now describe the types. As explained above, they are built recursively from basic types which are as follows. We use the letter T to designate any type; the symbolic names **t_XXX** correspond to the internal representations of the types.

type t_INT	$\mathbf{Z}$	Integers (with arbitrary precision)
type t_REAL	$\mathbf{R}$	Real numbers (with arbitrary precision)
type t_INTMOD	$\mathbf{Z}/n\mathbf{Z}$	Intmods (integers modulo n)
type t_FRAC	$\mathbf{Q}$	Rational numbers (in irreducible form)
type t_FFELT	$\mathbf{F}_q$	Finite field element
type t_COMPLEX	$T[i]$	Complex numbers
type t_PADIC	$\mathbf{Q}_p$	p -adic numbers
type t_QUAD	$\mathbf{Q}[w]$	Quadratic Numbers (where $[\mathbf{Z}[w] : \mathbf{Z}] = 2$)
type t_POLMOD	$T[X]/(P)$	Polmods (polynomials modulo $P \in T[X]$)

type <code>t_POL</code>	$T[X]$	Polynomials
type <code>t_SER</code>	$T((X))$	Power series (finite Laurent series)
type <code>t_RFRAC</code>	$T(X)$	Rational functions (in irreducible form)
type <code>t_VEC</code>	T^n	Row (i.e. horizontal) vectors
type <code>t_COL</code>	T^n	Column (i.e. vertical) vectors
type <code>t_MAT</code>	$\mathcal{M}_{m,n}(T)$	Matrices
type <code>t_LIST</code>	T^n	Lists
type <code>t_STR</code>		Character strings
type <code>t_CLOSURE</code>		Functions
type <code>t_ERROR</code>		Error messages
type <code>t_INFINITY</code>		$-\infty$ and $+\infty$

and where the types T in recursive types can be different in each component. The first nine basic types, from `t_INT` to `t_POLMOD`, are called scalar types because they essentially occur as coefficients of other more complicated objects. Type `t_POLMOD` is used to define algebraic extensions of a base ring, and as such is a scalar type.

In addition, there exist the type `t_QFB` for integral binary quadratic forms, and the internal type `t_VECSMALL`. The latter holds vectors of small integers, whose absolute value is bounded by 2^{31} (resp. 2^{63}) on 32-bit, resp. 64-bit, machines. They are used internally to represent permutations, polynomials or matrices over a small finite field, etc.

Every PARI object (called `GEN` in the sequel) belongs to one of these basic types. Let us have a closer look.

1.3.1 Integers and reals. They are of arbitrary and varying length (each number carrying in its internal representation its own length or precision) with the following mild restrictions (given for 32-bit machines, the restrictions for 64-bit machines being so weak as to be considered nonexistent): integers must be in absolute value less than $2^{536870815}$ (i.e. roughly 161614219 decimal digits). The precision of real numbers is also at most 161614219 significant decimal digits, and the binary exponent must be in absolute value less than 2^{29} , resp. 2^{61} , on 32-bit, resp. 64-bit machines.

Integers and real numbers are nonrecursive types.

1.3.2 Intmods, rational numbers, p -adic numbers, polmods, and rational functions. These are recursive, but in a restricted way.

For intmods or polmods, there are two components: the modulus, which must be of type integer (resp. polynomial), and the representative number (resp. polynomial).

For rational numbers or rational functions, there are also only two components: the numerator and the denominator, which must both be of type integer (resp. polynomial).

Finally, p -adic numbers have three components: the prime p , the “modulus” p^k , and an approximation to the p -adic number. Here $\mathbf{Z}_p$ is considered as the projective limit $\varprojlim \mathbf{Z}/p^k \mathbf{Z}$ via its finite quotients, and $\mathbf{Q}_p$ as its field of fractions. Like real numbers, the codewords contain an exponent, giving the p -adic valuation of the number, and also the information on the precision of the number, which is redundant with p^k , but is included for the sake of efficiency.

1.3.3 Finite field elements. The exact internal format depends of the finite field size, but it includes the field characteristic p , an irreducible polynomial $T \in \mathbf{F}_p[X]$ defining the finite field $\mathbf{F}_p[X]/(T)$ and the element expressed as a polynomial in (the class of) X .

1.3.4 Complex numbers and quadratic numbers. Quadratic numbers are numbers of the form $a + bw$, where w is such that $[\mathbf{Z}[w] : \mathbf{Z}] = 2$, and more precisely $w = \sqrt{d}/2$ when $d \equiv 0 \pmod{4}$, and $w = (1 + \sqrt{d})/2$ when $d \equiv 1 \pmod{4}$, where d is the discriminant of a quadratic order. Complex numbers correspond to the important special case $w = \sqrt{-1}$.

Complex numbers are partially recursive: the two components a and b can be of type `t_INT`, `t_REAL`, `t_INTMOD`, `t_FRAC`, or `t_PADIC`, and can be mixed, subject to the limitations mentioned above. For example, $a + bi$ with a and b p -adic is in $\mathbf{Q}_p[i]$, but this is equal to $\mathbf{Q}_p$ when $p \equiv 1 \pmod{4}$, hence we must exclude these p when one explicitly uses a complex p -adic type. Quadratic numbers are more restricted: their components may be as above, except that `t_REAL` is not allowed.

1.3.5 Polynomials, power series, vectors, matrices. They are completely recursive, over a commutative base ring: their components can be of any type, and types can be mixed (however beware when doing operations). Note in particular that a polynomial in two variables is simply a polynomial with polynomial coefficients. Polynomials or matrices over noncommutative rings are not supported.

In the present version 2.17.4 of PARI, it is not possible to handle conveniently power series of power series, i.e. power series in several variables. However power series of polynomials (which are power series in several variables of a special type) are OK. This is a difficult design problem: the mathematical problem itself contains some amount of imprecision, and it is not easy to design an intuitive generic interface for such beasts.

1.3.6 Strings. These contain objects just as they would be printed by the `gp` calculator.

1.3.7 Zero. What is zero? This is a crucial question in all computer systems. The answer we give in PARI is the following. For exact types, all zeros are equivalent and are exact, and thus are usually represented as an integer zero. The problem becomes nontrivial for imprecise types: there are infinitely many distinct zeros of each of these types! For p -adics and power series the answer is as follows: every such object, including 0, has an exponent e . This p -adic or X -adic zero is understood to be equal to $O(p^e)$ or $O(X^e)$ respectively.

Real numbers also have exponents and a real zero is in fact $O(2^e)$ where e is now usually a negative binary exponent. This of course is printed as usual for a floating point number ($0.00 \dots$ or $0.Exx$ depending on the output format) and not with a O symbol as with p -adics or power series. With respect to the natural ordering on the reals we make the following convention: whatever its exponent a real zero is smaller than any positive number, and any two real zeroes are equal.

1.4 The PARI philosophy.

The basic principles which govern PARI is that operations and functions should, firstly, give as exact a result as possible, and secondly, be permitted if they make any kind of sense.

In this respect, we make an important distinction between exact and inexact objects: by definition, types `t_REAL`, `t_PADIC` or `t_SER` are imprecise. A PARI object having one of these imprecise types anywhere in its tree is *inexact*, and *exact* otherwise. No loss of accuracy (rounding error) is involved when dealing with exact objects. Specifically, an exact operation between exact objects will yield an exact object. For example, dividing 1 by 3 does not give $0.333\dots$, but the rational number $(1/3)$. To get the result as a floating point real number, evaluate `1./3` or `0.+1/3`.

Conversely, the result of operations between imprecise objects, although inexact by nature, will be as precise as possible. Consider for example the addition of two real numbers x and y . The accuracy of the result is *a priori* unpredictable; it depends on the precisions of x and y , on their sizes, and also on the size of $x + y$. From this data, PARI works out the right precision for the result. Even if it is working in calculator mode `gp`, where there is a notion of default precision, its value is only used to convert exact types to inexact ones.

In particular, if an operation involves objects of different accuracies, some digits will be disregarded by PARI. It is a common source of errors to forget, for instance, that a real number is given as $r + 2^e \varepsilon$ where r is a rational approximation, e a binary exponent and ε is a nondescript real number less than 1 in absolute value. Hence, any number less than 2^e may be treated as an exact zero:

```
? 0.E-38 + 1.E-100
%1 = 0.E-38
? 0.E100 + 1
%2 = 0.E100
```

As an exercise, if `a = 2^(-100)`, why do `a + 0.` and `a * 1.` differ?

The second principle is that PARI operations are in general quite permissive. For instance taking the exponential of a vector should not make sense. However, it frequently happens that one wants to apply a given function to all elements in a vector. This is easily done using a loop, or using the `apply` built-in function, but in fact PARI assumes that this is exactly what you want to do when you apply a scalar function to a vector. Taking the exponential of a vector will do just that, so no work is necessary. Most transcendental functions work in the same way*.

In the same spirit, when objects of different types are combined they are first automatically mapped to a suitable ring, where the computation becomes meaningful:

```
? 1/3 + Mod(1,5)
%1 = Mod(3, 5)
? I + 0(5^9)
%2 = 2 + 5 + 2*5^2 + 5^3 + 3*5^4 + 4*5^5 + 2*5^6 + 3*5^7 + 0(5^9)
? Mod(1,15) + Mod(1,10)
%3 = Mod(2, 5)
```

The first example is straightforward: since 3 is invertible mod 5, $(1/3)$ is easily mapped to $\mathbf{Z}/5\mathbf{Z}$. In the second example, `I` stands for the customary square root of -1 ; we obtain a 5-adic number, 5-adically close to a square root of -1 . The final example is more problematic, but there are natural maps from $\mathbf{Z}/15\mathbf{Z}$ and $\mathbf{Z}/10\mathbf{Z}$ to $\mathbf{Z}/5\mathbf{Z}$, and the computation takes place there.

* An ambiguity arises with square matrices. PARI always considers that you want to do componentwise function evaluation in this context, hence to get for example the standard exponential of a square matrix you would need to implement a different function.

1.5 Operations and functions.

The available operations and functions in PARI are described in detail in Chapter 3. Here is a brief summary:

1.5.1 Standard arithmetic operations.

Of course, the four standard operators $+$, $-$, $*$, $/$ exist. We emphasize once more that division is, as far as possible, an exact operation: 4 divided by 3 gives $(4/3)$. In addition to this, operations on integers or polynomials, like $\backslash$ (Euclidean division), $\%$ (Euclidean remainder) exist; for integers, $\backslash/$ computes the quotient such that the remainder has smallest possible absolute value. There is also the exponentiation operator $^$, when the exponent is of type integer; otherwise, it is considered as a transcendental function. Finally, the logical operators $!$ (not prefix operator), $\&\&$ (and operator), $||$ (or operator) exist, giving as results 1 (true) or 0 (false).

1.5.2 Conversions and similar functions.

Many conversion functions are available to convert between different types. For example floor, ceiling, rounding, truncation, etc. . . . Other simple functions are included like real and imaginary part, conjugation, norm, absolute value, changing precision or creating an intmod or a polmod.

1.5.3 Transcendental functions.

They usually operate on any complex number, power series, and some also on p -adics. The list is ever-expanding and of course contains all the elementary functions (exp/log, trigonometric functions), plus many others (modular functions, Bessel functions, polylogarithms. . .). Recall that by extension, PARI usually allows a transcendental function to operate componentwise on vectors or matrices.

1.5.4 Arithmetic functions.

Apart from a few like the factorial function or the Fibonacci numbers, these are functions which explicitly use the prime factor decomposition of integers. The standard functions are included. A number of factoring methods are used by a rather sophisticated factoring engine (to name a few, Shanks's SQUFOF, Pollard's rho, Lenstra's ECM, the MPQS quadratic sieve). These routines output strong pseudoprimes, which may be certified by the APRCL test.

There is also a large package to work with algebraic number fields. All the usual operations on elements, ideals, prime ideals, etc. are available. More sophisticated functions are also implemented, like solving Thue equations, finding integral bases and discriminants of number fields, computing class groups and fundamental units, computing in relative number field extensions, Galois and class field theory, and also many functions dealing with elliptic curves over $\mathbf{Q}$ or over local fields.

1.5.5 Other functions.

Quite a number of other functions dealing with polynomials (e.g. finding complex or p -adic roots, factoring, etc), power series (e.g. substitution, reversion), linear algebra (e.g. determinant, characteristic polynomial, linear systems), and different kinds of recursions are also included. In addition, standard numerical analysis routines like univariate integration (using the double exponential method), real root finding (when the root is bracketed), polynomial interpolation, infinite series evaluation, and plotting are included.

And now, you should really have a look at the tutorial before proceeding.

Chapter 2:

The gp Calculator

2.1 Introduction.

Originally, `gp` was designed as a debugging device for the PARI system library. Over the years, it has become a powerful user-friendly stand-alone calculator. The mathematical functions available in PARI and `gp` are described in the next chapter. In the present one, we describe the specific use of the `gp` programmable calculator.

EMACS: If you have GNU Emacs and use the PariEmacs package, you can work in a special Emacs shell, described in Section 2.16. Specific features of this Emacs shell are indicated by an EMACS sign in the left margin.

We briefly mention at this point GNU TeXmacs (<https://www.texmacs.org/>), a free wysiwyg editing platform that allows to embed an entire `gp` session in a document, and provides a nice alternative to PariEmacs.

2.1.1 Startup.

To start the calculator, the general command line syntax is:

```
gp [-D key=val] [files]
```

where items within brackets are optional. The `[files]` argument is a list of files written in the GP scripting language, which will be loaded on startup. There can be any number of arguments of the form `-D key=val`, setting some internal parameters of `gp`, or *defaults*: each sets the default *key* to the value *val*. See Section 2.12 below for a list and explanation of all defaults. These defaults can be changed by adding parameters to the input line as above, or interactively during a `gp` session, or in a preferences file also known as `gprc`.

If a preferences file (to be discussed in Section 2.14) is found, `gp` then reads it and executes the commands it contains. This provides an easy way to customize `gp`. The *files* argument is processed right after the `gprc`.

A copyright banner then appears which includes the version number, and a lot of useful technical information. After the copyright, the computer writes the top-level help information, some initial defaults, and then waits after printing its prompt, which is `'? '` by default. Whether extended on-line help and line editing are available or not is indicated in this `gp` banner, between the version number and the copyright message. Consider investigating the matter with the person who installed `gp` if they are not. Do this as well if there is no mention of the GMP kernel.

2.1.2 Getting help.

To get help, type a `?` and hit return. A menu appears, describing the main categories of available functions and how to get more detailed help. If you now type `?n` with $n = 1, 2, \dots$, you get the list of commands corresponding to category n and simultaneously to Section 3. n of this manual. If you type `?functionname` where *functionname* is the name of a PARI function, you will get a short explanation of this function.

If extended help (see Section 2.13.1) is available on your system, you can double or triple the `?` sign to get much more: respectively the complete description of the function (e.g. `??sqrt`), or a list of `gp` functions relevant to your query (e.g. `??? "elliptic curve"` or `??? "quadratic field"`).

If `gp` was properly installed (see Appendix A), a line editor is available to correct the command line, get automatic completions, and so on. See Section 2.15 or `??readline` for a short summary of the line editor's commands.

If you type `?\` you will get a short description of the metacommands (keyboard shortcuts).

Finally, typing `?.` will return the list of available (pre-defined) member functions. These are functions attached to specific kind of objects, used to retrieve easily some information from complicated structures (you can define your own but they won't be shown here). We will soon describe these commands in more detail.

More generally, commands starting with the symbols `\` or `?`, are not computing commands, but are metacommands which allow you to exchange information with `gp`. The available metacommands can be divided into default setting commands (explained below) and simple commands (or keyboard shortcuts, to be dealt with in Section 2.13).

2.1.3 Input.

Just type in an instruction, e.g. `1 + 1`, or `Pi`. No action is undertaken until you hit the `<Return>` key. Then computation starts, and a result is eventually printed. To suppress printing of the result, end the expression with a `;` sign. Note that many systems use `;` to indicate end of input. Not so in `gp`: a final semicolon means the result should not be printed. (Which is certainly useful if it occupies several screens.)

2.1.4 Interrupt, Quit.

Typing `quit` at the prompt ends the session and exits `gp`. At any point you can type `Ctrl-C` (that is press simultaneously the `Control` and `C` keys): the current computation is interrupted and control given back to you at the `gp` prompt, together with a message like

```
*** at top-level: gcd(a,b)
***          ^-----
*** gcd: user interrupt after 236 ms.
```

telling you how much time elapsed since the last command was typed in and in which GP function the computation was aborted. It does not mean that that much time was spent in the function, only that the evaluator was busy processing that specific function when you stopped it.

2.2 The general gp input line.

The **gp** calculator uses a purely interpreted language GP. The structure of this language is reminiscent of LISP with a functional notation, $f(x,y)$ rather than $(f\ x\ y)$: all programming constructs, such as **if**, **while**, etc...are functions*, and the main loop does not really execute, but rather evaluates (sequences of) expressions. Of course, it is by no means a true LISP, and has been strongly influenced by C and Perl since then.

2.2.1 Introduction. User interaction with a **gp** session proceeds as follows. First, one types a sequence of characters at the **gp** prompt; see Section 2.15 for a description of the line editor. When you hit the <Return> key, **gp** gets your input, evaluates it, then prints the result and assigns it to an “history” array.

More precisely, the input is case-sensitive and, outside of character strings, blanks are completely ignored. Inputs are either metacommands or sequences of expressions. Metacommands are shortcuts designed to alter **gp**’s internal state, such as the working precision or general verbosity level; we shall describe them in Section 2.13, and ignore them for the time being.

The evaluation of a sequence of instructions proceeds in two phases: your input is first digested (byte-compiled) to a bytecode suitable for fast evaluation, in particular loop bodies are compiled only once but a priori evaluated many times; then the bytecode is evaluated.

An expression is formed by combining constants, variables, operator symbols, functions and control statements. It is evaluated using the conventions about operator priorities and left to right associativity. An expression always has a value, which can be any PARI object:

```
? 1 + 1
%1 = 2          \\ an ordinary integer
? x
%2 = x          \\ a polynomial of degree 1 in the unknown x
? print("Hello")
Hello           \\ void return value, 'Hello' printed as side effect
? f(x) = x^2
%4 = (x)->x^2    \\ a user function
```

In the third example, **Hello** is printed as a side effect, but is not the return value. The **print** command is a *procedure*, which conceptually returns nothing. But in fact procedures return a special **void** object, meant to be ignored (but which evaluates to 0 in a numeric context, and stored as 0 in the history or results). The final example assigns to the variable **f** the function $x \mapsto x^2$, the alternative form **f = x->x^2** achieving the same effect; the return value of a function definition is, unsurprisingly, a function object (of type **t_CLOSURE**).

Several expressions are combined on a single line by separating them with semicolons (**;**). Such an expression sequence will be called a *seq*. A *seq* also has a value, which is the value of the last expression in the sequence. Under **gp**, the value of the *seq*, and only this last value, becomes an history entry. The values of the other expressions in the *seq* are discarded after the execution of the *seq* is complete, except of course if they were assigned into variables. In addition, the value of the *seq* is printed if the line does not end with a semicolon **;**.

* Not exactly, since not all their arguments need be evaluated. For instance it would be stupid to evaluate both branches of an **if** statement: since only one will apply, only this one is evaluated.

2.2.2 The gp history of results.

This is not to be confused with the history of your *commands*, maintained by `readline`. The `gp` history contains the *results* they produced, in sequence.

The successive elements of the history array are called `%1`, `%2`, ... As a shortcut, the latest computed expression can also be called `%`, the previous one `%'`, the one before that `%''` and so on.

When you suppress the printing of the result with a semicolon, it is still stored in the history, but its history number will not appear either. It is a better idea to assign it to a variable for later use than to mentally recompute what its number is. Of course, on the next line, you may just use `%`.

The time used to compute that history entry is also stored as part of the entry and can be recovered using the `%#` operator: `%#1`, `%#2`, `%#'`; `%#` by itself returns the time needed to compute the last result (the one returned by `%`). The output is a vector with two components `[cpu, real]` where `cpu` is the CPU time and `real` is the wall clock time.

Remark. The history “array” is in fact better thought of as a queue: its size is limited to 5000 entries by default, after which `gp` starts forgetting the initial entries. So `%1` becomes unavailable as `gp` prints `%5001`. You can modify the history size using `histsize`.

2.2.3 Special editing characters. A GP program can of course have more than one line. Since your commands are executed as soon as you have finished typing them, there must be a way to tell `gp` to wait for the next line or lines of input before doing anything. There are three ways of doing this.

The first one is to use the backslash character `\` at the end of the line that you are typing, just before hitting `<Return>`. This tells `gp` that what you will write on the next line is the physical continuation of what you have just written. In other words, it makes `gp` forget your newline character. You can type a `\` anywhere. It is interpreted as above only if (apart from ignored whitespace characters) it is immediately followed by a newline. For example, you can type

```
? 3 + \  
4
```

instead of typing `3 + 4`.

The second one is a variation on the first, and is mostly useful when defining a user function (see Section 2.7): since an equal sign can never end a valid expression, `gp` disregards a newline immediately following an `=`.

```
? a =  
123  
%1 = 123
```

The third one is in general much more useful, and uses braces `{` and `}`. An opening brace `{` signals that you are typing a multi-line command, and newlines are ignored until you type a closing brace `}`. There are two important, but easily obeyed, restrictions: first, braces do not nest; second, inside an open brace-close brace pair, all input lines are concatenated, suppressing any newlines. Thus, all newlines should occur after a semicolon `;`, a comma `,` or an operator (for clarity's sake, never split an identifier over two lines in this way). For instance, the following program

```
{  
  a = b
```

```

    b = c
}

```

would silently produce garbage, since this is interpreted as `a=bb=c` which assigns the value of `c` to both `bb` and `a`. It should have been written

```

{
    a = b;
    b = c;
}

```

2.3 The PARI types.

We see here how to input values of the different data types known to PARI. Recall that blanks are ignored in any expression which is not a string (see below).

A note on efficiency. The following types are provided for convenience, not for speed: `t_INTMOD`, `t_FRAC`, `t_PADIC`, `t_QUAD`, `t_POLMOD`, `t_RFRAC`. Indeed, they always perform a reduction of some kind after each basic operation, even though it is usually more efficient to perform a single reduction at the end of some complex computation. For instance, in a convolution product $\sum_{i+j=n} x_i y_j$ in $\mathbf{Z}/N\mathbf{Z}$ — common when multiplying polynomials! —, it is quite wasteful to perform n reductions modulo N . In short, basic individual operations on these types are fast, but recursive objects with such components could be handled more efficiently: programming with `libpari` will save large constant factors here, compared to GP.

2.3.1 Integers (`t_INT`). After an (optional) leading `+` or `-`, type in the decimal digits of your integer. No decimal point!

```

? 1234567
%1 = 1234567
? -3
%2 = -3
? 1.          \\ oops, not an integer
%3 = 1.000000000000000000000000000000

```

Integers can be input in hexadecimal notation by prefixing them with `0x`; hexadecimal digits $(a, \dots, f)$ can be input either in lowercase or in uppercase:

```

? 0xF
%4 = 15
? 0x1abcd
%5 = 109517

```

Integers can also be input in binary by prefixing them with `0b`:

```

? 0b010101
%6 = 21

```

2.3.2 Real numbers (`t_REAL`).

Real numbers are represented (approximately) in a floating point system, internally in base 2, but converted to base 10 for input / output purposes. A `t_REAL` object has a given *bit accuracy* (or *bit precision*) $\ell \geq 0$; it comprises

- a sign s : $+1$, -1 or 0 ;
- a mantissa m : a multiprecision integer, $0 \leq m < 2^\ell$;
- an exponent e : a small integer in $[-2^B, 2^B[$, where $B = 31$ on a 32-bit machine and 63 otherwise.

This data may represent any real number x such that

$$|x - sm2^e| < 2^{e-\ell}.$$

We consider that a `t_REAL` with sign $s = 0$ has accuracy $\ell = 0$, so that its mantissa is useless, but it still has an exponent e and acts like a machine epsilon for all accuracies $< e$.

After an (optional) leading $+$ or $-$, type a number with a decimal point. Leading zeroes may be omitted, up to the decimal point, but trailing zeroes are important: your `t_REAL` is assigned an internal precision, which is the supremum of the input precision, one more than the number of decimal digits input, and the default `realprecision`. For example, if the default precision is 38 digits, typing `2.` yields a precision of 38 digits, but `2.0...0` with 45 zeros gives a number with internal decimal precision at least 45, although less may be printed.

You can also use scientific notation with the letter `E` or `e`. As usual, `en` is interpreted as $\times 10^n$ for all integers n . Since the result is converted to a `t_REAL`, you may often omit the decimal point in this case: `6.02 E 23` or `1e-5` are fine, but `e10` is not.

By definition, `0.E n` returns a real 0 of exponent n , whereas `0.` returns a real 0 “of default precision” (of exponent $-\text{realprecision}$), see Section 1.3.7, behaving like the machine epsilon for the current default accuracy: any float of smaller absolute value is indistinguishable from 0.

Note on output formats. A zero real number is printed in `e` format as `0.Exx` where xx is the (usually negative) *decimal* exponent of the number (cf. Section 1.3.7). This allows the user to check the accuracy of that particular zero.

When the integer part of a real number x is not known exactly because the exponent of x is greater than the internal precision, the real number is printed in `e` format.

Technical note. The internal *precision* is actually expressed in bits and can be viewed and manipulated globally in interactive use via `realprecision` (decimal digits, as explained above; shortcut `\p`) or `realbitprecision` (bits; shortcut `\pb`), the latter allowing finer granularity. See Section 3.11 for details. In programs we advise to leave this global variable alone and adapt precision locally for a given sequence of computations using `localbitprec`.

Note that most decimal floating point numbers cannot be converted exactly in binary, the (binary) number actually stored is a rounded version of the (decimal) number input. Analogously, a decimal output is rounded from the internal binary representation.

2.3.3 Intmods (`t_INTMOD`). To create the image of the integer a in $\mathbf{Z}/b\mathbf{Z}$ (for some nonzero integer b), type `Mod(a,b)`; *not* `a%b`. Internally, all operations are done on integer representatives belonging to $[0, b - 1]$.

Note that this type is available for convenience, not for speed: each elementary operation involves a reduction modulo b .

If x is a `t_INTMOD` `Mod(a,b)`, the following member function is defined:

`x.mod`: return the modulus b .

2.3.4 Rational numbers (`t_FRAC`). All fractions are automatically reduced to lowest terms, so it is impossible to work with reducible fractions. To enter n/m just type it as written. As explained in Section 3.5.8, floating point division is *not* performed, only reduction to lowest terms.

Note that rational computation are almost never the fastest method to proceed: in the PARI implementation, each elementary operation involves computing a gcd. It is generally a little more efficient to cancel denominators and work with integers only:

```
? P = Pol( vector(10^3,i, 1/i) ); \\ big polynomial with small rational coeffs
? P^2
time = 1,392 ms.
? c = content(P); c^2 * (P/c)^2; \\ same computation in integers
time = 1,116 ms.
```

And much more efficient (but harder to setup) to use homomorphic imaging schemes and modular computations. As the simple example below indicates, if you only need modular information, it is very worthwhile to work with `t_INTMOD`s directly, rather than deal with `t_FRAC`s all the way through:

```
? p = nextprime(10^7);
? sum(i=1, 10^5, 1/i) % p
time = 13,288 ms.
%1 = 2759492
? sum(i=1, 10^5, Mod(1/i, p))
time = 60 ms.
%2 = Mod(2759492, 10000019)
```

2.3.5 Finite field elements (`t_FFELT`). Let $T \in \mathbf{F}_p[X]$ be a monic irreducible polynomial defining your finite field over $\mathbf{F}_p$, for instance obtained using `ffinit`. Then the `ffgen` function creates a generator of the finite field as an $\mathbf{F}_p$ -algebra, namely the class of X in $\mathbf{F}_p[X]/(T)$, from which you can build all other elements. For instance, to create the field $\mathbf{F}_{2^8}$, we write

```
? T = ffinit(2, 8);
? y = ffgen(T, 'y);
? y^0 \\ the unit element in the field
%3 = 1
? y^8
%4 = y^6 + y^5 + y^4 + y^3 + y + 1
```

The second (optional) parameter to `ffgen` is only used to display the result; it is customary to use the name of the variable we assign the generator to. If g is a `t_FFELT`, the following member functions are defined:

g.pol: the polynomial (with reduced integer coefficients) expressing **g** in term of the field generator.

g.p: the characteristic of the finite field.

g.f: the dimension of the definition field over its prime field; the cardinality of the definition field is thus p^f .

g.mod: the minimal polynomial (with reduced integer coefficients) of the field generator.

2.3.6 Complex numbers (t_COMPLEX). To enter $x + iy$, type **x + I*y**. (That's I, *not* i!) The letter I stands for $\sqrt{-1}$. The “real” and “imaginary” parts x and y can be of type **t_INT**, **t_REAL**, **t_INTMOD**, **t_FRAC**, or **t_PADIC**.

2.3.7 p-adic numbers (t_PADIC):. Typing **0(p^k)**, where p is a prime and k is an integer, yields a p -adic 0 of accuracy k , representing any p -adic number whose valuation is $\geq k$. To input a general nonzero p -adic number, write a suitably precise rational or integer approximation and add **0(p^k)** to it. For example, you can type in the 7-adic number

2*7^(-1) + 3 + 4*7 + 2*7^2 + 0(7^3)

exactly as shown, or equivalently as **905/7 + 0(7^3)**.

Note that it is not checked whether p is indeed prime but results are undefined if this is not the case: you can try to work on 10-adics if you want, but disasters will happen as soon as you do something nontrivial. For instance:

```
? t = 2 * (1/10 + 0(10^5));
? lift(t)
%2 = 2/10  \\ not reduced (invalid t_FRAC)
? factor(x^2-t)
***   at top-level: factor(x^2-%1)
***                               ^-----
*** factor: impossible inverse in F1_inv: Mod(2, 10000).
```

Note that **0(25)** is not the same as **0(5^2)**; you want the latter!

If a is a **t_PADIC**, the following member functions are defined:

a.mod: returns the modulus p^k .

a.p: returns p .

Note that this type is available for convenience, not for speed: internally, **t_PADICs** are stored as p -adic units modulo some p^k . Each elementary operation involves updating p^k (multiplying or dividing by powers of p) and a reduction mod p^k . In particular, additions are slow.

```
? n = 1+0(2^20);   for (i=1,10^6, n++)
time = 841 ms.
? n = Mod(1,2^20); for (i=1,10^6, n++)
time = 441 ms.
? n = 1;           for (i=1,10^6, n++)
time = 328 ms.
```

The penalty attached to maintaining p^k decreases steeply as p increases (and updates become rare). But **t_INTMODs** remain at least 25% more efficient. (On the other hand, they do not allow denominators!)

2.3.8 Quadratic numbers (`t_QUAD`). This type is used to work in the quadratic order of *discriminant* d , where d is a nonsquare integer congruent to 0 or 1 (modulo 4). The command

```
w = quadgen(d, 'w)
```

assigns to w the “canonical” generator for the integer basis of the order of discriminant d , i.e. $w = \sqrt{d}/2$ if $d \equiv 0 \pmod{4}$, and $w = (1 + \sqrt{d})/2$ if $d \equiv 1 \pmod{4}$ and set its name to w . The name $'w$ is used for printing and we advise to store it in a variable of the same name. Beware, two `t_QUAD`s with different discriminants can be printed in the same way and not be equal; however, `gp` will refuse to add or multiply them for example, so use different names for different discriminants.

Since the order is $\mathbf{Z} + w\mathbf{Z}$, any other element can be input as $a = x + yw$ for some integers x and y . In fact, you may work in its fraction field $\mathbf{Q}(\sqrt{d})$ and use `t_FRAC` values for x and y .

The following member functions are defined:

`a.disc` retrieves the discriminant d ;

`a.mod`: returns the minimal polynomial T of w ;

`a.pol`: returns the `t_POL` $x + wy$. In particular `[x,y] = Vecrev(a.pol)` recovers x and y . The components x and y are also obtained via `real(a)` and `imag(a)` respectively.

2.3.9 Polmods (`t_POLMOD`). Exactly as for `intmods`, to enter $x \bmod y$ (where x and y are polynomials), type `Mod(x,y)`, not `x%y`. Note that when y is an irreducible polynomial in one variable, polmods whose modulus is y are simply algebraic numbers in the finite extension defined by the polynomial y . This allows us to work easily in number fields, finite extensions of the p -adic field $\mathbf{Q}_p$, or finite fields.

Note that this type is available for convenience, not for speed: each elementary operation involves a reduction modulo y . If p is a `t_POLMOD`, the following member functions are defined:

`p.pol`: return a representative of the polynomial class of minimal degree.

`p.mod`: return the modulus.

Important remark. Mathematically, the variables occurring in a polmod are not free variables. But internally, a congruence class in $R[t]/(y)$ is represented by its representative of lowest degree, which is a `t_POL` in $R[t]$, and computations occur with polynomials in the variable t . PARI will not recognize that `Mod(y, y^2 + 1)` is “the same” as `Mod(x, x^2 + 1)`, since x and y are different variables.

To avoid inconsistencies, polmods must use the same variable in internal operations (i.e. between polmods) and variables of lower priority for external operations, typically between a polynomial and a polmod. See Section 2.5.3 for a definition of “priority” and a discussion of (PARI’s idea of) multivariate polynomial arithmetic. For instance:

```
? Mod(x, x^2+ 1) + Mod(x, x^2 + 1)
%1 = Mod(2*x, x^2 + 1)    \\ 2i (or -2i), with i^2 = -1
? x + Mod(y, y^2 + 1)
%2 = x + Mod(y, y^2 + 1)  \\ in Q(i)[x]
? y + Mod(x, x^2 + 1)
%3 = Mod(x + y, x^2 + 1)  \\ in Q(y)[i]
```

The first two are straightforward, but the last one may not be what you want: y is treated here as a numerical parameter, not as a polynomial variable.

If the main variables are the same, it is allowed to mix `t_POL` and `t_POLMODs`. The result is the expected `t_POLMOD`. For instance

```
? x + Mod(x, x^2 + 1)
%1 = Mod(2*x, x^2 + 1)
```

2.3.10 Polynomials (`t_POL`). Type the polynomial in a natural way, not forgetting to put a “*” between a coefficient and a formal variable;

```
? 1 + 2*x + 3*x^2
%1 = 3*x^2 + 2*x + 1
```

This assumes that `x` is still a “free variable”.

```
? x = 1; 1 + 2*x + 3*x^2
%2 = 6
```

generates an integer, not a polynomial! It is good practice to never assign values to polynomial variables to avoid the above problem, but a foolproof construction is available using `'x` instead of `x`: `'x` is a constant evaluating to the free variable with name `x`, independently of the current value of `x`.

```
? x = 1; 1 + 2*'x + 3*'x^2
%3 = 1 + 2*x + 3*x^2
? x = 'x; 1 + 2*x + 3*x^2
%4 = 1 + 2*x + 3*x^2
```

You may also use the functions `Pol` or `Polrev`:

```
? Pol([1,2,3])          \\ Pol creates a polynomial in x by default
%1 = x^2 + 2*x + 3
? Polrev([1,2,3])
%2 = 3*x^2 + 2*x + 1
? Pol([1,2,3], 'y)      \\ we use 'y, safer than y
%3 = y^2 + 2*y + 3
```

The latter two are much more efficient constructors than an explicit summation (the latter is quadratic in the degree, the former linear):

```
? for (i=1, 10^4, Polrev( vector(100, i,i) ) )
time = 124ms
? for (i=1, 10^4, sum(i = 1, 100, (i+1) * 'x^i) )
time = 3,985ms
```

Polynomials are always printed as *univariate* polynomials over a commutative base ring, with monomials sorted by decreasing degree:

```
? (x+y+1)^2
%1 = x^2 + (2*y + 2)*x + (y^2 + 2*y + 1)
```

(Univariate polynomial in `x` whose coefficients are polynomials in `y`.) See Section 2.5 for valid variable names, and a discussion of multivariate polynomial rings. Polynomials over noncommutative rings are not supported.

2.3.11 Power series (t_SER). Typing $0(X^k)$, where k is an integer, yields an X -adic 0 of accuracy k , representing any power series in X whose valuation is $\geq k$. Of course, X can be replaced by any other variable name! To input a general nonzero power series, type in a polynomial or rational function (in X , say), and add $0(X^k)$ to it. The discussion in the t_POL section about variables remains valid; a constructor **Ser** replaces **Pol** and **Polrev**. Power series over noncommutative rings are not supported.

Caveat. Power series with inexact coefficients sometimes have a nonintuitive behavior: if k significant terms are requested, an inexact zero is counted as significant, even if it is the coefficient of lowest degree. This means that useful higher order terms may be disregarded.

If a series with a zero leading coefficient must be inverted, then as a desperation measure that coefficient is discarded, and a warning is issued:

```
? C = 0. + y + O(y^2);
? 1/C
*** _/_: Warning: normalizing a series with 0 leading term.
%2 = y^-1 + O(1)
```

The last output could be construed as a bug since it is a priori impossible to deduce such a result from the input (0. represents any sufficiently small real number). But it was thought more useful to try and go on with an approximate computation than to raise an early exception.

If the series precision is insufficient, errors may occur (mostly division by 0), which could have been avoided by a better global understanding of the computation:

[illegible]

2.3.12 Rational functions (t_RFRAC). As for fractions, all rational functions are automatically reduced to lowest terms. All that was said about fractions in Section 2.3.4 remains valid here.

2.3.13 Binary quadratic forms (τ_QFB). These are input using the function `Qfb`. For example, both `Qfb(1,2,3)` and `Qfb([1,2,3])` create the binary form $q = x^2 + 2xy + 3y^2$. It is imaginary since its discriminant $2^2 - 4 \times 3 = -8$ is negative. Although imaginary forms could be positive or negative definite, only positive definite forms are implemented.

The discriminant can be retrieved via `q.disc`. The individual components are obtained via either of

```
[a,b,c] = Vec(q);  
a = component(q,1);  
b = component(q,2);  
c = component(q,3);
```

See also the function `qfbprimeform` which creates a prime form of given discriminant.

2.3.14 Row and column vectors (t_VEC and t_COL). To enter a row vector, type the components separated by commas “,” and enclosed between brackets “[” and “]”, e.g. `[1,2,3]`. To enter a column vector, type the vector horizontally, and add a tilde “~” to transpose. `[ ]` yields the empty (row) vector. The function `Vec` can be used to transform any object into a vector (see Chapter 3). The construction `[i..j]`, where $i \leq j$ are two integers returns the vector $[i, i+1, \dots, j-1, j]$

```
? [1,2,3]
%1 = [1, 2, 3]
? [-2..3]
%2 = [-2, -1, 0, 1, 2, 3]
```

Let the variable v contain a (row or column) vector:

- `v[m]` refers to its m -th entry; you can assign any value to `v[m]`, i.e. write something like $v[m] = \text{expr}$.

- `v[i..j]`, where $i \leq j$, returns the vector slice containing elements $v[i], \dots, v[j]$; you can *not* assign a result to `v[i..j]`.

- `v[^i]` returns the vector whose i -th entry has been removed; you can *not* assign a result to `v[^i]`.

In the last two constructions `v[i..j]` and `v[^i]`, i and j are allowed to be negative integers, in which case, we start counting from the end of the vector: e.g., -1 is the index of the last element.

```
? v = [1,2,3,4];
? v[2..4]
%2 = [2, 3, 4]
? v[^3]
%3 = [1, 2, 4]
? v[^-1]
%3 = [1, 2, 3]
? v[-3..-1]
%4 = [2, 3, 4]
```

Remark. `vector` is the standard constructor for row vectors whose i -th entry is given by a simple function of i ; `vectorv` is similar for column vectors:

```
? vector(10, i, i^2+1)
%1 = [2, 5, 10, 17, 26, 37, 50, 65, 82, 101]
```

The functions `Vec` and `Col` convert objects to row and column vectors respectively (as well as `Vecrev` and `Colrev`, which revert the indexing):

```
? T = poltchebi(5)  \\ 5-th Chebyshev polynomial
%1 = 16*x^5 - 20*x^3 + 5*x
? Vec(T)
%2 = [16, 0, -20, 0, 5, 0]  \\ coefficients of T
? Vecrev(T)
%3 = [0, 5, 0, -20, 0, 16]  \\ ... in reverse order
```

Remark. For v a `t_VEC`, `t_COL`, `t_VECSMALL`, `t_LIST` or `t_MAT`, the alternative set-notations

```
[g(x) | x <- v, f(x)]
[x | x <- v, f(x)]
[g(x) | x <- v]
```

are available as shortcuts for

```
apply(g, select(f, Vec(v)))
select(f, Vec(v))
apply(g, Vec(v))
```

respectively, and may serve as `t_VEC` constructors:

```
? [ p | p <- primes(10), isprime(p+2) ]
%2 = [3, 5, 11, 17, 29]
```

returns the primes p (among the first 10 primes) such that $(p, p+2)$ is a twin pair;

```
? [ p^2 | p <- primes(10), p % 4 == 1 ]
%3 = [25, 169, 289, 841]
```

returns the squares of the primes congruent to 1 modulo 4, where p runs among the first 10 primes.

To iterate over more than one variable, you may separate clauses with `;` as in

```
? [ x+y | x <- [1..3]; y <- [1..2] ]
%4 = [2, 3, 3, 4, 4, 5]
? [ [x,y] | x <- [1..4], isprime(x); y <- [2..5], y % 3 == 1 ]
%5 = [[2, 4], [3, 4]]
```

2.3.15 Matrices (`t_MAT`). To enter a matrix, type the components row by row, the components being separated by commas “,”, the rows by semicolons “;”, and everything enclosed in brackets “[” and “]”, e.g. `[x,y; z,t; u,v]`. `[;]` yields an empty (0×0) matrix. The function `Mat` transforms any object into a matrix, and `matrix` creates matrices whose (i, j) -th entry is described by a function $f(i, j)$:

```
? Mat(1)
%1 =
[1]
? matrix(2,2, i,j, 2*i+j)
%2 =
[3 4]
[5 6]
```

Matrix multiplication assumes that the base ring containing the matrix entries is commutative.

Let the variable M contain a matrix, and let i, j, k, l denote four integers:

- $M[i, j]$ refers to its (i, j) -th entry; you can assign any result to $M[i, j]$.

- $M[i,]$ refers to its i -th row; you can assign a $\mathbf{t_VEC}$ of the right dimension to $M[i,]$.
- $M[, j]$ refers to its j -th column; you can assign a $\mathbf{t_COL}$ of the right dimension to $M[, j]$.

But $M[i]$ is meaningless and triggers an error. The “range” $i..j$ and “caret” $\wedge c$ notations are available as for vectors; you can not *assign* to any of these:

- $M[i..j, k..l]$, $i \leq j$, $k \leq l$, returns the submatrix built from the rows i to j and columns k to l of M .
- $M[i..j,]$ returns the submatrix built from the rows i to j of M .
- $M[, i..j]$ returns the submatrix built from the columns i to j of M .
- $M[i..j, \wedge k]$, $i \leq j$, returns the submatrix built from the rows i to j and column k removed.
- $M[\wedge k,]$ returns the submatrix with row k removed.
- $M[, \wedge k]$ returns the submatrix with column k removed.

Finally,

- $M[i..j, k]$ returns the $\mathbf{t_COL}$ built from the k -th column (entries i to j).
- $M[\wedge i, k]$ returns the $\mathbf{t_COL}$ built from the k -th column (entry i removed).
- $M[k, i..j]$ returns the $\mathbf{t_VEC}$ built from the k -th row (entries i to j).
- $M[k, \wedge i]$ returns the $\mathbf{t_VEC}$ built from the k -th row (entry i removed).

```
? M = [1,2,3;4,5,6;7,8,9];
? M[1..2, 2..3]
%2 =
[2 3]

[5 6]
? M[1..2, ]
%3 =
[1 2 3]

[4 5 6]
? M[, 2..3]
%4 =
[2 3]

[5 6]

[8 9]
```

All this is recursive, so if M is a matrix of matrices of $\dots$, an expression such as $M[1,1][,3][4] = 1$ is perfectly valid (and actually identical to $M[1,1][4,3] = 1$), assuming that all matrices along the way have compatible dimensions.

Technical note (design flaw). Matrices are internally represented as a vector of columns. All matrices with 0 columns are thus represented by the same object (internally, an empty vector), and there is no way to distinguish between them. Thus it is not possible to create or represent matrices with zero columns and an actual nonzero number of rows. The empty matrix `[]` is handled as though it had an arbitrary number of rows, exactly as many as needed for the current computation to make sense:

```
? [1,2,3; 4,5,6] * []
%1 = []
```

The empty matrix on the first line is understood as a 3×0 matrix, and the result as a 2×0 matrix. On the other hand, it is possible to create matrices with a given positive number of columns, each of which has zero rows, e.g. using `Mat` as above or using the `matrix` function.

Note that although the internal representation is essentially the same, a row vector of column vectors is *not* a matrix; for example, multiplication will not work in the same way. It is easy to go from one representation to the other using `Vec` / `Mat`, though:

```
? [1,2,3;4,5,6]
%1 =
[1 2 3]
[4 5 6]
? Vec(%)
%2 = [[1, 4]~, [2, 5]~, [3, 6]~]
? Mat(%)
%3 =
[1 2 3]
[4 5 6]
```

2.3.16 Lists (t_LIST). Lists can be input directly, as in `List([1,2,3,4])`; but in most cases, one creates an empty list, then appends elements using `listput`:

```
? L = List(); listput(~L,1); listput(~L,2);
? L
%2 = List([1, 2])
```

Note the `~L`: this means that the function is called with a *reference* to `L` and changes `L` in place. Elements can be accessed directly as with the vector types described above.

2.3.17 Strings (t_STR). To enter a string, enclose it between double quotes `"`, as in: `"this is a string"`. The function `Str` can be used to transform any object into a string.

2.3.18 Small vectors (`t_VECSMALL`). This type codes in an efficient way vectors containing only small integers, such as permutations. Most `gp` functions will refuse to operate on these objects, notable exceptions being `vecsort` and conversion functions such as `Vec`, but you can retrieve entries and assign to them as for ordinary vectors. You can also convert back and forth between `t_VECSMALL` and `t_VEC` objects using `Vec` and `Vecsmall`.

```
? v = Vecsmall([2, 4, 6])
%1 = Vecsmall([2, 4, 6])
? v[1]
%2 = 2
? v[1] = 3; v
%3 = Vecsmall([3, 2, 3])
? v[2..3]
%4 = Vecsmall([2, 3])
? v[^2]
%5 = Vecsmall([3, 3])
? Vec(v)
%6 = [3, 2, 3]
```

Allowed entries for a `t_VECSMALL` are signed integer x such that $|x| < 2^{31}$ on a 32-bit architecture, resp. $|x| < 2^{63}$ on a 64-bit architecture. Assigning a larger integer to a `t_VECSMALL` entry triggers an exception:

```
? v[1] = 2^63
*** at top-level: v[1]=2^63
***      ^-----
*** incorrect type in t_VECSMALL assignment (t_INT).
```

2.3.19 Functions (`t_CLOSURE`). We will explain this at length in Section 2.7. For the time being, suffice it to say that functions can be assigned to variables, as any other object, and the following equivalent basic forms are available to create new ones

```
f = (x,y) -> x^2 + y^2
f(x,y) = x^2 + y^2
```

2.3.20 Error contexts (`t_ERROR`). An object of this type is created whenever an error occurs: it contains some information about the error and the error context. Usually, an appropriate error is printed immediately, the computation is aborted, and GP enters the “break loop”:

```
? 1/0; 1 + 1
*** at top-level: 1/0;1+1
***      ^-----
*** _/_: division by a noninvertible object
*** Break loop: type 'break' to go back to the GP prompt
```

Here the computation is aborted as soon as we try to evaluate $1/0$, and $1 + 1$ is never executed. Exceptions can be trapped using `iferr`, however: we can evaluate some expression and either recover an ordinary result (no error occurred), or an exception (an error did occur).

```
? i = Mod(6,12); iferr(1/i, E, print(E)); 1 + 1
error("impossible inverse modulo: Mod(6, 12).")
%1 = 2
```

One can ignore the exception, print it as above, or extract non trivial information from the error context:

```
? i = Mod(6,12); iferr(1/i, E, print(component(E,1)));
Mod(6, 12)
```

We can also rethrow the exception: `error(E)`.

2.3.21 Infinity (`t_INFINITY`).

There are only two objects of this type `+oo` and `-oo`, representing $\pm\infty$. This type only contain only two elements `oo` and `-oo`. They are used in functions such as `intnum` or `polrootsreal`, to encode infinite real intervals. These objects can only be negated and compared to real numbers (`t_INT`, `t_REAL`, `t_FRAC`), but not included in any computation, i.e. `1+oo` is an error, not `oo` again.

2.4 GP operators.

Loosely speaking, an operator is a function, usually attached to basic arithmetic operations, whose name contains only nonalphanumeric characters. For instance `+` or `-`, but also `=` or `+=`, or even `[ ]` (the selection operator). As all functions, operators take arguments, and return a value; *assignment* operators also have side effects: besides returning a value, they change the value of some variable.

Each operator has a fixed and unchangeable priority, which means that, in a given expression, the operation with the highest priority is performed first. Unless mentioned otherwise, operators at the same priority level are left-associative (performed from left to right), unless they are assignments, in which case they are right-associative. Anything enclosed between parenthesis is considered a complete subexpression, and is resolved recursively, independently of the surrounding context. For instance,

<code>a + b + c</code>	<code>--></code>	<code>(a + b) + c</code>	<code>\\ left-associative</code>
<code>a = b = c</code>	<code>--></code>	<code>a = (b = c)</code>	<code>\\ right-associative</code>

Assuming that op_1 , op_2 , op_3 are binary operators with increasing priorities (think of `+`, `*`, `^`),

$$x \ op_1 \ y \ op_2 \ z \ op_2 \ x \ op_3 \ y$$

is equivalent to

$$x \ op_1 \ ((y \ op_2 \ z) \ op_2 \ (x \ op_3 \ y)).$$

GP contains many different operators, either unary (having only one argument) or binary, plus a few special selection operators. Unary operators are defined as either *prefix* or *postfix*, meaning that they respectively precede ($op \ x$) and follow ($x \ op$) their single argument. Some symbols are syntactically correct in both positions, like `!`, but then represent different operators: the `!` symbol represents the negation and factorial operators when in prefix and postfix position respectively. Binary operators all use the (infix) syntax $x \ op \ y$.

Most operators are standard (`+`, `%`, `=`), some are borrowed from the C language (`++`, `<<`), and a few are specific to GP (`\`, `#`). Beware that some GP operators differ slightly from their C counterparts. For instance, GP's postfix `++` returns the *new* value, like the prefix `++` of C, and the binary shifts `<<`, `>>` have a priority which is different from (higher than) that of their C counterparts. When in doubt, just surround everything by parentheses; besides, your code will be more legible.

Here is the list of available operators, ordered by decreasing priority, binary and left-associative unless mentioned otherwise. An expression is an *lvalue* if something can be assigned to it. (The name comes from left-value, to the left of a = operator; e.g. `x`, or `v[1]` are lvalues, but `x + 1` is not.)

- Priority 14

`:` as in `x:small`, is used to indicate to the GP2C compiler that the variable on the left-hand side always contains objects of the type specified on the right hand-side (here, a small integer) in order to produce more efficient or more readable C code. This is ignored by GP.

- Priority 13

`( )` is the function call operator. If f is a closure and $args$ is a comma-separated list of arguments (possibly empty), $f(args)$ evaluates f on those arguments.

- Priority 12

`++` and `--` (unary, postfix): if x is an lvalue, $x++$ assigns the value $x + 1$ to x , then returns the new value of x . This corresponds to the C statement `++x`: there is no prefix `++` operator in GP. $x--$ does the same with $x - 1$. These operators are not associative, i.e. `x++++` is invalid, since `x++` is not an lvalue.

- Priority 11

`.member` (unary, postfix): $x.member$ extracts *member* from structure x (see Section 2.8).

`[ ]` is the selection operator. $x[i]$ returns the i -th component of vector x ; $x[i,j]$, $x[,j]$ and $x[i,]$ respectively return the entry of coordinates (i,j) , the j -th column, and the i -th row of matrix x . If the assignment operator (`=`) immediately follows a sequence of selections, it assigns its right hand side to the selected component. E.g `x[1][1] = 0` is valid; but beware that `(x[1])[1] = 0` is not (because the parentheses force the complete evaluation of `x[1]`, and the result is not modifiable).

- Priority 10

`'` (unary, postfix): derivative with respect to the main variable. If f is a function (`t_CLOSURE`), f' is allowed and defines a new function, which will perform numerical derivation when evaluated at a scalar x ; this is defined as $(f(x + \epsilon) - f(x - \epsilon))/2\epsilon$ for a suitably small epsilon depending on current precision.

```
? (x^2 + y*x + y^2)'  \\ derive with respect to main variable x
%1 = 2*x + y
? SIN = cos'
%2 = cos'
? SIN(Pi/6)           \\ numerical derivation
%3 = -0.50000000000000000000000000000000
? cos'(Pi/6)          \\ works directly: no need for intermediate SIN
%4 = -0.50000000000000000000000000000000
```

`~` (unary, postfix): vector/matrix transpose.

`!` (unary, postfix): factorial. $x! = x(x - 1) \cdots 1$.

`#` (unary, postfix): primorial. For a non-negative integer n , $n\#$ is the product of all prime numbers less than or equal to n .

- Priority 9

`#` (unary, prefix): cardinality; `#x` returns `length(x)`.

! (unary, prefix): logical *not*. `!x` returns 1 if x is equal to 0 (specifically, if `gequal0(x)==1`), and 0 otherwise.

- Priority 8

^: powering. This operator is right associative: `2 ^ 3 ^ 4` is understood as `2 ^ (3 ^ 4)`.

- Priority 7

+, **-** (unary, prefix): **-** toggles the sign of its argument, **+** has no effect whatsoever.

- Priority 6

*****: multiplication.

/: exact division (`3/2` yields `3/2`, not `1.5`).

****, **%**: Euclidean quotient and remainder, i.e. if $x = qy + r$, then `x\y = q`, `x%y = r`. If x and y are scalars, then q is an integer and r satisfies $0 \leq r < |y|$; if x and y are polynomials, then q and r are polynomials such that $\deg r < \deg y$ and the leading terms of r and x have the same sign.

\|: rounded Euclidean quotient for integers (rounded towards $+\infty$ when the exact quotient would be a half-integer).

<<, **>>**: left and right binary shift. By definition, `x<<n = x * 2n` if $n > 0$, and `truncate(x2-n)` otherwise. Right shift is defined by `x>>n = x<<(-n)`.

- Priority 5

+, **-**: addition/subtraction.

- Priority 4

<, **>**, **<=**, **>=**: the usual comparison operators, returning 1 for **true** and 0 for **false**. For instance, `x<=1` returns 1 if $x \leq 1$ and 0 otherwise.

<>, **!=**: test for (exact) inequality.

==: test for (exact) equality.

===: test whether two objects are identical component-wise. This is stricter than **==**: for instance, the integer 0, a 0 polynomial or a vector with 0 entries, are all tested equal by **==**, but they are not identical.

- Priority 3

&&: logical *and*.

||: logical (inclusive) *or*. Any sequence of logical *or* and *and* operations is evaluated from left to right, and aborted as soon as the final truth value is known. Thus, for instance,

```
x == 0 || test(1/x)
```

will never produce an error since `test(1/x)` is not even evaluated when the first test is true (hence the final truth value is true). Similarly

```
type(p) == "t_INT" && isprime(p)
```

does not evaluate `isprime(p)` if `p` is not an integer.

- Priority 2

= (assignment, *lvalue = expr*). The result of `x = y` is the value of the expression y , which is also assigned to the variable `x`. This assignment operator is right-associative. This is *not* the equality test operator; a statement like `x = 1` is always true (i.e. nonzero), and sets `x` to 1; the

equality test would be `x == 1`. The right hand side of the assignment operator is evaluated before the left hand side.

It is crucial that the left hand-side be an *lvalue* there, it avoids ambiguities in expressions like `1 + x = 1`. The latter evaluates as `1 + (x = 1)`, not as `(1 + x) = 1`, even though the priority of `=` is lower than the priority of `+`: `1 + x` is not an lvalue.

If the expression cannot be parsed in a way where the left hand side is an lvalue, raise an error.

```
? x + 1 = 1
*** syntax error, unexpected '=', expecting $end or ';' : x+1=1
***                                     ^--
```

Assignment to all variables is a deep copy: after `x = y`, modifying a component of `y` will *not* change `x`. To globals it is a full copy to the heap. Space used by local objects in local variables is released when they go out of scope or when the value changes in local scope. Assigning a value to a vector or matrix entry allocates room for that entry only (on the heap).

`op=`, where `op` is any binary operator among `+`, `-`, `*`, `%`, `/`, `\`, `\ /`, `<<`, or `>>` (composed assignment *lvalue op= expr*). The expression `x op= y` assigns `(x op y)` to `x`, and returns the new value of `x`. The result is *not* an *lvalue*; thus

```
(x += 2) = 3
```

is invalid. These assignment operators are right-associative:

```
? x = 'x'; x += x *= 2
%1 = 3*x
```

- Priority 1

-> (function definition): `(vars)->expr` returns a function object, of type `t_CLOSURE`.

Remark. Use the `op=` operators as often as possible since they make complex assignments more legible. Compare

```
v[i+j-1] = v[i+j-1] + 1    -->    v[i+j-1]++
M[i,i+j] = M[i,i+j] * 2    -->    M[i,i+j] *= 2
```

Remark about efficiency. the operators `++` and `--` are usually a little more efficient than their expended counterpart:

```
? N = 10^7;
? i = 0; for(k = 1, N, i=i+1)
time = 949 ms.
? i = 0; for(k = 1, N, i++)
time = 933 ms.
```

On the other hand, this is not the case for the `op=` operators which may even be a little less efficient:

```
? i = 0; for(k = 1, N, i=i+10)
time = 949 ms.
? i = 0; for(k = 1, N, i+=10)
time = 1,064 ms.
```

2.5 Variables and symbolic expressions.

In this section we use *variable* in the standard mathematical sense, symbols representing algebraically independent elements used to build rings of polynomials and power series, and explain the all-important concept of *variable priority*. In the next Section 2.6, we shall no longer consider only free variables, but adopt the viewpoint of computer programming and assign values to these symbols: (bound) variables are names attached to values in a given scope.

2.5.1 Variable names. A valid name starts with a letter, followed by any number of keyword characters: `_` or alphanumeric characters (`[A-Za-z0-9]`). As a rule, the built-in function names are reserved and cannot be used; see the list with `\c`, including the constants `Pi`, `Euler`, `Catalan`, $I = \sqrt{-1}$ and `oo` = ∞ . Beware in particular of `gamma`, `omega`, `theta`, `sum` or `0`, none of which are free to use. (We shall see in Section 2.6 how this rule can be circumvented. It *is* possible to name a *lexical* variable `gamma`.)

GP names are case sensitive. For instance, the symbol `i` is perfectly safe to use, and will not be mistaken for $I = \sqrt{-1}$; analogously, `o` is not synonymous to `0`.

In GP you can use up to 16383 variable names (up to 65535 on 64-bit machines). If you ever need thousands of variables and this becomes a serious limitation, you should probably be using vectors instead: e.g. instead of variables `X1`, `X2`, `X3`, ..., you might equally well store their values in `X[1]`, `X[2]`, `X[3]`, ...

2.5.2 Variables and polynomials. The quote operator `'t` registers a new *free variable* with the interpreter, which will be written as `t`, and evaluates to a monomial of degree 1 in the said variable.

Caveat. For reasons of backward compatibility, there is no such thing as an “unbound” (uninitialized) variable in GP. If you use a valid variable name in an expression, `t` say, for the first time *before* assigning a value into it, it is interpreted as `'t` rather than raising an exception. One should not rely on this feature in serious programs, which would otherwise break if some unexpected assignment (e.g. `t = 1`) occurs: use `'t` directly or `t = 't` first, then `t`. A statement like `t = 't` in effect restores `t` as a free variable.

```
? t = 't; t^2 + 1
%1 = t^2 + 1
? t = 2; t^2 + 1
%2 = 5
? %1
%3 = t^2 + 1
? eval(%1)
%4 = 5
```

In the above, we initialize `t` to a monomial, then bind it to 2. Assigning a value to a polynomial variable does not affect previous expressions involving it; to take into account the new variable's value, one must force a new evaluation, using the function `eval` (see Section 3.9.6).

Caveat2. The use of an explicit quote operator avoids the following kind of problems:

```
? t = 't; p = t^2 + 1; subst(p, t, 2)
%1 = 5
? t = 2;
? subst(p, t, 3)    \\ t is no longer free: it evaluates to 2
***   at top-level: subst(p,t,3)
***               ^----
***   variable name expected.
? subst(p, 't, 3)    \\ OK
%3 = 10
```

2.5.3 Variable priorities, multivariate objects. A multivariate polynomial in PARI is just a polynomial (in one variable), whose coefficients are themselves polynomials, arbitrary but for the fact that they do not involve the main variable. (PARI currently has no sparse representation for polynomials, listing only nonzero monomials.) All computations are then done formally on the coefficients as if the polynomial was univariate.

This is not symmetrical. So if I enter 'x + 'y in a clean session, what happens? This is understood as

$$x^1 + (y^1 + 0 * y^0) * x^0 \in (\mathbf{Z}[y])[x]$$

but how do we know that x is “more important” than y ? Why not $y^1 + x * y^0$, which is the same mathematical entity after all?

The answer is that variables are ordered implicitly by the interpreter: when a new identifier (e.g x , or y as above) is input, the corresponding variable is registered as having a strictly lower priority than any variable in use at this point*. To see the ordering used by **gp** at any given time, type **variable()**.

Given such an ordering, multivariate polynomials are stored so that the variable with the highest priority is the main variable. And so on, recursively, until all variables are exhausted. A different storage pattern (which could only be obtained via **libpari** programming and low-level constructors) would produce an invalid object, and eventually a disaster.

In any case, if you are working with expressions involving several variables and want to have them ordered in a specific manner in the internal representation just described, the simplest is just to write down the variables one after the other under **gp** before starting any real computations. You may also define variables from your **gprc** to have a consistent ordering of common variable names in all your **gp** sessions, e.g read in a file **variables.gp** containing

```
'x; 'y; 'z; 't; 'a;
```

There is no way to change the priority of existing variables, but you may always create new ones with well-defined priorities using **varhigher** or **varlower**.

* This is not strictly true: the variables x and y are predefined, and satisfy $x > y$. Variables of higher priority than x can be created using **varhigher**.

Important note. PARI allows Euclidean division of multivariate polynomials, but assumes that the computation takes place in the fraction field of the coefficient ring (if it is not an integral domain, the result will a priori not make sense). This can become tricky. For instance assume x has highest priority, then y :

```
? x % y
%1 = 0
? y % x
%2 = y          \\ these two take place in  $\mathbf{Q}(y)[x]$ 
? x * Mod(1,y)
%3 = Mod(1, y)*x  \\ in  $(\mathbf{Q}(y)/y\mathbf{Q}(y))[x] \sim \mathbf{Q}[x]$ 
? Mod(x,y)
%4 = 0
```

In the last example, the division by y takes place in $\mathbf{Q}(y)[x]$, hence the `Mod` object is a coset in $(\mathbf{Q}(y)[x])/(y\mathbf{Q}(y)[x])$, which is the null ring since y is invertible! So be very wary of variable ordering when your computations involve implicit divisions and many variables. This also affects functions like `numerator/denominator` or `content`:

```
? denominator(x / y)
%1 = 1
? denominator(y / x)
%2 = x
? content(x / y)
%3 = 1/y
? content(y / x)
%4 = y
? content(2 / x)
%5 = 2
```

Can you see why? Hint: $x/y = (1/y) * x$ is in $\mathbf{Q}(y)[x]$ and `denominator` is taken with respect to $\mathbf{Q}(y)(x)$; $y/x = (y * x^0)/x$ is in $\mathbf{Q}(y)(x)$ so y is invertible in the coefficient ring. On the other hand, $2/x$ involves a single variable and the coefficient ring is simply $\mathbf{Z}$.

These problems arise because the variable ordering defines an *implicit* variable with respect to which division takes place. This is the price to pay to allow `%` and `/` operators on polynomials instead of requiring a more cumbersome `divrem(x, y, var)` (which also exists). Unfortunately, in some functions like `content` and `denominator`, there is no way to set explicitly a main variable like in `divrem` and remove the dependence on implicit orderings. This will hopefully be corrected in future versions.

2.5.4 Multivariate power series. Just like multivariate polynomials, power series are fundamentally single-variable objects. It is awkward to handle many variables at once, since PARI's implementation cannot handle multivariate error terms like $O(x^i y^j)$. (It can handle the polynomial $O(y^j) \times x^i$ which is a very different thing, see below.)

The basic assumption in our model is that if variable x has higher priority than y , then y does not depend on x : setting y to a function of x after some computations with bivariate power series does not make sense a priori. This is because implicit constants in expressions like $O(x^i)$ depend on y (whereas in $O(y^j)$ they can not depend on x). For instance

```
? 0(x) * y
```

```
%1 = 0(x)
? 0(y) * x
%2 = 0(y)*x
```

Here is a more involved example:

```
? A = 1/x^2 + 1 + 0(x); B = 1/x + 1 + 0(x^3);
? subst(z*A, z, B)
%2 = x^-3 + x^-2 + x^-1 + 1 + 0(x)
? B * A
%3 = x^-3 + x^-2 + x^-1 + 0(1)
? z * A
%4 = z*x^-2 + z + 0(x)
```

The discrepancy between %2 and %3 is surprising. Why does %2 contain a spurious constant term, which cannot be deduced from the input? Well, we ignored the rule that forbids to substitute an expression involving high-priority variables to a low-priority variable. The result %4 is correct according to our rules since the implicit constant in $O(x)$ may depend on z . It is obviously wrong if z is allowed to have negative valuation in x . Of course, the correct error term should be $O(xz)$, but this is not possible in PARI.

2.6 Variables and Scope.

This section is rather technical, and strives to explain potentially confusing concepts. Skip to the last subsection for practical advice, if the next discussion does not make sense to you. After learning about user functions, study the example in Section 2.7.7 then come back.

Definitions.

A *scope* is an enclosing context where names and values are attached. A user's function body, the body of a loop, an individual command line, all define scopes; the whole program defines the *global* scope. The argument of `eval` is evaluated in the enclosing scope.

Variables are bound to values within a given scope. This is traditionally implemented in two different ways:

- lexical (or static) scoping: the binding makes sense within a given block of program text. The value is private to the block and may not be accessed from outside. Where to find the value is determined at compile time.
- dynamic scoping: introducing a local variable, say x , pushes a new value on a stack attached to the name x (possibly empty at this point), which is popped out when the control flow leaves the scope. Evaluating x in any context, possibly outside of the given block, always yields the top value on this dynamic stack.

GP implements both lexical and dynamic scoping, using the keywords* `my` (lexical) and `local` (dynamic):

```
x = 0;
f() = x
g() = my(x = 1); f()
```

* The names are borrowed from the Perl scripting language.

```
h() = local(x = 1); f()
```

The function `g` returns 0 since the global `x` binding is unaffected by the introduction of a private variable of the same name in `g`. On the other hand, `h` returns 1; when it calls `f()`, the binding stack for the `x` identifier contains two items: the global binding to 0, and the binding to 1 introduced in `h`, which is still present on the stack since the control flow has not left `h` yet.

The rule mentioned in the previous section about built-in function names being reserved does *not* apply to lexically scoped variables. Those may temporarily shadow an existing function name:

```
my(gamma = 0) ; ...
```

Without the `my`, this would be invalid since `gamma` is the Γ function.

2.6.1 Scoping rules.

Named parameters in a function definition, as well as all loop indices**, have lexical scope within the function body and the loop body respectively.

```
p = 0;
forprime (p = 2, 11, print(p)); p  \\ prints 0 at the end
x = 0;
f(x) = x++;
f(1)  \\ returns 2, and leave global x unaffected (= 0)
```

If you exit the loop prematurely, e.g. using the `break` statement, you must save the loop index in another variable since its value prior the loop will be restored upon exit. For instance

```
for(i = 1, n,
  if (ok(i), break);
);
if (i > n, return(failure));
```

is incorrect, since the value of `i` tested by the `(i > n)` is quite unrelated to the loop index. One ugly workaround is

```
for(i = 1, n,
  if (ok(i), isave = i; break);
);
if (isave > n, return(failure));
```

But it is usually more natural to wrap the loop in a user function and use `return` instead of `break`:

```
try() =
{
  for(i = 1, n,
    if (ok(i), return (i));
  );
  0 \\ failure
}
```

A list of variables can be lexically or dynamically scoped (to the block between the declaration and the end of the innermost enclosing scope) using a `my` or `local` declaration:

** More generally, in all iterative constructs which use a variable name (`for`, `prod`, `sum`, `vector`, `matrix`, `plot`, etc.) the given variable is lexically scoped to the construct's body.

```

for (i = 1, 10,
    my(x, y, z, i2 = i^2); \\ temps needed within the loop body
    ...
)

```

Note how the declaration can include (optional) initial values, `i2 = i^2` in the above. Variables for which no explicit default value is given in the declaration are initialized to 0. It would be more natural to initialize them to free variables, but this would break backward compatibility. To obtain this behavior, you may explicitly use the quoting operator:

```
my(x = 'x, y = 'y, z = 'z);
```

A more complicated example:

```

for (i = 1, 3,
    print("main loop");
    my(x = i);          \\ local to the outermost loop
    for (j = 1, 3,
        my (y = x^2);   \\ local to the innermost loop
        print (y + y^2);
        x++;
    )
)

```

When we leave the loops, the values of `x`, `y`, `i`, `j` are the same as before they were started.

Note that `eval` is evaluated in the given scope, and can access values of lexical variables:

```

? x = 1;
? my(x = 0); eval("x")
%2 = 0    \\ we see the local x scoped to this command line, not the global one

```

Variables dynamically scoped using `local` should more appropriately be called *temporary values* since they are in fact local to the function declaring them *and* any subroutine called from within. In practice, you almost certainly want true private variables, hence should use almost exclusively `my`.

We strongly recommended to explicitly scope (lexically) all variables to the smallest possible block. Should you forget this, in expressions involving such “rogue” variables, the value used will be the one which happens to be on top of the value stack at the time of the call; which depends on the whole calling context in a nontrivial way. This is in general *not* what you want.

2.7 User defined functions.

User-defined functions are ordinary GP objects, bound to variables just like any other object. Those variables are subject to scoping rules as any other: while you can define all your functions in global scope, it is usually possible and cleaner to lexically scope your private helper functions to the block of text where they will be needed.

Whenever gp meets a construction of the form `expr(argument list)` and the expression `expr` evaluates to a function (an object of type `t_CLOSURE`), the function is called with the proper arguments. For instance, constructions like `funcs[i](x)` are perfectly valid, assuming `funcs` is an array of functions.

As regards argument passing conventions, GP functions support both

- call by value: the function operates on a copy of a variable, changes made to the argument in the function do not affect the original variable;
- and call by reference: the function receives a reference to the variable and original data is affected.

2.7.1 Defining a function.

A user function is defined as follows:

(list of formal variables) -> seq.

The list of formal variables is a comma-separated list of *distinct* variable names and allowed to be empty. If there is a single formal variable, the parentheses are optional. This list corresponds to the list of parameters you will supply to your function when calling it. By default, GP functions use call by value to pass arguments; a variable name may be prefixed by a tilde `~` to use instead a call by reference.

In most cases you want to assign a function to a variable immediately, as in

```
R = (x,y) -> sqrt( x^2+y^2 );  
sq = x -> x^2;  \\ or equivalently (x) -> x^2
```

but it is quite possible to define short-lived anonymous functions. The trailing semicolon is not part of the definition, but as usual prevents gp from printing the result of the evaluation, i.e. the function object. The construction

f(list of formal variables) = seq

is available as an alias for

f = (list of formal variables) -> seq

Using that syntax, it is not possible to define anonymous functions (obviously), and the above two examples become:

```
R(x,y) = sqrt( x^2+y^2 );  
sq(x) = x^2;
```

The semicolon serves the same purpose as above: preventing the printing of the resulting function object; compare

```
? sq(x) = x^2;  \\ no output  
? sq(x) = x^2  \\ print the result: a function object
```

```
%2 = (x)->x^2
```

Of course, the sequence *seq* can be arbitrarily complicated, in which case it will look better written on consecutive lines, with properly scoped variables:

```
{
f(x0, x1, ...) =
    my(t0, t1, ...); \\ variables lexically scoped to the function body
    ...
}
```

Note that the following variant would also work:

```
f(x0, x1, ...) =
{
    my(t0, t1, ...); \\ variables lexically scoped to the function body
    ...
}
```

(the first newline is disregarded due to the preceding = sign, and the others because of the enclosing braces). The *my* statements can actually occur anywhere within the function body, scoping the variables to more restricted blocks than the whole function body.

Formal parameters are lexically scoped to the function body. It is not allowed to use the same variable name for different parameters of your function:

```
? f(x,x) = 1
***      variable declared twice: f(x,x)=1
***                               ^----
```

Finishing touch. You can add a specific help message for your function using `addhelp`, but the online help system already handles it. By default `?name` will print the definition of the function *name*: the list of arguments, as well as their default values, the text of *seq* as you input it. Just as `\c` prints the list of all built-in commands, `\u` outputs the list of all user-defined functions.

2.7.2 Call by value, call by reference.

By default, arguments are passed by value, not as variables: modifying a function's argument in the function body is allowed, but does not modify its value in the calling scope. In fact, a *copy* of the actual parameter is assigned to the formal parameter when the function is called. (This is not literally true: a form of copy-on-write is implemented so an object is not duplicated unless modified in the function.) If an argument is prefixed by a tilde `~` in the function declaration *and* the call, it is passed by reference. (If either the declaration or the call is missing a tilde, we revert to a call by value.)

```
? x = [1];
? f(v)  = v[1]++;
? F(~v) = v[1]++;

? f(x)
%4 = 2
? x      \\ unchanged
%5 = [1]
? F(~x)
```

```

%6 = 2
? x      \\ incremented
%7 = [2]
? F(x)    \\ forgot the ~: call by value
%8 = 3
? x      \\ => contents of x did not change
%9 = [2]
? f(~x)    \\ adding a ~ in call, missing in declaration
%10 = 3
? x      \\ => call by value
%11 = [2]

```

Caveat. In GP, a call by reference means that the function accesses the value and may change the original variable content, but only if it is a container type (a vector, list or matrix), as shown above with the vector x . It will not alter its value in other cases !

```

? v = 1    \\ not a container
? F(~v) = v++
? F(~v)
%3 = 2
? v \\ components of v could be altered, not v itself
%4 = 1

```

2.7.3 Functions taking an unlimited number of arguments.

A function taking an unlimited number of arguments is called *variadic*. To create such a function, use the syntax

```
(list of formal variables, var[...]) -> seq
```

The parameter *var* is replaced by a vector containing all the remaining arguments. The name may not be prefixed by a tilde to (absurdly) indicate a call by reference.

```

? f(c[...]) = sum(i=1,#c,c[i]);
? f(1,2,3)
%1 = 6
? sep(s,v[...]) = for(i=1,#v-1,print1(v[i],s)); if (#v, print(v[#v]));
? sep(":", 1, 2, 3)
1:2:3

```

2.7.4 Backward compatibility (lexical scope). Lexically scoped variables were introduced in version 2.4.2. Before that, the formal parameters were dynamically scoped. If your script depends on this behavior, you may use the following trick: replace the initial $f(x) =$ by

```
f(x_orig) = local(x = x_orig)
```


<code>f()</code>	<code>f(0,0),</code>
<code>f(,3)</code>	<code>f(0,3).</code> (“Empty argument” trick)

This *implicit* default value of 0, is actually deprecated and setting

```
default(strictargs, 1)
```

allows to disable it (see Section 3.4.43).

The recommended practice is to *explicitly* set a default value: in the function definition, you can append `=expr` to a formal parameter, to give that variable a default value. The expression gets evaluated the moment the function is called, and may involve the preceding function parameters: a default value for x_i may involve x_j for $j < i$. For instance, after

```
f(x = 1, y = 2, z = y+1) = ....
```

typing in `f(3,4)` would give you `f(3,4,5)`. In the rare case when you want to set some far away argument, and leave the defaults in between as they stand, use the “empty argument” trick: `f(6,,1)` would yield `f(6,2,1)`. Of course, `f()` by itself yields `f(1,2,3)` as was to be expected.

In short, the argument list is filled with user supplied values, in order. A comma or closing parenthesis, where a value should have been, signals we must use a default value. When no input arguments are left, the defaults are used instead to fill in remaining formal parameters. A final example:

```
f(x, y=2, z=3) = print(x, ":", y, ":", z);
```

defines a function which prints its arguments (at most three of them), separated by colons.

```
? f(6,7)
6:7:3
? f(,5)
0:5:3
? f()
0:2:3
```

If `strictargs` is set (recommended), x is now a mandatory argument, and the above becomes:

```
? default(strictargs,1)
? f(6,7)
6:7:3
? f(,5)
*** at top-level: f(,5)
***      ^-----
*** in function f: x,y=2,z=3
***      ^-----
*** missing mandatory argument 'x' in user function.
```

Example. We conclude with an amusing example, intended to illustrate both user-defined functions and the power of the `sumalt` function. Although the Riemann zeta-function is included (as `zeta`) among the standard functions, let us assume that we want to check other implementations. Since we are highly interested in the critical strip, we use the classical formula

$$(2^{1-s} - 1)\zeta(s) = \sum_{n \geq 1} (-1)^n n^{-s}, \quad \Re s > 0.$$

The implementation is obvious:

```
ZETA(s) = sumalt(n=1, (-1)^n*n^(-s)) / (2^(1-s) - 1)
```

Note that `n` is automatically lexically scoped to the `sumalt` “loop”, so that it is unnecessary to add a `my(n)` declaration to the function body. Surprisingly, this gives very good accuracy in a larger region than expected:

```
? check = z -> ZETA(z) / zeta(z);
? check(2)
%1 = 1.000000000000000000000000000000
? check(200)
%2 = 1.000000000000000000000000000000
? check(0)
%3 = 0.999999999999999999999999999994
? check(-5)
%4 = 1.000000000000000000000000000007549266557
? check(-11)
%5 = 0.9999752641047824902660847745
? check(1/2+14.134*I) \\ very close to a nontrivial zero
%6 = 1.0000000000000000000000000000003747432 + 7.62329066 E-21*I
? check(-1+10*I)
%7 = 1.00000000000000000000000000000002511 + 2.989950968 E-24*I
```

Now wait a minute; not only are we summing a series which is certainly no longer alternating (it has complex coefficients), but we are also way outside of the region of convergence, and still get decent results! No programming mistake this time: `sumalt` is a “magic” function*, providing very good convergence acceleration; in effect, we are computing the analytic continuation of our original function. To convince ourselves that `sumalt` is a nontrivial implementation, let us try a simpler example:

```
? sum(n=1, 10^7, (-1)^n/n, 0.) / (-log(2)) \\ approximates the well-known formula
time = 7,417 ms.
%1 = 0.9999999278652515622893405457
? sumalt(n=1, (-1)^n/n) / (-log(2)) \\ accurate and fast
time = 0 ms.
%2 = 1.000000000000000000000000000000
```

No, we are not using a powerful simplification tool here, only numerical computations. Remember, PARI is not a generalist computer algebra system!

* `sumalt` is heuristic, but its use can be rigorously justified for a given function, in particular our $\zeta(s)$ formula. Indeed, Peter Borwein (*An efficient algorithm for the Riemann zeta function*, CMS Conf. Proc. **27** (2000), pp. 29–34) proved that the formula used in `sumalt` with n terms computes $(1 - 2^{1-s})\zeta(s)$ with a relative error of the order of $(3 + \sqrt{8})^{-n}|\Gamma(s)|^{-1}$.

2.7.7 Beware scopes. Be extra careful with the scopes of variables. What is wrong with the following definition?

```
FirstPrimeDiv(x) =
{ my(p);
  forprime(p=2, x, if (x%p == 0, break));
  p
}
? FirstPrimeDiv(10)
%1 = 0
```

Hint. The function body is equivalent to

```
{ my(newp = 0);
  forprime(p=2, x, if (x%p == 0, break));
  newp
}
```

Detailed explanation. The index p in the `forprime` loop is lexically scoped to the loop and is not visible to the outside world. Hence, it will not survive the `break` statement. More precisely, at this point the loop index is restored to its preceding value. The initial `my(p)`, although well-meant, adds to the confusion: it indeed scopes p to the function body, with initial value 0, but the `forprime` loop introduces *another* variable, unfortunately also called p , scoped to the loop body, which shadows the one we wanted. So we always return 0, since the value of the p scoped to the function body never changes and is initially 0.

To sum up, the routine returns the p declared local to it, not the one which was local to `forprime` and ran through consecutive prime numbers. Here is a corrected version:

```
? FirstPrimeDiv(x) = forprime(p=2, x, if (x%p == 0, return(p)))
```

2.7.8 Recursive functions. Recursive functions can easily be written as long as one pays proper attention to variable scope. Here is an example, used to retrieve the coefficient array of a multivariate polynomial (a nontrivial task due to PARI's unsophisticated representation for those objects):

```
coeffs(P, nvar) =
{ my (d = poldegree(P));
  if (d <= 0,
    P = simplify(P); for (i=1, nvar, P = [P]);
    return (P));
  vector(d + 1, i, coeffs(polcoef(P, i-1), nvar-1));
}
```

If P is a polynomial in k variables, show that after the assignment $\mathbf{v} = \text{coeffs}(P, \mathbf{k})$, the coefficient of $x_1^{n_1} \dots x_k^{n_k}$ in P is given by $\mathbf{v}[\mathbf{n}_1+1][\dots][\mathbf{n}_k+1]$, provided a monomial $x_1^{N_1} \dots x_k^{N_k}$ with $n \leq N$ (lexicographically) exists with a non-zero coefficient.

When the operating system allows querying the maximum size of the process stack, we automatically limit the recursion depth:

```
? dive(n) = dive(n+1)
```

```

? dive(0);
***   [...] at: dive(n+1)
***           ^-----
***   in function dive: dive(n+1)
***           ^-----
\\ (last 2 lines repeated 19 times)
***   deep recursion.

```

All Unix variants support this mechanism and the recursion limit may be different from one machine to the next; other systems may crash on deep recursion. There is no way to increase the limit from within `gp`. On a Unix system, you may increase it before launching `gp` with `ulimit` or `limit`, depending on your shell, and raise the process available stack space (increase `stacksize`).

2.7.9 Function which take functions as parameters. This is done as follows:

```

? calc(f, x) = f(x)
? calc(sin, Pi)
%2 = -5.04870979 E-29
? g(x) = x^2;
? calc(g, 3)
%4 = 9

```

If we do not need `g` elsewhere, we should use an anonymous function here, `calc(x->x^2, 3)`. Here is a variation:

```

? funs = [cos, sin, tan, x->x^3+1]; \\ an array of functions
? call(i, x) = funs[i](x)

```

evaluates the appropriate function on argument `x`, provided $1 \leq i \leq 4$. Finally, a more useful example:

```

APPLY(f, v) = vector(#v, i, f(v[i]))

```

applies the function `f` to every element in the vector `v`. (The built-in function `apply` is more powerful since it also applies to lists and matrices.)

2.7.10 Defining functions within a function. Defining a single function is easy:

```

init(x) = (add = y -> x+y);

```

Basically, we are defining a global variable `add` whose value is the function `y->x+y`. The parentheses were added for clarity and are not mandatory.

```

? init(5);
? add(2)
%2 = 7

```

A more refined approach is to avoid global variables and *return* the function:

```

init(x) = y -> x+y
add = init(5)

```

Then `add(2)` still returns 7, as expected! Of course, if `add` is in global scope, there is no gain, but we can lexically scope it to the place where it is useful:

```

my ( add = init(5) );

```

How about multiple functions then? We can use the last idea and return a vector of functions, but if we insist on global variables? The first idea

```
init(x) = add(y) = x+y; mul(y) = x*y;
```

does not work since in the construction `f() = seq`, the function body contains everything until the end of the expression. Hence executing `init` defines the wrong function `add` (itself defining a function `mul`). The way out is to use parentheses for grouping, so that enclosed subexpressions will be evaluated independently:

```
? init(x) = ( add(y) = x+y ); ( mul(y) = x*y );
? init(5);
? add(2)
%3 = 7
? mul(3)
%4 = 15
```

This defines two global functions which have access to the lexical variables private to `init`! The following would work in exactly the same way:

```
? init5() = my(x = 5); ( add(y) = x+y ); ( mul(y) = x*y );
```

2.7.11 Closures as Objects. Contrary to what you might think after the preceding examples, GP's closures may not be used to simulate true “objects”, with private and public parts and methods to access and manipulate them. In fact, closures indeed incorporate an existing context (they may access lexical variables that existed at the time of their definition), but then may not change it. More precisely, they access a copy, which they are welcome to change, but a further function call still accesses the original context, as it existed at the time the function was defined:

```
init() =
{ my(count = 0);
  (inc()=count++);
  (dec()=count--);
}
? init();
? inc()
%1 = 1
? inc()
%2 = 1
? dec()
%3 = -1
? dec()
%4 = -1
```

2.8 Member functions.

Member functions use the ‘dot’ notation to retrieve information from complicated structures. The built-in structures are *bid*, *ell*, *galois*, *ff*, *nf*, *bnf*, *bnr* and *prid*, which will be described at length in Chapter 3. The syntax `structure.member` is taken to mean: retrieve `member` from `structure`, e.g. `E.j` returns the *j*-invariant of the elliptic curve `E`, or outputs an error message if `E` is not a proper *ell* structure. To define your own member functions, use the syntax

```
var.member = seq,
```

where the formal variable *var* is scoped to the function body *seq*. This is of course reminiscent of a user function with a single formal variable *var*. For instance, the current implementation of the *ell* type is a vector, the *j*-invariant being the thirteenth component. It could be implemented as

```
x.j =  
{  
  if (type(x) != "t_VEC" || #x < 14, error("not an elliptic curve: " x));  
  x[13]  
}
```

As for user functions, you can redefine your member functions simply by typing new definitions. On the other hand, as a safety measure, you cannot redefine the built-in member functions, so attempting to redefine `x.j` as above would in fact produce an error; you would have to call it e.g. `x.myj` in order for `gp` to accept it.

Member functions use call by reference to pass arguments, your function may modify in place the contents of a variable (of container type).

Rationale. In most cases, member functions are simple accessors of the form

```
x.a = x[1];  
x.b = x[2];  
x.c = x[3];
```

where `x` is a vector containing relevant data. There are at least three alternative approaches to the above member functions: 1) hardcode `x[1]`, etc. in the program text, 2) define constant global variables `AINDEX = 1`, `BINDEX = 2` and hardcode `x[AINDEX]`, 3) user functions `a(x) = x[1]` and so on.

Even if 2) improves on 1), these solutions are neither elegant nor flexible, and they scale badly. 3) is a genuine possibility, but the main advantage of member functions is that their namespace is independent from the variables (and functions) namespace, hence we can use very short identifiers without risk. The *j*-invariant is a good example: it would clearly not be a good idea to define `j(E) = E[13]`, because clashes with loop indices are likely.

Beware that there is no guarantee that a built-in member function *is* a simple accessor and it could involve a computation. Thus you should not use them on a constant object in tight loops: store them in a variable before the loop.

Note. Typing `\um` will output all user-defined member functions.

Member function names. A valid name starts with a letter followed by any number of keyword characters: `_` or alphanumeric characters (`[A-Za-z0-9]`). The built-in member function names are reserved and cannot be used (see the list with `?.`). Finally, names starting with `e` or `E` followed by a digit are forbidden, due to a clash with the floating point exponent notation: we understand `1.e2` as `100.000...`, not as extracting member `e2` of object `1`.

2.9 Strings and Keywords.

2.9.1 Strings.

GP variables can hold values of type character string (internal type `t_STR`). This section describes how they are actually used, as well as some convenient tricks (automatic concatenation and expansion, keywords) valid in string context.

As explained above, the general way to input a string is to enclose characters between quotes ". This is the only input construct where whitespace characters are significant: the string will contain the exact number of spaces you typed in. Besides, you can “escape” characters by putting a `\` just before them; the translation is as follows

```
\e: <Escape>
\n: <Newline>
\t: <Tab>
```

For any other character x , `\x` is expanded to x . In particular, the only way to put a " into a string is to escape it. Thus, for instance, `"\"a\""` would produce the string whose content is “a”. This is definitely *not* the same thing as typing `"a"`, whose content is merely the one-letter string a.

You can concatenate two strings using the `concat` function. If either argument is a string, the other is automatically converted to a string if necessary (it will be evaluated first).

```
? concat("ex", 1+1)
%1 = "ex2"
? a = 2; b = "ex"; concat(b, a)
%2 = "ex2"
? concat(a, b)
%3 = "2ex"
```

Some functions expect strings for some of their arguments: `print` would be an obvious example, `Str` is a less obvious but useful one (see the end of this section for a complete list). While typing in such an argument, you will be said to be in *string context*. The rest of this section is devoted to special syntactical tricks which can be used with such arguments (and only here; you will get an error message if you try these outside of string context):

- Writing two strings alongside one another will just concatenate them, producing a longer string. Thus it is equivalent to type in `"a " "b"` or `"a b"`. A little tricky point in the first expression: the first whitespace is enclosed between quotes, and so is part of a string; while the second (before the `"b"`) is completely optional and `gp` actually suppresses it, as it would with any number of whitespace characters at this point (i.e. outside of any string).

- If you insert any expression when a string is expected, it gets “expanded”: it is evaluated as a standard GP expression, and the final result (as would have been printed if you had typed it by itself) is then converted to a string, as if you had typed it directly. For instance `"a" 1+1 "b"` is equivalent to `"a2b"`: three strings get created, the middle one being the expansion of `1+1`, and these are then concatenated according to the rule described above. Another tricky point here: assume you did not assign a value to `aaa` in a GP expression before. Then typing `aaa` by itself in a string context will actually produce the correct output (i.e. the string whose content is `aaa`), but in a fortuitous way. This `aaa` gets expanded to the monomial of degree one in the variable `aaa`, which is of course printed as `aaa`, and thus will expand to the three letters you were expecting.

Warning. Expression involving strings are not handled in a special way; even in string context, the largest possible expression is evaluated, hence `print("a"[1])` is incorrect since "a" is not an object whose first component can be extracted. On the other hand `print("a", [1])` is correct (two distinct argument, each converted to a string), and so is `print("a" 1)` (since "a"1 is not a valid expression, only "a" gets expanded, then 1, and the result is concatenated as explained above).

2.9.2 Keywords. Since there are cases where expansion is not desirable, we now distinguish between “Keywords” and “Strings”. String is what has been described so far. Keywords are special relatives of Strings which are automatically assumed to be quoted, whether you actually type in the quotes or not. Thus expansion is never performed on them. They get concatenated, though. The analyzer supplies automatically the quotes you have “forgotten” and treats Keywords just as normal strings otherwise. For instance, if you type `"a"b+b` in Keyword context, you will get the string whose contents are `ab+b`. In String context, on the other hand, you would get `a2*b`.

All GP functions have prototypes (described in Chapter 3 below) which specify the types of arguments they expect: either generic PARI objects (GEN), or strings, or keywords, or unevaluated expression sequences. In the keyword case, only a very small set of words will actually be meaningful (the `default` function is a prominent example).

Reference. The arguments of the following functions are processed in string context:

```
Str
addhelp (second argument)
default (second argument)
error
extern
plotstring (second argument)
plotterm (first argument)
read and readvec
system
all the printxxx functions
all the writexxx functions
```

The arguments of the following functions are processed as keywords:

```
alias
default (first argument)
install (all arguments but the last)
trap (first argument)
whatnow
```

2.9.3 Useful example. The function `Str` converts its arguments into strings and concatenate them. Coupled with `eval`, it is very powerful. The following example creates generic matrices:

```
? genmat(u,v,s="x") = matrix(u,v,i,j, eval( Str(s,i,j) ))
? genmat(2,3) + genmat(2,3,"m")
%1 =
[x11 + m11 x12 + m12 x13 + m13]
[x21 + m21 x22 + m22 x23 + m23]
```

2.10 Errors and error recovery.

2.10.1 Errors. Your input program is first compiled to a more efficient bytecode; then the latter is evaluated, calling appropriate functions from the PARI library. Accordingly, there are two kind of errors: syntax errors produced by the compiler, and runtime errors produced by the PARI library either by the evaluator itself, or in a mathematical function. Both kinds are fatal to your computation: **gp** will report the error and perform some cleanup (restore variables modified while evaluating the erroneous command, close open files, reclaim unused memory, etc.).

At this point, the default is to return to the usual prompt, but if the **recover** option (Section 3.4.38) is off then **gp** exits immediately. This can be useful for batch-mode operation to make untrapped errors fatal.

When reporting a *syntax error*, **gp** gives meaningful context by copying (part of) the expression it was trying to compile, indicating where the error occurred with a caret `^`-, as in

```
? factor()
***   too few arguments: factor()
***                               ^-
? 1+
***   syntax error, unexpected $end: 1+
***                               ^-
```

possibly enlarged to a full arrow given enough trailing context

```
? if (isprime(1+, do_something()))
***   syntax error, unexpected ',': if(isprime(1+,do_something()))
***                               ^-----
```

These error messages may be mysterious, because **gp** cannot guess what you were trying to do, and the error may occur once **gp** has been sidetracked. The first error is straightforward: **factor** has one mandatory argument, which is missing.

The other two are simple typos involving an ill-formed addition `1 +` missing its second operand. The error messages differ because the parsing context is slightly different: in the first case we reach the end of input (**\$end**) while still expecting a token, and in the second one, we received an unexpected token (the comma).

Here is a more complicated one:

```
? factor(x
***   syntax error, unexpected $end, expecting )-> or ', ' or ')': factor(x
***                               ^-
```

The error is a missing parenthesis, but from **gp**'s point of view, you might as well have intended to give further arguments to **factor** (this is possible and useful, see the description of the function). In fact **gp** expected either a closing parenthesis, or a second argument separated from the first by a comma. And this is essentially what the error message says: we reached the end of the input (**\$end**) while expecting a `)` or a `,`.

Actually, a third possibility is mentioned in the error message `)->`, which could never be valid in the above context, but a subexpression like `(x)->sin(x)`, defining an inline closure would be valid, and the parser is not clever enough to rule that out, so we get the same message as in

```
? (x
```

```
*** syntax error, unexpected $end, expecting )-> or ',' or ')': (x
***                                     ^_
```

where all three proposed continuations would be valid.

Runtime errors from the evaluator are nicer because they answer a correctly worded query, otherwise the bytecode compiler would have protested first; here is a slightly pathological case:

```
? if (siN(x) < eps, do_something())
*** at top-level: if(siN(x)<eps,do_someth
***             ^-----
*** not a function in function call
```

(no arrow!) The code is syntactically correct and compiled correctly, even though the `siN` function, a typo for `sin`, was not defined at this point. When trying to evaluate the bytecode, however, it turned out that `siN` is still undefined so we cannot evaluate the function call `siN(x)`.

Library runtime errors are even nicer because they have more mathematical content, which is easier to grasp than a parser's logic:

```
? 1/Mod(2,4)
*** at top-level: 1/Mod(2,4)
***             ^-----
*** _/_: impossible inverse in Fp_inv: Mod(2, 4).
```

telling us that a runtime error occurred while evaluating the binary `/` operator (the `_` surrounding the operator are placeholders), more precisely the `Fp_inv` library function was fed the argument `Mod(2,4)` and could not invert it. More context is provided if the error occurs deep in the call chain:

```
? f(x) = 1/x;
? g(N) = for(i = -N, N, f(i + 0(5)));
? g(10)
*** at top-level: g(10)
***             ^-----
*** in function g: for(i=-N,N,f(i))
***             ^-----
*** in function f: 1/x
***             ^--
*** _/_: impossible inverse in ginv: 0(5).
```

In this example, the debugger reports (at least) 3 enclosed frames: last (innermost) is the body of user function *f*, the body of *g*, and the top-level (global scope). In fact, the `for` loop in *g*'s body defines an extra frame, since there exist variables scoped to the loop body.

2.10.2 Error recovery.

It is annoying to wait for a program to finish and find out the hard way that there was a mistake in it (like the division by 0 above), sending you back to the prompt. First you may lose some valuable intermediate data. Also, correcting the error may not be obvious; you might have to change your program, adding a number of extra statements and tests to narrow down the problem.

A different situation, still related to error recovery, is when you actually foresee that some error may occur, are unable to prevent it, but quite capable of recovering from it, given the chance. Examples include lazy factorization, where you knowingly use a pseudo prime N as if it were prime; you may then encounter an “impossible” situation, but this would usually exhibit a factor of N , enabling you to refine the factorization and go on. Or you might run an expensive computation at low precision to guess the size of the output, hence the right precision to use. You can then encounter errors like “precision loss in truncation”, e.g when trying to convert 1E1000, known to 38 digits of accuracy, to an integer; or “division by 0”, e.g inverting 0E1000 when all accuracy has been lost, and no significant digit remains. It would be enough to restart part of the computation at a slightly higher precision.

We now describe *error trapping*, a useful mechanism which alleviates much of the pain in the first situation (the break loop debugger), and provides satisfactory ways out of the second one (the `iferr` exception handler).

2.10.3 Break loop.

A *break loop* is a special debugging mode that you enter whenever a user interrupt (**Control-C**) or runtime error occurs, freezing the `gp` state, and preventing cleanup until you get out of the loop. By runtime error, we mean an error from the evaluator, the library or a user error (from `error`), *not* syntax errors. When a break loop starts, a prompt is issued (**break>**). You can type in a `gp` command, which is evaluated when you hit the **<Return>** key, and the result is printed as during the main `gp` loop, except that no history of results is kept. Then the break loop prompt reappears and you can type further commands as long as you do not exit the loop. If you are using `readline`, the history of commands is kept, and line editing is available as usual. If you type in a command that results in an error, you are sent back to the break loop prompt: errors do *not* terminate the loop.

To get out of a break loop, you can use `next`, `break`, `return`, or type **C-d** (EOF), any of which will let `gp` perform its usual cleanup, and send you back to the `gp` prompt. Note that **C-d** is slightly dangerous, since typing it *twice* will not only send you back to the `gp` prompt, but to your shell prompt! (Since **C-d** at the `gp` prompt exits the `gp` session.)

If the break loop was started by a user interrupt **Control-C**, and not by an error, inputting an empty line, i.e hitting the **<Return>** key at the **break>** prompt, resumes the temporarily interrupted computation. A single empty line has no effect in case of a fatal error, to avoid getting get out of the loop prematurely, thereby losing valuable debugging data. Any of `next`, `break`, `return`, or **C-d** will abort the computation and send you back to the `gp` prompt as above.

Break loops are useful as a debugging tool. You may inspect the values of `gp` variables to understand why an error occurred, or change `gp`'s state in the middle of a computation (increase debugging level, start storing results in a log file, set variables to different values...): hit **C-c**, type in your modifications, then let the computation go on as explained above. A break loop looks like this:

```
? v = 0; 1/v
```

```

*** at top-level: v=0;1/v
***      ^__
*** _/: impossible inverse in gdiv: 0.
*** Break loop (type 'break' to go back to the GP prompt)
break>

```

So the standard error message is printed first. The `break>` at the bottom is a prompt, and hitting `v` then `<Return>`, we see:

```

break> v
0

```

explaining the problem. We could have typed any `gp` command, not only the name of a variable, of course. Lexically-scoped variables are accessible to the evaluator during the break loop:

```

? for(v = -2, 2, print(1/v))
-1/2
-1
*** at top-level: for(v=-2,2,print(1/v))
***      ^----
*** _/: impossible inverse in gdiv: 0.
*** Break loop (type 'break' to go back to the GP prompt)
break> v
0

```

Even though loop indices are automatically lexically scoped and no longer exist when the break loop is run, enough debugging information is retained in the bytecode to reconstruct the evaluation context. Of course, when the error occurs in a nested chain of user function calls, lexically scoped variables are available only in the corresponding frame:

```

? f(x) = 1/x;
? g(x) = for(i = 1, 10, f(x+i));
? for(j = -5,5, g(j))
*** at top-level: for(j=-5,5,g(j))
***      ^-----
*** in function g: for(i=1,10,f(x+i))
***      ^-----
*** in function f: 1/x
***      ^__
*** _/: impossible inverse in gdiv: 0.
*** Break loop: type 'break' to go back to GP prompt
break> [i,j,x]      \\ the x in f's body.
[i, j, 0]
break> dbg_up      \\ go up one frame
*** at top-level: for(j=-5,5,g(j))
***      ^-----
*** in function g: for(i=1,10,f(x+i))
***      ^-----
break> [i,j,x]      \\ the x in g's body, i in the for loop.
[5, j, -5]

```

The following GP commands are available during a break loop to help debugging:

`dbg_up(n)`: go up *n* frames, as seen above.
`dbg_down(n)`: go down *n* frames, cancelling previous `dbg_up`'s.
`dbg_x(t)`: examine *t*, as `\x` but more flexible.

`dbg_err()`: returns the current error context `t_ERROR`. The error components often provide useful additional information:

```
? 0(2) + 0(3)
***   at top-level: 0(2)+0(3)
***               ^-----
***   _+: inconsistent addition t_PADIC + t_PADIC.
***   Break loop: type 'break' to go back to GP prompt
break> E = dbg_err()
error("inconsistent addition t_PADIC + t_PADIC.")
break> Vec(E)
["e_OP", "+", 0(2), 0(3)]
```

Note. The debugger is enabled by default, and fires up as soon as a runtime error occurs. If you do not like this behavior, you may disable it by setting the default `breakloop` to 0 in for `gprc`. A runtime error will send you back to the prompt. Note that the break loop is automatically disabled when running `gp` in non interactive mode, i.e. when the program's standard input is not attached to a terminal.

Technical Note. When you enter a break loop due to a PARI stack overflow, the PARI stack is reset so that you can run commands. Otherwise the stack would immediately overflow again! Still, as explained above, you do not lose the value of any `gp` variable in the process.

2.10.4 Protecting code. The expression

`iferr(statements, ERR, recovery)`

evaluates and returns the value of *statements*, unless an error occurs during the evaluation in which case the value of *recovery* is returned. As in an if/else clause, with the difference that *statements* has been partially evaluated, with possible side effects. We shall give a lot more details about the `ERR` argument shortly; it is the name of a variable, lexically scoped to the *recovery* expression sequence, whose value is set by the exception handler to help the recovery code decide what to do about the error.

For instance one can define a fault tolerant inversion function as follows:

```
? inv(x) = iferr(1/x, ERR, "oo")    \\ ERR is unused...
? for (i=-1,1, print(inv(i)))
-1
oo
1
```

Protected codes can be nested without adverse effect. Let's now see how `ERR` can be used; as written, `inv` is too tolerant:

```
? inv("blah")
%2 = "oo"
```

Let's improve it by checking that we caught a "division by 0" exception, and not an unrelated one like the type error `1 / "blah"`.

```
? inv2(x) = {
    iferr(1/x,
        ERR, if (errname(ERR) != "e_INV", error(ERR)); "oo")
}
? inv2(0)
%3 = "oo"  \\ as before
? inv2("blah")
*** at top-level: inv2("blah")
***      ^-----
*** in function inv2: ...f(errname(ERR)!="e_INV",error(ERR));"oo"
***      ^-----
*** error: forbidden division t_INT / t_STR.
```

In the `inv2("blah")` example, the error type was not expected, so we rethrow the exception: `error(ERR)` triggers the original error that we mistakenly trapped. Since the recovery code should always check whether the error is the one expected, this construction is very common and can be simplified to

```
? inv3(x) = iferr(1/x,
    ERR, "oo",
    errname(ERR) == "e_INV")
```

More generally

```
iferr(statements, ERR, recovery, predicate)
```

only catches the exception if *predicate* (allowed to check various things about `ERR`, not only its name) is nonzero.

Rather than trapping everything, then rethrowing whatever we do not like, we advise to only trap errors of a specific kind, as above. Of course, sometimes, one just want to trap *everything* because we do not know what to expect. The following function check whether `install` works correctly in your `gp`:

```
broken_install() =
{ \\ can we install?
  iferr(install(addii,GG),
    ERR, return ("OS"));
  \\ can we use the installed function?
  iferr(if (addii(1,1) != 2, return("BROKEN")),
    ERR, return("USE"));
  return (0);
}
```

The function returns `OS` if the operating system does not support `install`, `USE` if using an installed function triggers an error, `BROKEN` if the installed function did not behave as expected, and `0` if everything works.

The `ERR` formal parameter contains more useful data than just the error name, which we recovered using `errname(ERR)`. In fact, a `t_ERROR` object usually has extra components, which can

be accessed as `component(ERR,1)`, `component(ERR,2)`, and so on. Or globally by casting the error to a `t_VEC`: `Vec(ERR)` returns the vector of all components at once. See Section 3.1.24 for the list of all exception types, and the corresponding contents of `ERR`.

2.11 Interfacing GP with other languages.

The PARI library was meant to be interfaced with C programs. This specific use is dealt with extensively in the *User's guide to the PARI library*. Of course, `gp` itself provides a convenient interpreter to execute rather intricate scripts (see Section 3.1).

Scripts, when properly written, tend to be shorter and clearer than C programs, and are certainly easier to write, maintain or debug. You don't need to deal with memory management, garbage collection, pointers, declarations, and so on. Because of their intrinsic simplicity, they are more robust as well. They are unfortunately somewhat slower. Thus their use will remain complementary: it is suggested that you test and debug your algorithms using scripts, before actually coding them in C if speed is paramount. The GP2C compiler often eases this part.

The `install` command (see Section 3.2.41) efficiently imports foreign functions for use under `gp`, which can of course be written using other libraries than PARI. Thus you may code only critical parts of your program in C, and still maintain most of the program as a GP script.

We are aware of three PARI-related Free Software packages to embed PARI in other languages. We *neither endorse nor support* any of them, but you may want to give them a try if you are familiar with the languages they are based on. The first is the Python-based SAGE system (<https://sagemath.org/>). The second is the `Math::Pari` Perl module (see any CPAN mirror), written by Ilya Zakharevich. Finally, Michael Stoll and Sam Steingold have integrated PARI into CLISP (<https://clisp.cons.org/>), a Common Lisp implementation.

These provide interfaces to `gp` functions for use in `python`, `perl`, or `Lisp` programs, respectively.

2.12 Defaults.

There are many internal variables in `gp`, defining how the system will behave in certain situations, unless a specific override has been given. Most of them are a matter of basic customization (colors, prompt) and will be set once and for all in your preferences file (see Section 2.14), but some of them are useful interactively (set timer on, increase precision, etc.).

The function used to manipulate these values is called `default`, which is described in Section 3.2.12. The basic syntax is

```
default(def, value),
```

which sets the default `def` to `value`. In interactive use, most of these can be abbreviated using `gp` metacommands (mostly, starting with `\`), which we shall describe in the next section.

Available defaults are described in the reference guide, Section 3.4, the most important one being `parisizemax`. Just be aware that typing `default` by itself will list all of them, as well as their current values (see `\d`).

Note. The suffixes **k**, **M**, **G** or **T** can be appended to a *value* which is a numeric argument, with the effect of multiplying it by 10^3 , 10^6 and 10^9 respectively. Case is not taken into account there, so for instance **30k** and **30K** both stand for 30000. This is mostly useful to modify or set the defaults **parisize** and **parisizemax** which typically involve a lot of trailing zeroes.

The suffixes **kB** or **KB**, **MB**, **GB**, **TB** can be appended to a *value* which is a numeric argument representing a memory size, with the usual meaning of counting in units of 2^{10} , 2^{20} , 2^{30} and 2^{40} bytes respectively. This allows to specify defaults such as **parisize** or **parisizemax** in customary units, such as gigabytes (or more properly gibibytes). For instance, **1k** represents 1000 bytes and **1kB** represents 1024 bytes.

(somewhat technical) Note. As we saw in Section 2.9.1, the second argument to **default** is subject to string context expansion, which means you can use run-time values. In other words, something like

```
a = 3;
default(logfile, "file" a ".log")
```

logs the output in **file3.log**.

Some special defaults, corresponding to file names and prompts, expand further the resulting value at the time they are set. Two kinds of expansions may be performed:

- **time expansion:** the string is sent through the library function **strftime**. This means that *%char* combinations have a special meaning, usually related to the time and date. For instance, **%H** = hour (24-hour clock) and **%M** = minute [00,59] (on a Unix system, you can try **man strftime** at your shell prompt to get a complete list). This is applied to **prompt** and **logfile**. For instance,

```
default(prompt, "(%H:%M) ? ")
```

will prepend the time of day, in the form (*hh:mm*) to **gp**'s usual prompt.

- **environment expansion:** When the string contains a sequence of the form **\$SOMEVAR**, e.g. **\$HOME**, the environment is searched and if *SOMEVAR* is defined, the sequence is replaced by the corresponding value. Also the **~** symbol has the same meaning as in many shells — **~** by itself stands for your home directory, and **~user** is expanded to **user**'s home directory. This is applied to all file names.

2.13 Simple metacommands.

Simple metacommands are meant as shortcuts and should not be used in GP scripts (see Section 3.1). Beware that these, as all of **gp** input, are *case sensitive*. For example, **\Q** is not identical to **\q**. Two kinds of arguments are allowed: numbers (denoted *n* below) and names (denoted *filename* below); braces are used to denote optional arguments, , e.g. **{n}** means that a numeric argument is expected but can be omitted. Names can be optionally surrounded by double quotes and in this case can contain whitespace, e.g. **"a b"** and are treated as ordinary character strings, see Section 2.9.1 for details.

Whitespace (or spaces) between the metacommand and its arguments and within unquoted arguments is optional. This can cause problems with **\w**, when you insist on having a file name whose first character is a digit, and with **\r** or **\w**, if the file name itself contains a space. In such cases, just quote filenames or use the underlying **read** or **write** function.

2.13.1 `?{command}`. The **gp** on-line help interface. If you type `?n` where n is a number from 1 to 11, you will get the list of functions in Section 3. n of the manual (the list of sections being obtained by simply typing `?`).

These names are in general not informative enough. More details can be obtained by typing `?function`, which gives a short explanation of the function's calling convention and effects. A help string is also attached to a symbolic operator, where arguments are replaced by a placeholder character `_`:

```
? ?sin
sin(x): sine of x.

? ?_*_
x*y: product of x and y.

? ?!_
!a: boolean operator "not".

? ?_!
n!: factorial of n.

? ? _^_
x^y: compute x to the power y.
```

Of course, to have complete information, read Chapter 3 of this manual. The source code is at your disposal as well, though a trifle less readable.

If the line before the copyright message indicates that extended help is available (this means `perl` is present on your system and the PARI distribution was correctly installed), you can add more `?` signs for extended functionality:

`?? keyword` yields the function description as it stands in this manual, usually in Chapter 2 or 3. If you're not satisfied with the default chapter chosen, you can impose a given chapter by ending the keyword with `@` followed by the chapter number, e.g. `?? Hello@2` will look in Chapter 2 for section heading `Hello` (which doesn't exist, by the way).

All operators (e.g. `+`, `&&`, etc.) are accepted by this extended help, as well as a few other keywords describing key **gp** concepts, e.g. `readline` (the line editor), `integer`, `nf` ("number field" as used in most algebraic number theory computations), `ell` (elliptic curves), etc.

In case of conflicts between *function* and *default* names (e.g `log`, `simplify`), the function has higher priority. To get the *default* help, use

```
?? default(log)
?? default(simplify)
```

`??? pattern` produces a list of sections in Chapter 3 of the manual related to your query. As before, if *pattern* ends by `@` followed by a chapter number, that chapter is searched instead; you also have the option to append a simple `@` (without a chapter number) to browse through the whole manual.

If your query contains dangerous characters (e.g `?` or blanks) it is advisable to enclose it within double quotes, as for GP strings (e.g `??? "elliptic curve"`).

Note that extended help is more powerful than the short help, since it knows about operators as well: you can type `?? *` or `?? &&`, whereas a single `?` would just yield a not too helpful

`&&`: unknown identifier.

message. Also, you can ask for extended help on section number n in Chapter 3, just by typing `?? n` (where $?n$ would yield merely a list of functions). Finally, a few key concepts in `gp` are documented in this way: metacommands (e.g. `?? "??"`), defaults (e.g. `?? default(log)`) not to be mistaken with `?? log` (the natural logarithm) and type names (e.g. `t_INT` or `integer`), as well as various miscellaneous keywords such as `edit` (short summary of line editor commands), `operator`, `member`, "user defined", `nf`, `ell`, ...

Last but not least: `??` without argument will open a dvi previewer (`xdvi` by default, `$GPXDVI` if it is defined in your environment) containing the full user's manual. `??tutorial` and `??refcard` do the same with the tutorial and reference card respectively.

Technical note. This functionality is provided by an external `perl` script that you are free to use outside any `gp` session (and modify to your liking, if you are perl-knowledgeable). It is called `gphelp`, lies in the `doc` subdirectory of your distribution (just make sure you run `Configure` first, see Appendix A) and is really two programs in one. The one which is used from within `gp` is `gphelp` which runs `TEX` on a selected part of this manual, then opens a previewer. `gphelp -detex` is a text mode equivalent, which looks often nicer especially on a colour-capable terminal (see `misc/gprc.dft` for examples). The default `help` selects which help program will be used from within `gp`. You are welcome to improve this help script, or write new ones (and we would like to know about it so that we may include them in future distributions). By the way, outside of `gp` you can give more than one keyword as argument to `gphelp`.

2.13.2 `/*...*/`. A comment. Everything between the stars is ignored by `gp`. These comments can span any number of lines.

2.13.3 `\\`. A one-line comment. The rest of the line is ignored by `gp`.

2.13.4 `\a {n}`. Prints the object number n (`%n`) in raw format (see `??output`). If the number n is omitted, print the latest computed object (`%`).

2.13.5 `\b {n}`. As `\a` using "beautified" (`prettymatrix`) format (see `??output`).

2.13.6 `\B {n}`. As `\b` using an external prettyprinter (see `??output` and `??prettyprinter`). If no prettyprinter is defined or available, this is identical to `\b`.

2.13.7 `\c`. Prints the list of all available hardcoded functions under `gp`, not including operators written as special symbols (see Section 2.4). More information can be obtained using the `?` metacommand (see above). For user-defined functions / member functions, see `\u` and `\um`.

2.13.8 `\d`. Prints the defaults as described in the previous section (shortcut for `default()`, see Section 3.2.12).

2.13.9 `\e {n}`. Switches the `echo` mode on (1) or off (0). If n is explicitly given, set `echo` to n .

2.13.10 `\g {n} {feature}`. Sets the debugging level `debug` to the nonnegative integer n . If *feature* is present (such as `bnf` or `qf111`), only set the debugging level for that feature, as by using `setdebug`.

2.13.11 `\g feature {n}`. Prints the debugging level for given *feature* (such as `bnf` or `qf111`, see `setdebug`). If the nonnegative integer n is present set the debugging level for that feature.

2.13.12 `\gf {n}`. Sets the "io" (or file usage) debugging level to the nonnegative integer n . This is a shortcut for `setdebug("io", n)`.

2.13.13 `\gm {n}`. Sets the memory debugging level `debugmem` to the nonnegative integer n .

2.13.14 `\h {m-n}`. Outputs some debugging info about the hashtable of identifiers used by the GP parser. If the argument is a number n , outputs the contents of cell n . Ranges can be given in the form $m-n$ (from cell m to cell n , $\$$ = last cell). If a function name is given instead of a number or range, outputs info on the internal structure of the hash cell this function occupies (a `struct entree` in C). If the range is reduced to a dash ('-'), outputs statistics about hash cell usage.

2.13.15 `\l {logfile}`. Switches `log` mode on and off. If a *logfile* argument is given, change the default logfile name to *logfile* and switch log mode on.

2.13.16 `\m`. As `\b`.

2.13.17 `\o {n}`. Sets output mode to n (0: raw, 1: prettymatrix, 3: external prettyprint). See `??output`

2.13.18 `\p {n}`. Sets `realprecision` to n decimal digits. Prints its current value if n is omitted.

2.13.19 `\pb {n}`. Sets `realbitprecision` to n bits. Prints its current value if n is omitted.

2.13.20 `\ps {n}`. Sets `seriesprecision` to n significant terms. Prints its current value if n is omitted.

2.13.21 `\q`. Quits the `gp` session and returns to the system. Shortcut for `quit()` (see Section 3.2.63).

2.13.22 `\r {filename}`. Reads into `gp` all the commands contained in the named file as if they had been typed from the keyboard, one line after the other. Can be used in combination with the `\w` command (see below). Related but not equivalent to the function `read` (see Section 3.2.64); in particular, if the file contains more than one line of input, there will be one history entry for each of them, whereas `read` would only record the last one. If *filename* is omitted, re-read the previously used input file (fails if no file has ever been successfully read in the current session). If a `gp` binary file (see Section 3.2.88) is read using this command, it is silently loaded, without cluttering the history.

Assuming `gp` figures how to decompress files on your machine, this command accepts compressed files in `compressed` (.Z) or `gzipped` (.gz or .z) format. They will be uncompressed on the fly as `gp` reads them, without changing the files themselves.

2.13.23 `\s`. Prints the state of the PARI *stack* and *heap*. This is used primarily as a debugging device for PARI.

2.13.24 `\t`. Prints the internal longword format of all the PARI types. The detailed bit or byte format of the initial codeword(s) is explained in Chapter 4, but its knowledge is not necessary for a `gp` user.

2.13.25 `\u`. Prints the definitions of all user-defined functions.

2.13.26 `\um`. Prints the definitions of all user-defined member functions.

2.13.27 `\uv`. Prints the definitions of all user-defined variables, closures being excluded.

2.13.28 `\v`. Prints the version number and implementation architecture (680x0, Sparc, Alpha, other) of the `gp` executable you are using.

2.13.29 `\w {n} {filename}`. Writes the object number n (`%n`) into the named file, in raw format. If the number n is omitted, writes the latest computed object (`%`). If *filename* is omitted, appends to `logfile` (the GP function `write` is a trifle more powerful, as you can have arbitrary file names).

2.13.30 `\x {n}`. Prints the complete tree with addresses and contents (in hexadecimal) of the internal representation of the object number n (`%n`). If the number n is omitted, uses the latest computed object in `gp`. As for `\s`, this is used primarily as a debugging device for PARI, and the format should be self-explanatory. The underlying GP function `dbg_x` is more versatile, since it can be applied to other objects than history entries.

2.13.31 `\y {n}`. Switches `simplify` on (1) or off (0). If n is explicitly given, set `simplify` to n .

2.13.32 `#`. Switches the `timer` on or off.

2.13.33 `##`. Prints the time taken by the latest computation. Useful when you forgot to turn on the `timer`.

2.14 The preferences file.

This file, called `gprc` in the sequel, is used to modify or extend `gp` default behavior, in all `gp` sessions: e.g. customize `default` values or load common user functions and aliases. `gp` opens the `gprc` file and processes the commands in there, *before* doing anything else, e.g. creating the PARI stack. If the file does not exist or cannot be read, `gp` will proceed to the initialization phase at once, eventually emitting a prompt. If any explicit command line switches are given, they override the values read from the preferences file.

2.14.1 Syntax. The syntax in the `gprc` file (and valid in this file only) is simple-minded, but should be sufficient for most purposes. The file is read line by line; as usual, white space is ignored unless surrounded by quotes and the standard multiline constructions using braces, `\`, or `=` are available (multiline comments between `/* ... */` are also recognized).

2.14.1.1 Preprocessor:. Two types of lines are first dealt with by a preprocessor:

- comments are removed. This applies to all text surrounded by `/* ... */` as well as to everything following `\\` on a given line.

- lines starting with `#if boolean` are treated as comments if *boolean* evaluates to `false`, and read normally otherwise. The condition can be negated using either `#if not`, `#ifnot` or `#if !`. If the rest of the current line is empty, the test applies to the next line (same behavior as `=` under `gp`). The following tests can be performed:

EMACS: true if `gp` is running in an Emacs or TeXmacs shell (see Section 2.16).

READL: true if `gp` is compiled with `readline` support (see Section 2.15).

`VERSION op number`: where *op* is in the set $\{>, <, <=, >=\}$, and *number* is a PARI version number of the form *Major.Minor.patch*, where the last two components can be omitted (i.e. 1 is understood as version 1.0.0). This is `true` if `gp`'s version number satisfies the required inequality.

`BITS_IN_LONG == number`: *number* is 32 (resp. 64). This is `true` if `gp` was built for a 32-bit (resp. 64-bit) architecture.

2.14.1.2 Commands: After preprocessing, the remaining lines are executed as sequence of expressions (as usual, separated by `;` if necessary). Only following kinds of expressions are recognized:

- `dft = value`, where *dft* is one of the available defaults (see Section 2.12), which will be set to *value* on actual startup. Don't forget the quotes around strings (e.g. for `prompt` or `help`).

- `default(dft, value)` as above.

- `setdebug(dom, value)` set debug level for domain *dom* to *value*.

- `read "some_GP_file"` where *some_GP_file* is a regular GP script this time, which will be read just before `gp` prompts you for commands, but after initializing the defaults. In particular, file input is delayed until the `gprc` has been fully loaded. This is the right place to input files containing `alias` commands, or your favorite macros.

For instance you could set your prompt in the following portable way:

```
\\ self modifying prompt looking like (18:03) gp >
prompt    = "(%H:%M) \e[1m\gp\e[m > "

\\ readline wants nonprinting characters to be braced between ^A/^B pairs
#if READL prompt = "(%H:%M) ^A\e[1m^Bgp^A\e[m^B > "

\\ escape sequences not supported under emacs
#if EMACS prompt = "(%H:%M) gp > "
```

Note that any of the last two lines could be broken in the following way

```
#if EMACS
    prompt = "(%H:%M) gp > "
```

since the preprocessor directive applies to the next line if the current one is empty.

A sample `gprc` file called `misc/gprc.dft` is provided in the standard distribution. It is a good idea to have a look at it and customize it to your needs. Since this file does not use multiline constructs, here is one (note the terminating `;` to separate the expressions):

```
#if VERSION > 2.2.3
{
    read "my_scripts";    \\ syntax errors in older versions
    new_galois_format = 1; \\ default introduced in 2.2.4
}
#endif ! EMACS
{
    colors = "9, 5, no, no, 4, 1, 2";
    help    = "gphelp -detex -ch 4 -cb 0 -cu 2";
}
```

2.14.2 The gprc location. When `gp` is started, it looks for a customization file, or `gprc` in the following places (in this order, only the first one found will be loaded):

- `gp` checks whether the environment variable `GPRC` is set. On Unix, this can be done with something like:

<code>GPRC=/my/dir/anyname; export GPRC</code>	in <code>sh</code> syntax (for instance in your <code>.profile</code>),
<code>setenv GPRC /my/dir/anyname</code>	in <code>csh</code> syntax (in your <code>.login</code> or <code>.cshrc</code> file).
<code>env GPRC=/my/dir/anyname gp</code>	on the command line launching <code>gp</code> .

If so, the file named by `$GPRC` is the `gprc`.

- If `GPRC` is not set, and if the environment variable `HOME` is defined, `gp` then tries

`$HOME/.gprc` on a Unix system

`$HOME\gprc.txt` on a DOS, OS/2, or Windows system.

- If no `gprc` was found among the user files mentioned above we look for `/etc/gprc` for a system-wide `gprc` file (you will need root privileges to set up such a file yourself).
- Finally, we look in `pari's datadir` for a file named

`.gprc` on a Unix system

`gprc.txt` on a DOS, OS/2, or Windows system. If you are using our Windows installer, this is where the default preferences file is written.

Note that on Unix systems, the `gprc's` default name starts with a `'.'` and thus is hidden to regular `ls` commands; you need to type `ls -a` to list it.

2.15 Using readline.

This very useful library provides line editing and contextual completion to `gp`. You are encouraged to read the `readline` user manual, but we describe basic usage here.

A (too) short introduction to readline. In the following, `C-` stands for “the **C**ontrol key combined with another” and the same for `M-` with the **M**eta key; generally `C-` combinations act on characters, while the `M-` ones operate on words. The **M**eta key might be called **A**lt on some keyboards, will display a black diamond on most others, and can safely be replaced by **E**sc in any case.

Typing any ordinary key inserts text where the cursor stands, the arrow keys enabling you to move in the line. There are many more movement commands, which will be familiar to the Emacs user, for instance `C-a/C-e` will take you to the start/end of the line, `M-b/M-f` move the cursor backward/forward by a word, etc. Just press the `<Return>` key at any point to send your command to `gp`.

All the commands you type at the `gp` prompt are stored in a history, a multiline command being saved as a single concatenated line. The Up and Down arrows (or `C-p/C-n`) will move you through the history, `M-</M->` sending you to the start/end of the history. `C-r/C-s` will start an incremental backward/forward search. You can kill text (`C-k` kills till the end of line, `M-d` to the end of current word) which you can then yank back using the `C-y` key (`M-y` will rotate the kill-ring). `C-_` will undo your last changes incrementally (`M-r` undoes all changes made to the current line). `C-t` and `M-t` will transpose the character (word) preceding the cursor and the one under the cursor.

Keeping the M- key down while you enter an integer (a minus sign meaning reverse behavior) gives an argument to your next readline command (for instance M-- C-k will kill text back to the start of line). If you prefer Vi-style editing, M-C-j will toggle you to Vi mode.

Of course you can change all these default bindings. For that you need to create a file named `.inputrc` in your home directory. For instance (notice the embedding conditional in case you would want specific bindings for `gp`):

```
$if Pari-GP
  set show-all-if-ambiguous
  "\C-h": backward-delete-char
  "\e\C-h": backward-kill-word
  "\C-xd": dump-functions
  (: "\C-v()\C-b"          # can be annoying when copy-pasting!
  [: "\C-v[]\C-b"
$endif
```

C-x C-r will re-read this init file, incorporating any changes made to it during the current session.

Note. By default, (and [are bound to the function `pari-matched-insert` which, if “electric parentheses” are enabled (default: off) will automatically insert the matching closure (respectively) and]). This behavior can be toggled on and off by giving the numeric argument `-2` to ((`M--2()`), which is useful if you want, e.g to copy-paste some text into the calculator. If you do not want a toggle, you can use `M--0` / `M--1` to specifically switch it on or off).

Note. In some versions of readline (2.1 for instance), the `Alt` or `Meta` key can give funny results (output 8-bit accented characters for instance). If you do not want to fall back to the `Esc` combination, put the following two lines in your `.inputrc`:

```
set convert-meta on
set output-meta off
```

Command completion and online help. Hitting `<TAB>` will complete words for you. This mechanism is context-dependent: `gp` will strive to only give you meaningful completions in a given context (it will fail sometimes, but only under rare and restricted conditions).

For instance, shortly after a `~`, we expect a user name, then a path to some file. Directly after `default(` has been typed, we would expect one of the `default` keywords. After a `'.`, we expect a member keyword. And generally of course, we expect any GP symbol which may be found in the hashing lists: functions (both yours and GP’s), and variables.

If, at any time, only one completion is meaningful, `gp` will provide it together with

- an ending comma if we are completing a default,
- a pair of parentheses if we are completing a function name. In that case hitting `<TAB>` again will provide the argument list as given by the online help. (Recall that you can always undo the effect of the preceding keys by hitting `C-;`; this applies here.)

Otherwise, hitting `<TAB>` once more will give you the list of possible completions. Just experiment with this mechanism as often as possible, you will probably find it very convenient. For instance, you can obtain `default(seriesprecision,10)`, just by hitting `def<TAB>se<TAB>10`, which saves 18 keystrokes (out of 27).

Hitting **M-h** will give you the usual short online help concerning the word directly beneath the cursor, **M-H** will yield the extended help corresponding to the **help** default program (usually opens a dvi previewer, or runs a primitive tex-to-ASCII program). None of these disturb the line you were editing.

2.16 GNU Emacs and PariEmacs.

If you install the PariEmacs package (see Appendix A), you may use **gp** as a subprocess in Emacs. You then need to include in your `.emacs` file the following lines:

```
(autoload 'gp-mode "pari" nil t)
(autoload 'gp-script-mode "pari" nil t)
(autoload 'gp "pari" nil t)
(autoload 'gpman "pari" nil t)
(setq auto-mode-alist
  (cons '("\\.gp$" . gp-script-mode) auto-mode-alist))
```

which autoloads functions from the PariEmacs package and ensures that file with the `.gp` suffix are edited in `gp-script` mode.

Once this is done, under GNU Emacs if you type **M-x gp** (where as usual **M** is the **Meta** key), a special shell will be started launching **gp** with the default stack size and prime limit. You can then work as usual under **gp**, but with all the facilities of an advanced text editor. See the PariEmacs documentation for customizations, menus, etc.

Chapter 3:

Functions and Operations Available in PARI and GP

The functions and operators available in PARI and in the GP/PARI calculator are numerous and ever-expanding. Here is a description of the ones available in version 2.17.4. It should be noted that many of these functions accept quite different types as arguments, but others are more restricted. The list of acceptable types will be given for each function or class of functions. Except when stated otherwise, it is understood that a function or operation which should make natural sense is legal.

On the other hand, many routines list explicit preconditions for some of their arguments, e.g. p is a prime number, or q is a positive definite quadratic form. For reasons of efficiency, all routines trust the user input and only perform minimal sanity checks. When a precondition is not satisfied, any of the following may occur: a regular exception is raised, the PARI stack overflows, a **SIGSEGV** or **SIGBUS** signal is generated, or we enter an infinite loop. The function can also quietly return a mathematically meaningless result: junk in, junk out. In the following, we document the results as *undefined* in this case.

In this chapter, we will describe the functions according to a rough classification. The general entry looks something like:

foo(x , {*flag* = 0}): short description.

The library syntax is **GEN foo**(GEN x , long *flag* = 0).

This means that the GP function **foo** has one mandatory argument x , and an optional one, *flag*, whose default value is 0. (The {} should not be typed, it is just a convenient notation that we will use throughout to denote optional arguments.) That is, you can type **foo**(x ,2), or **foo**(x), which is then understood to mean **foo**(x ,0). As well, a comma or closing parenthesis, where an optional argument should have been, signals to GP it should use the default. Thus, the syntax **foo**(x ,) is also accepted as a synonym for our last expression. When a function has more than one optional argument, the argument list is filled with user supplied values, in order. When none are left, the defaults are used instead. Thus, assuming that **foo**'s prototype had been

$$\mathbf{foo}(\{x = 1\}, \{y = 2\}, \{z = 3\}),$$

typing in **foo**(6,4) would give you **foo**(6,4,3). In the rare case when you want to set some far away argument, and leave the defaults in between as they stand, you can use the “empty arg” trick alluded to above: **foo**(6,,1) would yield **foo**(6,2,1). By the way, **foo**() by itself yields **foo**(1,2,3) as was to be expected.

In this rather special case of a function having no mandatory argument, you can even omit the (): a standalone **foo** would be enough (though we do not recommend it for your scripts, for the sake of clarity). In defining GP syntax, we strove to put optional arguments at the end of the argument list (of course, since they would not make sense otherwise), and in order of decreasing usefulness so that, most of the time, you will be able to ignore them.

Finally, an optional argument (between braces) followed by a star, like $\{x\}$ *, means that any number of such arguments (possibly none) can be given. This is in particular used by the various **print** routines.

Flags. A *flag* is an argument which, rather than conveying actual information to the routine, instructs it to change its default behavior, e.g. return more or less information. All such flags are optional, and will be called *flag* in the function descriptions to follow. There are two different kind of flags

- generic: all valid values for the flag are individually described (“If *flag* is equal to 1, then...”).
- binary: use customary binary notation as a compact way to represent many toggles with just one integer. Let $(p_0, \dots, p_n)$ be a list of switches (i.e. of properties which take either the value 0 or 1), the number $2^3 + 2^5 = 40$ means that p_3 and p_5 are set (that is, set to 1), and none of the others are (that is, they are set to 0). This is announced as “The binary digits of *flag* mean 1: p_0 , 2: p_1 , 4: p_2 ”, and so on, using the available consecutive powers of 2.

Mnemonics for binary flags. Numeric flags as mentioned above are obscure, error-prone, and quite rigid: should the authors want to adopt a new flag numbering scheme, it would break backward compatibility. The only advantage of explicit numeric values is that they are fast to type, so their use is only advised when using the calculator `gp`.

As an alternative, one can replace a binary flag by a character string containing symbolic identifiers (mnemonics). In the function description, mnemonics corresponding to the various toggles are given after each of them. They can be negated by prepending `no_` to the mnemonic, or by removing such a prefix. These toggles are grouped together using any punctuation character (such as `,` or `;`). For instance (taken from description of `plott(X = a, b, expr, {flag = 0}, {n = 0})`)

Binary digits of flags mean: 1 = `Parametric`, 2 = `Recursive`, ...

so that, instead of 1, one could use the mnemonic `"Parametric; no_Recursive"`, or simply `"Parametric"` since `Recursive` is unset by default (default value of *flag* is 0, i.e. everything unset). People used to the bit-or notation in languages like C may also use the form `"Parametric | no_Recursive"`.

Pointers. If a parameter in the function prototype is prefixed with a `&` sign, as in

`foo(x, &e)`

it means that, besides the normal return value, the function may assign a value to *e* as a side effect. When passing the argument, the `&` sign has to be typed in explicitly. As of version 2.17.4, this *pointer* argument is optional for all documented functions, hence the `&` will always appear between brackets as in `Z_issquare(x, {&e})`.

About library programming. The *library* function `foo`, as defined at the beginning of this section, is seen to have two mandatory arguments, *x* and *flag*: no function seen in the present chapter has been implemented so as to accept a variable number of arguments, so all arguments are mandatory when programming with the library (usually, variants are provided corresponding to the various flag values). We include an `= default value` token in the prototype to signal how a missing argument should be encoded. Most of the time, it will be a `NULL` pointer, or -1 for a variable number. Refer to the *User's Guide to the PARI library* for general background and details.

3.1 Programming in GP: control statements.

A number of control statements are available in GP. They are simpler and have a syntax slightly different from their C counterparts, but are quite powerful enough to write any kind of program. Some of them are specific to GP, since they are made for number theorists. As usual, X will denote any simple variable name, and *seq* will always denote a sequence of expressions, including the empty sequence.

Caveat. In constructs like

```
for (X = a,b, seq)
```

the variable X is lexically scoped to the loop, leading to possibly unexpected behavior:

```
n = 5;
for (n = 1, 10,
    if (something_nice(), break);
);
\\ at this point n is 5 !
```

If the sequence *seq* modifies the loop index, then the loop is modified accordingly:

```
? for (n = 1, 10, n += 2; print(n))
3
6
9
12
```

3.1.1 break($\{n = 1\}$). Interrupts execution of current *seq*, and immediately exits from the n innermost enclosing loops, within the current function call (or the top level loop); the integer n must be positive. If n is greater than the number of enclosing loops, all enclosing loops are exited.

3.1.2 breakpoint(). Interrupt the program and enter the breakloop. The program continues when the breakloop is exited.

```
? f(N,x)=my(z=x^2+1);breakpoint();gcd(N,z^2+1-z);
? f(221,3)
*** at top-level: f(221,3)
***      ^-----
*** in function f: my(z=x^2+1);breakpoint();gcd(N,z
***      ^-----
*** Break loop: type <Return> to continue; 'break' to go back to GP
break> z
10
break>
%2 = 13
```

3.1.3 dbg_down({n = 1}). (In the break loop) go down n frames. This allows to cancel a previous call to `dbg_up`.

```
? x = 0;
? g(x) = x-3;
? f(x) = 1 / g(x+1);
? for (x = 1, 5, f(x+1))
    *** at top-level: for(x=1,5,f(x+1))
    ***          ^-----
    *** in function f: 1/g(x+1)
    ***          ^-----
    *** _/_: impossible inverse in gdiv: 0.
    *** Break loop: type 'break' to go back to GP prompt
break> dbg_up(3) \\ go up 3 frames
    *** at top-level: for(x=1,5,f(x+1))
    ***          ^-----
break> x
0
break> dbg_down()
    *** at top-level: for(x=1,5,f(x+1))
    ***          ^-----
break> x
1
break> dbg_down()
    *** at top-level: for(x=1,5,f(x+1))
    ***          ^-----
break> x
1
break> dbg_down()
    *** at top-level: for(x=1,5,f(x+1))
    ***          ^-----
    *** in function f: 1/g(x+1)
    ***          ^-----
break> x
2
```

The above example shows that the notion of GP frame is finer than the usual stack of function calls (as given for instance by the GDB `backtrace` command): GP frames are attached to variable scopes and there are frames attached to control flow instructions such as a `for` loop above.

3.1.4 dbg_err(). In the break loop, return the error data of the current error, if any. See `iferr` for details about error data. Compare:

```
? iferr(1/(Mod(2,12019)^(6!)-1),E,Vec(E))
%1 = ["e_INV", "Fp_inv", Mod(119, 12019)]
? 1/(Mod(2,12019)^(6!)-1)
    *** at top-level: 1/(Mod(2,12019)^(6!)-
    ***          ^-----
    *** _/_: impossible inverse in Fp_inv: Mod(119, 12019).
    *** Break loop: type 'break' to go back to GP prompt
```

```
break> Vec(dbg_err())
["e_INV", "Fp_inv", Mod(119, 12019)]
```

3.1.5 dbg_up($\{n = 1\}$). (In the break loop) go up n frames, which allows to inspect data of the parent function. To cancel a `dbg_up` call, use `dbg_down`.

```
? x = 0;
? g(x) = x-3;
? f(x) = 1 / g(x+1);
? for (x = 1, 5, f(x+1))
  *** at top-level: for(x=1,5,f(x+1))
  ***      ^-----
  *** in function f: 1/g(x+1)
  ***      ^-----
  *** _/_: impossible inverse in gdiv: 0.
  *** Break loop: type 'break' to go back to GP prompt
break> x
2
break> dbg_up()
  *** at top-level: for(x=1,5,f(x+1))
  ***      ^-----
break> x
1
break> dbg_up()
  *** at top-level: for(x=1,5,f(x+1))
  ***      ^-----
break> x
1
break> dbg_up()
  *** at top-level: for(x=1,5,f(x+1))
  ***      ^-----
break> x
0
break> dbg_down()    \\ back up once
  *** at top-level: for(x=1,5,f(x+1))
  ***      ^-----
break> x
1
```

The above example shows that the notion of GP frame is finer than the usual stack of function calls (as given for instance by the GDB `backtrace` command): GP frames are attached to variable scopes and there are frames attached to control flow instructions such as a `for` loop above.

3.1.6 dbg_x($A, \{n\}$). Print the inner structure of A , complete if n is omitted, up to level n otherwise. This function is useful for debugging. It is similar to `\x` but does not require A to be a history entry. In particular, it can be used in the break loop.

3.1.7 for($X = a, b, seq$). Evaluates seq , where the formal variable X goes from a to b , where a and b must be in $\mathbf{R}$. Nothing is done if $a > b$. If b is set to $+\infty$, the loop will not stop; it is expected that the caller will break out of the loop itself at some point, using `break` or `return`.

3.1.8 forcomposite($n = a, \{b\}, seq$). Evaluates *seq*, where the formal variable *n* ranges over the composite numbers between the nonnegative real numbers *a* to *b*, including *a* and *b* if they are composite. Nothing is done if $a > b$.

```
? forcomposite(n = 0, 10, print(n))
4
6
8
9
10
```

Omitting *b* means we will run through all composites $\geq a$, starting an infinite loop; it is expected that the user will break out of the loop himself at some point, using **break** or **return**.

Note that the value of *n* cannot be modified within *seq*:

```
? forcomposite(n = 2, 10, n = [])
*** at top-level: forcomposite(n=2,10,n=[])
***                                     ^---
*** index read-only: was changed to [].
```

3.1.9 fordiv(*n*, *X*, *seq*). Evaluates *seq*, where the formal variable *X* ranges through the divisors of *n* (see **divisors**, which is used as a subroutine). It is assumed that **factor** can handle *n*, without negative exponents. Instead of *n*, it is possible to input a factorization matrix, i.e. the output of **factor**(*n*).

This routine uses **divisors** as a subroutine, then loops over the divisors. In particular, if *n* is an integer, divisors are sorted by increasing size.

To avoid storing all divisors, possibly using a lot of memory, the following (slower) routine loops over the divisors using essentially constant space:

```
FORDIV(N)=
{ my(F = factor(N), P = F[,1], E = F[,2]);
  forvec(v = vector(#E, i, [0,E[i]]), X = factorback(P, v));
}
? for(i=1, 10^6, FORDIV(i))
time = 11,180 ms.
? for(i=1, 10^6, fordiv(i, d, ))
time = 2,667 ms.
```

Of course, the divisors are no longer sorted by increasing size.

3.1.10 fordivfactored(n, X, seq). Evaluates seq , where the formal variable X ranges through $[d, \text{factor}(d)]$, where d is a divisors of n (see `divisors`, which is used as a subroutine). Note that such a pair is accepted as argument to all multiplicative functions.

It is assumed that `factor` can handle n , without negative exponents. Instead of n , it is possible to input a factorization matrix, i.e. the output of `factor(n)`. This routine uses `divisors(,1)` as a subroutine, then loops over the divisors. In particular, if n is an integer, divisors are sorted by increasing size.

This function is particularly useful when n is hard to factor and one must evaluate multiplicative function on its divisors: we avoid refactoring each divisor in turn. It also provides a small speedup when n is easy to factor; compare

```
? A = 10^8; B = A + 10^5;
? for (n = A, B, fordiv(n, d, eulerphi(d)));
time = 2,091 ms.
? for (n = A, B, fordivfactored(n, d, eulerphi(d)));
time = 1,298 ms. \\ avoid refactoring the divisors
? forfactored (n = A, B, fordivfactored(n, d, eulerphi(d)));
time = 1,270 ms. \\ also avoid factoring the consecutive n's !
```

3.1.11 foreach(V, X, seq). Evaluates seq , where the formal variable X ranges through the components of V (`t_VEC`, `t_COL`, `t_LIST` or `t_MAT`). A matrix argument is interpreted as a vector containing column vectors, as in `Vec(V)`.

3.1.12 forell($E, a, b, seq, \{flag = 0\}$). Evaluates seq , where the formal variable $E = [name, M, G]$ ranges through all elliptic curves of conductors from a to b . In this notation $name$ is the curve name in Cremona's elliptic curve database, M is the minimal model, G is a $\mathbf{Z}$ -basis of the free part of the Mordell-Weil group $E(\mathbf{Q})$. If $flag$ is nonzero, select only the first curve in each isogeny class.

```
? forell(E, 1, 500, my([name,M,G] = E); \
    if (#G > 1, print(name)))
389a1
433a1
446d1
? c = 0; forell(E, 1, 500, c++); c    \\ number of curves
%2 = 2214
? c = 0; forell(E, 1, 500, c++, 1); c \\ number of isogeny classes
%3 = 971
```

The `elldata` database must be installed and contain data for the specified conductors.

The library syntax is `forell(void *data, long (*f)(void*,GEN), long a, long b, long flag)`.

3.1.13 forfactored($N = a, b, seq$). Evaluates seq , where the formal variable N is $[n, \text{factor}(n)]$ and n goes from a to b ; a and b must be integers. Nothing is done if $a > b$.

This function is only implemented for $|a|, |b| < 2^{64}$ (2^{32} on a 32-bit machine). It uses a sieve and runs in time $O(\sqrt{b} + b - a)$. It should be at least 3 times faster than regular factorization as long as the interval length $b - a$ is much larger than $\sqrt{b}$ and get relatively faster as the bounds increase. The function slows down dramatically if `primelimit` $< \sqrt{b}$.

```
? B = 10^9;
? for (N = B, B+10^6, factor(N))
time = 4,538 ms.
? forfactored (N = B, B+10^6, [n,fan] = N)
time = 1,031 ms.

? B = 10^11;
? for (N = B, B+10^6, factor(N))
time = 15,575 ms.
? forfactored (N = B, B+10^6, [n,fan] = N)
time = 2,375 ms.

? B = 10^14;
? for (N = B, B+10^6, factor(N))
time = 1min, 4,948 ms.
? forfactored (N = B, B+10^6, [n,fan] = N)
time = 58,601 ms.
```

The last timing is with the default `primelimit` (500000) which is much less than $\sqrt{B + 10^6}$; it goes down to 26,750ms if `primelimit` gets bigger than that bound. In any case $\sqrt{B + 10^6}$ is much larger than the interval length 10^6 so `forfactored` gets relatively slower for that reason as well.

Note that all PARI multiplicative functions accept the `[n,fan]` argument natively:

```
? s = 0; forfactored(N = 1, 10^7, s += moebius(N)*eulerphi(N)); s
time = 6,001 ms.
%1 = 6393738650
? s = 0; for(N = 1, 10^7, s += moebius(N)*eulerphi(N)); s
time = 28,398 ms. \\ slower, we must factor N. Twice.
%2 = 6393738650
```

The following loops over the fundamental dicriminants less than X :

```
? X = 10^8;
? forfactored(d=1,X, if (isfundamental(d),));
time = 34,030 ms.
? for(d=1,X, if (isfundamental(d),))
time = 1min, 24,225 ms.
```

3.1.14 forpart($X = k, seq, \{a = k\}, \{n = k\}$). Evaluate *seq* over the partitions $X = [x_1, \dots, x_n]$ of the integer k , i.e. increasing sequences $x_1 \leq x_2 \leq \dots \leq x_n$ of sum $x_1 + \dots + x_n = k$. By convention, 0 admits only the empty partition and negative numbers have no partitions. A partition is given by a `t_VECSMALL`, where parts are sorted in nondecreasing order. The partitions are listed by increasing size and in lexicographic order when sizes are equal:

```
? forpart(X=4, print(X))
Vecsmall([4])
Vecsmall([1, 3])
Vecsmall([2, 2])
Vecsmall([1, 1, 2])
Vecsmall([1, 1, 1, 1])
```

Optional parameters n and a are as follows:

- $n = nmax$ (resp. $n = [nmin, nmax]$) restricts partitions to length less than $nmax$ (resp. length between $nmin$ and $nmax$), where the *length* is the number of nonzero entries.
- $a = amax$ (resp. $a = [amin, amax]$) restricts the parts to integers less than $amax$ (resp. between $amin$ and $amax$).

By default, parts are positive and we remove zero entries unless $amin \leq 0$, in which case we fix the size $\#X = nmax$:

```
\\ at most 3 nonzero parts, all <= 4
? forpart(v=5,print(Vec(v)), 4, 3)
[1, 4]
[2, 3]
[1, 1, 3]
[1, 2, 2]

\\ between 2 and 4 parts less than 5, fill with zeros
? forpart(v=5,print(Vec(v)), [0,5], [2,4])
[0, 0, 1, 4]
[0, 0, 2, 3]
[0, 1, 1, 3]
[0, 1, 2, 2]
[1, 1, 1, 2]

\\ no partitions of 1 with 2 to 4 nonzero parts
? forpart(v=1,print(v), [0,5], [2,4])
?
```

The behavior is unspecified if X is modified inside the loop.

The library syntax is `forpart(void *data, long (*call)(void*,GEN), long k, GEN a, GEN n)`.

3.1.15 forperm(a, p, seq). Evaluates seq , where the formal variable p goes through some permutations given by a `t_VECSMALL`. If a is a positive integer then P goes through the permutations of $\{1, 2, \dots, a\}$ in lexicographic order and if a is a small vector then p goes through the (multi)permutations lexicographically larger than or equal to a .

```
? forperm(3, p, print(p))
Vecsmall([1, 2, 3])
Vecsmall([1, 3, 2])
Vecsmall([2, 1, 3])
Vecsmall([2, 3, 1])
Vecsmall([3, 1, 2])
Vecsmall([3, 2, 1])
```

When a is itself a `t_VECSMALL` or a `t_VEC` then p iterates through multipermutations

```
? forperm([2,1,1,3], p, print(p))
Vecsmall([2, 1, 1, 3])
Vecsmall([2, 1, 3, 1])
Vecsmall([2, 3, 1, 1])
Vecsmall([3, 1, 1, 2])
Vecsmall([3, 1, 2, 1])
Vecsmall([3, 2, 1, 1])
```

3.1.16 forprime($p = a, \{b\}, seq$). Evaluates seq , where the formal variable p ranges over the prime numbers between the real numbers a to b , including a and b if they are prime. More precisely, the value of p is incremented to `nextprime(p + 1)`, the smallest prime strictly larger than p , at the end of each iteration. Nothing is done if $a > b$.

```
? forprime(p = 4, 10, print(p))
5
7
```

Setting b to `+oo` means we will run through all primes $\geq a$, starting an infinite loop; it is expected that the caller will break out of the loop itself at some point, using `break` or `return`.

Note that the value of p cannot be modified within seq :

```
? forprime(p = 2, 10, p = [])
*** at top-level: forprime(p=2,10,p=[])
*** ^---
*** prime index read-only: was changed to [].
```

3.1.17 forprimestep($p = a, b, q, seq$). Evaluates seq , where the formal variable p ranges over the prime numbers in an arithmetic progression in $[a, b]$: q is either an integer ($p \equiv a \pmod{q}$) or an integer $\text{Mod}(c, N)$ and we restrict to that congruence class. Nothing is done if $a > b$.

```
? forprimestep(p = 4, 30, 5, print(p))
19
29
? forprimestep(p = 4, 30, Mod(1,5), print(p))
11
```

Setting b to $+\infty$ means we will run through all primes $\geq a$, starting an infinite loop; it is expected that the caller will break out of the loop itself at some point, using **break** or **return**.

Note that the value of p cannot be modified within seq :

```
? forprimestep(p = 2, 10, 3, p = [])
*** at top-level: forprimestep(p=2,10,3,p=[])
***                                     ^---
*** prime index read-only: was changed to [].
```

3.1.18 forsquarefree($N = a, b, seq$). Evaluates seq , where the formal variable N is $[n, \text{factor}(n)]$ and n goes through squarefree integers from a to b ; a and b must be integers. Nothing is done if $a > b$.

```
? forsquarefree(N=-3,9,print(N))
[-3, [-1, 1; 3, 1]]
[-2, [-1, 1; 2, 1]]
[-1, Mat([-1, 1])]
[1, matrix(0,2)]
[2, Mat([2, 1])]
[3, Mat([3, 1])]
[5, Mat([5, 1])]
[6, [2, 1; 3, 1]]
[7, Mat([7, 1])]
```

This function is only implemented for $|a|, |b| < 2^{64}$ (2^{32} on a 32-bit machine). It uses a sieve and runs in time $O(\sqrt{b} + b - a)$. It should be at least 5 times faster than regular factorization as long as the interval length $b - a$ is much larger than $\sqrt{b}$ and get relatively faster as the bounds increase. The function slows down dramatically if $\text{primelimit} < \sqrt{b}$. It is comparable to **forfactored**, but about $\zeta(2) = \pi^2/6$ times faster due to the relative density of squarefree integers.

```
? B = 10^9;
? for (N = B, B+10^6, factor(N))
time = 2,463 ms.
? forfactored (N = B, B+10^6, [n,fan] = N)
time = 567 ms.
? forsquarefree (N = B, B+10^6, [n,fan] = N)
time = 343 ms.

? B = 10^11;
? for (N = B, B+10^6, factor(N))
time = 8,012 ms.
? forfactored (N = B, B+10^6, [n,fan] = N)
```

```

time = 1,293 ms.
? forsquarefree (N = B, B+10^6, [n,fan] = N)
time = 713 ms.

? B = 10^14;
? for (N = B, B+10^6, factor(N))
time = 41,283 ms.
? forsquarefree (N = B, B+10^6, [n,fan] = N)
time = 33,399 ms.

```

The last timing is with the default `primelimit` (500000) which is much less than $\sqrt{B+10^6}$; it goes down to 29,253ms if `primelimit` gets bigger than that bound. In any case $\sqrt{B+10^6}$ is much larger than the interval length 10^6 so `forsquarefree` gets relatively slower for that reason as well.

Note that all PARI multiplicative functions accept the `[n,fan]` argument natively:

```

? s = 0; forsquarefree(N = 1, 10^7, s += moebius(N)*eulerphi(N)); s
time = 2,003 ms.
%1 = 6393738650
? s = 0; for(N = 1, 10^7, s += moebius(N)*eulerphi(N)); s
time = 18,024 ms. \\ slower, we must factor N. Twice.
%2 = 6393738650

```

The following loops over the fundamental discriminants less than X :

```

? X = 10^8;
? for(d=1,X, if (isfundamental(d),))
time = 53,387 ms.
? forfactored(d=1,X, if (isfundamental(d),));
time = 13,861 ms.
? forsquarefree(d=1,X, D = quaddisc(d); if (D <= X, ));
time = 14,341 ms.

```

Note that in the last loop, the fundamental discriminants D are not evaluated in order (since `quaddisc(d)` for squarefree d is either d or $4d$) but the set of numbers we run through is the same. Not worth the complication since it's slower than testing `isfundamental`. A faster, more complicated approach uses two loops. For simplicity, assume X is divisible by 4:

```

? forsquarefree(d=1,X/4, D = quaddisc(d));
time = 3,642 ms.
? forsquarefree(d=X/4+1,X, if (d[1] % 4 == 1,));
time = 7,772 ms.

```

This is the price we pay for a faster evaluation,

We can run through negative fundamental discriminants in the same way:

```

? forfactored(d=-X,-1, if (isfundamental(d),));

```

3.1.19 forstep($X = a, b, s, seq$). Evaluates seq , where the formal variable X goes from a to b in increments of s . Nothing is done if $s > 0$ and $a > b$ or if $s < 0$ and $a < b$. The s can be

- a positive real number, preferably an integer: $X = a, a + s, a + 2s \dots$
- an intmod $\text{Mod}(c, N)$ (restrict to the corresponding arithmetic progression starting at the smallest integer $A \geq a$ and congruent to c modulo N): $X = A, A + N, \dots$
- a vector of steps $[s_1, \dots, s_n]$ (the successive steps in $\mathbf{R}^*$ are used in the order they appear in s): $X = a, a + s_1, a + s_1 + s_2, \dots$

```
? forstep(x=5, 10, 2, print(x))
5
7
9
? forstep(x=5, 10, Mod(1,3), print(x))
7
10
? forstep(x=5, 10, [1,2], print(x))
5
6
8
9
```

Setting b to $+\infty$ will start an infinite loop; it is expected that the caller will break out of the loop itself at some point, using **break** or **return**.

3.1.20 forsubgroup($H = G, \{bound\}, seq$). Evaluates seq for each subgroup H of the *abelian* group G (given in SNF form or as a vector of elementary divisors).

If $bound$ is present, and is a positive integer, restrict the output to subgroups of index less than $bound$. If $bound$ is a vector containing a single positive integer B , then only subgroups of index exactly equal to B are computed

The subgroups are not ordered in any obvious way, unless G is a p -group in which case Birkhoff's algorithm produces them by decreasing index. A subgroup is given as a matrix whose columns give its generators on the implicit generators of G . For example, the following prints all subgroups of index less than 2 in $G = \mathbf{Z}/2\mathbf{Z}g_1 \times \mathbf{Z}/2\mathbf{Z}g_2$:

```
? G = [2,2]; forsubgroup(H=G, 2, print(H))
[1; 1]
[1; 2]
[2; 1]
[1, 0; 1, 1]
```

The last one, for instance is generated by $(g_1, g_1 + g_2)$. This routine is intended to treat huge groups, when **subgrouplist** is not an option due to the sheer size of the output.

For maximal speed the subgroups have been left as produced by the algorithm. To print them in canonical form (as left divisors of G in HNF form), one can for instance use

```
? G = matdiagonal([2,2]); forsubgroup(H=G, 2, print(mathnf(concat(G,H))))
[2, 1; 0, 1]
[1, 0; 0, 2]
```

```
[2, 0; 0, 1]
[1, 0; 0, 1]
```

Note that in this last representation, the index $[G : H]$ is given by the determinant. See `galois-subcyclo` and `galoisfixedfield` for applications to Galois theory.

The library syntax is `for subgroup(void *data, long (*call)(void*,GEN), GEN G, GEN bound)`.

3.1.21 forsubset(nk, s, seq). If nk is an integer n , evaluates `seq`, where the formal variable s goes through all subsets of $\{1, 2, \dots, n\}$; if nk is a pair $[n, k]$ of integers, s goes through subsets of size k of $\{1, 2, \dots, n\}$. In both cases s goes through subsets in lexicographic order among subsets of the same size and smaller subsets come first.

```
? forsubset([5,3], s, print(s))
Vecsmall([1, 2, 3])
Vecsmall([1, 2, 4])
Vecsmall([1, 2, 5])
Vecsmall([1, 3, 4])
Vecsmall([1, 3, 5])
Vecsmall([1, 4, 5])
Vecsmall([2, 3, 4])
Vecsmall([2, 3, 5])
Vecsmall([2, 4, 5])
Vecsmall([3, 4, 5])

? forsubset(3, s, print(s))
Vecsmall([])
Vecsmall([1])
Vecsmall([2])
Vecsmall([3])
Vecsmall([1, 2])
Vecsmall([1, 3])
Vecsmall([2, 3])
Vecsmall([1, 2, 3])
```

The running time is proportional to the number of subsets enumerated, respectively 2^n and $\text{binomial}(n, k)$:

```
? c = 0; forsubset([40,35], s, c++); c
time = 128 ms.
%4 = 658008
? binomial(40,35)
%5 = 658008
```

The function is not useful for negative n or k , but these corner cases are supported nonetheless. When $n \leq 0$, we iterate through subsets of the empty set, i.e., the empty set; when $k < 0$ no subset of the required length exist and nothing happens.

3.1.22 forvec($X = v, seq, \{flag = 0\}$). Let v be an n -component vector (where n is arbitrary) of two-component vectors $[a_i, b_i]$ for $1 \leq i \leq n$, where all entries a_i, b_i are real numbers. This routine lets X vary over the n -dimensional box given by v with unit steps: X is an n -dimensional vector whose i -th entry $X[i]$ runs through $a_i, a_i + 1, a_i + 2, \dots$ stopping with the first value greater than b_i (note that neither a_i nor $b_i - a_i$ are required to be integers). The values of X are ordered lexicographically, like embedded **for** loops, and the expression seq is evaluated with the successive values of X . The type of X is the same as the type of v : **t_VEC** or **t_COL**.

If $flag = 1$, generate only nondecreasing vectors X , and if $flag = 2$, generate only strictly increasing vectors X .

```
? forvec (X=[[0,1],[-1,1]], print(X));
[0, -1]
[0, 0]
[0, 1]
[1, -1]
[1, 0]
[1, 1]
? forvec (X=[[0,1],[-1,1]], print(X), 1);
[0, 0]
[0, 1]
[1, 1]
? forvec (X=[[0,1],[-1,1]], print(X), 2)
[0, 1]
```

As a shortcut, a vector of the form $v = [[0, c_1 - 1], \dots, [0, c_n - 1]]$ can be abbreviated as $v = [c_1, \dots, c_n]$ and $flag$ is ignored in this case. More generally, if v is a vector of nonnegative integers c_i the loop runs over representatives of $\mathbf{Z}^n/v\mathbf{Z}^n$; and $flag$ is again ignored. The vector v may contain zero entries, in which case the loop spans an infinite lattice. The values are ordered lexicographically, graded by increasing L_1 -norm on free ($c_i = 0$) components.

This allows to iterate over elements of abelian groups using their **.cyc** vector.

```
? forvec (X=[2,3], print(X));
[0, 0]
[0, 1]
[0, 2]
[1, 0]
[1, 1]
[1, 2]
? my(i);forvec (X=[0,0], print(X); if (i++ > 10, break));
[0, 0]
[-1, 0]
[0, -1]
[0, 1]
[1, 0]
[-2, 0]
[-1, -1]
[-1, 1]
[0, -2]
[0, 2]
```

```

[1, -1]
? zn = znstar(36,1);
? forvec (chi = zn.cyc, if (chareval(zn,chi,5) == 5/6, print(chi)));
[1, 0]
[1, 1]
? bnrchar(zn, [5], [5/6]) \\ much more efficient in general
%5 = [[1, 1], [1, 0]]

```

3.1.23 `if(a, {seq1}, {seq2})`. Evaluates the expression sequence *seq1* if *a* is nonzero, otherwise the expression *seq2*. Of course, *seq1* or *seq2* may be empty:

`if (a, seq)` evaluates *seq* if *a* is not equal to zero (you don't have to write the second comma), and does nothing otherwise,

`if (a,, seq)` evaluates *seq* if *a* is equal to zero, and does nothing otherwise. You could get the same result using the `!` (not) operator: `if (!a, seq)`.

The value of an `if` statement is the value of the branch that gets evaluated: for instance

```
x = if(n % 4 == 1, y, z);
```

sets *x* to *y* if *n* is 1 modulo 4, and to *z* otherwise.

Successive 'else' blocks can be abbreviated in a single compound `if` as follows:

```

if (test1, seq1,
    test2, seq2,
    ...
    testn, seqn,
    seqdefault);

```

is equivalent to

```

if (test1, seq1
    , if (test2, seq2
        , ...
        if (testn, seqn, seqdefault)...));

```

For instance, this allows to write traditional switch / case constructions:

```

if (x == 0, do0(),
    x == 1, do1(),
    x == 2, do2(),
    dodefault());

```

Remark. The boolean operators `&&` and `||` are evaluated according to operator precedence as explained in Section 2.4, but, contrary to other operators, the evaluation of the arguments is stopped as soon as the final truth value has been determined. For instance

```
if (x != 0 && f(1/x), ...)
```

is a perfectly safe statement.

Remark. Functions such as `break` and `next` operate on *loops*, such as `forxxx`, `while`, `until`. The `if` statement is *not* a loop. (Obviously!)

3.1.24 `iferr(seq1, E, seq2, {pred})`. Evaluates the expression sequence `seq1`. If an error occurs, set the formal parameter `E` set to the error data. If `pred` is not present or evaluates to true, catch the error and evaluate `seq2`. Both `pred` and `seq2` can reference `E`. The error type is given by `errname(E)`, and other data can be accessed using the `component` function. The code `seq2` should check whether the error is the one expected. In the negative the error can be rethrown using `error(E)` (and possibly caught by an higher `iferr` instance). The following uses `iferr` to implement Lenstra's ECM factoring method

```
? ecm(N, B = 1000!, nb = 100)=
{
  for(a = 1, nb,
    iferr(ellmul(ellinit([a,1]*Mod(1,N)), [0,1]*Mod(1,N), B),
      E, return(gcd(lift(component(E,2)),N)),
      errname(E)=="e_INV" && type(component(E,2)) == "t_INTMOD"))
  }
? ecm(2^101-1)
%2 = 7432339208719
```

The return value of `iferr` itself is the value of `seq2` if an error occurs, and the value of `seq1` otherwise. We now describe the list of valid error types, and the attached error data `E`; in each case, we list in order the components of `E`, accessed via `component(E,1)`, `component(E,2)`, etc.

Internal errors, “system” errors.

- `"e_ARCH"`. A requested feature `s` is not available on this architecture or operating system. `E` has one component (`t_STR`): the missing feature name `s`.
- `"e_BUG"`. A bug in the PARI library, in function `s`. `E` has one component (`t_STR`): the function name `s`.
- `"e_FILE"`. Error while trying to open a file. `E` has two components, 1 (`t_STR`): the file type (input, output, etc.), 2 (`t_STR`): the file name.
- `"e_IMPL"`. A requested feature `s` is not implemented. `E` has one component, 1 (`t_STR`): the feature name `s`.
- `"e_PACKAGE"`. Missing optional package `s`. `E` has one component, 1 (`t_STR`): the package name `s`.

Syntax errors, type errors.

- **"e_DIM"**. The dimensions of arguments x and y submitted to function s does not match up. E.g., multiplying matrices of inconsistent dimension, adding vectors of different lengths, ... E has three components, 1 (**t_STR**): the function name s , 2: the argument x , 3: the argument y .

- **"e_FLAG"**. A flag argument is out of bounds in function s . E has one component, 1 (**t_STR**): the function name s .

- **"e_NOTFUNC"**. Generated by the PARI evaluator; tried to use a **GEN** x which is not a **t_CLOSURE** in a function call syntax (as in `f = 1; f(2);`). E has one component, 1: the offending **GEN** x .

- **"e_OP"**. Impossible operation between two objects than cannot be typecast to a sensible common domain for deeper reasons than a type mismatch, usually for arithmetic reasons. As in $0(2) + 0(3)$: it is valid to add two **t_PADICs**, provided the underlying prime is the same; so the addition is not forbidden a priori for type reasons, it only becomes so when inspecting the objects and trying to perform the operation. E has three components, 1 (**t_STR**): the operator name op , 2: first argument, 3: second argument.

- **"e_TYPE"**. An argument x of function s had an unexpected type. (As in `factor("blah")`.) E has two components, 1 (**t_STR**): the function name s , 2: the offending argument x .

- **"e_TYPE2"**. Forbidden operation between two objects than cannot be typecast to a sensible common domain, because their types do not match up. (As in `Mod(1,2) + Pi`.) E has three components, 1 (**t_STR**): the operator name op , 2: first argument, 3: second argument.

- **"e_PRIORITY"**. Object o in function s contains variables whose priority is incompatible with the expected operation. E.g. `Pol([x,1], 'y)`: this raises an error because it's not possible to create a polynomial whose coefficients involve variables with higher priority than the main variable. E has four components: 1 (**t_STR**): the function name s , 2: the offending argument o , 3 (**t_STR**): an operator op describing the priority error, 4 (**t_POL**): the variable v describing the priority error. The argument satisfies `variable(x) op variable(v)`.

- **"e_VAR"**. The variables of arguments x and y submitted to function s does not match up. E.g., considering the algebraic number `Mod(t,t^2+1)` in `nfini(x^2+1)`. E has three components, 1 (**t_STR**): the function name s , 2 (**t_POL**): the argument x , 3 (**t_POL**): the argument y .

Overflows.

- **"e_COMPONENT"**. Trying to access an inexistent component in a vector/matrix/list in a function: the index is less than 1 or greater than the allowed length. E has four components, 1 (**t_STR**): the function name, 2 (**t_STR**): an operator op ($<$ or $>$), 2 (**t_GEN**): a numerical limit l bounding the allowed range, 3 (**GEN**): the index x . It satisfies $x op l$.

- **"e_DOMAIN"**. An argument is not in the function's domain. E has five components, 1 (**t_STR**): the function name, 2 (**t_STR**): the mathematical name of the out-of-domain argument, 3 (**t_STR**): an operator op describing the domain error, 4 (**t_GEN**): the numerical limit l describing the domain error, 5 (**GEN**): the out-of-domain argument x . The argument satisfies $x op l$, which prevents it from belonging to the function's domain.

- **"e_MAXPRIME"**. A function using the precomputed list of prime numbers ran out of primes. E has one component, 1 (**t_INT**): the requested prime bound, which overflowed `primelimit` or 0 (bound is unknown).

- "e_MEM". A call to `pari_malloc` or `pari_realloc` failed. E has no component.
- "e_OVERFLOW". An object in function s becomes too large to be represented within PARI's hardcoded limits. (As in $2^{2^{2^{10}}}$ or `exp(1e100)`, which overflow in `lg` and `expo`.) E has one component, 1 (`t_STR`): the function name s .
- "e_PREC". Function s fails because input accuracy is too low. (As in `floor(1e100)` at default accuracy.) E has one component, 1 (`t_STR`): the function name s .
- "e_STACK". The PARI stack overflows. E has no component.

Errors triggered intentionally.

- "e_ALARM". A timeout, generated by the `alarm` function. E has one component (`t_STR`): the error message to print.
- "e_USER". A user error, as triggered by `error(g1, ..., gn)`. E has one component, 1 (`t_VEC`): the vector of n arguments given to `error`.

Mathematical errors.

- "e_CONSTPOL". An argument of function s is a constant polynomial, which does not make sense. (As in `galoisinit(Pol(1))`.) E has one component, 1 (`t_STR`): the function name s .
- "e_COPRIME". Function s expected coprime arguments, and did receive x, y , which were not. E has three component, 1 (`t_STR`): the function name s , 2: the argument x , 3: the argument y .
- "e_INV". Tried to invert a noninvertible object x in function s . E has two components, 1 (`t_STR`): the function name s , 2: the noninvertible x . If $x = \text{Mod}(a, b)$ is a `t_INTMOD` and a is not 0 mod b , this allows to factor the modulus, as `gcd(a, b)` is a nontrivial divisor of b .
- "e_IRREDPOL". Function s expected an irreducible polynomial, and did receive T , which was not. (As in `nfinit(x^2-1)`.) E has two component, 1 (`t_STR`): the function name s , 2 (`t_POL`): the polynomial x .
- "e_MISC". Generic uncategorized error. E has one component (`t_STR`): the error message to print.
- "e_MODULUS". moduli x and y submitted to function s are inconsistent. As in

$$\text{nfbalgtobasis}(\text{nfinit}(t^3-2), \text{Mod}(t, t^2+1))$$

E has three component, 1 (`t_STR`): the function s , 2: the argument x , 3: the argument y .

- "e_PRIME". Function s expected a prime number, and did receive p , which was not. (As in `idealprimedec(nf, 4)`.) E has two component, 1 (`t_STR`): the function name s , 2: the argument p .
- "e_ROOTS0". An argument of function s is a zero polynomial, and we need to consider its roots. (As in `polroots(0)`.) E has one component, 1 (`t_STR`): the function name s .
- "e_SQRTN". Trying to compute an n -th root of x , which does not exist, in function s . (As in `sqrt(Mod(-1,3))`.) E has two components, 1 (`t_STR`): the function name s , 2: the argument x .

3.1.25 next($\{n = 1\}$). Interrupts execution of current *seq*, resume the next iteration of the innermost enclosing loop, within the current function call (or top level loop). If *n* is specified, resume at the *n*-th enclosing loop. If *n* is bigger than the number of enclosing loops, all enclosing loops are exited.

3.1.26 return($\{x = 0\}$). Returns from current subroutine, with result *x*. If *x* is omitted, return the (void) value (return no result, like `print`).

3.1.27 setdebug($\{D\}, \{n\}$). Sets debug level for domain *D* to *n* ($0 \leq n \leq 20$). The domain *D* is a character string describing a Pari feature or code module, such as "bnf", "qflll" or "polgalois". This allows to selectively increase or decrease the diagnostics attached to a particular feature. If *n* is omitted, returns the current level for domain *D*. If *D* is omitted, returns a two-column matrix which lists the available domains with their levels. The `debug` default allows to reset all debug levels to a given value.

```
? setdebug()[,1] \\ list of all domains
["alg", "arith", "bern", "bnf", "bnr", "bnrclassfield", ..., "zetamult"]
? \g 1 \\ sets all debug levels to 1
  debug = 1
? setdebug("bnf", 0); \\ kills messages related to bnfinit and bnfisrincipal
```

3.1.28 until(*a*, *seq*). Evaluates *seq* until *a* is not equal to 0 (i.e. until *a* is true). If *a* is initially not equal to 0, *seq* is evaluated once (more generally, the condition on *a* is tested *after* execution of the *seq*, not before as in `while`).

3.1.29 while(*a*, *seq*). While *a* is nonzero, evaluates the expression sequence *seq*. The test is made *before* evaluating the *seq*, hence in particular if *a* is initially equal to zero the *seq* will not be evaluated at all.

3.2 Programming in GP: other specific functions.

In addition to the general PARI functions, it is necessary to have some functions which will be of use specifically for `gp`, though a few of these can be accessed under library mode. Before we start describing these, we recall the difference between *strings* and *keywords* (see Section 2.9.1): the latter don't get expanded at all, and you can type them without any enclosing quotes. The former are dynamic objects, where everything outside quotes gets immediately expanded.

3.2.1 Strchr(*x*). Deprecated alias for `strchr`.

The library syntax is GEN `pari_strchr`(GEN *x*).

3.2.2 Strexpand($\{x\}*$). Deprecated alias for `strexpand`

The library syntax is GEN `strexpand`(GEN *vec_x*).

3.2.3 Strprintf(*fmt*, $\{x\}*$). Deprecated alias for `strprintf`.

The library syntax is GEN `strprintf`(const char **fmt*, GEN *vec_x*).

3.2.4 Strtex($\{x\}^*$). Deprecatd alias for strtex.

The library syntax is `GEN strtex(GEN vec_x)`.

3.2.5 addhelp(*sym*, *str*). Changes the help message for the symbol **sym**. The string *str* is expanded on the spot and stored as the online help for **sym**. It is recommended to document global variables and user functions in this way, although **gp** will not protest if you don't.

You can attach a help text to an alias, but it will never be shown: aliases are expanded by the `? help` operator and we get the help of the symbol the alias points to. Nothing prevents you from modifying the help of built-in PARI functions. But if you do, we would like to hear why you needed it!

Without `addhelp`, the standard help for user functions consists of its name and definition.

```
gp> f(x) = x^2;
gp> ?f
f =
      (x)->x^2
```

Once `addhelp` is applied to f , the function code is no longer included. It can still be consulted by typing the function name:

```
gp> addhelp(f, "Square")
gp> ?f
Square

gp> f
%2 = (x)->x^2
```

The library syntax is `void addhelp(const char *sym, const char *str)`.

3.2.6 alarm($\{s = 0\}, \{code\}$). If *code* is omitted, trigger an *e_ALARM* exception after *s* seconds (wall-clock time), cancelling any previously set alarm; stop a pending alarm if *s* = 0 or is omitted.

Otherwise, if *s* is positive, the function evaluates *code*, aborting after *s* seconds. The return value is the value of *code* if it ran to completion before the alarm timeout, and a `t_ERROR` object otherwise.

[illegible]

Here is a more involved example: the function `timefact(N,sec)` below tries to factor N and gives up after *sec* seconds, returning a partial factorization.

\\ Time-bounded partial factorization

```
default(factor_add_primes,1);
timefact(N,sec)=
{
    F = alarm(sec, factor(N));
    if (type(F) == "t_ERROR", factor(N, 2^24), F);
}
```

We either return the factorization directly, or replace the `t_ERROR` result by a simple bounded factorization `factor(N, 224)`. Note the `factor_add_primes` trick: any prime larger than 2^{24} discovered while attempting the initial factorization is stored and remembered. When the alarm rings, the subsequent bounded factorization finds it right away.

Caveat. It is not possible to set a new alarm *within* another `alarm` code: the new timer erases the parent one.

Caveat2. In a parallel-enabled `gp`, if the *code* involves parallel subtasks, then `alarm` may not return right away: it will prevent new tasks from being launched but will not interrupt previously launched secondary threads. This avoids leaving the system in an inconsistent state.

The library syntax is `GEN gp_alarm(long s, GEN code = NULL)`.

3.2.7 alias(*newsym*, *sym*). Defines the symbol *newsym* as an alias for the symbol *sym*:

```
? alias("det", "matdet");
? det([1,2;3,4])
%1 = -2
```

You are not restricted to ordinary functions, as in the above example: to alias (from/to) member functions, prefix them with ‘`_.`’; to alias operators, use their internal name, obtained by writing `_` in lieu of the operators argument: for instance, `_!` and `!_` are the internal names of the factorial and the logical negation, respectively.

```
? alias("mod", "_mod");
? alias("add", "_+_");
? alias("_sin", "sin");
? mod(Mod(x,x^4+1))
%2 = x^4 + 1
? add(4,6)
%3 = 10
? Pi.sin
%4 = 0.E-37
```

Alias expansion is performed directly by the internal GP compiler. Note that since alias is performed at compilation-time, it does not require any run-time processing, however it only affects GP code compiled *after* the alias command is evaluated. A slower but more flexible alternative is to use variables. Compare

[illegible]

Note that changing either `parisize` or `parisizemax` at the break loop prompt would interrupt the computation, contrary to the above.

In most cases, `parisize` will increase automatically (up to `parisizemax`) and there is no need to perform the above maneuvers. But that the garbage collector is sufficiently efficient that a given computation can still run without increasing the stack size, albeit very slowly due to the frequent garbage collections.

Deprecated: when `parisizemax` is unset. This is currently still the default behavior in order not to break backward compatibility. The rest of this section documents the behavior of `allocatemem` in that (deprecated) situation: it becomes a synonym for `default(parisize,...)`. In that case, there is no notion of a virtual stack, and the stack size is always equal to `parisize`. If more memory is needed, the PARI stack overflows, aborting the computation.

Thus, increasing `parisize` via `allocatemem` or `default(parisize,...)` before a big computation is important. Unfortunately, either must be typed at the `gp` prompt in interactive usage, or left by itself at the start of batch files. They cannot be used meaningfully in loop-like constructs, or as part of a larger expression sequence, e.g

```
allocatemem(); x = 1;    \\ This will not set x!
```

In fact, all loops are immediately exited, user functions terminated, and the rest of the sequence following `allocatemem()` is silently discarded, as well as all pending sequences of instructions. We just go on reading the next instruction sequence from the file we are in (or from the user). In particular, we have the following possibly unexpected behavior: in

```
read("file.gp"); x = 1
```

were `file.gp` contains an `allocatemem` statement, the `x = 1` is never executed, since all pending instructions in the current sequence are discarded.

The reason for these unfortunate side-effects is that, with `parisizemax` disabled, increasing the stack size physically moves the stack, so temporary objects created during the current expression evaluation are not correct anymore. (In particular byte-compiled expressions, which are allocated on the stack.) To avoid accessing obsolete pointers to the old stack, this routine ends by a `longjmp`.

The library syntax is `void gp_allocatemem(GEN s = NULL)`.

3.2.9 `apply(f, A)`. Apply the `t_CLOSURE` `f` to the entries of `A`.

- If `A` is a scalar, return `f(A)`.
- If `A` is a polynomial or power series $\sum a_i x^i (+O(x^N))$, apply `f` on all coefficients and return $\sum f(a_i) x^i (+O(x^N))$.
- If `A` is a vector or list $[a_1, \dots, a_n]$, return the vector or list $[f(a_1), \dots, f(a_n)]$. If `A` is a matrix, return the matrix whose entries are the $f(A[i, j])$.

```
? apply(x->x^2, [1,2,3,4])
%1 = [1, 4, 9, 16]
? apply(x->x^2, [1,2;3,4])
%2 =
[1 4]
[9 16]
? apply(x->x^2, 4*x^2 + 3*x + 2)
```

```
%3 = 16*x^2 + 9*x + 4
? apply(sign, 2 - 3* x + 4*x^2 + 0(x^3))
%4 = 1 - x + x^2 + 0(x^3)
```

Note that many functions already act componentwise on vectors or matrices, but they almost never act on lists; in this case, `apply` is a good solution:

```
? L = List([Mod(1,3), Mod(2,4)]);
? lift(L)
*** at top-level: lift(L)
***      ^-----
*** lift: incorrect type in lift.
? apply(lift, L);
%2 = List([1, 2])
```

Remark. For v a `t_VEC`, `t_COL`, `t_VECSMALL`, `t_LIST` or `t_MAT`, the alternative set-notations

```
[g(x) | x <- v, f(x)]
[x | x <- v, f(x)]
[g(x) | x <- v]
```

are available as shortcuts for

```
apply(g, select(f, Vec(v)))
select(f, Vec(v))
apply(g, Vec(v))
```

respectively:

```
? L = List([Mod(1,3), Mod(2,4)]);
? [ lift(x) | x<-L ]
%2 = [1, 2]
```

The library syntax is `genapply(void *E, GEN (*fun)(void*,GEN), GEN a)`.

3.2.10 `arity(C)`. Return the arity of the closure C , i.e., the number of its arguments.

```
? f1(x,y=0)=x+y;
? arity(f1)
%1 = 2
? f2(t,s[.])=print(t,":",s);
? arity(f2)
%2 = 2
```

Note that a variadic argument, such as s in `f2` above, is counted as a single argument.

The library syntax is `GEN arity0(GEN C)`.

3.2.11 call(f, A). $A = [a_1, \dots, a_n]$ being a vector and f being a function, returns the evaluation of $f(a_1, \dots, a_n)$. f can also be the name of a built-in GP function. If $\#A = 1$, $\text{call}(f, A) = \text{apply}(f, A)[1]$. If f is variadic (has a variable number of arguments), then the variadic arguments are grouped in a vector in the last component of A .

This function is useful

- when writing a variadic function, to call another one:

```
fprintf(file,format,args[..]) = write(file, call(strprintf,[format,args]))
```

- when dealing with function arguments with unspecified arity.

The function below implements a global memoization interface:

```
memo=Map();
memoize(f,A[..])=
{
  my(res);
  if(!mapisdefined(memo, [f,A], &res),
    res = call(f,A);
    mapput(memo,[f,A],res));
  res;
}
```

for example:

```
? memoize(factor,2^128+1)
%3 = [59649589127497217,1;5704689200685129054721,1]
? ##
*** last result computed in 76 ms.
? memoize(factor,2^128+1)
%4 = [59649589127497217,1;5704689200685129054721,1]
? ##
*** last result computed in 0 ms.
? memoize(ffinit,3,3)
%5 = Mod(1,3)*x^3+Mod(1,3)*x^2+Mod(1,3)*x+Mod(2,3)
? fibo(n)=if(n==0,0,n==1,1,memoize(fibo,n-2)+memoize(fibo,n-1));
? fibo(100)
%7 = 354224848179261915075
```

- to call operators through their internal names without using alias

```
matnbelts(M) = call("_*_",matsize(M))
```

The library syntax is GEN call0(GEN f, GEN A).

3.2.12 default($\{key\}, \{val\}$). Returns the default corresponding to keyword key . If val is present, sets the default to val first (which is subject to string expansion first). Typing `default()` (or `\d`) yields the complete default list as well as their current values. See Section 2.12 for an introduction to GP defaults, Section 3.4 for a list of available defaults, and Section 2.13 for some shortcut alternatives. Note that the shortcuts are meant for interactive use and usually display more information than `default`.

The library syntax is GEN default0(const char *key = NULL, const char *val = NULL)

3.2.13 `errname(E)`. Returns the type of the error message `E` as a string.

```
? iferr(1 / 0, E, print(errname(E)))
e_INV
? ?? e_INV
[...]
* "e_INV". Tried to invert a noninvertible object x in function s.
[...]
```

The library syntax is `GEN errname(GEN E)`.

3.2.14 `error({str}*)`. Outputs its argument list (each of them interpreted as a string), then interrupts the running `gp` program, returning to the input prompt. For instance

```
error("n = ", n, " is not squarefree!")
```

The library syntax is `void error0(GEN vec_str)`.

The variadic version `void pari_err(e_USER, ...)` is usually preferable.

3.2.15 `export(x{= ...},...,z{= ...})`. Export the variables $x, \dots, z$ to the parallel world. Such variables are visible inside parallel sections in place of global variables, but cannot be modified inside a parallel section. `export(a)` set the variable a in the parallel world to current value of a . `export(a=z)` set the variable a in the parallel world to z , without affecting the current value of a .

```
? fun(x)=x^2+1;
? parvector(10,i,fun(i))
*** mt: please use export(fun).
? export(fun)
? parvector(10,i,fun(i))
%4 = [2,5,10,17,26,37,50,65,82,101]
```

3.2.16 `exportall()`. Declare all current dynamic variables as exported variables. Such variables are visible inside parallel sections in place of global variables.

```
? fun(x)=x^2+1;
? parvector(10,i,fun(i))
*** mt: please use export(fun).
? exportall()
? parvector(10,i,fun(i))
%4 = [2,5,10,17,26,37,50,65,82,101]
```

The library syntax is `void exportall()`.

3.2.17 `extern(str)`. The string `str` is the name of an external command (i.e. one you would type from your UNIX shell prompt). This command is immediately run and its output fed into `gp`, just as if read from a file.

The library syntax is `GEN gpextern(const char *str)`.

3.2.18 externstr(*str*). The string *str* is the name of an external command (i.e. one you would type from your UNIX shell prompt). This command is immediately run and its output is returned as a vector of GP strings, one component per output line.

The library syntax is GEN externstr(const char *str).

3.2.19 fileclose(*n*). Close the file descriptor *n*, created via fileopen or fileextern. Finitely many files can be opened at a given time, closing them recycles file descriptors and avoids running out of them:

```
? n = 0; while(n++, fileopen("/tmp/test", "w"))
***   at top-level: n=0;while(n++,fileopen("/tmp/test","w"))
***                                     ^-----
*** fileopen: error opening requested file: '/tmp/test'.
***   Break loop: type 'break' to go back to GP prompt
break> n
65533
```

This is a limitation of the operating system and does not depend on PARI: if you open too many files in gp without closing them, the operating system will also prevent unrelated applications from opening files. Independently, your operating system (e.g. Windows) may prevent other applications from accessing or deleting your file while it is opened by gp. Quitting gp implicitly calls this function on all opened file descriptors.

On files opened for writing, this function also forces a write of all buffered data to the file system and completes all pending write operations. This function is implicitly called for all open file descriptors when exiting gp but it is cleaner and safer to call it explicitly, for instance in case of a gp crash or general system failure, which could cause data loss.

```
? n = fileopen("./here");
? while(1 = fileread(n), print(1));
? fileclose(n);

? n = fileopen("./there", "w");
? for (i = 1, 100, filewrite(n, i^2+1))
? fileclose(n)
```

Until a fileclose, there is no guarantee that the file on disk contains all the expected data from previous filewrites. (And even then the operating system may delay the actual write to hardware.)

Closing a file twice raises an exception:

```
? n = fileopen("/tmp/test");
? fileclose(n)
? fileclose(n)
***   at top-level: fileclose(n)
***                                     ^-----
*** fileclose: invalid file descriptor 0
```

The library syntax is void gp_fileclose(long n).

3.2.20 fileextern(*str*). The string *str* is the name of an external command, i.e. one you would type from your UNIX shell prompt. This command is immediately run and the function returns a file descriptor attached to the command output as if it were read from a file.

```
? n = fileextern("ls -l");
? while(1 = filereadstr(n), print(1))
? fileclose(n)
```

If the `secure` default is set, this function will raise an exception.

The library syntax is `long gp_fileextern(const char *str)`.

3.2.21 fflush(*{n}*). Flushes the file descriptor *n*, created via `fileopen` or `fileextern`. On files opened for writing, this function forces a write of all buffered data to the file system and completes all pending write operations. This function is implicitly called by `fileclose` but you may want to call it explicitly at synchronization points, for instance after writing a large result to file and before printing diagnostics on screen. (In order to be sure that the file contains the expected content on inspection.)

If *n* is omitted, flush all descriptors to output streams.

```
? n = fileopen("./here", "w");
? for (i = 1, 10^5, \
      filewrite(n, i^2+1); \
      if (i % 10000 == 0, fflush(n)))
```

Until a `fflush` or `fileclose`, there is no guarantee that the file contains all the expected data from previous `filewrites`.

The library syntax is `void gp_fflush0(GEN n = NULL)`. But the direct and more specific variant `void gp_fflush(long n)` is also available.

3.2.22 fileopen(*path, mode*). Open the file pointed to by '*path*' and return a file descriptor which can be used with other file functions.

The mode can be

- "`r`": (default): open for reading; allow `fileread` and `filereadstr`.
- "`w`": open for writing, discarding existing content; allow `filewrite`, `filewrite1`.
- "`a`": open for writing, appending to existing content; same operations allowed as "`w`".

Eventually, the file should be closed and the descriptor recycled using `fileclose`.

```
? n = fileopen("./here"); \\ "r" by default
? while (1 = filereadstr(n), print(1)) \\ print successive lines
? fileclose(n) \\ done
```

In *read* mode, raise an exception if the file does not exist or the user does not have read permission. In *write* mode, raise an exception if the file cannot be written to. Trying to read or write to a file that was not opened with the right mode raises an exception.

```
? n = fileopen("./read", "r");
? filewrite(n, "test") \\ not open for writing
*** at top-level: filewrite(n,"test")
*** ^-----
```

*** fwrite: invalid file descriptor 0

The library syntax is `long gp_fileopen(const char *path, const char *mode)`.

3.2.23 fileread(*n*). Read a logical line from the file attached to the descriptor *n*, opened for reading with `fileopen`. Return 0 at end of file.

A logical line is a full command as it is prepared by gp's preprocessor (skipping blanks and comments or assembling multiline commands between braces) before being fed to the interpreter. The function `filereadstr` would read a *raw* line exactly as input, up to the next carriage return `\n`.

Compare raw lines

```
? n = fileopen("examples/bench.gp");
? while(1 = filereadstr(n), print(1));
{
  u=v=p=q=1;
  for (k=1, 2000,
    [u,v] = [v,u+v];
    p *= v; q = lcm(q,v);
    if (k%50 == 0,
      print(k, " ", log(p)/log(q))
    )
  )
}
```

and logical lines

```
? n = fileopen("examples/bench.gp");
? while(1 = fileread(n), print(1));
u=v=p=q=1;for(k=1,2000,[u,v]=[v,u+v];p*=v;q=lcm(q,v);[...])
```

The library syntax is GEN `gp_fileread(long n)`.

3.2.24 filereadstr(*n*). Read a raw line from the file attached to the descriptor *n*, opened for reading with `fileopen`, discarding the terminating newline. In other words the line is read exactly as input, up to the next carriage return `\n`. By comparison, `fileread` would read a logical line, as assembled by gp's preprocessor (skipping blanks and comments for instance).

The library syntax is GEN `gp_filereadstr(long n)`.

3.2.25 fwrite(*n*, *s*). Write the string *s* to the file attached to descriptor *n*, ending with a newline. The file must have been opened with `fileopen` in "w" or "a" mode. There is no guarantee that *s* is completely written to disk until `fileclose(n)` is executed, which is automatic when quitting gp.

If the newline is not desired, use `fwrite1`.

Variant. The high-level function `write` is expensive when many consecutive writes are expected because it cannot use buffering. The low-level interface `fileopen` / `filewrite` / `fileclose` is more efficient:

```
? f = "/tmp/bigfile";
? for (i = 1, 10^5, write(f, i^2+1))
time = 240 ms.

? v = vector(10^5, i, i^2+1);
time = 10 ms. \\ computing the values is fast
? write("/tmp/bigfile2",v)
time = 12 ms. \\ writing them in one operation is fast

? n = fileopen("/tmp/bigfile", "w");
? for (i = 1, 10^5, filewrite(n, i^2+1))
time = 24 ms. \\ low-level write is ten times faster
? fileclose(n);
```

In the final example, the file needs not be in a consistent state until the ending `fileclose` is evaluated, e.g. some lines might be half-written or not present at all even though the corresponding `filewrite` was executed already. Both a single high-level `write` and a succession of low-level `filewrites` achieve the same efficiency, but the latter is often more natural. In fact, concatenating naively the entries to be written is quadratic in the number of entries, hence much more expensive than the original write operations:

```
? v = []; for (i = 1, 10^5, v = concat(v,i))
time = 1min, 41,456 ms.
```

The library syntax is `void gp_filewrite(long n, const char *s)`.

3.2.26 filewrite1(*n*, *s*). Write the string *s* to the file attached to descriptor *n*. The file must have been opened with `fileopen` in "w" or "a" mode.

If you want to append a newline at the end of *s*, you can use `Str(s, "\n")` or `filewrite`.

The library syntax is `void gp_filewrite1(long n, const char *s)`.

3.2.27 fold(*f*, *A*). Apply the `t_CLOSURE` *f* of arity 2 to the entries of *A*, in order to return `f(...f(f(A[1],A[2]),A[3])... ,A[#A])`.

```
? fold((x,y)->x*y, [1,2,3,4])
%1 = 24
? fold((x,y)->[x,y], [1,2,3,4])
%2 = [[1, 2], 3], 4]
? fold((x,f)->f(x), [2,sqr,sqr,sqr])
%3 = 256
? fold((x,y)->(x+y)/(1-x*y), [1..5])
%4 = -9/19
? bestappr(tan(sum(i=1,5,atan(i))))
%5 = -9/19
```

The library syntax is `GEN fold0(GEN f, GEN A)`. Also available is `GEN genfold(void *E, GEN (*fun)(void*, GEN, GEN), GEN A)`.

3.2.28 getabstime(). Returns the CPU time (in milliseconds) elapsed since `gp` startup. This provides a reentrant version of `gettime`:

```
my (t = getabstime());
...
print("Time: ", strtime(getabstime() - t));
```

For a version giving wall-clock time, see `getwalltime`.

The library syntax is `long getabstime()`.

3.2.29 getcache(). Returns information about various auto-growing caches. For each resource, we report its name, its size, the number of cache misses (since the last extension), the largest cache miss and the size of the cache in bytes.

The caches are initially empty, then set automatically to a small inexpensive default value, then grow on demand up to some maximal value. Their size never decreases, they are only freed on exit.

The current caches are

- Hurwitz class numbers $H(D)$ for $|D| \leq N$, computed in time $O(N^{3/2})$ using $O(N)$ space.
- Factorizations of small integers up to N , computed in time $O(N^{1+\varepsilon})$ using $O(N \log N)$ space.
- Divisors of small integers up to N , computed in time $O(N^{1+\varepsilon})$ using $O(N \log N)$ space.
- Coredisc's of negative integers down to $-N$, computed in time $O(N^{1+\varepsilon})$ using $O(N)$ space.
- Primitive dihedral forms of weight 1 and level up to N , computed in time $O(N^{2+\varepsilon})$ and space $O(N^2)$.

```
? getcache()  \\ on startup, all caches are empty
%1 =
[ "Factors" 0 0 0 0]
[ "Divisors" 0 0 0 0]
[      "H" 0 0 0 0]
["CorediscF" 0 0 0 0]
[ "Dihedral" 0 0 0 0]
? mfdim([500,1,0],0); \\ nontrivial computation
time = 540 ms.
? getcache()
%3 =
[ "Factors" 50000 0      0 4479272]
[ "Divisors" 50000 1 100000 5189808]
[      "H" 50000 0      0 400008]
[ "Dihedral" 1000 0      0 2278208]
```

The library syntax is `GEN getcache()`.

3.2.30 getenv(s). Return the value of the environment variable `s` if it is defined, otherwise return 0.

The library syntax is `GEN gp_getenv(const char *s)`.

3.2.31 getheap(). Returns a two-component row vector giving the number of objects on the heap and the amount of memory they occupy in long words. Useful mainly for debugging purposes.

The library syntax is GEN `getheap()`.

3.2.32 getlocalbitprec(). Returns the current dynamic bit precision.

3.2.33 getlocalprec(). Returns the current dynamic precision, in decimal digits.

3.2.34 getrand(). Returns the current value of the seed used by the pseudo-random number generator `random`. Useful mainly for debugging purposes, to reproduce a specific chain of computations. The returned value is technical (reproduces an internal state array), and can only be used as an argument to `setrand`.

The library syntax is GEN `getrand()`.

3.2.35 getstack(). Returns the current value of `top - avma`, i.e. the number of bytes used up to now on the stack. Useful mainly for debugging purposes.

The library syntax is long `getstack()`.

3.2.36 gettime(). Returns the CPU time (in milliseconds) used since either the last call to `gettime`, or to the beginning of the containing GP instruction (if inside `gp`), whichever came last.

For a reentrant version, see `getabstime`.

For a version giving wall-clock time, see `getwalltime`.

The library syntax is long `gettime()`.

3.2.37 getwalltime(). Returns the time (in milliseconds) elapsed since 00:00:00 UTC Thursday 1, January 1970 (the Unix epoch).

```
my (t = getwalltime());  
...  
print("Time: ", strtime(getwalltime() - t));
```

The library syntax is GEN `getwalltime()`.

3.2.38 global(*list of variables*). Obsolete. Scheduled for deletion.

3.2.39 inline(*x, ..., z*). Declare *x, ..., z* as inline variables. Such variables behave like lexically scoped variable (see `my()`) but with unlimited scope. It is however possible to exit the scope by using `uninline()`. When used in a GP script, it is recommended to call `uninline()` before the script's end to avoid inline variables leaking outside the script. DEPRECATED, use `export`.

3.2.40 input(). Reads a string, interpreted as a GP expression, from the input file, usually standard input (i.e. the keyboard). If a sequence of expressions is given, the result is the result of the last expression of the sequence. When using this instruction, it is useful to prompt for the string by using the `print1` function. Note that in the present version 2.19 of `pari.el`, when using `gp` under GNU Emacs (see Section 2.16) one *must* prompt for the string, with a string which ends with the same prompt as any of the previous ones (a "`?` " will do for instance).

The library syntax is GEN `gp_input()`.

3.2.41 install(*name*, *code*, {*gpname*}, {*lib*}). Loads from dynamic library *lib* the function *name*. Assigns to it the name *gpname* in this **gp** session, with *prototype code* (see below). If *gpname* is omitted, uses *name*. If *lib* is omitted, all symbols known to **gp** are available: this includes the whole of **libpari.so** and possibly others (such as **libc.so**).

Most importantly, **install** gives you access to all nonstatic functions defined in the PARI library. For instance, the function

```
GEN addii(GEN x, GEN y)
```

adds two PARI integers, and is not directly accessible under **gp** (it is eventually called by the **+** operator of course):

```
? install("addii", "GG")
? addii(1, 2)
%1 = 3
```

It also allows to add external functions to the **gp** interpreter. For instance, it makes the function **system** obsolete:

```
? install(system, vs, sys,/*omitted*/)
? sys("ls gp*")
gp.c          gp.h          gp_rl.c
```

This works because **system** is part of **libc.so**, which is linked to **gp**. It is also possible to compile a shared library yourself and provide it to **gp** in this way: use **gp2c**, or do it manually (see the **modules.build** variable in **pari.cfg** for hints).

Re-installing a function will print a warning and update the prototype code if needed. However, it will not reload a symbol from the library, even if the latter has been recompiled.

Prototype. We only give a simplified description here, covering most functions, but there are many more possibilities. The full documentation is available in **libpari.dvi**, see

??prototype

- First character *i*, *l*, *u*, *v* : return type **int** / **long** / **ulong** / **void**. (Default: **GEN**)
- One letter for each mandatory argument, in the same order as they appear in the argument list: **G** (**GEN**), **&** (**GEN***), **L** (**long**), **U** (**ulong**), **s** (**char ***), **n** (variable).
- **p** to supply **realprecision** (usually **long prec** in the argument list), **b** to supply **realbit-precision** (usually **long bitprec**), **P** to supply **seriesprecision** (usually **long precdl**).

We also have special constructs for optional arguments and default values:

- **DG** (optional **GEN**, **NULL** if omitted),
- **D&** (optional **GEN***, **NULL** if omitted),
- **Dn** (optional variable, **-1** if omitted),

For instance the prototype corresponding to

```
long issquareall(GEN x, GEN *n = NULL)
```

is **lGD&**.

Caution. This function may not work on all systems, especially when `gp` has been compiled statically. In that case, the first use of an installed function will provoke a Segmentation Fault (this should never happen with a dynamically linked executable). If you intend to use this function, please check first on some harmless example such as the one above that it works properly on your machine.

The library syntax is `void gpinstall(const char *name, const char *code, const char *gpname, const char *lib)`.

3.2.42 `kill(sym)`. Restores the symbol `sym` to its “undefined” status, and deletes any help messages attached to `sym` using `addhelp`. Variable names remain known to the interpreter and keep their former priority: you cannot make a variable “less important” by killing it!

```
? z = y = 1; y
%1 = 1
? kill(y)
? y          \\ restored to ‘‘undefined’’ status
%2 = y
? variable()
%3 = [x, y, z] \\ but the variable name y is still known, with y > z !
```

For the same reason, killing a user function (which is an ordinary variable holding a `t_CLOSURE`) does not remove its name from the list of variable names.

If the symbol is attached to a variable — user functions being an important special case —, one may use the quote operator `a = 'a` to reset variables to their starting values. However, this will not delete a help message attached to `a`, and is also slightly slower than `kill(a)`.

```
? x = 1; addhelp(x, "foo"); x
%1 = 1
? x = 'x; x    \\ same as 'kill', except we don't delete help.
%2 = x
? ?x
foo
```

On the other hand, `kill` is the only way to remove aliases and installed functions.

```
? alias(fun, sin);
? kill(fun);

? install(addii, GG);
? kill(addii);
```

The library syntax is `void kill0(const char *sym)`.

3.2.43 `listcreate({n})`. This function is obsolete, use `List`.

Creates an empty list. This routine used to have a mandatory argument, which is now ignored (for backward compatibility).

3.2.44 listinsert($\sim L, x, n$). Inserts the object x at position n in L (which must be of type `t_LIST`). This has complexity $O(\#L - n + 1)$: all the remaining elements of $list$ (from position $n + 1$ onwards) are shifted to the right. If n is greater than the list length, appends x .

```
? L = List([1,2,3]);
? listput(~L, 4); L \\ listput inserts at end
%4 = List([1, 2, 3, 4])
? listinsert(~L, 5, 1); L \\insert at position 1
%5 = List([5, 1, 2, 3, 4])
? listinsert(~L, 6, 1000); L \\ trying to insert beyond position #L
%6 = List([5, 1, 2, 3, 4, 6]) \\ ... inserts at the end
```

Note the $\sim L$: this means that the function is called with a *reference* to L and changes L in place.

The library syntax is `GEN listinsert0(GEN ~L, GEN x, long n)`.

3.2.45 listkill($\sim L$). Obsolete, retained for backward compatibility. Just use `L = List()` instead of `listkill(L)`. In most cases, you won't even need that, e.g. local variables are automatically cleared when a user function returns.

The library syntax is `void listkill(GEN ~L)`.

3.2.46 listpop($\sim list, \{n\}$). Removes the n -th element of the list $list$ (which must be of type `t_LIST`). If n is omitted, or greater than the list current length, removes the last element. If the list is already empty, do nothing. This runs in time $O(\#L - n + 1)$.

```
? L = List([1,2,3,4]);
? listpop(~L); L \\ remove last entry
%2 = List([1, 2, 3])
? listpop(~L, 1); L \\ remove first entry
%3 = List([2, 3])
```

Note the $\sim L$: this means that the function is called with a *reference* to L and changes L in place.

The library syntax is `void listpop0(GEN ~list, long n)`.

3.2.47 listput($\sim list, x, \{n\}$). Sets the n -th element of the list $list$ (which must be of type `t_LIST`) equal to x . If n is omitted, or greater than the list length, appends x .

```
? L = List();
? listput(~L, 1)
? listput(~L, 2)
? L
%4 = List([1, 2])
```

Note the $\sim L$: this means that the function is called with a *reference* to L and changes L in place.

You may put an element into an occupied cell (not changing the list length), but it is easier to use the standard `list[n] = x` construct.

```
? listput(~L, 3, 1) \\ insert at position 1
? L
%6 = List([3, 2])
? L[2] = 4 \\ simpler
%7 = List([3, 4])
```

```

? L[10] = 1  \\ can't insert beyond the end of the list
***   at top-level: L[10]=1
***               ^-----
***   nonexistent component: index > 2
? listput(L, 1, 10) \\ but listput can
? L
%9 = List([3, 2, 1])

```

This function runs in time $O(\#L)$ in the worst case (when the list must be reallocated), but in time $O(1)$ on average: any number of successive `listputs` run in time $O(\#L)$, where $\#L$ denotes the list *final* length.

The library syntax is `GEN listput0(GEN ~list, GEN x, long n)`.

3.2.48 `listsort`($\sim L, \{flag = 0\}$). Sorts the `t_LIST` *list* in place, with respect to the (somewhat arbitrary) universal comparison function `cmp`. In particular, the ordering is the same as for sets and `setsearch` can be used on a sorted list. No value is returned. If *flag* is nonzero, suppresses all repeated coefficients.

```

? L = List([1,2,4,1,3,-1]); listsort(~L); L
%1 = List([-1, 1, 1, 2, 3, 4])
? setsearch(L, 4)
%2 = 6
? setsearch(L, -2)
%3 = 0
? listsort(~L, 1); L \\ remove duplicates
%4 = List([-1, 1, 2, 3, 4])

```

Note the `~L`: this means that the function is called with a *reference* to `L` and changes `L` in place: this is faster than the `vecsrt` command since the list is sorted in place and we avoid unnecessary copies.

```

? v = vector(100,i,random); L = List(v);
? for(i=1,10^4, vecsort(v))
time = 162 ms.
? for(i=1,10^4, vecsort(L))
time = 162 ms.
? for(i=1,10^4, listsort(~L))
time = 63 ms.

```

The library syntax is `void listsort(GEN ~L, long flag)`.

3.2.49 `localbitprec`(*p*). Set the real precision to *p* bits in the dynamic scope. All computations are performed as if `realbitprecision` was *p*: transcendental constants (e.g. `Pi`) and conversions from exact to floating point inexact data use *p* bits, as well as iterative routines implicitly using a floating point accuracy as a termination criterion (e.g. `solve` or `intnum`). But `realbitprecision` itself is unaffected and is “unmasked” when we exit the dynamic (*not* lexical) scope. In effect, this is similar to

```

my(bit = default(realbitprecision));
default(realbitprecision,p);
...

```

```
default(realbitprecision, bit);
```

but is both less cumbersome, cleaner (no need to manipulate a global variable, which in fact never changes and is only temporarily masked) and more robust: if the above computation is interrupted or an exception occurs, `realbitprecision` will not be restored as intended.

Such `localbitprec` statements can be nested, the innermost one taking precedence as expected. Beware that `localbitprec` follows the semantic of `local`, not `my`: a subroutine called from `localbitprec` scope uses the local accuracy:

```
? f()=bitprecision(1.0);
? f()
%2 = 128
? localbitprec(1000); f()
%3 = 1024
```

Note that the bit precision of *data* (1.0 in the above example) increases by steps of 64 (32 on a 32-bit machine) so we get 1024 instead of the expected 1000; `localbitprec` bounds the relative error exactly as specified in functions that support that granularity (e.g. `lfun`), and rounded to the next multiple of 64 (resp. 32) everywhere else.

Warning. Changing `realbitprecision` or `realprecision` in programs is deprecated in favor of `localbitprec` and `localprec`. Think about the `realprecision` and `realbitprecision` defaults as interactive commands for the `gp` interpreter, best left out of GP programs. Indeed, the above rules imply that mixing both constructs yields surprising results:

```
? \p38
? localprec(19); default(realprecision,1000); Pi
%1 = 3.141592653589793239
? \p
realprecision = 1001 significant digits (1000 digits displayed)
```

Indeed, `realprecision` itself is ignored within `localprec` scope, so `Pi` is computed to a low accuracy. And when we leave the `localprec` scope, `realprecision` only regains precedence, it is not “restored” to the original value.

3.2.50 `localprec(p)`. Set the real precision to p in the dynamic scope. All computations are performed as if `realprecision` was p : transcendental constants (e.g. `Pi`) and conversions from exact to floating point inexact data use p decimal digits, as well as iterative routines implicitly using a floating point accuracy as a termination criterion (e.g. `solve` or `intnum`). But `realprecision` itself is unaffected and is “unmasked” when we exit the dynamic (*not* lexical) scope. In effect, this is similar to

```
my(prec = default(realprecision));
default(realprecision,p);
...
default(realprecision, prec);
```

but is both less cumbersome, cleaner (no need to manipulate a global variable, which in fact never changes and is only temporarily masked) and more robust: if the above computation is interrupted or an exception occurs, `realprecision` will not be restored as intended.

Such `localprec` statements can be nested, the innermost one taking precedence as expected. Beware that `localprec` follows the semantic of `local`, not `my`: a subroutine called from `localprec` scope uses the local accuracy:

```
? f()=precision(1.);
? f()
%2 = 38
? localprec(19); f()
%3 = 19
```

Warning. Changing `realprecision` itself in programs is now deprecated in favor of `localprec`. Think about the `realprecision` default as an interactive command for the `gp` interpreter, best left out of GP programs. Indeed, the above rules imply that mixing both constructs yields surprising results:

```
? \p38
? localprec(19); default(realprecision,100); Pi
%1 = 3.141592653589793239
? \p
    realprecision = 115 significant digits (100 digits displayed)
```

Indeed, `realprecision` itself is ignored within `localprec` scope, so `Pi` is computed to a low accuracy. And when we leave `localprec` scope, `realprecision` only regains precedence, it is not “restored” to the original value.

3.2.51 mapapply($\sim M, x, f, \{u\}$). Applies the closure f to the image y of x by the map M and returns the evaluation $f(y)$. The closure f is allowed to modify the components of y in place. If M is not defined at x , and the optional argument u (for *undefined*) is present and is a closure of arity 0, return the evaluation $u()$.

To apply f to *all* entries (values) of M , use `apply( $f, M$ )` instead. There are two main use-cases:

- performing a computation on a value directly, without using `mapget`, avoiding a copy:

```
? M = Map(); mapput(~M, "a", mathilbert(2000));
? matsize(mapget(M, "a"))  \\ Slow because mapget(M, "a") copies the value
%2 = [2000, 2000]
time = 101 ms.
? mapapply(~M, "a", matsize)  \\ Fast
time = 0 ms.
%3 = [2000, 2000]
```

- modifying a value in place, for example to append an element to a value in a map of lists.

This requires to use `~` in the function declaration. In the following `maplistput`, M is a map of lists and we append v to the list `mapget(M,k)`, except this is done in place ! When the map is undefined at k , we use the `u(ndefined)` argument `()->List(v)` to convert v to a list then insert it in the map:

```
? maplistput(~M, k, v) = mapapply(~M, k, (~y)->listput(~y,v), ()->List(v));
? M = Map();
%2 = Map([;])
? maplistput(~M, "a", 1); M
%3 = Map(["a", List([1])])
```

```
? maplistput(~M, "a", 2); M
%4 = Map(["a", List([1, 2])])
? maplistput(~M, "b", 3); M
%5 = Map(["a", List([1, 2]); "b", List([3])])
? maplistput(~M, "a", 4); M
%6 = Map(["a", List([1, 2, 4]); "b", List([])])
```

The library syntax is GEN mapapply(GEN ~M, GEN x, GEN f, GEN u = NULL).

3.2.52 mapdelete($\sim M, x$). Removes x from the domain of the map M .

```
? M = Map(["a",1; "b",3; "c",7]);
? mapdelete(M,"b");
? Mat(M)
["a" 1]
["c" 7]
```

The library syntax is void mapdelete(GEN ~M, GEN x).

3.2.53 mapget(M, x). Returns the image of x by the map M .

```
? M=Map(["a",23;"b",43]);
? mapget(M,"a")
%2 = 23
? mapget(M,"b")
%3 = 43
```

Raises an exception when the key x is not present in M .

```
? mapget(M,"c")
*** at top-level: mapget(M,"c")
*** ^-----
*** mapget: nonexistent component in mapget: index not in map
```

The library syntax is GEN mapget(GEN M, GEN x).

3.2.54 mapisdefined($M, x, \{&z\}$). Returns true (1) if x has an image by the map M , false (0) otherwise. If z is present, set z to the image of x , if it exists.

```
? M1 = Map([1, 10; 2, 20]);
? mapisdefined(M1,3)
%1 = 0
? mapisdefined(M1, 1, &z)
%2 = 1
? z
%3 = 10

? M2 = Map(); N = 19;
? for (a=0, N-1, mapput(M2, a^3%N, a));
? {for (a=0, N-1,
    if (mapisdefined(M2, a, &b),
        printf("%d is the cube of %d mod %d\n",a,b,N));}
0 is the cube of 0 mod 19
```

```

1 is the cube of 11 mod 19
7 is the cube of 9 mod 19
8 is the cube of 14 mod 19
11 is the cube of 17 mod 19
12 is the cube of 15 mod 19
18 is the cube of 18 mod 19

```

The library syntax is `int mapisdefined(GEN M, GEN x, GEN *z = NULL)`.

3.2.55 mapput($\sim M, x, y$). Associates x to y in the map M . The value y can be retrieved with `mapget`.

```

? M = Map();
? mapput(~M, "foo", 23);
? mapput(~M, 7718, "bill");
? mapget(M, "foo")
%4 = 23
? mapget(M, 7718)
%5 = "bill"
? Vec(M) \\ keys
%6 = [7718, "foo"]
? Mat(M)
%7 =
[ 7718 "bill"
["foo"      23]

```

The library syntax is `void mapput(GEN ~M, GEN x, GEN y)`.

3.2.56 print($\{str\}*$). Outputs its arguments in raw format ending with a newline. The arguments are converted to strings following the rules in Section 2.9.1.

```

? m = matid(2);
? print(m) \\ raw format
[1, 0; 0, 1]
? printp(m) \\ prettymatrix format
[1 0]
[0 1]

```

The library syntax is `void print(GEN vec_str)`.

3.2.57 print1($\{str\}*$). Outputs its arguments in raw format, without ending with a newline. Note that you can still embed newlines within your strings, using the `\n` notation ! The arguments are converted to strings following the rules in Section 2.9.1.

The library syntax is `void print1(GEN vec_str)`.

3.2.58 printf(*fmt*, {*x*}*). This function is based on the C library command of the same name. It prints its arguments according to the format *fmt*, which specifies how subsequent arguments are converted for output. The format is a character string composed of zero or more directives:

- ordinary characters (not %), printed unchanged,
- conversions specifications (% followed by some characters) which fetch one argument from the list and prints it according to the specification.

More precisely, a conversion specification consists in a %, one or more optional flags (among #, 0, -, +, ' '), an optional decimal digit string specifying a minimal field width, an optional precision in the form of a period ('.') followed by a decimal digit string, and the conversion specifier (among d, i, o, u, x, X, p, e, E, f, g, G, s).

The flag characters. The character % is followed by zero or more of the following flags:

- #: the value is converted to an “alternate form”. For o conversion (octal), a 0 is prefixed to the string. For x and X conversions (hexa), respectively 0x and 0X are prepended. For other conversions, the flag is ignored.
- 0: the value should be zero padded. For d, i, o, u, x, X, e, E, f, F, g, and G conversions, the value is padded on the left with zeros rather than blanks. (If the 0 and - flags both appear, the 0 flag is ignored.)
- -: the value is left adjusted on the field boundary. (The default is right justification.) The value is padded on the right with blanks, rather than on the left with blanks or zeros. A - overrides a 0 if both are given.
- ' ' (a space): a blank is left before a positive number produced by a signed conversion.
- +: a sign (+ or -) is placed before a number produced by a signed conversion. A + overrides a space if both are used.

The field width. An optional decimal digit string (whose first digit is nonzero) specifying a *minimum* field width. If the value has fewer characters than the field width, it is padded with spaces on the left (or right, if the left-adjustment flag has been given). In no case does a small field width cause truncation of a field; if the value is wider than the field width, the field is expanded to contain the conversion result. Instead of a decimal digit string, one may write * to specify that the field width is given in the next argument.

The precision. An optional precision in the form of a period ('.') followed by a decimal digit string. This gives the number of digits to appear after the radix character for e, E, f, and F conversions, the maximum number of significant digits for g and G conversions, and the maximum number of characters to be printed from an s conversion. Instead of a decimal digit string, one may write * to specify that the field width is given in the next argument.

The length modifier. This is ignored under gp, but necessary for libpari programming. Description given here for completeness:

- l: argument is a long integer.
- P: argument is a GEN.

The conversion specifier. A character that specifies the type of conversion to be applied.

- **d, i**: a signed integer.
- **o, u, x, X**: an unsigned integer, converted to unsigned octal (**o**), decimal (**u**) or hexadecimal (**x** or **X**) notation. The letters **abcdef** are used for **x** conversions; the letters **ABCDEF** are used for **X** conversions.
- **e, E**: the (real) argument is converted in the style `[ -]d.ddd e[ -]dd`, where there is one digit before the decimal point, and the number of digits after it is equal to the precision; if the precision is missing, use the current **realprecision** for the total number of printed digits. If the precision is explicitly 0, no decimal-point character appears. An **E** conversion uses the letter **E** rather than **e** to introduce the exponent.
- **f, F**: the (real) argument is converted in the style `[ -]ddd.ddd`, where the number of digits after the decimal point is equal to the precision; if the precision is missing, use the current **realprecision** for the total number of printed digits. If the precision is explicitly 0, no decimal-point character appears. If a decimal point appears, at least one digit appears before it.
- **g, G**: the (real) argument is converted in style **e** or **f** (or **E** or **F** for **G** conversions) `[ -]ddd.ddd`, where the total number of digits printed is equal to the precision; if the precision is missing, use the current **realprecision**. If the precision is explicitly 0, it is treated as 1. Style **e** is used when the decimal exponent is < -4 , to print `0.`, or when the integer part cannot be decided given the known significant digits, and the **f** format otherwise.
- **c**: the integer argument is converted to an unsigned char, and the resulting character is written.
- **s**: convert to a character string. If a precision is given, no more than the specified number of characters are written.
- **p**: print the address of the argument in hexadecimal (as if by `%#x`).
- **%**: a `%` is written. No argument is converted. The complete conversion specification is `%%`.

Examples:

```
? printf("floor: %d, field width 3: %3d, with sign: %+3d\n", Pi, 1, 2);
floor: 3, field width 3:   1, with sign:  +2

? printf("%.5g %.5g %.5g\n",123,123/456,123456789);
123.00 0.26974 1.2346 e8

? printf("%-2.5s:%2.5s:%2.5s\n", "P", "PARI", "PARIGP");
P :PARI:PARIG

\\ min field width and precision given by arguments
? x = 23; y=-1/x; printf("x=%+06.2f y=%+0*.*f\n", x, 6, 2, y);
x=+23.00 y=-00.04

\\ minimum fields width 5, pad left with zeroes
? for (i = 2, 5, printf("%05d\n", 10^i))
00100
01000
10000
100000  \\ don't truncate fields whose length is larger than the minimum width
? printf("%.2f  |%06.2f|", Pi,Pi)
```

3.14 | 3.14|

All numerical conversions apply recursively to the entries of complex numbers, vectors and matrices:

```
? printf("%4d", [1,2,3]);
[  1,   2,   3]
? printf("%5.2f", mathilbert(3));
[ 1.00  0.50  0.33]
[ 0.50  0.33  0.25]
[ 0.33  0.25  0.20]
? printf("%.3g", Pi+I)
3.14+1.00I
```

Technical note. Our implementation of `printf` deviates from the C89 and C99 standards in a few places:

- whenever a precision is missing, the current `realprecision` is used to determine the number of printed digits (C89: use 6 decimals after the radix character).
- in conversion style `e`, we do not impose that the exponent has at least two digits; we never write a `+` sign in the exponent; 0 is printed in a special way, always as `0.Exp`.
- in conversion style `f`, we switch to style `e` if the exponent is greater or equal to the precision.
- in conversion `g` and `G`, we do not remove trailing zeros from the fractional part of the result; nor a trailing decimal point; 0 is printed in a special way, always as `0.Exp`.

The library syntax is `void printf0(const char *fmt, GEN vec_x)`.

The variadic version `void pari_printf(const char *fmt, ...)` is usually preferable.

3.2.59 `printp`(`{str}*`). Outputs its arguments in prettymatrix format, ending with a newline. The arguments are converted to strings following the rules in Section 2.9.1.

```
? m = matid(2);
? print(m) \\ raw format
[1, 0; 0, 1]
? printp(m) \\ prettymatrix format
[1 0]
[0 1]
```

The library syntax is `void printp(GEN vec_str)`.

3.2.60 `printsep`(`sep, {str}*`). Outputs its arguments in raw format, ending with a newline. The arguments are converted to strings following the rules in Section 2.9.1. Successive entries are separated by `sep`:

```
? printsep(":", 1,2,3,4)
1:2:3:4
```

The library syntax is `void printsep(const char *sep, GEN vec_str)`.

3.2.61 printsep1(*sep*, {*str*}*). Outputs its arguments in raw format, without ending with a newline. The arguments are converted to strings following the rules in Section 2.9.1. Successive entries are separated by *sep*:

```
? printsep1(":", 1,2,3,4);print("|")
1:2:3:4|
```

The library syntax is `void printsep1(const char *sep, GEN vec_str)`.

3.2.62 printtex({*str*}*). Outputs its arguments in $\text{\TeX}$ format. This output can then be used in a $\text{\TeX}$ manuscript, see `strtex` for details. The arguments are converted to strings following the rules in Section 2.9.1. The printing is done on the standard output. If you want to print it to a file you should use `writetex` (see there).

Another possibility is to enable the `log` default (see Section 2.12). You could for instance do:

```
default(logfile, "new.tex");
default(log, 1);
printtex(result);
```

The library syntax is `void printtex(GEN vec_str)`.

3.2.63 quit({*status* = 0}). Exits `gp` and return to the system with exit status *status*, a small integer. A nonzero exit status normally indicates abnormal termination. (Note: the system actually sees only *status* mod 256, see your man pages for `exit(3)` or `wait(2)`).

3.2.64 read({*filename*}). Reads in the file *filename* (subject to string expansion). If *filename* is omitted, re-reads the last file that was fed into `gp`. The return value is the result of the last expression evaluated.

If a GP `binary file` is read using this command (see Section 3.2.88), the file is loaded and the last object in the file is returned.

In case the file you read in contains an `allocatemem` statement (to be generally avoided), you should leave `read` instructions by themselves, and not part of larger instruction sequences.

Variants. `readvec` allows to read a whole file at once; `fileopen` followed by either `fileread` (evaluated lines) or `filereadstr` (lines as nonevaluated strings) allows to read a file one line at a time.

The library syntax is `GEN gp_read_file(const char *filename)`.

3.2.65 readstr({*filename*}). Reads in the file *filename* and return a vector of GP strings, each component containing one line from the file. If *filename* is omitted, re-reads the last file that was fed into `gp`.

The library syntax is `GEN readstr(const char *filename)`.

3.2.66 readvec(*{filename}*). Reads in the file *filename* (subject to string expansion). If *filename* is omitted, re-reads the last file that was fed into `gp`. The return value is a vector whose components are the evaluation of all sequences of instructions contained in the file. For instance, if *file* contains

```
1
2
3
```

then we will get:

```
? \r a
%1 = 1
%2 = 2
%3 = 3
? read(a)
%4 = 3
? readvec(a)
%5 = [1, 2, 3]
```

In general a sequence is just a single line, but as usual braces and `\` may be used to enter multiline sequences.

The library syntax is `GEN gp_readvec_file(const char *filename)`. The underlying library function `GEN gp_readvec_stream(FILE *f)` is usually more flexible.

3.2.67 select(*f*, *A*, *{flag = 0}*). We first describe the default behavior, when *flag* is 0 or omitted. Given a vector or list *A* and a `t_CLOSURE` *f*, `select` returns the elements *x* of *A* such that *f*(*x*) is nonzero. In other words, *f* is seen as a selection function returning a boolean value.

```
? select(x->isprime(x), vector(50,i,i^2+1))
%1 = [2, 5, 17, 37, 101, 197, 257, 401, 577, 677, 1297, 1601]
? select(x->(x<100), %)
%2 = [2, 5, 17, 37]
```

returns the primes of the form $i^2 + 1$ for some $i \leq 50$, then the elements less than 100 in the preceding result. The `select` function also applies to a matrix *A*, seen as a vector of columns, i.e. it selects columns instead of entries, and returns the matrix whose columns are the selected ones.

Remark. For *v* a `t_VEC`, `t_COL`, `t_VECSMALL`, `t_LIST` or `t_MAT`, the alternative set-notations

```
[g(x) | x <- v, f(x)]
[x | x <- v, f(x)]
[g(x) | x <- v]
```

are available as shortcuts for

```
apply(g, select(f, Vec(v)))
select(f, Vec(v))
apply(g, Vec(v))
```

respectively:

```
? [ x | x <- vector(50,i,i^2+1), isprime(x) ]
%1 = [2, 5, 17, 37, 101, 197, 257, 401, 577, 677, 1297, 1601]
```

If *flag* = 1, this function returns instead the *indices* of the selected elements, and not the elements themselves (indirect selection):

```
? V = vector(50,i,i^2+1);
? select(x->isprime(x), V, 1)
%2 = Vecsmall([1, 2, 4, 6, 10, 14, 16, 20, 24, 26, 36, 40])
? vecextract(V, %)
%3 = [2, 5, 17, 37, 101, 197, 257, 401, 577, 677, 1297, 1601]
```

The following function lists the elements in $(\mathbf{Z}/N\mathbf{Z})^*$:

```
? invertibles(N) = select(x->gcd(x,N) == 1, [1..N])
```

Finally

```
? select(x->x, M)
```

selects the nonzero entries in M. If the latter is a `t_MAT`, we extract the matrix of nonzero columns. Note that *removing* entries instead of selecting them just involves replacing the selection function *f* with its negation:

```
? select(x->!isprime(x), vector(50,i,i^2+1))
```

The library syntax is `genselect(void *E, long (*fun)(void*,GEN), GEN a)`. Also available is `GEN genindexselect(void *E, long (*fun)(void*, GEN), GEN a)`, corresponding to *flag* = 1.

3.2.68 self(). Return the calling function or closure as a `t_CLOSURE` object. This is useful for defining anonymous recursive functions.

```
? (n -> if(n==0,1,n*self()(n-1)))(5)
%1 = 120 \\ 5!
? (n -> if(n<=1, n, self()(n-1)+self()(n-2)))(20)
%2 = 6765 \\ Fibonacci(20)
```

The library syntax is `GEN pari_self()`.

3.2.69 setrand(*n*). Reseeds the random number generator using the seed *n*. No value is returned. The seed is a small positive integer $0 < n < 2^{64}$ used to generate deterministically a suitable state array. All gp session start by an implicit `setrand(1)`, so resetting the seed to this value allows to replay all computations since the session start. Alternatively, running a randomized computation starting by `setrand(n)` twice with the same *n* will generate the exact same output.

In the other direction, including a call to `setrand(getwalltime())` from your gprc will cause GP to produce different streams of random numbers in each session. (Unix users may want to use `/dev/urandom` instead of `getwalltime()`.)

For debugging purposes, one can also record a particular random state using `getrand` (the value is encoded as a huge integer) and feed it to `setrand`:

```
? state = getrand(); \\ record seed
...
? setrand(state); \\ we can now replay the exact same computations
```

The library syntax is `void setrand(GEN n)`.

3.2.70 strchr(x). Converts integer or vector of integers x to a string, translating each integer (in the range $[1, 255]$) into a character using ASCII encoding.

```
? strchr(97)
%1 = "a"
? Vecsmall("hello world")
%2 = Vecsmall([104, 101, 108, 108, 111, 32, 119, 111, 114, 108, 100])
? strchr(%)
%3 = "hello world"
```

The library syntax is GEN `pari_strchr`(GEN x).

3.2.71 strexand($\{x\}^*$). Converts its argument list into a single character string (type `t_STR`, the empty string if x is omitted). Then perform environment expansion, see Section 2.12. This feature can be used to read environment variable values.

```
? strexand("$HOME/doc")
%1 = "/home/pari/doc"
? module = "aprcl"; n = 10;
? strexand("$HOME/doc/", module, n, ".tex")
%3 = "/home/pari/doc/aprcl10.tex"
```

The individual arguments are read in string context, see Section 2.9.1.

The library syntax is GEN `strexand`(GEN vec_x).

3.2.72 strjoin($v, \{p = ""\}$). Joins the strings in vector v , separating them with delimiter p . The reverse operation is `strsplit`.

```
? v = ["abc", "def", "ghi"]
? strjoin(v, "/")
%2 = "abc/def/ghi"
? strjoin(v)
%3 = "abcdefghi"
```

The library syntax is GEN `strjoin`(GEN v , GEN $p = \text{NULL}$).

3.2.73 sprintf($f_{\text{mt}}, \{x\}^*$). Returns a string built from the remaining arguments according to the format f_{mt} . The format consists of ordinary characters (not `%`), printed unchanged, and conversions specifications. See `printf`.

```
? dir = "/home/pari"; file = "aprcl"; n = 10;
? sprintf("%s/%s%ld.tex", dir, file, n)
%2 = "/home/pari/aprcl10.tex"
```

The library syntax is GEN `sprintf`(const char $*f_{\text{mt}}$, GEN vec_x).

The variadic version `char * pari_sprintf(const char  $*f_{\text{mt}}$ , ...)` is usually preferable.

3.2.74 strsplit(*s*, {*p* = ""}). Splits the string *s* into a vector of strings, with *p* acting as a delimiter. If *p* is empty or omitted, split the string into characters.

```
? strsplit("abc::def::ghi", "::")
%1 = ["abc", "def", "ghi"]
? strsplit("abc", "")
%2 = ["a", "b", "c"]
? strsplit("aba", "a")
```

If *s* starts (resp. ends) with the pattern *p*, then the first (resp. last) entry in the vector is the empty string:

```
? strsplit("aba", "a")
%3 = ["", "b", ""]
```

The library syntax is GEN `strsplit(GEN s, GEN p = NULL)`.

3.2.75 strtex({*x*}*). Translates its arguments to TeX format, and concatenates the results into a single character string (type `t_STR`, the empty string if *x* is omitted).

The individual arguments are read in string context, see Section 2.9.1.

```
? v = [1, 2, 3]
%1 [1, 2, 3]
? strtex(v)
%2 = "\\pmatrix{ 1&2&3\\cr}\\n"
```

TeX-nical notes. The TeX output engine was originally written for plain TeX and designed for maximal portability. Unfortunately later LaTeX packages have obsoleted valid TeX primitives, leading us to replace TeX's `\over` by LaTeX's `\frac` in PARI's TeX output. We have decided not to update further our TeX markup and let the users of various LaTeX engines customize their preambles. The following documents the precise changes you may need to include in your style files to incorporate PARI TeX output verbatim:

- if you enabled bit 4 in `TeXstyle` default, you must define `\PARIbreak`; see `??TeXstyle`;
- if you use plain TeX only: you must define `\frac` as follows

```
\def\frac#1#2{{#1\over#2}}
```

• if you use LaTeX and `amsmath`, `\pmatrix` is obsoleted in favor of the `pmatrix` environment; see `examples/parigp.sty` for how to re-enable the deprecated construct.

The library syntax is GEN `strtex(GEN vec_x)`.

3.2.76 strtime(*t*). Return a string describing the time *t* in milliseconds in the format used by the GP timer.

```
? print(strtime(12345678))
3h, 25min, 45,678 ms
? {
  my(t=getabstime());
  F=factor(2^256+1);t=getabstime()-t;
  print("factor(2^256+1) took ",strtime(t));
}
factor(2^256+1) took 1,320 ms
```

The library syntax is GEN `strtime(long t)`.

3.2.77 system(*str*). *str* is a string representing a system command. This command is executed, its output written to the standard output (this won't get into your logfile), and control returns to the PARI system. This simply calls the C `system` command. Return the shell return value (which is system-dependent). Beware that UNIX shell convention for boolean is opposite to GP, true is 0 and false is non-0.

```
? system("test -d /") \\ test if '/' is a directory (true)
%1 = 0
? system("test -f /") \\ test if '/' is a file (false)
%2 = 1
```

The library syntax is `long gpsystem(const char *str)`.

3.2.78 trap(*{e}*, *{rec}*, *seq*). This function is obsolete, use `iferr`, which has a nicer and much more powerful interface. For compatibility's sake we now describe the *obsolete* function `trap`.

This function tries to evaluate *seq*, trapping runtime error *e*, that is effectively preventing it from aborting computations in the usual way; the recovery sequence *rec* is executed if the error occurs and the evaluation of *rec* becomes the result of the command. If *e* is omitted, all exceptions are trapped. See Section 2.10.2 for an introduction to error recovery under `gp`.

```
? \\ trap division by 0
? inv(x) = trap (e_INV, INFINITY, 1/x)
? inv(2)
%1 = 1/2
? inv(0)
%2 = INFINITY
```

Note that *seq* is effectively evaluated up to the point that produced the error, and the recovery sequence is evaluated starting from that same context, it does not "undo" whatever happened in the other branch (restore the evaluation context):

```
? x = 1; trap (, /* recover: */ x, /* try: */ x = 0; 1/x)
%1 = 0
```

Note. The interface is currently not adequate for trapping individual exceptions. In the current version 2.17.4, the following keywords are recognized, but the name list will be expanded and changed in the future (all library mode errors can be trapped: it's a matter of defining the keywords to `gp`):

`e_ALARM`: alarm time-out

`e_ARCH`: not available on this architecture or operating system

`e_STACK`: the PARI stack overflows

`e_INV`: impossible inverse

`e_IMPL`: not yet implemented

`e_OVERFLOW`: all forms of arithmetic overflow, including length or exponent overflow (when a larger value is supplied than the implementation can handle).

`e_SYNTAX`: syntax error

`e_MISC`: miscellaneous error

e_TYPE: wrong type

e_USER: user error (from the `error` function)

The library syntax is `GEN trap0(const char *e = NULL, GEN rec = NULL, GEN seq = NULL)`.

3.2.79 type(*x*). This is useful only under `gp`. Returns the internal type name of the PARI object *x* as a string. Check out existing type names with the metacommand `\t`. For example `type(1)` will return `"t_INT"`.

The library syntax is `GEN type0(GEN x)`. The macro `typ` is usually simpler to use since it returns a `long` that can easily be matched with the symbols `t_*`. The name `type` was avoided since it is a reserved identifier for some compilers.

3.2.80 unexport(*x*, ..., *z*). Remove *x*, ..., *z* from the list of variables exported to the parallel world. See `export`.

3.2.81 unexportall(). Empty the list of variables exported to the parallel world.

The library syntax is `void unexportall()`.

3.2.82 uninline(). Exit the scope of all current `inline` variables. DEPRECATED, use `export / unexport`.

3.2.83 version(). Returns the current version number as a `t_VEC` with three integer components (major version number, minor version number and patchlevel); if your sources were obtained through our version control system, this will be followed by further more precise arguments, including e.g. a `git commit hash`.

This function is present in all versions of PARI following releases 2.3.4 (stable) and 2.4.3 (testing).

Unless you are working with multiple development versions, you probably only care about the 3 first numeric components. In any case, the `lex` function offers a clever way to check against a particular version number, since it will compare each successive vector entry, numerically or as strings, and will not mind if the vectors it compares have different lengths:

```
if (lex(version(), [2,3,5]) >= 0,
    \\ code to be executed if we are running 2.3.5 or more recent.
    ,
    \\ compatibility code
);
```

On a number of different machines, `version()` could return either of

```
%1 = [2, 3, 4]    \\ released version, stable branch
%1 = [2, 4, 3]    \\ released version, testing branch
%1 = [2, 6, 1, 15174, "505ab9b"] \\ development
```

In particular, if you are only working with released versions, the first line of the `gp` introductory message can be emulated by

```
[M,m,p] = version();
printf("GP/PARI CALCULATOR Version %s.%s.%s", M,m,p);
```

If you *are* working with many development versions of PARI/GP, the 4th and/or 5th components can be profitably included in the name of your logfiles, for instance.

Technical note. For development versions obtained via `git`, the 4th and 5th components are liable to change eventually, but we document their current meaning for completeness. The 4th component counts the number of reachable commits in the branch (analogous to `svn`'s revision number), and the 5th is the `git` commit hash. In particular, `lex` comparison still orders correctly development versions with respect to each others or to released versions (provided we stay within a given branch, e.g. `master`)!

The library syntax is `GEN pari_version()`.

3.2.84 `warning`(*{str}**). Outputs the message “user warning” and the argument list (each of them interpreted as a string). If colors are enabled, this warning will be in a different color, making it easy to distinguish.

```
warning(n, " is very large, this might take a while.")
```

The library syntax is `void warning0(GEN vec_str)`.

3.2.85 `whatnow`(*key*). If keyword *key* is the name of a function that was present in GP version 1.39.15, outputs the new function name and syntax, if it changed at all. Functions that were introduced since then, then modified are also recognized.

```
? whatnow("mu")
New syntax: mu(n) ==> moebius(n)
moebius(x): Moebius function of x.
? whatnow("sin")
This function did not change
```

When a function was removed and the underlying functionality is not available under a compatible interface, no equivalent is mentioned:

```
? whatnow("buchfu")
This function no longer exists
```

(The closest equivalent would be to set `K = bnfinit(T)` then access `K.fu`.)

3.2.86 `write`(*filename*, *{str}**). Writes (appends) to *filename* the remaining arguments, and appends a newline (same output as `print`).

Variant. The high-level function `write` is expensive when many consecutive writes are expected because it cannot use buffering. The low-level interface `fileopen` / `filewrite` / `fileclose` is more efficient. It also allows to truncate existing files and replace their contents.

The library syntax is `void write0(const char *filename, GEN vec_str)`.

3.2.87 `write1`(*filename*, *{str}**). Writes (appends) to *filename* the remaining arguments without a trailing newline (same output as `print1`).

The library syntax is `void write1(const char *filename, GEN vec_str)`.

3.2.88 writebin(*filename*, {*x*}). Writes (appends) to *filename* the object *x* in binary format. This format is not human readable, but contains the exact internal structure of *x*, and is much faster to save/load than a string expression, as would be produced by **write**. The binary file format includes a magic number, so that such a file can be recognized and correctly input by the regular **read** or **\r** function. If saved objects refer to polynomial variables that are not defined in the new session, they will be displayed as **tn** for some integer *n* (the attached variable number). Installed functions and history objects can not be saved via this function.

If *x* is omitted, saves all user variables from the session, together with their names. Reading such a “named object” back in a **gp** session will set the corresponding user variable to the saved value. E.g after

```
x = 1; writebin("log")
```

reading **log** into a clean session will set **x** to 1. The relative variables priorities (see Section 2.5.3) of new variables set in this way remain the same (preset variables retain their former priority, but are set to the new value). In particular, reading such a session log into a clean session will restore all variables exactly as they were in the original one.

Just as a regular input file, a binary file can be compressed using **gzip**, provided the file name has the standard **.gz** extension.

In the present implementation, the binary files are architecture dependent and compatibility with future versions of **gp** is not guaranteed. Hence binary files should not be used for long term storage (also, they are larger and harder to compress than text files).

The library syntax is `void gpwritebin(const char *filename, GEN x = NULL)`.

3.2.89 writetex(*filename*, {*str*}*). As **write**, in **T_EX** format. See **strtex** for details: this function is essentially equivalent to calling **strtex** on remaining arguments and writing them to file.

The library syntax is `void writetex(const char *filename, GEN vec_str)`.

3.3 Parallel programming.

These function are only available if PARI was configured using **Configure --mt=...**. Two multithread interfaces are supported:

- POSIX threads
- Message passing interface (MPI)

As a rule, POSIX threads are well-suited for single systems, while MPI is used by most clusters. However the parallel GP interface does not depend on the chosen multithread interface: a properly written GP program will work identically with both.

3.3.1 parapply(f, x). Parallel evaluation of f on the elements of x . The function f must not access global variables or variables declared with `local()`, and must be free of side effects.

```
parapply(factor, [2^256 + 1, 2^193 - 1])
```

factors $2^{256} + 1$ and $2^{193} - 1$ in parallel.

```
{
  my(E = ellinit([1,3]), V = vector(12,i,randomprime(2^200)));
  parapply(p->ellcard(E,p), V)
}
```

computes the order of $E(\mathbf{F}_p)$ for 12 random primes of 200 bits.

The library syntax is `GEN parapply(GEN f, GEN x)`.

3.3.2 pareval(x). Parallel evaluation of the elements of x , where x is a vector of closures. The closures must be of arity 0, must not access global variables or variables declared with `local` and must be free of side effects.

Here is an artificial example explaining the MOV attack on the elliptic discrete log problem (by reducing it to a standard discrete log over a finite field):

```
{
  my(q = 2^30 + 3, m = 40 * q, p = 1 + m^2); \\ p, q are primes
  my(E = ellinit([0,0,0,1,0] * Mod(1,p)));
  my([P, Q] = ellgenerators(E));
  \\ E(F_p) ~ Z/m P + Z/m Q and the order of the
  \\ Weil pairing <P,Q> in (Z/p)^* is m
  my(F = [m,factor(m)], e = random(m), R, wR, wQ);
  R = ellpow(E, Q, e);
  wR = ellweilpairing(E,P,R,m);
  wQ = ellweilpairing(E,P,Q,m); \\ wR = wQ^e
  pareval([()->znlog(wR,wQ,F), ()->elllog(E,R,Q), ()->e])
}
```

Note the use of `my` to pass "arguments" to the functions we need to evaluate while satisfying the listed requirements: closures of arity 0 and no global variables (another possibility would be to use `export`). As a result, the final three statements satisfy all the listed requirements and are run in parallel. (Which is silly for this computation but illustrates the use of `pareval`.) The function `parfor` is more powerful but harder to use.

The library syntax is `GEN pareval(GEN x)`.

3.3.3 parfor($i = a, \{b\}, \text{expr1}, \{r\}, \{\text{expr2}\}$). Evaluates in parallel the expression **expr1** in the formal argument i running from a to b . If b is set to $+\infty$, the loop runs indefinitely. If r and **expr2** are present, the expression **expr2** in the formal variables r and i is evaluated with r running through all the different results obtained for **expr1** and i takes the corresponding argument.

The computations of **expr1** are *started* in increasing order of i ; otherwise said, the computation for $i = c$ is started after those for $i = 1, \dots, c-1$ have been started, but before the computation for $i = c+1$ is started. Notice that the order of *completion*, that is, the order in which the different r become available, may be different; **expr2** is evaluated sequentially on each r as it appears.

The following example computes the sum of the squares of the integers from 1 to 10 by computing the squares in parallel and is equivalent to **parsum** (i=1, 10, i^2):

```
? s=0;
? parfor (i=1, 10, i^2, r, s=s+r)
? s
%3 = 385
```

More precisely, apart from a potentially different order of evaluation due to the parallelism, the line containing **parfor** is equivalent to

```
? my (r); for (i=1, 10, r=i^2; s=s+r)
```

The sequentiality of the evaluation of **expr2** ensures that the variable **s** is not modified concurrently by two different additions, although the order in which the terms are added is nondeterministic.

It is allowed for **expr2** to exit the loop using **break/next/return**. If that happens for $i = c$, then the evaluation of **expr1** and **expr2** is continued for all values $i < c$, and the return value is the one obtained for the smallest i causing an interruption in **expr2** (it may be undefined if this is a **break/next**). In that case, using side-effects in **expr2** may lead to undefined behavior, as the exact number of values of i for which it is executed is nondeterministic. The following example computes **nextprime**(1000) in parallel:

```
? parfor (i=1000, , isprime (i), r, if (r, return (i)))
%1 = 1009
```

3.3.4 parforeach($V, x, \text{expr1}, \{r\}, \{\text{expr2}\}$). Evaluates in parallel the expression **expr1** in the formal argument x , where x runs through all components of V . If r and **expr2** are present, evaluate sequentially the expression **expr2**, in which the formal variables x and r are replaced by the successive arguments and corresponding values. The sequential evaluation ordering is not specified:

```
? parforeach([50..100], x,isprime(x), r, if(r,print(x)))
53
67
71
79
83
89
97
73
59
61
```

3.3.5 parforprime($p = a, \{b\}, \text{expr1}, \{r\}, \{\text{expr2}\}$). Behaves exactly as **parfor**, but loops only over prime values p . Precisely, the functions evaluates in parallel the expression **expr1** in the formal argument p running through the primes from a to b . If b is set to $+\infty$, the loop runs indefinitely. If r and **expr2** are present, the expression **expr2** in the formal variables r and p is evaluated with r running through all the different results obtained for **expr1** and p takes the corresponding argument.

It is allowed fo **expr2** to exit the loop using **break/next/return**; see the remarks in the documentation of **parfor** for details.

3.3.6 parforprimestep($p = a, \{b\}, q, \text{expr1}, \{r\}, \{\text{expr2}\}$). Behaves exactly as **parfor**, but loops only over prime values p in an arithmetic progression Precisely, the functions evaluates in parallel the expression **expr1** in the formal argument p running through the primes from a to b in an arithmetic progression of the form $a + kq$. ($p \equiv a \pmod{q}$) or an intmod $\text{Mod}(c, N)$. If b is set to $+\infty$, the loop runs indefinitely. If r and **expr2** are present, the expression **expr2** in the formal variables r and p is evaluated with r running through all the different results obtained for **expr1** and p takes the corresponding argument.

It is allowed fo **expr2** to exit the loop using **break/next/return**; see the remarks in the documentation of **parfor** for details.

3.3.7 parforstep($i = a, \{b\}, s, \text{expr1}, \{r\}, \{\text{expr2}\}$). Evaluates in parallel the expression **expr1** in the formal argument i running from a to b in steps of s (can be a positive real number, an intmod for an arithmetic progression, or finally a vector of steps, see **forstep**). If r and **expr2** are present, the expression **expr2** in the formal variables r and i is evaluated with r running through all the different results obtained for **expr1** and i takes the corresponding argument.

```
? parforstep(i=3,8,2,2*i,x,print([i,x]))
[3, 6]
[5, 10]
[7, 14]
? parforstep(i=3,8,Mod(1,3),2*i,x,print([i,x]))
[4, 8]
[7, 14]
? parforstep(i=3,10,[1,3],2*i,x,print([i,x]))
[3, 6]
[4, 8]
[7, 14]
[8, 16]
```

The library syntax is `void parforstep0(GEN i, GEN b = NULL, GEN s, GEN expr1, GEN r = NULL)`.

3.3.8 parforvec($X = v, \text{expr1}, \{j\}, \{\text{expr2}\}, \{\text{flag}\}$). Evaluates the sequence **expr2** (dependent on X and j) for X as generated by **forvec**, in random order, computed in parallel. Substitute for j the value of **expr1** (dependent on X).

It is allowed fo **expr2** to exit the loop using **break/next/return**, however in that case, **expr2** will still be evaluated for all remaining value of p less than the current one, unless a subsequent **break/next/return** happens.

3.3.9 parselect($f, A, \{flag = 0\}$). Selects elements of A according to the selection function f , done in parallel. If $flag$ is 1, return the indices of those elements (indirect selection) The function f must not access global variables or variables declared with `local()`, and must be free of side effects.

The library syntax is `GEN parselect(GEN f, GEN A, long flag)`.

3.3.10 parsum($i = a, b, expr$). Sum of expression $expr$, the formal parameter going from a to b , evaluated in parallel in random order. The expression $expr$ must not access global variables or variables declared with `local()`, and must be free of side effects.

```
? parsum(i=1,1000,ispseudoprime(2^prime(i)-1))
cpu time = 1min, 26,776 ms, real time = 5,854 ms.
%1 = 20
```

returns the number of prime numbers among the first 1000 Mersenne numbers.

Note. This function is only useful when summing entries that are too large to fit in memory simultaneously. To sum a small number of values, using `vecsum(parvector())` is likely to be more efficient; the summation order also becomes deterministic.

3.3.11 parvector($N, i, expr$). As `vector(N, i, expr)` but the evaluations of $expr$ are done in parallel. The expression $expr$ must not access global variables or variables declared with `local()`, and must be free of side effects.

```
parvector(10,i,quadclassunit(2^(100+i)+1).no)
```

computes the class numbers in parallel.

3.4 GP defaults.

This section documents the GP defaults, that can be set either by the GP function `default` or in your GPRC. Be sure to check out `parisize` and `parisizemax` !

3.4.1 TeXstyle. The bits of this default allow `gp` to use less rigid TeX formatting commands in the logfile. This default is only taken into account when `log = 3`. The bits of `TeXstyle` have the following meaning

2: insert `\right / \left` pairs where appropriate.

4: insert discretionary breaks in polynomials, to enhance the probability of a good line break. You *must* then define `\PARIbreak` as follows:

```
\def\PARIbreak{\hskip 0pt plus \hsize\relax\discretionary{}{}{}}
```

The default value is 0.

3.4.2 breakloop. If true, enables the “break loop” debugging mode, see Section [2.10.3](#).

The default value is 1 if we are running an interactive `gp` session, and 0 otherwise.

3.4.3 colors. This default is only usable if `gp` is running within certain color-capable terminals. For instance `rxvt`, `color_xterm` and modern versions of `xterm` under X Windows, or standard Linux/DOS text consoles. It causes `gp` to use a small palette of colors for its output. With `xterms`, the colormap used corresponds to the resources `Xterm*colorn` where n ranges from 0 to 15 (see the file `misc/color.dft` for an example). Accepted values for this default are strings " $a_1, \dots, a_k$ " where $k \leq 7$ and each a_i is either

- the keyword `no` (use the default color, usually black on transparent background)
- an integer between 0 and 15 corresponding to the aforementioned colormap
- a triple $[c_0, c_1, c_2]$ where c_0 stands for foreground color, c_1 for background color, and c_2 for attributes (0 is default, 1 is bold, 4 is underline).

The output objects thus affected are respectively error messages, history numbers, prompt, input line, output, help messages, timer (that's seven of them). If $k < 7$, the remaining a_i are assumed to be `no`. For instance

```
default(colors, "9, 5, no, no, 4")
```

typesets error messages in color 9, history numbers in color 5, output in color 4, and does not affect the rest.

A set of default colors for dark (reverse video or PC console) and light backgrounds respectively is activated when `colors` is set to `darkbg`, resp. `lightbg` (or any proper prefix: `d` is recognized as an abbreviation for `darkbg`). A bold variant of `darkbg`, called `boldfg`, is provided if you find the former too pale.

EMACS: In the present version, this default is incompatible with `PariEmacs`. Changing it will just fail silently (the alternative would be to display escape sequences as is, since Emacs will refuse to interpret them). You must customize color highlighting from the `PariEmacs` side, see its documentation.

The default value is "" (no colors).

3.4.4 compatible. Obsolete. This default is now a no-op.

3.4.5 datadir. The name of directory containing the optional data files. For now, this includes the `elldata`, `galdata`, `galpol`, `seadata` packages.

The default value is `/usr/local/share/pari`, or the override specified via `Configure --datadir=`.

Windows-specific note. On Windows operating systems, the special value `@` stands for “the directory where the `gp` binary is installed”. This is the default value.

3.4.6 debug. Debugging level. If it is nonzero, some extra messages may be printed, according to what is going on (see `\g`). To turn on and off diagnostics attached to a specific feature (such as the LLL algorithm), use `setdebug`.

The default value is 0 (no debugging messages).

3.4.7 debugfiles. This is a deprecated alias for `setdebug("io",)`. If nonzero, `gp` will print information on file descriptors in use and I/O operations (see `\gf`).

The default value is 0 (no debugging messages).

3.4.8 debugmem. Memory debugging level (see `\gm`). If this is nonzero, `gp` will print increasingly precise notifications about memory use:

- `debugmem > 0`, notify when `parisize` changes (within the boundaries set by `parisizemax`);
- `debugmem > 1`, indicate any important garbage collection and the function it is taking place in;
- `debugmem > 2`, indicate the creation/destruction of “blocks” (or clones); expect lots of messages.

Important Note: if you are running a version compiled for debugging (see Appendix A) and `debugmem > 1`, `gp` will further regularly print information on memory usage, notifying whenever stack usage goes up or down by 1 MByte. This functionality is disabled on non-debugging builds as it noticeably slows down the performance.

The default value is 1.

3.4.9 echo. This default can be 0 (off), 1 (on) or 2 (on, raw). When `echo` mode is on, each command is reprinted before being executed. This can be useful when reading a file with the `\r` or `read` commands. For example, it is turned on at the beginning of the test files used to check whether `gp` has been built correctly (see `\e`). When `echo` is set to 1 the input is cleaned up, removing white space and comments and uniting multi-line input. When set to 2 (raw), the input is written as-is, without any pre-processing.

The default value is 0 (no echo).

3.4.10 factor_add_primes. This toggle is either 1 (on) or 0 (off). If on, the integer factorization machinery calls `addprimes` on prime factors that were difficult to find (larger than 2^{24}), so they are automatically tried first in other factorizations. If a routine is performing (or has performed) a factorization and is interrupted by an error or via Control-C, this lets you recover the prime factors already found. The downside is that a huge `addprimes` table unrelated to the current computations will slow down arithmetic functions relying on integer factorization; one should then empty the table using `removeprimes`.

The default value is 0.

3.4.11 factor_proven. This toggle is either 1 (on) or 0 (off). By default, the factors output by the integer factorization machinery are only pseudo-primes, not proven primes. If this toggle is set, a primality proof is done for each factor and all results depending on integer factorization are fully proven. This flag does not affect partial factorization when it is explicitly requested. It also does not affect the private table managed by `addprimes`: its entries are included as is in factorizations, without being tested for primality.

The default value is 0.

3.4.12 factorlimit. `gp` precomputes a list of all primes less than `primelimit` at initialization time (and can quickly generate more primes on demand, up to the square of that bound). Let N be an integer. The command `factor(N)` factors the integer, starting by trial division by all primes up to some bound (which depends on the size of N and less than 2^{19} is any case), then moving on to more advanced algorithms. When additionally D is an integer, `factor(N, D)` uses *only* trial division by primes less than D . In both case, trial division is sped up by precomputations involving primes up to another bound called `factorlimit`. Trial division up to a larger bound is possible, but will be slower than for bounds lower than `factorlimit` and will slow down factorization on average. If `factorlimit` is larger than `primelimit`, then `primelimit` is increased to match `factorlimit`.

In the present version, precomputations are only used on startup and changing either `primelimit` or `factorlimit` will not recompute new tables. Changing `primelimit` has no effect, while changing `factorlimit` affects the behavior in factorizations.

The default value is 2^{20} , which is the default `primelimit`. This default is only used on startup: changing it will not recompute a new table.

Note that the precomputations are expensive both in terms of time and space, although softly linear in the bound, and the ones attached to `factorlimit` more so. So neither should be taken too large. Here are sample timings: in the first column are the increasing values of `primelimit`, in the second column is the startup time keeping `factorlimit` at its default value, and the third column is the startup time with `factorlimit = primelimit`.

2^{20} :	40 ms	40 ms
2^{23} :	40 ms	230 ms
2^{26} :	140 ms	2,410 ms
2^{29} :	810 ms	27,240 ms
2^{32} :	6,040 ms	293,660 ms

The final 2^{32} for `factorlimit` requires a 10GB stack. On the other hand, here are timings trying `factor(p, D)` for some random 1000-bit prime (so we are in the worst case of performing trial division in a setting where it cannot succeed) and increasing values of D . We use a `primelimit` of 2^{32} ; the first column corresponds to the values of D , the second to the times for the default `factorlimit` and the third to fifth for `factorlimit` matching D , $D/2$ and $D/4$ respectively.

2^{20} :	1 ms	1 ms	6 ms	18 ms
2^{23} :	72 ms	18 ms	21 ms	63 ms
2^{26} :	296 ms	50 ms	176 ms	233 ms
2^{29} :	1,911 ms	266 ms	1,023 ms	1,404 ms
2^{32} :	15,505 ms	2,406 ms	6,954 ms	15,264 ms

As expected, matching `factorlimit`'s fast trial division to the desired trial division bound D is optimal if we do not take precomputation time into account. But this data also shows that if you need to often trial divide above 4 `factorlimit`, then you should not bother and can just as well stick with the default value: the extra efficiency up to `factorlimit` is negligible compared to the naive trial division that will follow. Whereas the increase in memory usage and startup time are *very* noticeable.

The default value is 2^{20} .

3.4.13 format. Of the form `x.n`, where `x` (conversion style) is a letter in `{e, f, g}`, and `n` (precision) is an integer; this affects the way real numbers are printed:

- If the conversion style is `e`, real numbers are printed in scientific format, always with an explicit exponent, e.g. `3.3 E-5`.
- In style `f`, real numbers are generally printed in fixed floating point format without exponent, e.g. `0.000033`. A large real number, whose integer part is not well defined (not enough significant digits), is printed in style `e`. For instance `10.^100` known to ten significant digits is always printed in style `e`.
- In style `g`, nonzero real numbers are printed in `f` format, except when their decimal exponent is < -4 , in which case they are printed in `e` format. Real zeroes (of arbitrary exponent) are printed in `e` format.

The precision `n` is the number of significant digits printed for real numbers, except if $n < 0$ where all the significant digits will be printed (initial default is 38 decimal digits). For more powerful formatting possibilities, see `printf` and `strprintf`.

The default value is `"g.38"`.

3.4.14 graphcolormap. A vector of colors, to be used by hi-res graphing routines. Its length is arbitrary, but it must contain at least 3 entries: the first 3 colors are used for background, frame/ticks and axes respectively. All colors in the colormap may be freely used in `plotcolor` calls.

A color is either given as in the default by character strings or by an RGB code. For valid color names, see the standard `rgb.txt` file in X11 distributions, where we restrict to lowercase letters and remove all whitespace from color names. An RGB code is a vector with 3 integer entries between 0 and 255 or a `#` followed by 6 hexadecimal digits. For instance `[250, 235, 215]`, `"#faebd7"` and `"antiquewhite"` all represent the same color.

The default value is `["white", "black", "blue", "violetred", "red", "green", "grey", "gainsboro"]`.

The colormap elements can not be changed individually as in a vector (you must either leave the colormap alone or change it globally). All color functions allow you either to hardcode a color given its descriptive name or RGB code, or to use a relative color scheme by changing the colormap and referring to an index in that table: for historical and compatibility reasons, the indexing is 0-based (as in C) and not 1-based as would be expected in a GP vector. This means that the index 0 in the default colormap represents `"white"`, 1 is `"black"`, and so on.

3.4.15 graphcolors. Entries in the `graphcolormap` that will be used to plot multi-curves. The successive curves are drawn in colors whose index in `graphcolormap` are the non-negative integers

`graphcolors[1], graphcolors[2], ...`

cycling when the `graphcolors` list is exhausted. Beware that for historical and compatibility reasons, `graphcolormap` is 0-based.

The default value is `[4,5]`. With factory settings for `graphcolormap`, this corresponds to `"red"` then `"green"`.

3.4.16 help. Name of the external help program to use from within `gp` when extended help is invoked, usually through a `??` or `???` request (see Section 2.13.1), or M-H under readline (see Section 2.15).

Windows-specific note. On Windows operating systems, if the first character of `help` is `@`, it is replaced by “the directory where the `gp` binary is installed”.

The default value is the path to the `gphelp` script we install.

3.4.17 histfile. Name of a file where `gp` will keep a history of all *input* commands (results are omitted). If this file exists when the value of `histfile` changes, it is read in and becomes part of the session history. Thus, setting this default in your `gprc` saves your readline history between sessions. Setting this default to the empty string `""` changes it to `<undefined>`. Note that, by default, the number of history entries saved is not limited: set `history-size` in readline’s `.inputrc` to limit the file size.

The default value is `<undefined>` (no history file).

3.4.18 histsize. `gp` keeps a history of the last `histsize` results computed so far, which you can recover using the `%` notation (see Section 2.13.4). When this number is exceeded, the oldest values are erased. Tampering with this default is the only way to get rid of the ones you do not need anymore.

The default value is 5000.

3.4.19 lines. If set to a positive value, `gp` prints at most that many lines from each result, terminating the last line shown with `[+++]` if further material has been suppressed. The various `print` commands (see Section 3.2) are unaffected, so you can always type `print(%)` or `\a` to view the full result. If the actual screen width cannot be determined, a “line” is assumed to be 80 characters long.

The default value is 0.

3.4.20 linewrap. If set to a positive value, `gp` wraps every single line after printing that many characters.

The default value is 0 (unset).

3.4.21 log. This can be either 0 (off) or 1, 2, 3 (on, see below for the various modes). When logging mode is turned on, `gp` opens a log file, whose exact name is determined by the `logfile` default. Subsequently, all the commands and results will be written to that file (see `\l`). In case a file with this precise name already existed, it will not be erased: your data will be *appended* at the end.

The specific positive values of `log` have the following meaning

- 1: plain logfile
- 2: emit color codes to the logfile (if `colors` is set).
- 3: write LaTeX output to the logfile (can be further customized using `TeXstyle`).

The default value is 0.

Note. Logging starts as soon as `log` is set to a nonzero value. In particular, when `log` is set in `gprc`, warnings and errors triggered from the rest of the file will be written in the logfile. For instance, on clean startup, the logfile will start by `Done.` (from the `Reading GPRC:...Done.` diagnostic printed when starting `gp`), then the `gp` header and prompt.

3.4.22 logfile. Name of the log file to be used when the `log` toggle is on. Environment and time expansion are performed.

The default value is `"pari.log"`.

3.4.23 nbthreads. This default is specific to the *parallel* version of PARI and `gp` (built via `Configure --mt=pthread` or `mpi`) and is ignored otherwise. In parallel mode, it governs the number of threads to use for parallel computing. The exact meaning and default value depend on the `mt` engine used:

- `single`: not used (always a single thread).
- `pthread`: number of threads (unlimited, default: number of cores)
- `mpi`: number of MPI processes to use (limited to the number allocated by `mpirun`, default: use all allocated processes).

See also `threadsize` and `threadsize_max`.

3.4.24 new_galois_format. This toggle is either 1 (on) or 0 (off). If on, the `polgalois` command will use a different, more consistent, naming scheme for Galois groups. This default is provided to ensure that scripts can control this behavior and do not break unexpectedly.

The default value is 0. This value will change to 1 (set) in the next major version.

3.4.25 output. There are three possible values: 0 (= *raw*), 1 (= *prettymatrix*), or 3 (= *external prettyprint*). This means that, independently of the default `format` for reals which we explained above, you can print results in three ways:

- *raw format*, i.e. a format which is equivalent to what you input, including explicit multiplication signs, and everything typed on a line instead of two dimensional boxes. This can have several advantages, for instance it allows you to pick the result with a mouse or an editor, and to paste it somewhere else.

- *prettymatrix format*: this is identical to raw format, except that matrices are printed as boxes instead of horizontally. This is prettier, but takes more space and cannot be used for input. Column vectors are still printed horizontally.

- *external prettyprint*: pipes all `gp` output in TeX format to an external prettyprinter, according to the value of `prettyprinter`. The default script (`tex2mail`) converts its input to readable two-dimensional text.

Independently of the setting of this default, an object can be printed in any of the three formats at any time using the commands `\a` and `\m` and `\B` respectively.

The default value is 1 (*prettymatrix*).

3.4.26 `parisize`. `gp`, and in fact any program using the PARI library, needs a *stack* in which to do its computations; `parisize` is the stack size, in bytes. It is recommended to increase this default using a `gprc`, to the value you believe PARI should be happy with, given your typical computation. We strongly recommend to also set `parisizemax` to a much larger value in your `gprc`, about what you believe your machine can stand: PARI will then try to fit its computations within about `parisize` bytes, but will increase the stack size if needed (up to `parisizemax`). PARI will restore the stack size to the originally requested `parisize` once we get back to the user's prompt.

If `parisizemax` is unset, this command has a very unintuitive behaviour since it must abort pending operations, see `??allocatemem`.

The default value is 8M.

3.4.27 `parisizemax`. `gp`, and in fact any program using the PARI library, needs a *stack* in which to do its computations. If nonzero, `parisizemax` is the maximum size the stack can grow to, in bytes. If zero, the stack will not automatically grow, and will be limited to the value of `parisize`.

When `parisizemax` is set, PARI tries to fit its computations within about `parisize` bytes, but will increase the stack size if needed, roughly doubling it each time (up to `parisizemax` of course!) and printing a message such as **Warning: increasing stack size to *some value***. Once the memory intensive computation is over, PARI will restore the stack size to the originally requested `parisize` without printing further messages.

We *strongly* recommend to set `parisizemax` permanently to a large nonzero value in your `gprc`, about what you believe your machine can stand. It is possible to increase or decrease `parisizemax` inside a running `gp` session, just use `default` as usual.

The default value is 0, for backward compatibility reasons.

3.4.28 `path`. This is a list of directories, separated by colons `:` (semicolons `;` in the DOS world, since colons are preempted for drive names). When asked to read a file whose name is not given by an absolute path (does not start with `/`, `./` or `../`), `gp` will look for it in these directories, in the order they were written in `path`. Here, as usual, `.` means the current directory, and `..` its immediate parent. Environment expansion is performed.

The default value is `".:~::~/gp"` on UNIX systems, `".;C:\;C:\GP"` on DOS, OS/2 and Windows, and `"."` otherwise.

3.4.29 `plotsizes`. If the graphic driver allows it, the array contains the size of the terminal, the size of the font, the size of the ticks.

3.4.30 `prettyprinter`. The name of an external prettyprinter to use when `output` is 3 (alternate prettyprinter). Note that the default `tex2mail` looks much nicer than the built-in "beautified format" (`output` = 2).

The default value is `"tex2mail -TeX -noindent -ragged -by_par"`.

3.4.31 primelimit. `gp` precomputes a list of all primes less than `primelimit` at initialization time, and can build fast sieves on demand to quickly iterate over primes up to the *square* of `primelimit`. These are used by functions looping over consecutive small primes. A related default is `factorlimit`, setting an upper bound for the small primes that can be quickly detected through fast trial division; you can still trial divide far above `factorlimit`, through `factor(N, B)` with large B but a slow algorithm will be used above `factorlimit`. If `primelimit` is set to a lower value than `factorlimit`, it is silently increased to match `factorlimit`.

The default value is 2^{20} . Since almost all arithmetic functions eventually require some table of prime numbers, PARI guarantees that the first 6547 primes, up to and including $65557 = 2^{16} + 21$, are precomputed, even if `primelimit` is 1.

A value of 2^{32} allows to quickly iterate over consecutive primes up to 2^{64} , and is the upper range of what is generally useful. (Allow for a startup time of about 6 seconds.) On the other hand, `factorlimit` is more expensive: it must build a product tree of all primes up to the bound, which can considerably increase startup time. A `factorlimit` of 2^{32} will increase startup time to about 5 minutes; and is only useful if you intend to call `factor(N, D)` *many* times with values of D about 2^{32} or 2^{33} .

This default is only used on startup: changing it will not recompute a new table. Here are sample timings for startup using increasing values of `primelimit`:

```
2^20:      40 ms
2^23:     230 ms
2^26:    2,410 ms
2^29:   27,240 ms
2^32:  293,660 ms
```

Deprecated feature. `factorlimit` was used in some situations by algebraic number theory functions using the `nf_PARTIALFACT` flag (`nfbasis`, `nfdisc`, `nfinit`, ...): this assumes that all primes $p > \text{factorlimit}$ have a certain property (the equation order is p -maximal). This is never done by default, and must be explicitly set by the user of such functions. Nevertheless, these functions now provide a more flexible interface, and their use of the global default `factorlimit` is deprecated.

Deprecated feature. `factor(N, 0)` is used to partially factor integers by removing all prime factors $\leq \text{factorlimit}$. Don't use this, supply an explicit bound: `factor(N, bound)`, which avoids relying on an unpredictable global variable.

The default value is $2^{20} = 1048576$.

3.4.32 prompt. A string that will be printed as prompt. Note that most usual escape sequences are available there: `\e` for Esc, `\n` for Newline, ..., `\\` for `\`. Time expansion is performed.

This string is sent through the library function `strftime` (on a Unix system, you can try `man strftime` at your shell prompt). This means that `%` constructs have a special meaning, usually related to the time and date. For instance, `%H` = hour (24-hour clock) and `%M` = minute [00,59] (use `%%` to get a real `%`).

If you use `readline`, escape sequences in your prompt will result in display bugs. If you have a relatively recent `readline` (see the comment at the end of Section 3.4.3), you can brace them with special sequences (`\[` and `\]`), and you will be safe. If these just result in extra spaces in

your prompt, then you'll have to get a more recent `readline`. See the file `misc/gprc.dft` for an example.

EMACS: **Caution:** PariEmacs needs to know about the prompt pattern to separate your input from previous `gp` results, without ambiguity. It is not a trivial problem to adapt automatically this regular expression to an arbitrary prompt (which can be self-modifying!). See PariEmacs's documentation.

The default value is `"? "`.

3.4.33 `prompt_cont`. A string that will be printed to prompt for continuation lines (e.g. in between braces, or after a line-terminating backslash). Everything that applies to `prompt` applies to `prompt_cont` as well.

The default value is `""`.

3.4.34 `psfile`. This default is obsolete, use one of `plotexport`, `plotexport` or `plotdrawexport` functions and write the result to file.

3.4.35 `readline`. Switches `readline` line-editing facilities on and off. This may be useful if you are running `gp` in a Sun `cmdtool`, which interacts badly with `readline`. Of course, until `readline` is switched on again, advanced editing features like automatic completion and editing history are not available.

The default value is `1`.

3.4.36 `realbitprecision`. The number of significant bits used to convert exact inputs given to transcendental functions (see Section 3.11), or to create absolute floating point constants (input as `1.0` or `Pi` for instance). Unless you tamper with the `format` default, this is also the number of significant bits used to print a `t_REAL` number; `format` will override this latter behavior, and allow you to have a large internal precision while outputting few digits for instance.

Note that most PARI's functions currently handle precision on a word basis (by increments of 32 or 64 bits), hence bit precision may be a little larger than the number of bits you expected. For instance to get 10 bits of precision, you need one word of precision which, on a 64-bit machine, correspond to 64 bits. To make things even more confusing, this internal bit accuracy is converted to decimal digits when printing floating point numbers: now 64 bits correspond to 19 printed decimal digits ($19 < \log_{10}(2^{64}) < 20$).

The value returned when typing `default(realbitprecision)` is the internal number of significant bits, not the number of printed decimal digits:

```
? default(realbitprecision, 10)
? \pb
      realbitprecision = 64 significant bits
? default(realbitprecision)
%1 = 64
? \p
      realprecision = 3 significant digits
? default(realprecision)
%2 = 19
```

Note that `realprecision` and `\p` allow to view and manipulate the internal precision in decimal digits.

The default value is 128 bits.

3.4.37 realprecision. The number of significant digits used to convert exact inputs given to transcendental functions (see Section 3.11), or to create absolute floating point constants (input as 1.0 or Pi for instance). Unless you tamper with the `format` default, this is also the number of significant digits used to print a `t_REAL` number; `format` will override this latter behavior, and allow you to have a large internal precision while outputting few digits for instance.

Note that PARI's internal precision works on a word basis (by increments of 32 or 64 bits), hence may be a little larger than the number of decimal digits you expected. For instance to get 2 decimal digits you need one word of precision which, on a 64-bit machine, actually gives you 19 digits ($19 < \log_{10}(2^{64}) < 20$). The value returned when typing `default(realprecision)` is the internal number of significant digits, not the number of printed digits:

```
? default(realprecision, 2)
      realprecision = 19 significant digits (2 digits displayed)
? default(realprecision)
%1 = 19
```

The default value is 38 decimal digits.

3.4.38 recover. This toggle is either 1 (on) or 0 (off). If you change this to 0, any error becomes fatal and causes the gp interpreter to exit immediately. Can be useful in batch job scripts.

The default value is 1.

3.4.39 secure. This toggle is either 1 (on) or 0 (off). If on, the `system` and `extern` command are disabled. These two commands are potentially dangerous when you execute foreign scripts since they let gp execute arbitrary UNIX commands. gp will ask for confirmation before letting you (or a script) unset this toggle.

The default value is 0.

3.4.40 seriesprecision. Number of significant terms when converting a polynomial or rational function to a power series (see `\ps`).

The default value is 16.

3.4.41 simplify. This toggle is either 1 (on) or 0 (off). When the PARI library computes something, the type of the result is not always the simplest possible. The only type conversions which the PARI library does automatically are rational numbers to integers (when they are of type `t_FRAC` and equal to integers), and similarly rational functions to polynomials (when they are of type `t_RFRAC` and equal to polynomials). This feature is useful in many cases, and saves time, but can be annoying at times. Hence you can disable this and, whenever you feel like it, use the function `simplify` (see Chapter 3) which allows you to simplify objects to the simplest possible types recursively (see `\y`).

The default value is 1.

3.4.42 `sopath`. This is a list of directories, separated by colons ':' (semicolons ';' in the DOS world, since colons are preempted for drive names). When asked to **install** an external symbol from a shared library whose name is not given by an absolute path (does not start with /, ./ or ../), **gp** will look for it in these directories, in the order they were written in **sopath**. Here, as usual, . means the current directory, and .. its immediate parent. Environment expansion is performed.

The default value is "", corresponding to an empty list of directories: **install** will use the library name as input (and look in the current directory if the name is not an absolute path).

3.4.43 `strictargs`. This toggle is either 1 (on) or 0 (off). If on, all arguments to *new* user functions are mandatory unless the function supplies an explicit default value. Otherwise arguments have the default value 0.

In this example,

```
fun(a,b=2)=a+b
```

a is mandatory, while **b** is optional. If **strictargs** is on:

```
? fun()
*** at top-level: fun()
***           ^-----
*** in function fun: a,b=2
***           ^-----
*** missing mandatory argument 'a' in user function.
```

This applies to functions defined while **strictargs** is on. Changing **strictargs** does not affect the behavior of previously defined functions.

The default value is 0.

3.4.44 `strictmatch`. Obsolete. This toggle is now a no-op.

3.4.45 `threadsize`. This default is specific to the *parallel* version of PARI and gp (built via `Configure --mt=pthread` or `mpi`) and is ignored otherwise. In parallel mode, each thread allocates its own private *stack* for its computations, see **parisize**. This value determines the size in bytes of the stacks of each thread, so the total memory allocated will be **parisize** + **nbthreads** × **threadsize**.

If set to 0, the value used is the same as **parisize**. It is not easy to estimate reliably a sufficient value for this parameter because PARI itself will parallelize computations and we recommend to not set this value explicitly unless it solves a specific problem for you. For instance if you see frequent messages of the form

```
*** Warning: not enough memory, new thread stack 10000002048
```

(Meaning that **threadsize** had to be temporarily increased.) On the other hand we strongly recommend to set **parisize**max and **threadsize**max to a nonzero value.

The default value is 0.

3.4.46 threadsize_{max}. This default is specific to the *parallel* version of PARI and gp (built via `Configure --mt=pthread` or `mpi`) and is ignored otherwise. In parallel mode, each threads allocates its own private *stack* for its computations, see `parisize` and `parisizemax`. The values of `threadsize` and `threadsizemax` determine the usual and maximal size in bytes of the stacks of each thread, so the total memory allocated will be between `parisize + nbthreads × threadsize`. and `parisizemax + nbthreads × threadsizemax`.

If set to 0, the value used is the same as `threadsize`. We strongly recommend to set both `parisizemax` and `threadsizemax` to a nonzero value.

The default value is 0.

3.4.47 timer. This toggle is either 1 (on) or 0 (off). Every instruction sequence in the gp calculator (anything ended by a newline in your input) is timed, to some accuracy depending on the hardware and operating system. When `timer` is on, each such timing is printed immediately before the output as follows:

```
? factor(2^2^7+1)
time = 108 ms.      \\ this line omitted if 'timer' is 0
%1 =
[      59649589127497217 1]
[5704689200685129054721 1]
```

(See also `#` and `##`.)

The time measured is the user CPU time, not including the time for printing the results. If the time is negligible (< 1 ms.), nothing is printed: in particular, no timing should be printed when defining a user function or an alias, or installing a symbol from the library.

If you are using a parallel version of gp, the output is more complex, such as

```
? isprime( 10^300 + 331 )
cpu time = 3,206 ms, real time = 1,289 ms. \\ omitted if 'timer' is 0
%1 = 1
```

Now, `real time` is the wallclock time, and `cpu time` is the sum of the CPU times spent by the different threads.

The default value is 0 (off).

3.5 Standard monadic or dyadic operators.

3.5.1 Boolean operators.

Any nonzero value is interpreted as *true* and any zero as *false* (this includes empty vectors or matrices). The standard boolean operators `||` (inclusive or), `&&` (and) and `!` in prefix notation (not) are available. Their value is 1 (true) or 0 (false):

```
? a && b    \\ 1 iff a and b are nonzero
? a || b    \\ 1 iff a or b is nonzero
? !a        \\ 1 iff a is zero
```

3.5.2 Comparison. The standard real comparison operators $\leq$, $<$, $\geq$, $>$, are available in GP. The result is 1 if the comparison is true, 0 if it is false. These operators allow to compare integers ($\mathbf{t_INT}$), rational ($\mathbf{t_FRAC}$) or real ($\mathbf{t_REAL}$) numbers, real quadratic numbers ($\mathbf{t_QUAD}$ of positive discriminant) and infinity (∞ , $\mathbf{t_INFINITY}$).

By extension, two character strings ($\mathbf{t_STR}$) are compared using the standard lexicographic order. Comparing a string to an object of a different type raises an exception. See also the `cmp` universal comparison function.

3.5.3 Equality. Two operators allow to test for equality: `==` (equality up to type coercion) and `===` (identity). The result is 1 if equality is decided, else 0.

The operator `===` is strict: objects of different type or length are never identical, polynomials in different variables are never identical, even if constant. On the contrary, `==` is very liberal: $a == b$ decides whether there is a natural map sending a to the domain of b or sending b to the domain of a , such that the comparison makes sense and equality holds. For instance

```
? 4 == Mod(1,3) \\ equal
%1 = 1
? 4 === Mod(1,3) \\ but not identical
%2 = 0
? 'x == 'y \\ not equal (nonconstant and different variables)
%3 = 0
? Pol(0,'x) == Pol(0,'y) \\ equal (constant: ignore variable)
%4 = 1
? Pol(0,'x) === Pol(0,'y) \\ not identical
%5 = 0
? 0 == Pol(0) \\ equal (not identical)
%6 = 1
? [0] == 0 \\ equal (not identical)
%7 = 1
? [0, 0] == 0 \\ equal (not identical)
%8 = 1
? [0] == [0,0] \\ not equal
%9 = 0
```

In particular `==` is not transitive in general. The operator `===` is transitive. The `==` operator allows two equivalent negated forms: `!=` or `<>`; there is no negated form for `===`.

Do not mistake `=` for `==`: the former is the assignment statement.

3.5.4 +/-. The expressions $+x$ and $-x$ refer to monadic operators: the first does nothing, the second negates x .

The library syntax is `GEN gneg(GEN x)` for $-x$.

3.5.5 +. The expression $x + y$ is the sum of x and y . Addition between a scalar type x and a $\mathbf{t_COL}$ or $\mathbf{t_MAT}$ y returns respectively $[y[1] + x, y[2], \dots]$ and $y + x\text{Id}$. Other additions between a scalar type and a vector or a matrix, or between vector/matrices of incompatible sizes are forbidden.

The library syntax is `GEN gadd(GEN x, GEN y)`.

3.5.6 -. The expression $x - y$ is the difference of x and y . Subtraction between a scalar type x and a `t_COL` or `t_MAT` y returns respectively $[y[1] - x, y[2], \dots]$ and $y - x\text{Id}$. Other subtractions between a scalar type and a vector or a matrix, or between vector/matrices of incompatible sizes are forbidden.

The library syntax is `GEN gsub(GEN x, GEN y)` for $x - y$.

3.5.7 *. The expression $x * y$ is the product of x and y . Among the prominent impossibilities are multiplication between vector/matrices of incompatible sizes, between a `t_INTMOD` or `t_PADIC`. Restricted to scalars, $*$ is commutative; because of vector and matrix operations, it is not commutative in general.

Multiplication between two `t_VECs` or two `t_COLs` is not allowed; to take the scalar product of two vectors of the same length, transpose one of the vectors (using the operator `~` or the function `mattranspose`, see Section 3.10) and multiply a row vector by a column vector:

```
? a = [1,2,3];
? a * a
***   at top-level: a*a
***               ^--
*** _*_: forbidden multiplication t_VEC * t_VEC.
? a * a~
%2 = 14
```

If x, y are binary quadratic forms, compose them; see also `qfbnucomp` and `qfbnupow`. If x, y are `t_VECSMALL` of the same length, understand them as permutations and compose them.

The library syntax is `GEN gmul(GEN x, GEN y)` for $x * y$. Also available is `GEN gsqr(GEN x)` for $x * x$.

3.5.8 /. The expression x / y is the quotient of x and y . In addition to the impossibilities for multiplication, note that if the divisor is a matrix, it must be an invertible square matrix, and in that case the result is $x*y^{-1}$. Furthermore note that the result is as exact as possible: in particular, division of two integers always gives a rational number (which may be an integer if the quotient is exact) and *not* the Euclidean quotient (see $x \setminus y$ for that), and similarly the quotient of two polynomials is a rational function in general. To obtain the approximate real value of the quotient of two integers, add `0.` to the result; to obtain the approximate p -adic value of the quotient of two integers, add `O(p^k)` to the result; finally, to obtain the Taylor series expansion of the quotient of two polynomials, add `O(X^k)` to the result or use the `taylor` function (see Section 3.9.63).

The library syntax is `GEN gdiv(GEN x, GEN y)` for x / y .

3.5.9 \. The expression $x \setminus y$ is the Euclidean quotient of x and y . If y is a real scalar, this is defined as `floor(x/y)` if $y > 0$, and `ceil(x/y)` if $y < 0$ and the division is not exact. Hence the remainder $x - (x \setminus y)*y$ is in $[0, |y|]$.

Note that when y is an integer and x a polynomial, y is first promoted to a polynomial of degree 0. When x is a vector or matrix, the operator is applied componentwise.

The library syntax is `GEN gdivent(GEN x, GEN y)` for $x \setminus y$.

3.5.10 $\backslash/$. The expression $x \backslash/ y$ evaluates to the rounded Euclidean quotient of x and y . This is the same as $x \setminus y$ except for scalar division: the quotient is such that the corresponding remainder is smallest in absolute value and in case of a tie the quotient closest to $+\infty$ is chosen (hence the remainder would belong to $[-|y|/2, |y|/2]$).

When x is a vector or matrix, the operator is applied componentwise.

The library syntax is `GEN gdivround(GEN x, GEN y)` for $x \backslash/ y$.

3.5.11 $\%$. The expression $x \% y$ evaluates to the modular Euclidean remainder of x and y , which we now define. When x or y is a nonintegral real number, $x\%y$ is defined as $x - (x\setminus y)*y$. Otherwise, if y is an integer, this is the smallest nonnegative integer congruent to x modulo y . (This actually coincides with the previous definition if and only if x is an integer.) If y is a polynomial, this is the polynomial of smallest degree congruent to x modulo y . For instance:

```
? (1/2) % 3
%1 = 2
? 0.5 % 3
%2 = 0.50000000000000000000000000000000
? (1/2) % 3.0
%3 = 1/2
```

Note that when y is an integer and x a polynomial, y is first promoted to a polynomial of degree 0. When x is a vector or matrix, the operator is applied componentwise.

The library syntax is `GEN gmod(GEN x, GEN y)` for $x \% y$.

3.5.12 $!$. The expression $n!$ is the factorial of the non-negative integer n .

The library syntax is `GEN mpfact(long n)`

3.5.13 $\#$. The expression $n\#$ is the primorial of the non-negative integer n , that is the product of all prime numbers less than or equal to x .

The library syntax is `GEN mpprimorial(long n)`

3.5.14 $\text{op}=\text{}$. When op is a binary arithmetic operator among $+$, $-$, $*$, $\%$, $/$, $\setminus$ or $\backslash/$, the construct $x \text{op} = y$ is a shortcut for $x = x \text{op} y$.

```
? v[1] += 10 \ \ increment v[1] by 10
? a /= 2 \ \ divide a by 2
```

3.5.15 $++$. $x++$ is a shortcut for $x = x + 1$ and for $x += 1$.

3.5.16 $--$. $x--$ is a shortcut for $x = x - 1$ and for $x -= 1$.

3.5.17 $\wedge$. The expression x^n is powering.

- If the exponent n is an integer, then exact operations are performed using binary (left-shift) powering techniques. By definition, x^0 is (an empty product interpreted as) an exact 1 in the underlying prime ring:

```
? 0.0 ^ 0
%1 = 1
? (1 + 0(2^3)) ^ 0
%2 = 1
? (1 + 0(x)) ^ 0
%3 = 1
? Mod(2,4)^0
%4 = Mod(1,4)
? Mod(x,x^2)^0
%5 = Mod(1, x^2)
```

If x is a p -adic number, its precision will increase if $v_p(n) > 0$ and $n \neq 0$. Powering a binary quadratic form (type `t_QFB`) returns a representative of the class, which is reduced if the input was. (In particular, x^1 returns x itself, whether it is reduced or not.)

PARI rewrites the multiplication $x * x$ of two *identical* objects as x^2 . Here, identical means the operands are reference the same chunk of memory; no equality test is performed. This is no longer true when more than two arguments are involved.

```
? a = 1 + 0(2); b = a;
? a * a \\ = a^2, precision increases
%2 = 1 + 0(2^3)
? a * b \\ not rewritten as a^2
%3 = 1 + 0(2)
? a*a*a \\ not rewritten as a^3
%4 = 1 + 0(2)
```

- If the exponent is a rational number p/q the behaviour depends on x . If x is a complex number, return $\exp(n \log x)$ (principal branch), in an exact form if possible:

```
? 4^(1/2) \\ 4 being a square, this is exact
%1 = 2
? 2^(1/2) \\ now inexact
%2 = 1.4142135623730950488016887242096980786
? (-1/4)^(1/2) \\ exact again
%3 = 1/2*I
? (-1)^(1/3)
%4 = 0.500...+ 0.866...*I
```

Note that even though -1 is an exact cube root of -1 , it is not $\exp(\log(-1)/3)$; the latter is returned.

Otherwise return a solution y of $y^q = x^p$ if it exists; beware that this is defined up to q -th roots of 1 in the base field. Intmods modulo composite numbers are not supported.

```
? Mod(7,19)^(1/2)
%1 = Mod(11, 19) \\ is any square root
? sqrt(Mod(7,19))
```

```
%2 = Mod(8, 19)  \\ is the smallest square root
? Mod(1,4)^(1/2)
*** at top-level: Mod(1,4)^(1/2)
*** ^-----
*** _^_: not a prime number in gpow: 4.
```

• If the exponent is a negative integer or rational number, an inverse must be computed. For noninvertible $t_INTMOD\ x$, this will fail and (for n an integer) implicitly exhibit a factor of the modulus:

```
? Mod(4,6)^(-1)
*** at top-level: Mod(4,6)^(-1)
*** ^-----
*** _^_: impossible inverse modulo: Mod(2, 6).
```

Here, a factor 2 is obtained directly. In general, take the gcd of the representative and the modulus. This is most useful when performing complicated operations modulo an integer N whose factorization is unknown. Either the computation succeeds and all is well, or a factor d is discovered and the computation may be restarted modulo d or N/d .

For noninvertible $t_POLMOD\ x$, the behavior is the same:

```
? Mod(x^2, x^3-x)^(-1)
*** at top-level: Mod(x^2,x^3-x)^(-1)
*** ^-----
*** _^_: impossible inverse in RgXQ_inv: Mod(x^2, x^3 - x).
```

Note that the underlying algorithm (subresultant) assumes that the base ring is a domain:

```
? a = Mod(3*y^3+1, 4); b = y^6+y^5+y^4+y^3+y^2+y+1; c = Mod(a,b);
? c^(-1)
*** at top-level: Mod(a,b)^(-1)
*** ^-----
*** _^_: impossible inverse modulo: Mod(2, 4).
```

In fact c is invertible, but $\mathbf{Z}/4\mathbf{Z}$ is not a domain and the algorithm fails. It is possible for the algorithm to succeed in such situations and any returned result will be correct, but chances are that an error will occur first. In this specific case, one should work with 2-adics. In general, one can also try the following approach

```
? inversemod(a, b) =
{ my(m, v = variable(b));
  m = polysylvestermatrix(polrecip(a), polrecip(b));
  m = matinverseimage(m, matid(#m)[,1]);
  Polrev(m[1..poldegree(b)], v);
}
? inversemod(a,b)
%2 = Mod(2,4)*y^5 + Mod(3,4)*y^3 + Mod(1,4)*y^2 + Mod(3,4)*y + Mod(2,4)
```

This is not guaranteed to work either since `matinverseimage` must also invert pivots. See Section 3.10.

For a $t_MAT\ x$, the matrix is expected to be square and invertible, except in the special case $x^{(-1)}$ which returns a left inverse if one exists (rectangular x with full column rank).

```
? x = Mat([1;2])
%1 =
[1]
[2]
? x^(-1)
%2 =
[1 0]
```

• Finally, if the exponent n is not an rational number, powering is treated as the transcendental function $\exp(n \log x)$, although it will be more precise than the latter when n and x are exact:

```
? s = 1/2 + 10^14 * I
? localprec(200); z = 2^s  \\ for reference
? exponent(2^s - z)
%3 = -127  \\ perfect
? exponent(exp(s * log(2)) - z)
%4 = -84  \\ not so good
```

The second computation is less precise because $\log(2)$ is first computed to 38 decimal digits, then multiplied by s , which has a huge imaginary part amplifying the error.

In this case, $x \mapsto x^n$ is treated as a transcendental function and in particular acts componentwise on vector or matrices, even square matrices ! (See Section 3.11.) If x is 0 and n is an inexact 0, this will raise an exception:

```
? 4 ^ 1.0
%1 = 4.0000000000000000000000000000000000000000000000000000000
? 0^ 0.0
*** at top-level: 0^0.0
*** ^-----
*** _^_: domain error in gpow(0,n): n <= 0
```

The library syntax is GEN gpow(GEN x, GEN n, long prec) for x^n .

3.5.18 cmp(x, y). Gives the result of a comparison between arbitrary objects x and y (as -1 , 0 or 1). The underlying order relation is transitive, the function returns 0 if and only if $x == y$. It has no mathematical meaning but satisfies the following properties when comparing entries of the same type:

- two `t_INTs` compare as usual (i.e. $\text{cmp}(x, y) < 0$ if and only if $x < y$);
- two `t_VECSMALLs` of the same length compare lexicographically;
- two `t_STRs` compare lexicographically.

In case all components are equal up to the smallest length of the operands, the more complex is considered to be larger. More precisely, the longest is the largest; when lengths are equal, we have matrix $>$ vector $>$ scalar. For example:

```
? cmp(1, 2)
%1 = -1
? cmp(2, 1)
%2 = 1
? cmp(1, 1.0)  \\ note that 1 == 1.0, but (1==1.0) is false.
```

```
%3 = -1
? cmp(x + Pi, [])
%4 = -1
```

This function is mostly useful to handle sorted lists or vectors of arbitrary objects. For instance, if v is a vector, the construction `vecsor(v, cmp)` is equivalent to `Set(v)`.

The library syntax is `int cmp_universal(GEN x, GEN y)`.

3.5.19 divrem($x, y, \{v\}$). Creates a column vector with two components, the first being the Euclidean quotient ($x \setminus y$), the second the Euclidean remainder ($x - (x \setminus y) * y$), of the division of x by y . This avoids the need to do two divisions if one needs both the quotient and the remainder. If v is present, and x, y are multivariate polynomials, divide with respect to the variable v .

Beware that `divrem(x, y) [2]` is in general not the same as $x \% y$; no GP operator corresponds to it:

```
? divrem(1/2, 3) [2]
%1 = 1/2
? (1/2) % 3
%2 = 2
? divrem(Mod(2,9), 3) [2]
*** at top-level: divrem(Mod(2,9), 3) [2]
***      ^-----
*** forbidden division t_INTMOD \ t_INT.
? Mod(2,9) % 6
%3 = Mod(2,3)
```

The library syntax is `GEN divrem(GEN x, GEN y, long v = -1)` where v is a variable number. Also available is `GEN gdiventres(GEN x, GEN y)` when v is not needed.

3.5.20 lex(x, y). Gives the result of a lexicographic comparison between x and y (as $-1, 0$ or 1). This is to be interpreted in quite a wide sense: it is admissible to compare objects of different types (scalars, vectors, matrices), provided the scalars can be compared, as well as vectors/matrices of different lengths; finally, when comparing two scalars, a complex number $a + I * b$ is interpreted as a vector $[a, b]$ and a real number a as $[a, 0]$. The comparison is recursive.

In case all components are equal up to the smallest length of the operands, the more complex is considered to be larger. More precisely, the longest is the largest; when lengths are equal, we have `matrix > vector > scalar`. For example:

```
? lex([1,3], [1,2,5])
%1 = 1
? lex([1,3], [1,3,-1])
%2 = -1
? lex([1], [[1]])
%3 = -1
? lex([1], [1]~)
%4 = 0
? lex(2 - I, 1)
%5 = 1
? lex(2 - I, 2)
```

`%6 = -1`

The library syntax is `int lexcmp(GEN x, GEN y)`.

3.5.21 `max(x, y)`. Creates the maximum of x and y when they can be compared.

The library syntax is `GEN gmax(GEN x, GEN y)`.

3.5.22 `min(x, y)`. Creates the minimum of x and y when they can be compared.

The library syntax is `GEN gmin(GEN x, GEN y)`.

3.5.23 `shift(x, n)`. Shifts x componentwise left by n bits if $n \geq 0$ and right by $|n|$ bits if $n < 0$. May be abbreviated as $x \ll n$ or $x \gg (-n)$. A left shift by n corresponds to multiplication by 2^n . A right shift of an integer x by $|n|$ corresponds to a Euclidean division of x by $2^{|n|}$ with a remainder of the same sign as x , hence is not the same (in general) as $x \setminus 2^n$.

The library syntax is `GEN gshift(GEN x, long n)`.

3.5.24 `shiftmul(x, n)`. Multiplies x by 2^n . The difference with `shift` is that when $n < 0$, ordinary division takes place, hence for example if x is an integer the result may be a fraction, while for shifts Euclidean division takes place when $n < 0$ hence if x is an integer the result is still an integer.

The library syntax is `GEN gmul2n(GEN x, long n)`.

3.5.25 `sign(x)`. `sign` (0, 1 or -1) of x , which must be of type integer, real or fraction; `t_QUAD` with positive discriminants and `t_INFINITY` are also supported.

The library syntax is `int gsigne(GEN x)`.

3.5.26 `vecmax(x, {&v})`. If x is a list, vector or matrix, returns the largest entry of x , otherwise returns a copy of x . Error if x is empty. Here, largest refers to the ordinary real ordering ($\leq$).

If v is given, set it to the index of a largest entry (indirect maximum), when x is a vector or list. If x is a matrix, set v to coordinates $[i, j]$ such that $x[i, j]$ is a largest entry. This argument v is ignored for other types. When the vector has equal largest entries, the first occurrence is chosen; in a matrix, the smallest j is chosen first, then the smallest i . vector or matrix.

```
? vecmax([10, 20, -30, 40])
%1 = 40
? vecmax([10, 20, -30, 40], &v); v
%2 = 4
? vecmax([10, 20; -30, 40], &v); v
%3 = [2, 2]
```

The library syntax is `GEN vecmax0(GEN x, GEN *v = NULL)`. When v is not needed, the function `GEN vecmax(GEN x)` is also available.

3.5.27 vecmin($x, \{\&v\}$). If x is a list, vector or matrix, returns the smallest entry of x , otherwise returns a copy of x . Error if x is empty. Here, smallest refers to the ordinary real ordering ($\leq$).

If v is given, set it to the index of a smallest entry (indirect minimum), when x is a vector or list. If x is a matrix, set v to coordinates $[i, j]$ such that $x[i, j]$ is a smallest entry. This argument v is ignored for other types. When a vector has equal smallest entries, the first occurrence is chosen; in a matrix, the smallest j is chosen first, then the smallest i .

```
? vecmin([10, 20, -30, 40])
%1 = -30
? vecmin([10, 20, -30, 40], &v); v
%2 = 3
? vecmin([10, 20; -30, 40], &v); v
%3 = [2, 1]
? vecmin([1,0;0,0], &v); v
%3 = [2, 1]
```

The library syntax is GEN `vecmin0(GEN x, GEN *v = NULL)`. When v is not needed, the function GEN `vecmin(GEN x)` is also available.

3.6 Conversions and similar elementary functions or commands.

Many of the conversion functions are rounding or truncating operations. In this case, if the argument is a rational function, the result is the Euclidean quotient of the numerator by the denominator, and if the argument is a vector or a matrix, the operation is done componentwise. This will not be restated for every function.

3.6.1 Col($x, \{n\}$). Transforms the object x into a column vector. The dimension of the resulting vector can be optionally specified via the extra parameter n .

If n is omitted or 0, the dimension depends on the type of x ; the vector has a single component, except when x is

- a vector or a quadratic form (in which case the resulting vector is simply the initial object considered as a row vector),
- a polynomial or a power series. In the case of a polynomial, the coefficients of the vector start with the leading coefficient of the polynomial, while for power series only the significant coefficients are taken into account, but this time by increasing order of degree. In this last case, **Vec** is the reciprocal function of **Pol** and **Ser** respectively,
- a matrix (the column of row vector comprising the matrix is returned),
- a character string (a vector of individual characters is returned).

In the last two cases (matrix and character string), n is meaningless and must be omitted or an error is raised. Otherwise, if n is given, 0 entries are appended at the end of the vector if $n > 0$, and prepended at the beginning if $n < 0$. The dimension of the resulting vector is $|n|$.

See `??Vec` for examples and further details.

The library syntax is GEN `gtocol0(GEN x, long n)`. GEN `gtocol(GEN x)` is also available.

3.6.2 Colrev($x, \{n\}$). As **Col**($x, -n$), then reverse the result. In particular, **Colrev** is the reciprocal function of **Polrev**: the coefficients of the vector start with the constant coefficient of the polynomial and the others follow by increasing degree.

The library syntax is **GEN gtocolrev0**(**GEN x**, **long n**). **GEN gtocolrev**(**GEN x**) is also available.

3.6.3 List($\{x = []\}$). Transforms a (row or column) vector x into a list, whose components are the entries of x . Similarly for a list, but rather useless in this case. For other types, creates a list with the single element x .

The library syntax is **GEN gtolist**(**GEN x = NULL**). The variant **GEN mklist**(**void**) creates an empty list.

3.6.4 Map($\{x\}$). A “Map” is an associative array, or dictionary: a data type composed of a collection of (*key*, *value*) pairs, such that each key appears just once in the collection. This function converts the matrix $[a_1, b_1; a_2, b_2; \dots; a_n, b_n]$ to the map $a_i \mapsto b_i$.

```
? M = Map(factor(13!));
? mapget(M, 3)
%2 = 5
? P = Map(matreduce(primes([1,20])))
%3 = Map([2,1;3,1;5,1;7,1;11,1;13,1;17,1;19,1])
? select(i->mapisdefined(P,i), [1..20])
%4 = [2, 3, 5, 7, 11, 13, 17, 19]
```

If the argument x is omitted, creates an empty map, which may be filled later via **mapput**.

The library syntax is **GEN gtomap**(**GEN x = NULL**).

3.6.5 Mat($\{x = []\}$). Transforms the object x into a matrix. If x is already a matrix, a copy of x is created. If x is a row (resp. column) vector, this creates a 1-row (resp. 1-column) matrix, *unless* all elements are column (resp. row) vectors of the same length, in which case the vectors are concatenated sideways and the attached big matrix is returned. If x is a binary quadratic form, creates the attached 2×2 matrix. Otherwise, this creates a 1×1 matrix containing x .

```
? Mat(x + 1)
%1 =
[x + 1]
? Vec( matid(3) )
%2 = [[1, 0, 0]~, [0, 1, 0]~, [0, 0, 1]~]
? Mat(%)
%3 =
[1 0 0]
[0 1 0]
[0 0 1]
? Col( [1,2; 3,4] )
%4 = [[1, 2], [3, 4]]~
? Mat(%)
%5 =
[1 2]
```

```

[3 4]
? Mat(Qfb(1,2,3))
%6 =
[1 1]
[1 3]

```

The library syntax is GEN `gtomat(GEN x = NULL)`.

3.6.6 Mod(a, b). In its basic form, create an `intmod` or a `polmod` ($a \bmod b$); b must be an integer or a polynomial. We then obtain a `t_INTMOD` and a `t_POLMOD` respectively:

```

? t = Mod(2,17); t^8
%1 = Mod(1, 17)
? t = Mod(x,x^2+1); t^2
%2 = Mod(-1, x^2+1)

```

If $a \% b$ makes sense and yields a result of the appropriate type (`t_INT` or `scalar/t_POL`), the operation succeeds as well:

```

? Mod(1/2, 5)
%3 = Mod(3, 5)
? Mod(7 + O(3^6), 3)
%4 = Mod(1, 3)
? Mod(Mod(1,12), 9)
%5 = Mod(1, 3)
? Mod(1/x, x^2+1)
%6 = Mod(-x, x^2+1)
? Mod(exp(x), x^4)
%7 = Mod(1/6*x^3 + 1/2*x^2 + x + 1, x^4)

```

If a is a complex object, “base change” it to $\mathbf{Z}/b\mathbf{Z}$ or $K[x]/(b)$, which is equivalent to, but faster than, multiplying it by `Mod(1,b)`:

```

? Mod([1,2;3,4], 2)
%8 =
[Mod(1, 2) Mod(0, 2)]
[Mod(1, 2) Mod(0, 2)]
? Mod(3*x+5, 2)
%9 = Mod(1, 2)*x + Mod(1, 2)
? Mod(x^2 + y*x + y^3, y^2+1)
%10 = Mod(1, y^2 + 1)*x^2 + Mod(y, y^2 + 1)*x + Mod(-y, y^2 + 1)

```

This function is not the same as $x \% y$, the result of which has no knowledge of the intended modulus y . Compare

```

? x = 4 % 5; x + 1
%11 = 5
? x = Mod(4,5); x + 1
%12 = Mod(0,5)

```

Note that such “modular” objects can be lifted via `lift` or `centerlift`. The modulus of a `t_INTMOD` or `t_POLMOD` z can be recovered via `z.mod`.

The library syntax is GEN `gmodulo(GEN a, GEN b)`.

3.6.7 Pol($t, \{v = 'x\}$). Transforms the object t into a polynomial with main variable v . If t is a scalar, this gives a constant polynomial. If t is a power series with nonnegative valuation or a rational function, the effect is similar to **truncate**, i.e. we chop off the $O(X^k)$ or compute the Euclidean quotient of the numerator by the denominator, then change the main variable of the result to v .

The main use of this function is when t is a vector: it creates the polynomial whose coefficients are given by t , with $t[1]$ being the leading coefficient (which can be zero). It is much faster to evaluate **Pol** on a vector of coefficients in this way, than the corresponding formal expression $a_n X^n + \dots + a_0$, which is evaluated naively exactly as written (linear versus quadratic time in n). **Polrev** can be used if one wants $x[1]$ to be the constant coefficient:

```
? Pol([1,2,3])
%1 = x^2 + 2*x + 3
? Polrev([1,2,3])
%2 = 3*x^2 + 2*x + 1
```

The reciprocal function of **Pol** (resp. **Polrev**) is **Vec** (resp. **Vecrev**).

```
? Vec(Pol([1,2,3]))
%1 = [1, 2, 3]
? Vecrev( Polrev([1,2,3]) )
%2 = [1, 2, 3]
```

Warning. This is *not* a substitution function. It will not transform an object containing variables of higher priority than v .

```
? Pol(x + y, y)
***   at top-level: Pol(x+y,y)
***           ^-----
*** Pol: variable must have higher priority in gtopoly.
```

The library syntax is GEN **gtopoly**(GEN t , long $v = -1$) where v is a variable number.

3.6.8 Polrev($t, \{v = 'x\}$). Transform the object t into a polynomial with main variable v . If t is a scalar, this gives a constant polynomial. If t is a power series, the effect is identical to **truncate**, i.e. it chops off the $O(X^k)$.

The main use of this function is when t is a vector: it creates the polynomial whose coefficients are given by t , with $t[1]$ being the constant term. **Pol** can be used if one wants $t[1]$ to be the leading coefficient:

```
? Polrev([1,2,3])
%1 = 3*x^2 + 2*x + 1
? Pol([1,2,3])
%2 = x^2 + 2*x + 3
```

The reciprocal function of **Pol** (resp. **Polrev**) is **Vec** (resp. **Vecrev**).

The library syntax is GEN **gtopolyrev**(GEN t , long $v = -1$) where v is a variable number.

3.6.9 Qfb($a, \{b\}, \{c\}$). Creates the binary quadratic form $ax^2 + bxy + cy^2$. Negative definite forms are not implemented, use their positive definite counterpart instead. The syntax **Qfb**(V) is also allowed with V being either a **t_VEC** $[a, b, c]$, a **t_POL** $ax^2 + bx + c$ or a **t_MAT** $[a, b_0; b_1, c]$ with $b_0 + b_1 = b$.

The library syntax is **GEN Qfb0**(**GEN a**, **GEN b** = **NULL**, **GEN c** = **NULL**).

3.6.10 Ser($s, \{v = 'x\}, \{d = \text{seriesprecision}\}$). Transforms the object s into a power series with main variable v (x by default) and precision (number of significant terms) equal to $d \geq 0$ ($d = \text{seriesprecision}$ by default). If s is a scalar, this gives a constant power series in v with precision d . If s is a polynomial, the polynomial is truncated to d terms if needed

```
? \ps
  seriesprecision = 16 significant terms
? Ser(1)  \\ 16 terms by default
%1 = 1 + 0(x^16)
? Ser(1, 'y, 5)
%2 = 1 + 0(y^5)
? Ser(x^2,, 5)
%3 = x^2 + 0(x^7)
? T = polcyclo(100)
%4 = x^40 - x^30 + x^20 - x^10 + 1
? Ser(T, 'x, 11)
%5 = 1 - x^10 + 0(x^11)
```

The function is more or less equivalent with multiplication by $1 + O(v^d)$ in these cases, only faster.

For the remaining types, vectors and power series, we first explain what occurs if d is omitted. In this case, the function uses exactly the amount of information given in the input:

- If s is already a power series in v , we return it verbatim;
- If s is a vector, the coefficients of the vector are understood to be the coefficients of the power series starting from the constant term (as in **Polrev**(x)); in other words we convert **t_VEC** / **t_COL** to the power series whose significant terms are exactly given by the vector entries.

On the other hand, if d is explicitly given, we abide by its value and return a series, truncated or extended with zeros as needed, with d significant terms.

```
? v = [1,2,3];
? Ser(v, t)  \\ 3 terms: seriesprecision is ignored!
%7 = 1 + 2*t + 3*t^2 + 0(t^3)
? Ser(v, t, 7)  \\ 7 terms as explicitly requested
%8 = 1 + 2*t + 3*t^2 + 0(t^7)
? s = 1+x+0(x^2);
? Ser(s)
%10 = 1 + x + 0(x^2)  \\ 2 terms: seriesprecision is ignored
? Ser(s, x, 7)  \\ extend to 7 terms
%11 = 1 + x + 0(x^7)
? Ser(s, x, 1)  \\ truncate to 1 term
%12 = 1 + 0(x)
```

The warning given for **Pol** also applies here: this is not a substitution function.

The library syntax is `GEN Ser0(GEN s, long v = -1, GEN d = NULL, long precd1)` where `v` is a variable number.

3.6.11 Set($\{x = []\}$). Converts x into a set, i.e. into a row vector, with strictly increasing entries with respect to the (somewhat arbitrary) universal comparison function `cmp`. Standard container types `t_VEC`, `t_COL`, `t_LIST` and `t_VECSMALL` are converted to the set with corresponding elements. All others are converted to a set with one element.

```
? Set([1,2,4,2,1,3])
%1 = [1, 2, 3, 4]
? Set(x)
%2 = [x]
? Set(Vecsmall([1,3,2,1,3]))
%3 = [1, 2, 3]
```

The library syntax is `GEN gtoset(GEN x = NULL)`.

3.6.12 Str($\{x\}$). Converts its argument list into a single character string (type `t_STR`, the empty string if x is omitted). To recover an ordinary `GEN` from a string, apply `eval` to it. The arguments of `Str` are evaluated in string context, see Section 2.9.1.

```
? x2 = 0; i = 2; Str(x, i)
%1 = "x2"
? eval(%)
%2 = 0
```

This function is mostly useless in library mode. Use the pair `strtoGEN`/`GENtostr` to convert between `GEN` and `char*`. The latter returns a malloced string, which should be freed after usage.

The library syntax is `GEN Str(GEN vec_x)`.

3.6.13 Vec($x, \{n\}$). Transforms the object x into a row vector. The dimension of the resulting vector can be optionally specified via the extra parameter n . If n is omitted or 0, the dimension depends on the type of x ; the vector has a single component, except when x is

- a vector or a quadratic form: returns the initial object considered as a row vector,
- a polynomial or a power series: returns a vector consisting of the coefficients. In the case of a polynomial, the coefficients of the vector start with the leading coefficient of the polynomial, while for power series only the significant coefficients are taken into account, but this time by increasing order of degree. In particular the valuation is ignored (which makes the function useful for series of negative valuation):

```
? Vec(3*x^2 + x)
%1 = [3, 1, 0]
? Vec(x^2 + 3*x^3 + 0(x^5))
%2 = [1, 3, 0]
? Vec(x^-2 + 3*x^-1 + 0(x))
%3 = [1, 3, 0]
```

`Vec` is the reciprocal function of `Pol` for a polynomial and of `Ser` for power series of valuation 0.

- a matrix: returns the vector of columns comprising the matrix,

```
? m = [1,2,3;4,5,6]
```

```
%4 =
[1 2 3]
[4 5 6]
? Vec(m)
%5 = [[1, 4]~, [2, 5]~, [3, 6]~]
```

- a character string: returns the vector of individual characters (as strings of length 1),

```
? Vec("PARI")
%6 = ["P", "A", "R", "I"]
```

- a map: returns the vector of the domain of the map,
- an error context (`t_ERROR`): returns the error components, see `iferr`.

In the last four cases (matrix, character string, map, error), n is meaningless and must be omitted or an error is raised. Otherwise, if n is given, 0 entries are appended at the end of the vector if $n > 0$, and prepended at the beginning if $n < 0$. The dimension of the resulting vector is $|n|$. If the original object had fewer than $|n|$ components, it is truncated from the right if $n > 0$ and from the left if $n < 0$:

```
? v = [1,2,3,4];
? forstep(i=5, 2, -1, print(Vec(v, i)));
[1, 2, 3, 4, 0]
[1, 2, 3, 4]
[1, 2, 3]  \\ truncated from the right
[1, 2]

? forstep(i=5, 2, -1, print(Vec(v, -i)));
[0, 1, 2, 3, 4]
[1, 2, 3, 4]
[2, 3, 4]  \\ truncated from the left
[3, 4]
```

These rules allow to write a conversion function for series that takes positive valuations into account:

```
? serVec(s) = Vec(s, -serprec(s,variable(s)));
? serVec(x^2 + 3*x^3 + 0(x^5))
%2 = [0, 0, 1, 3, 0]
```

(That function is not intended for series of negative valuation.)

The library syntax is GEN `gtovect0(GEN x, long n)`. GEN `gtovect(GEN x)` is also available.

3.6.14 Vecrev($x, \{n\}$). As `Vec`($x, -n$), then reverse the result. In particular, `Vecrev` is the reciprocal function of `Polrev`: the coefficients of the vector start with the constant coefficient of the polynomial and the others follow by increasing degree.

The library syntax is GEN `gtovectrev0(GEN x, long n)`. GEN `gtovectrev(GEN x)` is also available.

3.6.15 Vecsmall($x, \{n\}$). Transforms the object x into a row vector of type `t_VECSMALL`. The dimension of the resulting vector can be optionally specified via the extra parameter n .

This acts as `Vec`(x, n), but only on a limited set of objects: the result must be representable as a vector of small integers. If x is a character string, a vector of individual characters in ASCII encoding is returned (`strchr` yields back the character string).

The library syntax is `GEN gtovecsmall0(GEN x, long n)`. `GEN gtovecsmall(GEN x)` is also available.

3.6.16 binary(x). Outputs the vector of the binary digits of $|x|$. Here x can be an integer, a real number (in which case the result has two components, one for the integer part, one for the fractional part) or a vector/matrix.

```
? binary(10)
%1 = [1, 0, 1, 0]
? binary(3.14)
%2 = [[1, 1], [0, 0, 1, 0, 0, 0, [...]]]
? binary([1,2])
%3 = [[1], [1, 0]]
```

For integer $x \geq 1$, the number of bits is `logint(x, 2) + 1`. By convention, 0 has no digits:

```
? binary(0)
%4 = []
```

The library syntax is `GEN binaire(GEN x)`.

3.6.17 bitand(x, y). Bitwise **and** of two integers x and y , that is the integer

$$\sum_i (x_i \text{ and } y_i) 2^i$$

Negative numbers behave 2-adically, i.e. the result is the 2-adic limit of `bitand`(x_n, y_n), where x_n and y_n are nonnegative integers tending to x and y respectively. (The result is an ordinary integer, possibly negative.)

```
? bitand(5, 3)
%1 = 1
? bitand(-5, 3)
%2 = 3
? bitand(-5, -3)
%3 = -7
```

The library syntax is `GEN gbitand(GEN x, GEN y)`. Also available is `GEN ibitand(GEN x, GEN y)`, which returns the bitwise *and* of $|x|$ and $|y|$, two integers.

3.6.18 bitneg($x, \{n = -1\}$). bitwise negation of an integer x , truncated to n bits, $n \geq 0$, that is the integer

$$\sum_{i=0}^{n-1} \text{not}(x_i) 2^i.$$

The special case $n = -1$ means no truncation: an infinite sequence of leading 1 is then represented as a negative number.

See Section 3.6.17 for the behavior for negative arguments.

The library syntax is GEN gbitneg(GEN x, long n).

3.6.19 bitnegimply(x, y). Bitwise negated imply of two integers x and y (or $\text{not}(x \Rightarrow y)$), that is the integer

$$\sum (x_i \text{ andnot}(y_i)) 2^i$$

See Section 3.6.17 for the behavior for negative arguments.

The library syntax is GEN gbitnegimply(GEN x, GEN y). Also available is GEN ibitnegimply(GEN x, GEN y), which returns the bitwise negated imply of $|x|$ and $|y|$, two integers.

3.6.20 bitor(x, y). bitwise (inclusive) or of two integers x and y , that is the integer

$$\sum (x_i \text{ or } y_i) 2^i$$

See Section 3.6.17 for the behavior for negative arguments.

The library syntax is GEN gbitor(GEN x, GEN y). Also available is GEN ibitor(GEN x, GEN y), which returns the bitwise or of $|x|$ and $|y|$, two integers.

3.6.21 bitprecision($x, \{n\}$). The function behaves differently according to whether n is present or not. If n is missing, the function returns the (floating point) precision in bits of the PARI object x .

If x is an exact object, the function returns +oo.

```
? bitprecision(exp(1e-100))
%1 = 512                \\ 512 bits
? bitprecision( [ exp(1e-100), 0.5 ] )
%2 = 128                \\ minimal accuracy among components
? bitprecision(2 + x)
%3 = +oo                \\ exact object
```

Use getlocalbitprec() to retrieve the working bit precision (as modified by possible localbitprec statements).

If n is present and positive, the function creates a new object equal to x with the new bit-precision roughly n . In fact, the smallest multiple of 64 (resp. 32 on a 32-bit machine) larger than or equal to n .

For x a vector or a matrix, the operation is done componentwise; for series and polynomials, the operation is done coefficientwise. For real x , n is the number of desired significant *bits*. If n

is smaller than the precision of x , x is truncated, otherwise x is extended with zeros. For exact or non-floating-point types, no change.

```
? bitprecision(Pi, 10)    \\ actually 64 bits ~ 19 decimal digits
%1 = 3.141592653589793239
? bitprecision(1, 10)
%2 = 1
? bitprecision(1 + O(x), 10)
%3 = 1 + O(x)
? bitprecision(2 + O(3^5), 10)
%4 = 2 + O(3^5)
```

The library syntax is GEN bitprecision00(GEN x, GEN n = NULL).

3.6.22 bittest(x, n). Outputs the n^{th} bit of x starting from the right (i.e. the coefficient of 2^n in the binary expansion of x). The result is 0 or 1. For $x \geq 1$, the highest 1-bit is at $n = \text{logint}(x)$ (and bigger n gives 0).

```
? bittest(7, 0)
%1 = 1 \\ the bit 0 is 1
? bittest(7, 2)
%2 = 1 \\ the bit 2 is 1
? bittest(7, 3)
%3 = 0 \\ the bit 3 is 0
```

See Section 3.6.17 for the behavior at negative arguments.

The library syntax is GEN gbittest(GEN x, long n). For a t_INT x , the variant long bittest(GEN x, long n) is generally easier to use, and if furthermore $n \geq 0$ the low-level function ulong int_bit(GEN x, long n) returns bittest(abs(x), n).

3.6.23 bitxor(x, y). Bitwise (exclusive) or of two integers x and y , that is the integer

$$\sum (x_i \text{ xor } y_i) 2^i$$

See Section 3.6.17 for the behavior for negative arguments.

The library syntax is GEN gbitxor(GEN x, GEN y). Also available is GEN ibitxor(GEN x, GEN y), which returns the bitwise xor of $|x|$ and $|y|$, two integers.

3.6.24 ceil(x). Ceiling of x . When x is in $\mathbf{R}$, the result is the smallest integer greater than or equal to x . Applied to a rational function, ceil(x) returns the Euclidean quotient of the numerator by the denominator.

The library syntax is GEN gceil(GEN x).

3.6.25 centerlift($x, \{v\}$). Same as **lift**, except that **t_INTMOD** and **t_PADIC** components are lifted using centered residues:

- for a **t_INTMOD** $x \in \mathbf{Z}/n\mathbf{Z}$, the lift y is such that $-n/2 < y \leq n/2$.
- a **t_PADIC** x is lifted in the same way as above (modulo $p^{\text{padicprec}(x)}$) if its valuation v is nonnegative; if not, returns the fraction $p^v \text{centerlift}(xp^{-v})$; in particular, rational reconstruction is not attempted. Use **bestappr** for this.

For backward compatibility, **centerlift**(**x**, 'v) is allowed as an alias for **lift**(**x**, 'v).

The library syntax is **centerlift**(GEN **x**).

3.6.26 characteristic(x). Returns the characteristic of the base ring over which x is defined (as defined by **t_INTMOD** and **t_FFELT** components). The function raises an exception if incompatible primes arise from **t_FFELT** and **t_PADIC** components.

```
? characteristic(Mod(1,24)*x + Mod(1,18)*y)
%1 = 6
```

The library syntax is GEN **characteristic**(GEN **x**).

3.6.27 component(x, n). Extracts the n^{th} -component of x . This is to be understood as follows: every PARI type has one or two initial code words. The components are counted, starting at 1, after these code words. In particular if x is a vector, this is indeed the n^{th} -component of x , if x is a matrix, the n^{th} column, if x is a polynomial, the n^{th} coefficient (i.e. of degree $n - 1$), and for power series, the n^{th} significant coefficient.

For polynomials and power series, one should rather use **polcoef**, and for vectors and matrices, the **[]** operator. Namely, if x is a vector, then **x[n]** represents the n^{th} component of x . If x is a matrix, **x[m,n]** represents the coefficient of row **m** and column **n** of the matrix, **x[m,]** represents the m^{th} row of x , and **x[,n]** represents the n^{th} column of x .

Using of this function requires detailed knowledge of the structure of the different PARI types, and thus it should almost never be used directly. Some useful exceptions:

```
? x = 3 + O(3^5);
? component(x, 2)
%2 = 81    \\ p^(p-adic accuracy)
? component(x, 1)
%3 = 3     \\ p
? q = Qfb(1,2,3);
? component(q, 1)
%5 = 1
```

The library syntax is GEN **compo**(GEN **x**, long **n**).

3.6.28 conj(x). Conjugate of x . The meaning of this is clear, except that for real quadratic numbers, it means conjugation in the real quadratic field. This function has no effect on integers, reals, intmods, fractions or p -adics. The only forbidden type is **polmod** (see **conjvec** for this).

The library syntax is GEN **gconj**(GEN **x**).

3.6.29 conjvec(z). Conjugate vector representation of z . If z is a polmod, equal to $\text{Mod}(a, T)$, this gives a vector of length $\text{degree}(T)$ containing:

- the complex embeddings of z if T has rational coefficients, i.e. the $a(r[i])$ where $r = \text{polroots}(T)$;
- the conjugates of z if T has some intmod coefficients;

if z is a finite field element, the result is the vector of conjugates $[z, z^p, z^{p^2}, \dots, z^{p^{n-1}}]$ where $n = \text{degree}(T)$.

If z is an integer or a rational number, the result is z . If z is a (row or column) vector, the result is a matrix whose columns are the conjugate vectors of the individual elements of z .

The library syntax is `GEN conjvec(GEN z, long prec)`.

3.6.30 denominator($f, \{D\}$). Denominator of f . The meaning of this is clear when f is a rational number or function. If f is an integer or a polynomial, it is treated as a rational number or function, respectively, and the result is equal to 1. For polynomials, you probably want to use

`denominator( content(f) )`

instead. As for modular objects, `t_INTMOD` and `t_PADIC` have denominator 1, and the denominator of a `t_POLMOD` is the denominator of its lift.

If f is a recursive structure, for instance a vector or matrix, the lcm of the denominators of its components (a common denominator) is computed. This also applies for `t_COMPLEXs` and `t_QUADS`.

Warning. Multivariate objects are created according to variable priorities, with possibly surprising side effects (x/y is a polynomial, but y/x is a rational function). See Section 2.5.3.

The optional argument D allows to control over which ring we compute the denominator and get a more predictable behaviour:

- 1: we only consider the underlying $\mathbf{Q}$ -structure and the denominator is a (positive) rational integer
- a simple variable, say 'x': all entries as rational functions in $K(x)$ and the denominator is a polynomial in x .

```
? f = x + 1/y + 1/2;
? denominator(f) \\ a t_POL in x
%2 = 1
? denominator(f, 1) \\ Q-denominator
%3 = 2
? denominator(f, x) \\ as a t_POL in x, seen above
%4 = 1
? denominator(f, y) \\ as a rational function in y
%5 = 2*y
```

The library syntax is `GEN denominator(GEN f, GEN D = NULL)`. Also available are `GEN denom(GEN x)` which implements the not very useful default behaviour (D is `NULL`) and `GEN Q_denom(GEN x)` ($D = 1$).

3.6.31 digits($x, \{b\}$). Outputs the vector of the digits of x in base b , where x and b are integers ($b = 10$ by default), from most significant down to least significant, the digits being the integers $0, 1, \dots, |b| - 1$. If $b > 0$ and $x < 0$, return the digits of $|x|$.

For $x \geq 1$ and $b > 0$, the number of digits is `logint( $x, b$ ) + 1`. See `fromdigits` for the reverse operation.

We also allow x an integral p -adic in which case b should be omitted or equal to p . Digits are still ordered from most significant to least significant in the p -adic sense (meaning we start from $x \bmod p$); trailing zeros are truncated.

```
? digits(1230)
%1 = [1, 2, 3, 0]

? digits(10, 2) \\ base 2
%2 = [1, 0, 1, 0]
```

By convention, 0 has no digits:

```
? digits(0)
%3 = []

? digits(10,-2) \\ base -2
%4 = [1, 1, 1, 1, 0] \\ 10 = -2 + 4 - 8 + 16
? 1105 + 0(5^5)
%5 = 5 + 4*5^2 + 3*5^3 + 0(5^5)
? digits(%)
%6 = [0, 1, 4, 3]
```

The library syntax is `GEN digits(GEN x, GEN b = NULL)`.

3.6.32 exponent(x). When x is a `t_REAL`, the result is the binary exponent e of x . For a nonzero x , this is the unique integer e such that $2^e \leq |x| < 2^{e+1}$. For a real 0, this returns the PARI exponent e attached to x (which may represent any floating-point number less than 2^e in absolute value).

```
? exponent(Pi)
%1 = 1
? exponent(4.0)
%2 = 2
? exponent(0.0)
%3 = -128
? default(realbitprecision)
%4 = 128
```

This definition extends naturally to nonzero integers, and the exponent of an exact 0 is $-\infty$ by convention.

For convenience, we *define* the exponent of a `t_FRAC` a/b as the difference of `exponent( $a$ )` and `exponent( $b$ )`; note that, if e' denotes the exponent of $a/b * 1.0$, then the exponent e we return is either e' or $e' + 1$, thus 2^{e+1} is an upper bound for $|a/b|$.

```
? [ exponent(9), exponent(10), exponent(9/10), exponent(9/10*1.) ]
%5 = [3, 3, 0, -1]
```

For a PARI object of type `t_COMPLEX`, `t_POL`, `t_SER`, `t_VEC`, `t_COL`, `t_MAT` this returns the largest exponent found among the components of x . Hence 2^{e+1} is a quick upper bound for the sup norm of real matrices or polynomials; and $2^{e+(3/2)}$ for complex ones.

```
? exponent(3*x^2 + 15*x - 100)
%5 = 6
? exponent(0)
%6 = -oo
```

The library syntax is `GEN gpexponent(GEN x)`.

Also available is `long gexpo(GEN x)`.

3.6.33 floor(x). Floor of x . When x is in $\mathbf{R}$, the result is the largest integer smaller than or equal to x . Applied to a rational function, `floor( $x$ )` returns the Euclidean quotient of the numerator by the denominator.

The library syntax is `GEN gffloor(GEN x)`.

3.6.34 frac(x). Fractional part of x . Identical to $x - \text{floor}(x)$. If x is real, the result is in $[0, 1[$.

The library syntax is `GEN gfrac(GEN x)`.

3.6.35 fromdigits($x, \{b = 10\}$). Gives the integer formed by the elements of x seen as the digits of a number in base b ($b = 10$ by default); b must be an integer satisfying $|b| > 1$. This is the reverse of `digits`:

```
? digits(1234, 5)
%1 = [1,4,4,1,4]
? fromdigits([1,4,4,1,4],5)
%2 = 1234
```

By convention, 0 has no digits:

```
? fromdigits([])
%3 = 0
```

This function works with x a `t_VECSMALL`; and also with $b < 0$ or $x[i]$ not an actual digit in base b (i.e., $x[i] < 0$ or $x[i] \geq b$): if x has length n , we return $\sum_{i=1}^n x[i]b^{n-i}$.

The library syntax is `GEN fromdigits(GEN x, GEN b = NULL)`.

3.6.36 imag(x). Imaginary part of x . When x is a quadratic number, this is the coefficient of ω in the “canonical” integral basis $(1, \omega)$.

```
? imag(3 + I)
%1 = 1
? x = 3 + quadgen(-23);
? imag(x) \\ as a quadratic number
%3 = 1
? imag(x * 1.) \\ as a complex number
%4 = 2.3979157616563597707987190320813469600
```

The library syntax is `GEN gimag(GEN x)`.

3.6.37 length(x). Length of x ; $\#x$ is a shortcut for `length( $x$ )`. This is mostly useful for

- vectors: dimension (0 for empty vectors),
- lists: number of entries (0 for empty lists),
- maps: number of entries (0 for empty maps),
- matrices: number of columns,
- character strings: number of actual characters (without trailing `\0`, should you expect it from `C char*`).

```
? # "a string"
%1 = 8
? # [3,2,1]
%2 = 3
? # []
%3 = 0
? #matrix(2,5)
%4 = 5
? L = List([1,2,3,4]); #L
%5 = 4
? M = Map([a,b; c,d; e,f]); #M
%6 = 3
```

The routine is in fact defined for arbitrary GP types, but is awkward and useless in other cases: it returns the number of non-code words in x , e.g. the effective length minus 2 for integers since the `t_INT` type has two code words.

The library syntax is `long glength(GEN x)`.

Also available is `long gtranslength(GEN x)` which return the length of x , that is the number of lines of matrices.

3.6.38 lift($x, \{v\}$). If v is omitted, lifts intmods from $\mathbf{Z}/n\mathbf{Z}$ in $\mathbf{Z}$, p -adics from $\mathbf{Q}_p$ to $\mathbf{Q}$ (as `truncate`), and polmods to polynomials. Otherwise, lifts only polmods whose modulus has main variable v . `t_FFELT` are not lifted, nor are List elements: you may convert the latter to vectors first, or use `apply(lift,L)`. More generally, components for which such lifts are meaningless (e.g. character strings) are copied verbatim.

```
? lift(Mod(5,3))
%1 = 2
? lift(3 + O(3^9))
%2 = 3
? lift(Mod(x,x^2+1))
%3 = x
? lift(Mod(x,x^2+1))
%4 = x
```

Lifts are performed recursively on an object components, but only by *one level*: once a `t_POLMOD` is lifted, the components of the result are *not* lifted further.

```
? lift(x * Mod(1,3) + Mod(2,3))
%4 = x + 2
```

```

? lift(x * Mod(y,y^2+1) + Mod(2,3))
%5 = y*x + Mod(2, 3)    \\ do you understand this one?
? lift(x * Mod(y,y^2+1) + Mod(2,3), 'x)
%6 = Mod(y, y^2 + 1)*x + Mod(Mod(2, 3), y^2 + 1)
? lift(%, y)
%7 = y*x + Mod(2, 3)

```

To recursively lift all components not only by one level, but as long as possible, use `liftall`. To lift only `t_INTMODs` and `t_PADICs` components, use `liftint`. To lift only `t_POLMODs` components, use `liftpol`. Finally, `centerlift` allows to lift `t_INTMODs` and `t_PADICs` using centered residues (lift of smallest absolute value).

The library syntax is `GEN lift0(GEN x, long v = -1)` where `v` is a variable number. Also available is `GEN lift(GEN x)` corresponding to `lift0(x,-1)`.

3.6.39 liftall(x). Recursively lift all components of x from $\mathbf{Z}/n\mathbf{Z}$ to $\mathbf{Z}$, from $\mathbf{Q}_p$ to $\mathbf{Q}$ (as `truncate`), and polmods to polynomials. `t_FFELT` are not lifted, nor are List elements: you may convert the latter to vectors first, or use `apply(liftall,L)`. More generally, components for which such lifts are meaningless (e.g. character strings) are copied verbatim.

```

? liftall(x * (1 + O(3)) + Mod(2,3))
%1 = x + 2
? liftall(x * Mod(y,y^2+1) + Mod(2,3)*Mod(z,z^2))
%2 = y*x + 2*z

```

The library syntax is `GEN liftall(GEN x)`.

3.6.40 liftint(x). Recursively lift all components of x from $\mathbf{Z}/n\mathbf{Z}$ to $\mathbf{Z}$ and from $\mathbf{Q}_p$ to $\mathbf{Q}$ (as `truncate`). `t_FFELT` are not lifted, nor are List elements: you may convert the latter to vectors first, or use `apply(liftint,L)`. More generally, components for which such lifts are meaningless (e.g. character strings) are copied verbatim.

```

? liftint(x * (1 + O(3)) + Mod(2,3))
%1 = x + 2
? liftint(x * Mod(y,y^2+1) + Mod(2,3)*Mod(z,z^2))
%2 = Mod(y, y^2 + 1)*x + Mod(Mod(2*z, z^2), y^2 + 1)

```

The library syntax is `GEN liftint(GEN x)`.

3.6.41 liftpol(x). Recursively lift all components of x which are polmods to polynomials. `t_FFELT` are not lifted, nor are List elements: you may convert the latter to vectors first, or use `apply(liftpol,L)`. More generally, components for which such lifts are meaningless (e.g. character strings) are copied verbatim.

```

? liftpol(x * (1 + O(3)) + Mod(2,3))
%1 = (1 + O(3))*x + Mod(2, 3)
? liftpol(x * Mod(y,y^2+1) + Mod(2,3)*Mod(z,z^2))
%2 = y*x + Mod(2, 3)*z

```

The library syntax is `GEN liftpol(GEN x)`.

3.6.42 norm(x). Algebraic norm of x , i.e. the product of x with its conjugate (no square roots are taken), or conjugates for polmods. For vectors and matrices, the norm is taken componentwise and hence is not the L^2 -norm (see `norml2`). Note that the norm of an element of $\mathbf{R}$ is its square, so as to be compatible with the complex norm.

The library syntax is `GEN gnorm(GEN x)`.

3.6.43 numerator($f, \{D\}$). Numerator of f . This is defined as $f * \text{denominator}(f, D)$, see `denominator` for details. The optional argument D allows to control over which ring we compute the denominator:

- 1: we only consider the underlying $\mathbf{Q}$ -structure and the denominator is a (positive) rational integer

- a simple variable, say 'x': all entries as rational functions in $K(x)$ and the denominator is a polynomial in x .

```
? f = x + 1/y + 1/2;
? numerator(f) \\ a t_POL in x
%2 = x + ((y + 2)/(2*y))
? numerator(f, 1) \\ Q-denominator is 2
%3 = x + ((y + 2)/y)
? numerator(f, y) \\ as a rational function in y
%5 = 2*y*x + (y + 2)
```

The library syntax is `GEN numerator(GEN f, GEN D = NULL)`. Also available are `GEN numerator(GEN x)` which implements the not very useful default behaviour (D is `NULL`) and `GEN Q_remove_denom(GEN x, GEN *ptd)` ($D = 1$) and also returns the denominator (coding 1 as `NULL`).

3.6.44 oo. Returns an object meaning $+\infty$, for use in functions such as `intnum`. It can be negated ($-\infty$ represents $-\infty$), and compared to real numbers (`t_INT`, `t_FRAC`, `t_REAL`), with the expected meaning: $+\infty$ is greater than any real number and $-\infty$ is smaller.

The library syntax is `GEN mkoo()`.

3.6.45 padicprec(x, p). Returns the absolute p -adic precision of the object x ; this is the minimum precision of the components of x . The result is `+oo` if x is an exact object (as a p -adic):

```
? padicprec((1 + 0(2^5)) * x + (2 + 0(2^4)), 2)
%1 = 4
? padicprec(x + 2, 2)
%2 = +oo
? padicprec(2 + x + 0(x^2), 2)
%3 = +oo
```

The function raises an exception if it encounters an object incompatible with p -adic computations:

```
? padicprec(0(3), 2)
*** at top-level: padicprec(0(3),2)
*** ^-----
*** padicprec: inconsistent moduli in padicprec: 3 != 2
? padicprec(1.0, 2)
```

```

***   at top-level: padicprec(1.0,2)
***               ^-----
*** padicprec: incorrect type in padicprec (t_REAL).

```

The library syntax is GEN `gppadicprec`(GEN `x`, GEN `p`). Also available is the function `long padicprec`(GEN `x`, GEN `p`), which returns `LONG_MAX` if $x = 0$ and the p -adic precision as a long integer.

3.6.46 `precision`($x, \{n\}$). The function behaves differently according to whether n is present or not. If n is missing, the function returns the floating point precision in decimal digits of the PARI object x . If x has no floating point component, the function returns `+oo`.

```

? precision(exp(1e-100))
%1 = 154                \\ 154 significant decimal digits
? precision(2 + x)
%2 = +oo                \\ exact object
? precision(0.5 + O(x))
%3 = 38                 \\ floating point accuracy, NOT series precision
? precision( [ exp(1e-100), 0.5 ] )
%4 = 38                 \\ minimal accuracy among components

```

Using `getlocalprec()` allows to retrieve the working precision (as modified by possible `localprec` statements).

If n is present, the function creates a new object equal to x with a new floating point precision n : n is the number of desired significant *decimal* digits. If n is smaller than the precision of a `t_REAL` component of x , it is truncated, otherwise it is extended with zeros. For non-floating-point types, no change.

The library syntax is GEN `precision00`(GEN `x`, GEN `n = NULL`). Also available are GEN `gprec`(GEN `x`, long `n`) and `long precision`(GEN `x`). In both, the accuracy is expressed in *words* (32-bit or 64-bit depending on the architecture).

3.6.47 `random`($\{N = 2^{31}\}$). Returns a random element in various natural sets depending on the argument N .

- `t_INT`: let $n = |N| - 1$; if $N > 0$ returns an integer uniformly distributed in $[0, n]$; if $N < 0$ returns an integer uniformly distributed in $[-n, n]$. Omitting the argument is equivalent to `random(2^31)`.

- `t_REAL`: returns a real number in $[0, 1[$ with the same accuracy as N (whose mantissa has the same number of significant words).

- `t_INTMOD`: returns a random intmod for the same modulus.

- `t_FFELT`: returns a random element in the same finite field.

- `t_VEC` of length 2, $N = [a, b]$: returns an integer uniformly distributed between a and b .

- `t_VEC` generated by `ellinit` over a finite field k (coefficients are `t_INTMODs` modulo a prime or `t_FFELTs`): returns a “random” k -rational *affine* point on the curve. More precisely if the curve has a single point (at infinity!) we return it; otherwise we return an affine point by drawing an abscissa uniformly at random until `ellordinate` succeeds. Note that this is definitely not a uniform distribution over $E(k)$, but it should be good enough for applications.

- `t_POL` return a random polynomial of degree at most the degree of N . The coefficients are drawn by applying `random` to the leading coefficient of N .

```
? random(10)
%1 = 9
? random(Mod(0,7))
%2 = Mod(1, 7)
? a = ffgen(ffinit(3,7), 'a'); random(a)
%3 = a^6 + 2*a^5 + a^4 + a^3 + a^2 + 2*a
? E = ellinit([3,7]*Mod(1,109)); random(E)
%4 = [Mod(103, 109), Mod(10, 109)]
? E = ellinit([1,7]*a^0); random(E)
%5 = [a^6 + a^5 + 2*a^4 + 2*a^2, 2*a^6 + 2*a^4 + 2*a^3 + a^2 + 2*a]
? random(Mod(1,7)*x^4)
%6 = Mod(5, 7)*x^4 + Mod(6, 7)*x^3 + Mod(2, 7)*x^2 + Mod(2, 7)*x + Mod(5, 7)
```

These variants all depend on a single internal generator, and are independent from your operating system's random number generators. A random seed may be obtained via `getrand`, and reset using `setrand`: from a given seed, and given sequence of `randoms`, the exact same values will be generated. The same seed is used at each startup, reseed the generator yourself if this is a problem. Note that internal functions also call the random number generator; adding such a function call in the middle of your code will change the numbers produced.

Technical note. Up to version 2.4 included, the internal generator produced pseudo-random numbers by means of linear congruences, which were not well distributed in arithmetic progressions. We now use Brent's XORGEN algorithm, based on Feedback Shift Registers, see <https://www.maths.anu.edu.au/~brent/random.html>. The generator has period $2^{4096} - 1$, passes the Crush battery of statistical tests of L'Ecuyer and Simard, but is not suitable for cryptographic purposes: one can reconstruct the state vector from a small sample of consecutive values, thus predicting the entire sequence.

Parallelism. In multi-threaded programs, each thread has a separate generator. They all start in the same `setrand(1)` state, so will all produce the same sequence of pseudo-random numbers although the various states are not shared. To avoid this, use `setrand` to provide a different starting state to each thread:

```
\\ with 8 threads
? parvector(8, i, random()) \\ all 8 threads return the same number
%1 = [1546275796, 1546275796, ... , 1546275796]
? parvector(8, i, random()) \\ ... and again since they are restarted
%2 = [1546275796, 1546275796, ... , 1546275796]

? s = [1..8]; \\ 8 random seeds; we could use vector(8,i,random())
? parvector(8, i, setrand(s[i]); random())
\\ now we get 8 different numbers
```

The library syntax is `GEN genrand(GEN N = NULL)`.

Also available: `GEN ellrandom(GEN E)` and `GEN ffrandom(GEN a)`.

3.6.48 $\text{real}(x)$. Real part of x . When x is a quadratic number, this is the coefficient of 1 in the “canonical” integral basis $(1, \omega)$.

```
? real(3 + I)
%1 = 3
? x = 3 + quadgen(-23);
? real(x) \\ as a quadratic number
%3 = 3
? real(x * 1.) \\ as a complex number
%4 = 3.5000000000000000000000000000000000000000000000000
```

The library syntax is `GEN greal(GEN x)`.

3.6.49 round($x, \{e\}$). If x is in $\mathbf{R}$, rounds x to the nearest integer (rounding to $+\infty$ in case of ties), then sets e to the number of error bits, that is the binary exponent of the difference between the original and the rounded value (the “fractional part”). If the exponent of x is too large compared to its precision (i.e. $e > 0$), the result is undefined and an error occurs if e was not given.

Important remark. Contrary to the other truncation functions, this function operates on every coefficient at every level of a PARI object. For example

$$\text{truncate}\left(\frac{2.4 * X^2 - 1.7}{X}\right) = 2.4 * X,$$

whereas

$$\text{round}\left(\frac{2.4 * X^2 - 1.7}{X}\right) = \frac{2 * X^2 - 2}{X}.$$

An important use of `round` is to get exact results after an approximate computation, when theory tells you that the coefficients must be integers.

The library syntax is `GEN round0(GEN x, GEN *e = NULL)`. Also available are `GEN grndtoi(GEN x, long *e)` and `GEN ground(GEN x)`.

3.6.50 serchop($s, \{n = 0\}$). Remove all terms of degree strictly less than n in series s . When the series contains no terms of degree $< n$, return $O(x^n)$.

```
? s = 1/x + x + 2*x^2 + O(x^3);
? serchop(s)
%2 = x + 2*x^3 + O(x^3)
? serchop(s, 2)
%3 = 2*x^2 + O(x^3)
? serchop(s, 100)
%4 = O(x^100)
```

The library syntax is `GEN serchop(GEN s, long n)`.

3.6.51 serprec(x, v). Returns the absolute precision of x with respect to power series in the variable v ; this is the minimum precision of the components of x . The result is $+\infty$ if x is an exact object (as a series in v):

```
? serprec(x + O(y^2), y)
%1 = 2
? serprec(x + 2, x)
%2 = +oo
? serprec(2 + x + O(x^2), y)
%3 = +oo
```

The library syntax is `GEN gpserprec(GEN x, long v)` where v is a variable number. Also available is `long serprec(GEN x, GEN p)`, which returns `LONG_MAX` if $x = 0$, otherwise the series precision as a long integer.

3.6.52 simplify(x). This function simplifies x as much as it can. Specifically, a complex or quadratic number whose imaginary part is the integer 0 (i.e. not `Mod(0,2)` or `0.E-28`) is converted to its real part, and a polynomial of degree 0 is converted to its constant term. Simplifications occur recursively.

This function is especially useful before using arithmetic functions, which expect integer arguments:

```
? x = 2 + y - y
%1 = 2
? isprime(x)
*** at top-level: isprime(x)
*** ^-----
*** isprime: not an integer argument in an arithmetic function
? type(x)
%2 = "t_POL"
? type(simplify(x))
%3 = "t_INT"
```

Note that GP results are simplified as above before they are stored in the history. (Unless you disable automatic simplification with `\y`, that is.) In particular

```
? type(%1)
%4 = "t_INT"
```

The library syntax is `GEN simplify(GEN x)`.

3.6.53 sizebyte(x). Outputs the total number of bytes occupied by the tree representing the PARI object x .

The library syntax is `long gsizebyte(GEN x)`. Also available is `long gsizeword(GEN x)` returning a number of *words*.

3.6.54 sizedigit(x). This function is DEPRECATED, essentially meaningless, and provided for backwards compatibility only. Don't use it!

outputs a quick upper bound for the number of decimal digits of (the components of) x , off by at most 1. More precisely, for a positive integer x , it computes (approximately) the ceiling of

$$\text{floor}(1 + \log_2 x) \log_{10} 2,$$

To count the number of decimal digits of a positive integer x , use `#digits(x)`. To estimate (recursively) the size of x , use `normlp(x)`.

The library syntax is `long sizedigit(GEN x)`.

3.6.55 truncate($x, \{&e\}$). Truncates x and sets e to the number of error bits. When x is in $\mathbf{R}$, this means that the part after the decimal point is chopped away, e is the binary exponent of the difference between the original and the truncated value (the “fractional part”). If the exponent of x is too large compared to its precision (i.e. $e > 0$), the result is undefined and an error occurs if e was not given. The function applies componentwise on vector / matrices; e is then the maximal number of error bits. If x is a rational function, the result is the “integer part” (Euclidean quotient of numerator by denominator) and e is not set.

Note a very special use of `truncate`: when applied to a power series, it transforms it into a polynomial or a rational function with denominator a power of X , by chopping away the $O(X^k)$. Similarly, when applied to a p -adic number, it transforms it into an integer or a rational number by chopping away the $O(p^k)$.

The library syntax is `GEN trunc0(GEN x, GEN *e = NULL)`. The following functions are also available: `GEN gtrunc(GEN x)` and `GEN gcvtol(GEN x, long *e)`.

3.6.56 valuation($x, \{p\}$). Computes the highest exponent of p dividing x . If p is of type integer, x must be an integer, an intmod whose modulus is divisible by p , a fraction, a q -adic number with $q = p$, or a polynomial or power series in which case the valuation is the minimum of the valuation of the coefficients.

If p is of type polynomial, x must be of type polynomial or rational function, and also a power series if x is a monomial. Finally, the valuation of a vector, complex or quadratic number is the minimum of the component valuations.

If $x = 0$, the result is `+oo` if x is an exact object. If x is a p -adic numbers or power series, the result is the exponent of the zero. Any other type combinations gives an error.

Finally, p can be omitted if x is a `t_PADIC` (taken to be the underlying prime), a `t_SER` or a `t_POL` (taken to be the main variable).

The library syntax is `GEN gpvaluation(GEN x, GEN p = NULL)`. Also available is `long gvaluation(GEN x, GEN p)`, which returns `LONG_MAX` if $x = 0$ and the valuation as a long integer.

3.6.57 varhigher(*name*, {*v*}). Return a variable *name* whose priority is higher than the priority of *v* (of all existing variables if *v* is omitted). This is a counterpart to **varlower**.

```
? Pol([x,x], t)
***   at top-level: Pol([x,x],t)
***           ^-----
*** Pol: incorrect priority in gtopoly: variable x <= t
? t = varhigher("t", x);
? Pol([x,x], t)
%3 = x*t + x
```

This routine is useful since new GP variables directly created by the interpreter always have lower priority than existing GP variables. When some basic objects already exist in a variable that is incompatible with some function requirement, you can now create a new variable with a suitable priority instead of changing variables in existing objects:

```
? K = nfinit(x^2+1);
? rnfequation(K,y^2-2)
***   at top-level: rnfequation(K,y^2-2)
***           ^-----
*** rnfequation: incorrect priority in rnfequation: variable y >= x
? y = varhigher("y", x);
? rnfequation(K, y^2-2)
%3 = y^4 - 2*y^2 + 9
```

Caution 1. The *name* is an arbitrary character string, only used for display purposes and need not be related to the GP variable holding the result, nor to be a valid variable name. In particular the *name* can not be used to retrieve the variable, it is not even present in the parser's hash tables.

```
? x = varhigher("#");
? x^2
%2 = #^2
```

Caution 2. There are a limited number of variables and if no existing variable with the given display name has the requested priority, the call to **varhigher** uses up one such slot. Do not create new variables in this way unless it's absolutely necessary, reuse existing names instead and choose sensible priority requirements: if you only need a variable with higher priority than *x*, state so rather than creating a new variable with highest priority.

```
\\ quickly use up all variables
? n = 0; while(1,varhigher("tmp"); n++)
***   at top-level: n=0;while(1,varhigher("tmp");n++)
***           ^-----
*** varhigher: no more variables available.
*** Break loop: type 'break' to go back to GP prompt
break> n
65510
\\ infinite loop: here we reuse the same 'tmp'
? n = 0; while(1,varhigher("tmp", x); n++)
```

The library syntax is GEN varhigher(const char *name, long v = -1) where v is a variable number.

3.6.58 variable($\{x\}$). Gives the main variable of the object x (the variable with the highest priority used in x), and p if x is a p -adic number. Return 0 if x has no variable attached to it.

```
? variable(x^2 + y)
%1 = x
? variable(1 + 0(5^2))
%2 = 5
? variable([x,y,z,t])
%3 = x
? variable(1)
%4 = 0
```

The construction

```
if (!variable(x),...)
```

can be used to test whether a variable is attached to x .

If x is omitted, returns the list of user variables known to the interpreter, by order of decreasing priority. (Highest priority is initially x , which come first until **varhigher** is used.) If **varhigher** or **varlower** are used, it is quite possible to end up with different variables (with different priorities) printed in the same way: they will then appear multiple times in the output:

```
? varhigher("y");
? varlower("y");
? variable()
%4 = [y, x, y]
```

Using $v = \text{variable}()$ then $v[1]$, $v[2]$, etc. allows to recover and use existing variables.

The library syntax is `GEN gpolve(GEN x = NULL)`. However, in library mode, this function should not be used for x non-NULL, since **gvar** is more appropriate. Instead, for x a p -adic (type `t_PADIC`), p is `gel(x, 2)`; otherwise, use `long gvar(GEN x)` which returns the variable number of x if it exists, `NO_VARIABLE` otherwise, which satisfies the property `varncmp(NO_VARIABLE, v) > 0` for all valid variable number v , i.e. it has lower priority than any variable.

3.6.59 variables($\{x\}$). Returns the list of all variables occurring in object x sorted by decreasing priority. If x is omitted, return all polynomial variables known to the interpreter (this will include x and y , which are always defined on startup); user variables which do not occur in `t_POL` or `t_SER` constructions are *not* included. To see all user variables, use `\uv`.

```
? variables([x^2 + y*z + 0(t), a+x])
%1 = [x, y, z, t, a]
```

The construction

```
if (!variables(x),...)
```

can be used to test whether a variable is attached to x .

If **varhigher** or **varlower** are used, it is quite possible to end up with different variables (having different priorities) printed in the same way. They will then appear multiple times in the output:

```
? y1 = varhigher("y"); y2 = varlower("y");
? variables(y*y1*y2)
```

```
%2 = [y, y, y]
```

The library syntax is `GEN variables_vec(GEN x = NULL)`.

Also available is `GEN variables_vecsmall(GEN x)` which returns the (sorted) variable numbers instead of the attached monomials of degree 1.

3.6.60 varlower(*name*, {*v*}). Return a variable *name* whose priority is lower than the priority of *v* (of all existing variables if *v* is omitted). This is a counterpart to **varhigher**.

New GP variables directly created by the interpreter always have lower priority than existing GP variables, but it is not easy to check whether an identifier is currently unused, so that the corresponding variable has the expected priority when it's created! Thus, depending on the session history, the same command may fail or succeed:

```
? t; z; \\ now t > z
? rnfequation(t^2+1,z^2-t)
*** at top-level: rnfequation(t^2+1,z^
*** ^-----
*** rnfequation: incorrect priority in rnfequation: variable t >= t
```

Restart and retry:

```
? z; t; \\ now z > t
? rnfequation(t^2+1,z^2-t)
%2 = z^4 + 1
```

It is quite annoying for package authors, when trying to define a base ring, to notice that the package may fail for some users depending on their session history. The safe way to do this is as follows:

```
? z; t; \\ In new session: now z > t
...
? t = varlower("t", 'z);
? rnfequation(t^2+1,z^2-2)
%2 = z^4 - 2*z^2 + 9
? variable()
%3 = [x, y, z, t]

? t; z; \\ In new session: now t > z
...
? t = varlower("t", 'z); \\ create a new variable, still printed "t"
? rnfequation(t^2+1,z^2-2)
%2 = z^4 - 2*z^2 + 9
? variable()
%3 = [x, y, t, z, t]
```

Now both constructions succeed. Note that in the first case, **varlower** is essentially a no-op, the existing variable *t* has correct priority. While in the second case, two different variables are displayed as *t*, one with higher priority than *z* (created in the first line) and another one with lower priority (created by **varlower**).

Caution 1. The *name* is an arbitrary character string, only used for display purposes and need not be related to the GP variable holding the result, nor to be a valid variable name. In particular the *name* can not be used to retrieve the variable, it is not even present in the parser's hash tables.

```
? x = varlower("#");
? x^2
%2 = #^2
```

Caution 2. There are a limited number of variables and if no existing variable with the given display name has the requested priority, the call to `varlower` uses up one such slot. Do not create new variables in this way unless it's absolutely necessary, reuse existing names instead and choose sensible priority requirements: if you only need a variable with higher priority than *x*, state so rather than creating a new variable with highest priority.

```
\\ quickly use up all variables
? n = 0; while(1,varlower("x"); n++)
*** at top-level: n=0;while(1,varlower("x");n++)
*** ^-----
*** varlower: no more variables available.
*** Break loop: type 'break' to go back to GP prompt
break> n
65510
\\ infinite loop: here we reuse the same 'tmp'
? n = 0; while(1,varlower("tmp", x); n++)
```

The library syntax is GEN `varlower(const char *name, long v = -1)` where *v* is a variable number.

3.7 Combinatorics.

Permutations are represented in gp as `t_VECSMALLs` and can be input directly as `Vecsmall([1,3,2,4])` or obtained from the iterator `forperm`:

```
? forperm(3, p, print(p)) \\ iterate through S_3
Vecsmall([1, 2, 3])
Vecsmall([1, 3, 2])
Vecsmall([2, 1, 3])
Vecsmall([2, 3, 1])
Vecsmall([3, 1, 2])
Vecsmall([3, 2, 1])
```

Permutations can be multiplied via `*`, raised to some power using `^`, inverted using `^(-1)`, conjugated as `p * q * p^(-1)`. Their order and signature are available via `permorder` and `perm-sign`.

3.7.1 bernfrac(n). Bernoulli number B_n , where $B_0 = 1$, $B_1 = -1/2$, $B_2 = 1/6, \dots$, expressed as a rational number. The argument n should be a nonnegative integer. The function **bernvec** creates a cache of successive Bernoulli numbers which greatly speeds up later calls to **bernfrac**:

```
? bernfrac(20000);
time = 107 ms.
? bernvec(10000); \\ cache B_0, B_2, ..., B_20000
time = 35,957 ms.
? bernfrac(20000); \\ now instantaneous
?
```

The library syntax is **GEN bernfrac(long n)**.

3.7.2 bernpol($n, \{a = 'x\}$). Bernoulli polynomial B_n evaluated at a ('x' by default), defined by

$$\sum_{n=0}^{\infty} B_n(x) \frac{T^n}{n!} = \frac{T e^{xT}}{e^T - 1}.$$

```
? bernpol(1)
%1 = x - 1/2
? bernpol(3)
%2 = x^3 - 3/2*x^2 + 1/2*x
? bernpol(3, 2)
%3 = 3
```

Note that evaluation at a is only provided for convenience and uniformity of interface: contrary to, e.g., **polcyclo**, computing the evaluation is no faster than

```
B = bernpol(k); subst(B, 'x, a)
```

and the latter allows to reuse B to evaluate B_k at different values.

The library syntax is **GEN bernpol_eval(long n, GEN a = NULL)**. The variant **GEN bernpol(long k, long v)** returns the k -th Bernoulli polynomial in variable v .

3.7.3 bernreal(n). Bernoulli number B_n , as **bernfrac**, but B_n is returned as a real number (with the current precision). The argument n should be a nonnegative integer. The function slows down as the precision increases:

```
? \p1000
? bernreal(200000);
time = 5 ms.
? \p10000
? bernreal(200000);
time = 18 ms.
? \p100000
? bernreal(200000);
time = 84 ms.
```

The library syntax is **GEN bernreal(long n, long prec)**.

3.7.4 bernvec(n). Returns a vector containing, as rational numbers, the Bernoulli numbers $B_0, B_2, \dots, B_{2n}$:

```
? bernvec(5) \\ B_0, B_2..., B_10
%1 = [1, 1/6, -1/30, 1/42, -1/30, 5/66]
? bernfrac(10)
%2 = 5/66
```

This routine uses a lot of memory but is much faster than repeated calls to **bernfrac**:

```
? forstep(n = 2, 10000, 2, bernfrac(n))
time = 18,245 ms.
? bernvec(5000);
time = 1,338 ms.
```

The computed Bernoulli numbers are stored in an incremental cache which makes later calls to **bernfrac** and **bernreal** instantaneous in the cache range: re-running the same previous **bernfracs** after the **bernvec** call gives:

```
? forstep(n = 2, 10000, 2, bernfrac(n))
time = 1 ms.
```

The time and space complexity of this function are $\tilde{O}(n^2)$; in the feasible range $n \leq 10^5$ (requires about two hours), the practical time complexity is closer to $\tilde{O}(n^{\log_2 6})$.

The library syntax is **GEN bernvec(long n)**.

3.7.5 binomial($n, \{k\}$). binomial coefficient $\binom{n}{k}$. Here k must be an integer, but n can be any PARI object. For nonnegative k , $\binom{n}{k} = (n)_k/k!$ is polynomial in n , where $(n)_k = n(n-1)\dots(n-k+1)$ is the Pochhammer symbol used by combinatorists (which is different from the one used by analysts).

```
? binomial(4,2)
%1 = 6
? n = 4; vector(n+1, k, binomial(n,k-1))
%2 = [1, 4, 6, 4, 1]
? binomial('x, 2)
%3 = 1/2*x^2 - 1/2*x
```

When n is a negative integer and k is negative, we use Daniel Loeb's extension,

$$\lim_{t \rightarrow 1} \Gamma(n+t)/\Gamma(k+t)/\Gamma(n-k+t).$$

(Sets with a negative number of elements, *Adv. Math.* **91** (1992), no. 1, 64–74. See also <https://arxiv.org/abs/1105.3689>.) This way the symmetry relation $\binom{n}{k} = \binom{n}{n-k}$ becomes valid for all integers n and k , and the binomial theorem holds for all complex numbers a, b, n with $|b| < |a|$:

$$(a+b)^n = \sum_{k \geq 0} \binom{n}{k} a^{n-k} b^k.$$

Beware that this extension is incompatible with another traditional extension ($\binom{n}{k} := 0$ if $k < 0$); to enforce the latter, use

```
BINOMIAL(n, k) = if (k >= 0, binomial(n, k));
```

The argument k may be omitted if n is a nonnegative integer; in this case, return the vector with $n + 1$ components whose $k + 1$ -th entry is `binomial( $n, k$ )`

```
? binomial(4)
%4 = [1, 4, 6, 4, 1]
```

The library syntax is `GEN binomial0(GEN n, GEN k = NULL)`.

3.7.6 eulerfrac(n). Euler number E_n , where $E_0 = 1$, $E_1 = 0$, $E_2 = -1$, $\dots$, are integers such that

$$\frac{1}{\cosh t} = \sum_{n \geq 0} \frac{E_n}{n!} t^n.$$

The argument n should be a nonnegative integer.

```
? vector(10,i,eulerfrac(i))
%1 = [0, -1, 0, 5, 0, -61, 0, 1385, 0, -50521]
? eulerfrac(20000);
? sizedigit(%)
%3 = 73416
```

The library syntax is `GEN eulerfrac(long n)`.

3.7.7 eulerianpol($n, \{v = 'x\}$). Eulerian polynomial A_n in variable v defined by

$$\sum_{n=0}^{\infty} A_n(x) \frac{T^n}{n!} = \frac{x-1}{x - e^{(x-1)T}}.$$

```
? eulerianpol(2)
%1 = x + 1
? eulerianpol(5, 't)
%2 = t^4 + 26*t^3 + 66*t^2 + 26*t + 1
```

The library syntax is `GEN eulerianpol(long n, long v = -1)` where v is a variable number.

3.7.8 eulerpol($n, \{v = 'x\}$). Euler polynomial E_n in variable v defined by

$$\sum_{n=0}^{\infty} E_n(x) \frac{T^n}{n!} = \frac{2e^{xT}}{e^T + 1}.$$

```
? eulerpol(1)
%1 = x - 1/2
? eulerpol(3)
%2 = x^3 - 3/2*x^2 + 1/4
```

The library syntax is `GEN eulerpol(long n, long v = -1)` where v is a variable number.

3.7.9 eulerreal(n). Euler number E_n , where $E_0 = 1$, $E_1 = 0$, $E_2 = -1$, $\dots$, are integers such that

$$\frac{1}{\cosh t} = \sum_{n \geq 0} \frac{E_n}{n!} t^n.$$

The argument n should be a nonnegative integer. Return E_n as a real number (with the current precision).

```
? sizedigit(eulerfrac(20000))
%1 = 73416
? eulerreal(20000);
%2 = 9.2736664576330851823546169139003297830 E73414
```

The library syntax is GEN `eulerreal(long n, long prec)`.

3.7.10 eulervec(n). Returns a vector containing the nonzero Euler numbers $E_0, E_2, \dots, E_{2n}$:

```
? eulervec(5) \\ E_0, E_2..., E_10
%1 = [1, -1, 5, -61, 1385, -50521]
? eulerfrac(10)
%2 = -50521
```

This routine uses more memory but is faster than repeated calls to `eulerfrac`:

```
? forstep(n = 2, 8000, 2, eulerfrac(n))
time = 27,3801ms.
? eulervec(4000);
time = 8,430 ms.
```

The computed Euler numbers are stored in an incremental cache which makes later calls to `eulerfrac` and `eulerreal` instantaneous in the cache range: re-running the same previous `eulerfracs` after the `eulervec` call gives:

```
? forstep(n = 2, 10000, 2, eulerfrac(n))
time = 0 ms.
```

The library syntax is GEN `eulervec(long n)`.

3.7.11 fibonacci(x). x^{th} Fibonacci number.

The library syntax is GEN `fibo(long x)`.

3.7.12 hammingweight(x). If x is a `t_INT`, return the binary Hamming weight of $|x|$. Otherwise x must be of type `t_POL`, `t_VEC`, `t_COL`, `t_VECSMALL`, or `t_MAT` and the function returns the number of nonzero coefficients of x .

```
? hammingweight(15)
%1 = 4
? hammingweight(x^100 + 2*x + 1)
%2 = 3
? hammingweight([Mod(1,2), 2, Mod(0,3)])
%3 = 2
? hammingweight(matid(100))
%4 = 100
```

The library syntax is long `hammingweight(GEN x)`.

3.7.13 harmonic($n, \{r = 1\}$). Generalized harmonic number of index $n \geq 0$ in power r , as a rational number. If $r = 1$ (or omitted), this is the harmonic number

$$H_n = \sum_{i=1}^n \frac{1}{i}.$$

In general, this is

$$H_{n,r} = \sum_{i=1}^n \frac{1}{i^r}.$$

The function runs in time $\tilde{O}(rn)$, essentially linear in the size of the output.

```
? harmonic(0)
%1 = 0
? harmonic(1)
%2 = 1
? harmonic(10)
%3 = 7381/2520
? harmonic(10, 2)
%4 = 1968329/1270080
? harmonic(10, -2)
%5 = 385
```

Note that the numerator and denominator are of order $\exp((r + o(1))n)$ and this will overflow for large n . To obtain H_n as a floating point number, use $H_n = \text{psi}(n + 1) + \text{Euler}$, or for $r > 1$,

```
H(n,r) = (-1)^(r+1)*(psi(n+1,r-1)/(r-1)!)+if(r==1,Euler,zeta(r))
```

The library syntax is `GEN harmonic0(ulong n, GEN r = NULL)`. Also available is `GEN harmonic(ulong n)` for $r = 1$.

3.7.14 numbpart(n). Gives the number of unrestricted partitions of n , usually called $p(n)$ in the literature; in other words the number of nonnegative integer solutions to $a + 2b + 3c + \dots = n$. n must be of type integer and $n < 10^{15}$ (with trivial values $p(n) = 0$ for $n < 0$ and $p(0) = 1$). The algorithm uses the Hardy-Ramanujan-Rademacher formula. To explicitly enumerate them, see `partitions`.

The library syntax is `GEN numbpart(GEN n)`.

3.7.15 numtoperm(n, k). Generates the k -th permutation (as a row vector of length n) of the numbers 1 to n . The number k is taken modulo $n!$, i.e. inverse function of `permtonum`. The numbering used is the standard lexicographic ordering, starting at 0.

The library syntax is `GEN numtoperm(long n, GEN k)`.

3.7.16 partitions($k, \{a = k\}, \{n = k\}$). Returns the vector of partitions of the integer k as a sum of positive integers (parts); for $k < 0$, it returns the empty set `[]`, and for $k = 0$ the trivial partition (no parts). A partition is given by a `t_VECSMALL`, where parts are sorted in nondecreasing order:

```
? partitions(3)
%1 = [Vecsmall([3]), Vecsmall([1, 2]), Vecsmall([1, 1, 1])]
```

correspond to 3, $1 + 2$ and $1 + 1 + 1$. The number of (unrestricted) partitions of k is given by `numbpart`:

```
? #partitions(50)
%1 = 204226
? numbpart(50)
%2 = 204226
```

Optional parameters n and a are as follows:

- $n = nmax$ (resp. $n = [nmin, nmax]$) restricts partitions to length less than $nmax$ (resp. length between $nmin$ and $nmax$), where the *length* is the number of nonzero entries.

- $a = amax$ (resp. $a = [amin, amax]$) restricts the parts to integers less than $amax$ (resp. between $amin$ and $amax$).

```
? partitions(4, 2) \\ parts bounded by 2
%1 = [Vecsmall([2, 2]), Vecsmall([1, 1, 2]), Vecsmall([1, 1, 1, 1])]
? partitions(4,, 2) \\ at most 2 parts
%2 = [Vecsmall([4]), Vecsmall([1, 3]), Vecsmall([2, 2])]
? partitions(4,[0,3], 2) \\ at most 2 parts
%3 = [Vecsmall([1,3]), Vecsmall([2,2])]
```

By default, parts are positive and we remove zero entries unless $amin \leq 0$, in which case $nmin$ is ignored and we fix $\#X = nmax$:

```
? partitions(4, [0,3]) \\ parts between 0 and 3
%1 = [Vecsmall([0, 0, 1, 3]), Vecsmall([0, 0, 2, 2]), \
      Vecsmall([0, 1, 1, 2]), Vecsmall([1, 1, 1, 1])]
? partitions(1, [0,3], [2,4]) \\ no partition with 2 to 4 nonzero parts
%2 = []
```

The library syntax is `GEN partitions(long k, GEN a = NULL, GEN n = NULL)`.

3.7.17 permcycles(x). Given a permutation x on n elements, return the orbits of $\{1, \dots, n\}$ under the action of x as cycles.

```
? permcycles(Vecsmall([1,2,3]))
%1 = [Vecsmall([1]),Vecsmall([2]),Vecsmall([3])]
? permcycles(Vecsmall([2,3,1]))
%2 = [Vecsmall([1,2,3])]
? permcycles(Vecsmall([2,1,3]))
%3 = [Vecsmall([1,2]),Vecsmall([3])]
```

The library syntax is `GEN permcycles(GEN x)`.

3.7.18 permorder(x). Given a permutation x on n elements, return its order.

```
? p = Vecsmall([3,1,4,2,5]);
? p^2
%2 = Vecsmall([4,3,2,1,5])
? p^4
%3 = Vecsmall([1,2,3,4,5])
? permorder(p)
%4 = 4
```

The library syntax is `GEN permorder(GEN x)`.

3.7.19 permsign(x). Given a permutation x on n elements, return its signature.

```
? p = Vecsmall([3,1,4,2,5]);
? permsign(p)
%2 = -1
? permsign(p^2)
%3 = 1
```

The library syntax is `long permsign(GEN x)`.

3.7.20 permtonum(x). Given a permutation x on n elements, gives the number k such that $x = \text{numtoperm}(n, k)$, i.e. inverse function of `numtoperm`. The numbering used is the standard lexicographic ordering, starting at 0.

The library syntax is `GEN permtonum(GEN x)`.

3.7.21 stirling($n, k, \{flag = 1\}$). Stirling number of the first kind $s(n, k)$ ($flag = 1$, default) or of the second kind $S(n, k)$ ($flag = 2$), where n, k are nonnegative integers. The former is $(-1)^{n-k}$ times the number of permutations of n symbols with exactly k cycles; the latter is the number of ways of partitioning a set of n elements into k nonempty subsets. Note that if all $s(n, k)$ are needed, it is much faster to compute

$$\sum_k s(n, k) x^k = x(x-1) \dots (x-n+1).$$

Similarly, if a large number of $S(n, k)$ are needed for the same k , one should use

$$\sum_n S(n, k) x^n = \frac{x^k}{(1-x) \dots (1-kx)}.$$

(Should be implemented using a divide and conquer product.) Here are simple variants for n fixed:

```
/* list of s(n,k), k = 1..n */
vecstirling(n) = Vec( factorback(vector(n-1,i,1-i*'x')) )

/* list of S(n,k), k = 1..n */
vecstirling2(n) =
{ my(Q = x^(n-1), t);
  vector(n, i, t = divrem(Q, x-i); Q=t[1]; simplify(t[2]));
}

/* Bell numbers, B_n = B[n+1] = sum(k = 0, n, S(n,k)), n = 0..N */
```

```
vecbell(N)=
{ my (B = vector(N+1));
  B[1] = B[2] = 1;
  for (n = 2, N,
    my (C = binomial(n-1));
    B[n+1] = sum(k = 1, n, C[k]*B[k]);
  ); B;
}
```

The library syntax is GEN `stirling(long n, long k, long flag)`. Also available are GEN `stirling1(ulong n, ulong k)` (*flag* = 1) and GEN `stirling2(ulong n, ulong k)` (*flag* = 2).

3.8 Arithmetic functions.

These functions are by definition functions whose natural domain of definition is either $\mathbf{Z}$ (or $\mathbf{Z}_{>0}$). The way these functions are used is completely different from transcendental functions in that there are no automatic type conversions: in general only integers are accepted as arguments. An integer argument N can be given in the following alternate formats:

- `t_MAT`: its factorization `fa = factor(N)`,
- `t_VEC`: a pair `[N, fa]` giving both the integer and its factorization.

This allows to compute different arithmetic functions at a given N while factoring the latter only once.

```
? N = 10!; faN = factor(N);
? eulerphi(N)
%2 = 829440
? eulerphi(faN)
%3 = 829440
? eulerphi(S = [N, faN])
%4 = 829440
? sigma(S)
%5 = 15334088
```

3.8.1 Arithmetic functions and the factoring engine. All arithmetic functions in the narrow sense of the word — Euler’s totient function, the Moebius function, the sums over divisors or powers of divisors etc.— call, after trial division by small primes, the same versatile factoring machinery described under `factorint`. It includes Shanks SQUFOF, Pollard Rho, ECM and MPQS stages, and has an early exit option for the functions `moebius` and (the integer function underlying) `issquarefree`. This machinery relies on a fairly strong probabilistic primality test, see `ispseudoprime`, but you may also set

```
default(factor_proven, 1)
```

to ensure that all tentative factorizations are fully proven. This should not slow down PARI too much, unless prime numbers with hundreds of decimal digits occur frequently in your application.

The following functions compute the order of an element in a finite group: `ellorder` (the rational points on an elliptic curve defined over a finite field), `fforder` (the multiplicative group of a finite field), `znorder` (the invertible elements in $\mathbf{Z}/n\mathbf{Z}$). The following functions compute discrete logarithms in the same groups (whenever this is meaningful) `elllog`, `fflog`, `znlog`.

- `t_INT`: the integer N ,
- `t_MAT`: the factorization `fa = factor(N)`,
- `t_VEC`: this is the preferred format and provides both the integer N and its factorization in a two-component vector `[N, fa]`.

```
? p = nextprime(10^30);  
? v = [p-1, factor(p-1)]; \\ data for discrete log & order computations  
? znorder(Mod(2,p), v)  
%3 = 5000000000000000000000000000028  
? g = znprimroot(p);  
? znlog(2, g, v)  
%5 = 543038070904014908801878611374
```

The finite abelian group $G = (\mathbf{Z}/N\mathbf{Z})^*$ can be written $G = \oplus_{i \leq n} (\mathbf{Z}/d_i\mathbf{Z})g_i$, with $d_n \mid \dots \mid d_2 \mid d_1$ (SNF condition), all $d_i > 0$, and $\prod_i d_i = \phi(N)$.

• a *character* on the abelian group $\oplus_j (\mathbf{Z}/d_j\mathbf{Z})g_j$ is given by a row vector $\chi = [a_1, \dots, a_n]$ of integers $0 \leq a_i < d_i$ such that $\chi(g_j) = e(a_j/d_j)$ for all j , with the standard notation $e(x) := \exp(2i\pi x)$. In other words, $\chi(\prod_j g_j^{n_j}) = e(\sum_j a_j n_j/d_j)$.

178

- The *column vector* of the m_j , $0 \leq m_j < D_j$ is called the *Conrey logarithm* of m (discrete logarithm in terms of the Conrey generators). Note that discrete logarithms in PARI/GP are always expressed as `t_COLs`.

- The attached character is called the *Conrey character* attached to m .

To sum up a Dirichlet character can be defined by a `t_INTMOD` $\text{Mod}(m, N)$, a `t_INT` lift (the Conrey label m), a `t_COL` (the Conrey logarithm of m , in terms of the Conrey generators) or a `t_VEC` (in terms of the SNF generators). The `t_COL` format, i.e. Conrey logarithms, is the preferred (fastest) representation.

Concretely, this works as follows:

`G = znstar(N, 1)` initializes $(\mathbf{Z}/N\mathbf{Z})^*$, which must be given as first arguments to all functions handling Dirichlet characters.

`znconreychar` transforms `t_INT`, `t_INTMOD` and `t_COL` to a SNF character.

`znconreylog` transforms `t_INT`, `t_INTMOD` and `t_VEC` to a Conrey logarithm.

`znconreyexp` transforms `t_VEC` and `t_COL` to a Conrey label.

Also available are `charconj`, `chardiv`, `charmulo`, `charker`, `chareval`, `charorder`, `zncharinduce`, `znconreyconductor` (also computes the primitive character attached to the input character). The prefix `char` indicates that the function applies to all characters, the prefix `znchar` that it is specific to Dirichlet characters (on $(\mathbf{Z}/N\mathbf{Z})^*$) and the prefix `znconrey` that it is specific to Conrey representation.

3.8.4 addprimes($\{x = []\}$). Adds the integers contained in the vector x (or the single integer x) to a special table of “user-defined primes”, and returns that table. Whenever `factor` is subsequently called, it will trial divide by the elements in this table. If x is empty or omitted, just returns the current list of extra primes.

```
? addprimes(37975227936943673922808872755445627854565536638199)
? factor(15226050279225333605356183781326374297180681149613806\
      88657908494580122963258952897654000350692006139)
%2 =
[37975227936943673922808872755445627854565536638199 1]
[40094690950920881030683735292761468389214899724061 1]
? ##
*** last result computed in 0 ms.
```

The entries in x must be primes: there is no internal check, even if the `factor_proven` default is set. To remove primes from the list use `removeprimes`.

The library syntax is `GEN addprimes(GEN x = NULL)`.

3.8.5 bestappr($x, \{B\}$). Using variants of the extended Euclidean algorithm, returns a rational approximation a/b to x , whose denominator is limited by B , if present. If B is omitted, returns the best approximation affordable given the input accuracy; if you are looking for true rational numbers, presumably approximated to sufficient accuracy, you should first try that option. Otherwise, B must be a positive real scalar (impose $0 < b \leq B$).

- If x is a `t_REAL` or a `t_FRAC`, this function uses continued fractions.

```
? bestappr(Pi, 100)
%1 = 22/7
? bestappr(0.1428571428571428571428571429)
%2 = 1/7
? bestappr([Pi, sqrt(2) + 'x], 10^3)
%3 = [355/113, x + 1393/985]
```

By definition, a/b is the best rational approximation to x if $|bx - a| < |vx - u|$ for all integers (u, v) with $0 < v \leq B$. (Which implies that n/d is a convergent of the continued fraction of x .)

- If x is a `t_INTMOD` modulo N or a `t_PADIC` of precision $N = p^k$, this function performs rational modular reconstruction modulo N . The routine then returns the unique rational number a/b in coprime integers $|a| < N/2B$ and $b \leq B$ which is congruent to x modulo N . Omitting B amounts to choosing it of the order of $\sqrt{N/2}$. If rational reconstruction is not possible (no suitable a/b exists), returns `[]`.

```
? bestappr(Mod(18526731858, 11^10))
%1 = 1/7
? bestappr(Mod(18526731858, 11^20))
%2 = []
? bestappr(3 + 5 + 3*5^2 + 5^3 + 3*5^4 + 5^5 + 3*5^6 + 0(5^7))
%2 = -1/3
```

In most concrete uses, B is a prime power and we performed Hensel lifting to obtain x .

The function applies recursively to components of complex objects (polynomials, vectors, ...). If rational reconstruction fails for even a single entry, returns `[]`.

The library syntax is `GEN bestappr(GEN x, GEN B = NULL)`.

3.8.6 bestapprPade($x, \{B\}, \{Q\}$). Using variants of the extended Euclidean algorithm (Padé approximants), returns a rational function approximation a/b to x , whose denominator is limited by B , if present. If B is omitted, return the best approximation affordable given the input accuracy; if you are looking for true rational functions, presumably approximated to sufficient accuracy, you should first try that option. Otherwise, B must be a nonnegative real (impose $0 \leq \text{degree}(b) \leq B$).

- If x is a `t_POLMOD` modulo N this function performs rational modular reconstruction modulo N . The routine then returns the unique rational function a/b in coprime polynomials, with $\text{degree}(b) \leq B$ and $\text{degree}(a)$ minimal, which is congruent to x modulo N . Omitting B amounts to choosing it equal to the floor of $\text{degree}(N)/2$. If rational reconstruction is not possible (no suitable a/b exists), returns `[]`.

```
? T = Mod(x^3 + x^2 + x + 3, x^4 - 2);
? bestapprPade(T)
%2 = (2*x - 1)/(x - 1)
? U = Mod(1 + x + x^2 + x^3 + x^5, x^9);
```

```
? bestapprPade(U)  \\ internally chooses B = 4
%3 = []
? bestapprPade(U, 5)  \\ with B = 5, a solution exists
%4 = (2*x^4 + x^3 - x - 1)/(-x^5 + x^3 + x^2 - 1)
```

- If x is a `t_SER`, we implicitly convert the input to a `t_POLMOD` modulo $N = t^k$ where k is the series absolute precision.

```
? T = 1 + t + t^2 + t^3 + t^4 + t^5 + t^6 + O(t^7);  \\ mod t^7
? bestapprPade(T)
%1 = 1/(-t + 1)
```

- If x is a `t_SER` and both B and Q are nonnegative, returns a rational function approximation a/b to x , with a of degree at most B and b of degree at most Q , with $x - a/b = O(t^{B+Q+1+v})$ if t is the variable, where v is the valuation of x , the empty vector if not possible.

- If x is a `t_RFRAC`, we implicitly convert the input to a `t_POLMOD` modulo $N = t^k$ where $k = 2B + 1$. If B was omitted, we return x :

```
? T = (4*t^2 + 2*t + 3)/(t+1)^10;
? bestapprPade(T,1)
%2 = []  \\ impossible
? bestapprPade(T,2)
%3 = 27/(337*t^2 + 84*t + 9)
? bestapprPade(T,3)
%4 = (4253*t - 3345)/(-39007*t^3 - 28519*t^2 - 8989*t - 1115)
```

The function applies recursively to components of complex objects (polynomials, vectors, ...). If rational reconstruction fails for even a single entry, return `[]`.

The library syntax is `GEN bestapprPade0(GEN x, long B, long Q)`.

`GEN bestapprPade(GEN x, long B)` as `bestapprPade0` when Q is omitted.

3.8.7 `bezout(x, y)`. Deprecated alias for `gcdext`

The library syntax is `GEN gcdext0(GEN x, GEN y)`.

3.8.8 `bigomega(x)`. Number of prime divisors of the integer $|x|$ counted with multiplicity:

```
? factor(392)
%1 =
[2 3]
[7 2]
? bigomega(392)
%2 = 5;  \\ = 3+2
? omega(392)
%3 = 2;  \\ without multiplicity
```

The library syntax is `long bigomega(GEN x)`.

3.8.9 charconj(*cyc*, *chi*). Let *cyc* represent a finite abelian group by its elementary divisors, i.e. (d_j) represents $\sum_{j \leq k} \mathbf{Z}/d_j \mathbf{Z}$ with $d_k \mid \dots \mid d_1$; any object which has a `.cyc` method is also allowed, e.g. the output of `znstar` or `bnrinit`. A character on this group is given by a row vector $\chi = [a_1, \dots, a_n]$ such that $\chi(\prod g_j^{n_j}) = \exp(2\pi i \sum_j a_j n_j / d_j)$, where g_j denotes the generator (of order d_j) of the j -th cyclic component. This function returns the conjugate character.

```
? cyc = [15,5]; chi = [1,1];
? charconj(cyc, chi)
%2 = [14, 4]
? bnf = bnfinit(x^2+23);
? bnf.cyc
%4 = [3]
? charconj(bnf, [1])
%5 = [2]
```

For Dirichlet characters (when *cyc* is `znstar(q,1)`), characters in Conrey representation are available, see Section 3.8.3 or `??character`:

```
? G = znstar(16, 1); \\ (Z/16Z)^*
? charconj(G, 3) \\ Conrey label
%2 = [1, 1]~
? znconreyexp(G, %)
%3 = 11 \\ attached Conrey label; indeed 11 = 3^(-1) mod 16
? chi = znconreylog(G, 3);
? charconj(G, chi) \\ Conrey logarithm
%5 = [1, 1]~
```

The library syntax is `GEN charconj0(GEN cyc, GEN chi)`. Also available is `GEN charconj(GEN cyc, GEN chi)`, when *cyc* is known to be a vector of elementary divisors and *chi* a compatible character (no checks).

3.8.10 chardiv(*cyc*, *a*, *b*). Let *cyc* represent a finite abelian group by its elementary divisors, i.e. (d_j) represents $\sum_{j \leq k} \mathbf{Z}/d_j \mathbf{Z}$ with $d_k \mid \dots \mid d_1$; any object which has a `.cyc` method is also allowed, e.g. the output of `znstar` or `bnrinit`. A character on this group is given by a row vector $a = [a_1, \dots, a_n]$ such that $\chi(\prod g_j^{n_j}) = \exp(2\pi i \sum_j a_j n_j / d_j)$, where g_j denotes the generator (of order d_j) of the j -th cyclic component.

Given two characters *a* and *b*, return the character $a/b = a\bar{b}$.

```
? cyc = [15,5]; a = [1,1]; b = [2,4];
? chardiv(cyc, a,b)
%2 = [14, 2]
? bnf = bnfinit(x^2+23);
? bnf.cyc
%4 = [3]
? chardiv(bnf, [1], [2])
%5 = [2]
```

For Dirichlet characters on $(\mathbf{Z}/N\mathbf{Z})^*$, additional representations are available (Conrey labels, Conrey logarithm), see Section 3.8.3 or `??character`. If the two characters are in the same format, the result is given in the same format, otherwise a Conrey logarithm is used.

```
? G = znstar(100, 1);
```

```

? G.cyc
%2 = [20, 2]
? a = [10, 1]; \\ usual representation for characters
? b = 7; \\ Conrey label;
? c = znconreylog(G, 11); \\ Conrey log
? chardiv(G, b,b)
%6 = 1 \\ Conrey label
? chardiv(G, a,b)
%7 = [0, 5]~ \\ Conrey log
? chardiv(G, a,c)
%7 = [0, 14]~ \\ Conrey log

```

The library syntax is `GEN chardiv0(GEN cyc, GEN a, GEN b)`. Also available is `GEN chardiv(GEN cyc, GEN a, GEN b)`, when `cyc` is known to be a vector of elementary divisors and a, b are compatible characters (no checks).

3.8.11 chareval($G, chi, x, \{z\}$). Let G be an abelian group structure affording a discrete logarithm method, e.g. $G = \text{znstar}(N, 1)$ for $(\mathbf{Z}/N\mathbf{Z})^*$ or a `bnr` structure, let x be an element of G and let chi be a character of G (see the note below for details). This function returns the value of chi at x .

Note on characters. Let K be some field. If G is an abelian group, let $\chi : G \rightarrow K^*$ be a character of finite order and let o be a multiple of the character order such that $\chi(n) = \zeta^{c(n)}$ for some fixed $\zeta \in K^*$ of multiplicative order o and a unique morphism $c : G \rightarrow (\mathbf{Z}/o\mathbf{Z}, +)$. Our usual convention is to write

$$G = (\mathbf{Z}/o_1\mathbf{Z})g_1 \oplus \cdots \oplus (\mathbf{Z}/o_d\mathbf{Z})g_d$$

for some generators (g_i) of respective order d_i , where the group has exponent $o := \text{lcm}_i o_i$. Since $\zeta^o = 1$, the vector (c_i) in $\prod_i (\mathbf{Z}/o_i\mathbf{Z})$ defines a character χ on G via $\chi(g_i) = \zeta^{c_i(o/o_i)}$ for all i . Classical Dirichlet characters have values in $K = \mathbf{C}$ and we can take $\zeta = \exp(2i\pi/o)$.

Note on Dirichlet characters. In the special case where bid is attached to $G = (\mathbf{Z}/q\mathbf{Z})^*$ (as per `G = znstar(q, 1)`), the Dirichlet character chi can be written in one of the usual 3 formats: a `t_VEC` in terms of `bid.gen` as above, a `t_COL` in terms of the Conrey generators, or a `t_INT` (Conrey label); see Section 3.8.3 or `??character`.

The character value is encoded as follows, depending on the optional argument z :

- If z is omitted: return the rational number $c(x)/o$ for x coprime to q , where we normalize $0 \leq c(x) < o$. If x can not be mapped to the group (e.g. x is not coprime to the conductor of a Dirichlet or Hecke character) we return the sentinel value -1 .
- If z is an integer o , then we assume that o is a multiple of the character order and we return the integer $c(x)$ when x belongs to the group, and the sentinel value -1 otherwise.
- z can be of the form $[zeta, o]$, where $zeta$ is an o -th root of 1 and o is a multiple of the character order. We return $\zeta^{c(x)}$ if x belongs to the group, and the sentinel value 0 otherwise. (Note that this coincides with the usual extension of Dirichlet characters to $\mathbf{Z}$, or of Hecke characters to general ideals.)
- Finally, z can be of the form $[vzeta, o]$, where $vzeta$ is a vector of powers $\zeta^0, \dots, \zeta^{o-1}$ of some o -th root of 1 and o is a multiple of the character order. As above, we return $\zeta^{c(x)}$ after a table lookup. Or the sentinel value 0.

The library syntax is `GEN chareval(GEN G, GEN chi, GEN x, GEN z = NULL)`.

3.8.12 chargalois(*cyc*, {*ORD*}). Let *cyc* represent a finite abelian group by its elementary divisors (any object which has a `.cyc` method is also allowed, i.e. the output of `znstar` or `bnrinit`). Return a list of representatives for the Galois orbits of complex characters of G . If *ORD* is present, select characters depending on their orders:

- if *ORD* is a `t_INT`, restrict to orders less than this bound;
- if *ORD* is a `t_VEC` or `t_VECSMALL`, restrict to orders in the list.

```
? G = znstar(96);
? #chargalois(G) \\ 16 orbits of characters mod 96
%2 = 16
? #chargalois(G,4) \\ order less than 4
%3 = 12
? chargalois(G,[1,4]) \\ order 1 or 4; 5 orbits
%4 = [[0, 0, 0], [2, 0, 0], [2, 1, 0], [2, 0, 1], [2, 1, 1]]
```

Given a character χ , of order n (`charorder(G,chi)`), the elements in its orbit are the $\phi(n)$ characters χ^i , $(i, n) = 1$.

The library syntax is `GEN chargalois(GEN cyc, GEN ORD = NULL)`.

3.8.13 charker(*cyc*, *chi*). Let *cyc* represent a finite abelian group by its elementary divisors, i.e. (d_j) represents $\sum_{j \leq k} \mathbf{Z}/d_j \mathbf{Z}$ with $d_k \mid \dots \mid d_1$; any object which has a `.cyc` method is also allowed, e.g. the output of `znstar` or `bnrinit`. A character on this group is given by a row vector $\chi = [a_1, \dots, a_n]$ such that $\chi(\prod g_j^{n_j}) = \exp(2\pi i \sum_j a_j n_j / d_j)$, where g_j denotes the generator (of order d_j) of the j -th cyclic component.

This function returns the kernel of χ , as a matrix K in HNF which is a left-divisor of `matdiagonal(d)`. Its columns express in terms of the g_j the generators of the subgroup. The determinant of K is the kernel index.

```
? cyc = [15,5]; chi = [1,1];
? charker(cyc, chi)
%2 =
[15 12]
[ 0  1]
? bnf = bnfinit(x^2+23);
? bnf.cyc
%4 = [3]
? charker(bnf, [1])
%5 =
[3]
```

Note that for Dirichlet characters (when *cyc* is `znstar(q, 1)`), characters in Conrey representation are available, see Section 3.8.3 or `??character`.

```
? G = znstar(8, 1); \\ (Z/8Z)^*
? charker(G, 1) \\ Conrey label for trivial character
%2 =
[1 0]
[0 1]
```

The library syntax is `GEN charker0(GEN cyc, GEN chi)`. Also available is `GEN charker(GEN cyc, GEN chi)`, when `cyc` is known to be a vector of elementary divisors and `chi` a compatible character (no checks).

3.8.14 charmul(*cyc*, *a*, *b*). Let *cyc* represent a finite abelian group by its elementary divisors, i.e. (d_j) represents $\sum_{j \leq k} \mathbf{Z}/d_j \mathbf{Z}$ with $d_k \mid \dots \mid d_1$; any object which has a `.cyc` method is also allowed, e.g. the output of `znstar` or `bnrinit`. A character on this group is given by a row vector $\chi = [a_1, \dots, a_n]$ such that $\chi(\prod g_j^{n_j}) = \exp(2\pi i \sum_j a_j n_j / d_j)$, where g_j denotes the generator (of order d_j) of the j -th cyclic component.

Given two characters *a* and *b*, return the product character *ab*.

```
? cyc = [15,5]; a = [1,1]; b = [2,4];
? charmul(cyc, a,b)
%2 = [3, 0]
? bnf = bnfinit(x^2+23);
? bnf.cyc
%4 = [3]
? charmul(bnf, [1], [2])
%5 = [0]
```

For Dirichlet characters on $(\mathbf{Z}/N\mathbf{Z})^*$, additional representations are available (Conrey labels, Conrey logarithm), see Section 3.8.3 or `??character`. If the two characters are in the same format, their product is given in the same format, otherwise a Conrey logarithm is used.

```
? G = znstar(100, 1);
? G.cyc
%2 = [20, 2]
? a = [10, 1]; \\ usual representation for characters
? b = 7; \\ Conrey label;
? c = znconreylog(G, 11); \\ Conrey log
? charmul(G, b,b)
%6 = 49 \\ Conrey label
? charmul(G, a,b)
%7 = [0, 15]~ \\ Conrey log
? charmul(G, a,c)
%7 = [0, 6]~ \\ Conrey log
```

The library syntax is `GEN charmul0(GEN cyc, GEN a, GEN b)`. Also available is `GEN charmul(GEN cyc, GEN a, GEN b)`, when `cyc` is known to be a vector of elementary divisors and *a*, *b* are compatible characters (no checks).

3.8.15 charorder(*cyc*, *chi*). Let *cyc* represent a finite abelian group by its elementary divisors, i.e. (d_j) represents $\sum_{j \leq k} \mathbf{Z}/d_j \mathbf{Z}$ with $d_k \mid \dots \mid d_1$; any object which has a `.cyc` method is also allowed, e.g. the output of `znstar` or `bnrinit`. A character on this group is given by a row vector $\chi = [a_1, \dots, a_n]$ such that $\chi(\prod g_j^{n_j}) = \exp(2\pi i \sum_j a_j n_j / d_j)$, where g_j denotes the generator (of order d_j) of the j -th cyclic component.

This function returns the order of the character *chi*.

```
? cyc = [15,5]; chi = [1,1];
? charorder(cyc, chi)
```

```

%2 = 15
? bnf = bnfinit(x^2+23);
? bnf.cyc
%4 = [3]
? charorder(bnf, [1])
%5 = 3

```

For Dirichlet characters (when `cyc` is `znstar(q, 1)`), characters in Conrey representation are available, see Section 3.8.3 or `??character`:

```

? G = znstar(100, 1); \\ (Z/100Z)^*
? charorder(G, 7)    \\ Conrey label
%2 = 4

```

The library syntax is `GEN charorder0(GEN cyc, GEN chi)`. Also available is `GEN charorder(GEN cyc, GEN chi)`, when `cyc` is known to be a vector of elementary divisors and `chi` a compatible character (no checks).

3.8.16 charpow(*cyc*, *a*, *n*). Let *cyc* represent a finite abelian group by its elementary divisors, i.e. (d_j) represents $\sum_{j \leq k} \mathbf{Z}/d_j \mathbf{Z}$ with $d_k \mid \dots \mid d_1$; any object which has a `.cyc` method is also allowed, e.g. the output of `znstar` or `bnrinit`. A character on this group is given by a row vector $\chi = [a_1, \dots, a_n]$ such that $\chi(\prod g_j^{n_j}) = \exp(2\pi i \sum_j a_j n_j / d_j)$, where g_j denotes the generator (of order d_j) of the j -th cyclic component.

Given $n \in \mathbf{Z}$ and a character a , return the character a^n .

```

? cyc = [15,5]; a = [1,1];
? charpow(cyc, a, 3)
%2 = [3, 3]
? charpow(cyc, a, 5)
%2 = [5, 0]
? bnf = bnfinit(x^2+23);
? bnf.cyc
%4 = [3]
? charpow(bnf, [1], 3)
%5 = [0]

```

For Dirichlet characters on $(\mathbf{Z}/N\mathbf{Z})^*$, additional representations are available (Conrey labels, Conrey logarithm), see Section 3.8.3 or `??character` and the output uses the same format as the input.

```

? G = znstar(100, 1);
? G.cyc
%2 = [20, 2]
? a = [10, 1]; \\ standard representation for characters
? b = 7; \\ Conrey label;
? c = znconreylog(G, 11); \\ Conrey log
? charpow(G, a, 3)
%6 = [10, 1]    \\ standard representation
? charpow(G, b, 3)
%7 = 43        \\ Conrey label
? charpow(G, c, 3)

```

```
%8 = [1, 8]~ \\ Conrey log
```

The library syntax is `GEN charpow0(GEN cyc, GEN a, GEN n)`. Also available is `GEN charpow(GEN cyc, GEN a, GEN n)`, when `cyc` is known to be a vector of elementary divisors (no check).

3.8.17 chinese($x, \{y\}$). If x and y are both intmods or both polmods, creates (with the same type) a z in the same residue class as x and in the same residue class as y , if it is possible.

```
? chinese(Mod(1,2), Mod(2,3))
%1 = Mod(5, 6)
? chinese(Mod(x,x^2-1), Mod(x+1,x^2+1))
%2 = Mod(-1/2*x^2 + x + 1/2, x^4 - 1)
```

This function also allows vector and matrix arguments, in which case the operation is recursively applied to each component of the vector or matrix.

```
? chinese([Mod(1,2),Mod(1,3)], [Mod(1,5),Mod(2,7)])
%3 = [Mod(1, 10), Mod(16, 21)]
```

For polynomial arguments in the same variable, the function is applied to each coefficient; if the polynomials have different degrees, the high degree terms are copied verbatim in the result, as if the missing high degree terms in the polynomial of lowest degree had been `Mod(0,1)`. Since the latter behavior is usually *not* the desired one, we propose to convert the polynomials to vectors of the same length first:

```
? P = x+1; Q = x^2+2*x+1;
? chinese(P*Mod(1,2), Q*Mod(1,3))
%4 = Mod(1, 3)*x^2 + Mod(5, 6)*x + Mod(3, 6)
? chinese(Vec(P,3)*Mod(1,2), Vec(Q,3)*Mod(1,3))
%5 = [Mod(1, 6), Mod(5, 6), Mod(4, 6)]
? Pol(%)
%6 = Mod(1, 6)*x^2 + Mod(5, 6)*x + Mod(4, 6)
```

If y is omitted, and x is a vector, `chinese` is applied recursively to the components of x , yielding a residue belonging to the same class as all components of x .

Finally `chinese`(x, x) = x regardless of the type of x ; this allows vector arguments to contain other data, so long as they are identical in both vectors.

The library syntax is `GEN chinese(GEN x, GEN y = NULL)`. `GEN chinese1(GEN x)` is also available.

3.8.18 content($x, \{D\}$). Computes the gcd of all the coefficients of x , when this gcd makes sense. This is the natural definition if x is a polynomial (and by extension a power series) or a vector/matrix. This is in general a weaker notion than the *ideal* generated by the coefficients:

```
? content(2*x+y)
%1 = 1 \\ = gcd(2,y) over Q[y]
```

If x is a scalar, this simply returns the absolute value of x if x is rational (`t_INT` or `t_FRAC`), and either 1 (inexact input) or x (exact input) otherwise; the result should be identical to `gcd(x, 0)`.

The content of a rational function is the ratio of the contents of the numerator and the denominator. In recursive structures, if a matrix or vector *coefficient* x appears, the gcd is taken not with x , but with its content:

```
? content([ [2], 4*matid(3) ])
%1 = 2
```

The content of a `t_VECSMALL` is computed assuming the entries are signed integers.

The optional argument D allows to control over which ring we compute and get a more predictable behaviour:

- 1: we only consider the underlying $\mathbf{Q}$ -structure and the denominator is a (positive) rational number
- a simple variable, say 'x': all entries are considered as rational functions in $K(x)$ for some field K and the content is an element of K .

```
? f = x + 1/y + 1/2;
? content(f) \\ as a t_POL in x
%2 = 1/(2*y)
? content(f, 1) \\ Q-content
%3 = 1/2
? content(f, y) \\ as a rational function in y
%4 = 1/2
? g = x^2*y + y^2*x;
? content(g, x)
%6 = y
? content(g, y)
%7 = x
```

The library syntax is `GEN content0(GEN x, GEN D = NULL)`.

3.8.19 contfrac($x, \{b\}, \{nmax\}$). Returns the row vector whose components are the partial quotients of the continued fraction expansion of x . In other words, a result $[a_0, \dots, a_n]$ means that $x \approx a_0 + 1/(a_1 + \dots + 1/a_n)$. The output is normalized so that $a_n \neq 1$ (unless we also have $n = 0$).

The number of partial quotients $n + 1$ is limited by `nmax`. If `nmax` is omitted, the expansion stops at the last significant partial quotient.

```
? \p19
  realprecision = 19 significant digits
? contfrac(Pi)
%1 = [3, 7, 15, 1, 292, 1, 1, 1, 2, 1, 3, 1, 14, 2, 1, 1, 2, 2]
? contfrac(Pi,, 3) \\ n = 2
%2 = [3, 7, 15]
```

x can also be a rational function or a power series.

If a vector b is supplied, the numerators are equal to the coefficients of b , instead of all equal to 1 as above; more precisely, $x \approx (1/b_0)(a_0 + b_1/(a_1 + \dots + b_n/a_n))$; for a numerical continued fraction (x real), the a_i are integers, as large as possible; if x is a rational function, they are polynomials with $\deg a_i = \deg b_i + 1$. The length of the result is then equal to the length of b , unless the next partial quotient cannot be reliably computed, in which case the expansion stops. This happens

when a partial remainder is equal to zero (or too small compared to the available significant digits for x a `t_REAL`).

A direct implementation of the numerical continued fraction `contfrac(x,b)` described above would be

```
\\ "greedy" generalized continued fraction
cf(x, b) =
{ my( a= vector(#b), t );
  x *= b[1];
  for (i = 1, #b,
    a[i] = floor(x);
    t = x - a[i]; if (!t || i == #b, break);
    x = b[i+1] / t;
  ); a;
}
```

There is some degree of freedom when choosing the a_i ; the program above can easily be modified to derive variants of the standard algorithm. In the same vein, although no builtin function implements the related Engel expansion (a special kind of Egyptian fraction decomposition: $x = 1/a_1 + 1/(a_1a_2) + \dots$), it can be obtained as follows:

```
\\ n terms of the Engel expansion of x
engel(x, n = 10) =
{ my( u = x, a = vector(n) );
  for (k = 1, n,
    a[k] = ceil(1/u);
    u = u*a[k] - 1;
    if (!u, break);
  ); a;
}
```

Obsolete hack. (don't use this): if b is an integer, $nmax$ is ignored and the command is understood as `contfrac(x, b)`.

The library syntax is `GEN contfrac0(GEN x, GEN b = NULL, long nmax)`. Also available are `GEN gboundcf(GEN x, long nmax)`, `GEN gcf(GEN x)` and `GEN gcf2(GEN b, GEN x)`.

3.8.20 contfracpnqn($x, \{n = -1\}$). When x is a vector or a one-row matrix, x is considered as the list of partial quotients $[a_0, a_1, \dots, a_n]$ of a rational number, and the result is the 2 by 2 matrix $[p_n, p_{n-1}; q_n, q_{n-1}]$ in the standard notation of continued fractions, so $p_n/q_n = a_0 + 1/(a_1 + \dots + 1/a_n)$. If x is a matrix with two rows $[b_0, b_1, \dots, b_n]$ and $[a_0, a_1, \dots, a_n]$, this is then considered as a generalized continued fraction and we have similarly $p_n/q_n = (1/b_0)(a_0 + b_1/(a_1 + \dots + b_n/a_n))$. Note that in this case one usually has $b_0 = 1$.

If $n \geq 0$ is present, returns all convergents from p_0/q_0 up to p_n/q_n . (All convergents if x is too small to compute the $n + 1$ requested convergents.)

```
? a = contfrac(Pi,10)
%1 = [3, 7, 15, 1, 292, 1, 1, 1, 3]
? allpnqn(x) = contfracpnqn(x,#x) \\ all convergents
? allpnqn(a)
```

```

%3 =
[3 22 333 355 103993 104348 208341 312689 1146408]

[1 7 106 113 33102 33215 66317 99532 364913]
? contfracpnqn(a) \\ last two convergents
%4 =
[1146408 312689]

[ 364913 99532]

? contfracpnqn(a,3) \\ first three convergents
%5 =
[3 22 333 355]

[1 7 106 113]

```

The library syntax is `GEN contfracpnqn(GEN x, long n)`. also available is `GEN pnqn(GEN x)` for $n = -1$.

3.8.21 `core`($n, \{flag = 0\}$). If n is an integer written as $n = df^2$ with d squarefree, returns d . If $flag$ is nonzero, returns the two-element row vector $[d, f]$. By convention, we write $0 = 0 \times 1^2$, so `core(0, 1)` returns $[0, 1]$.

The library syntax is `GEN core0(GEN n, long flag)`. Also available are `GEN core(GEN n)` ($flag = 0$) and `GEN core2(GEN n)` ($flag = 1$)

3.8.22 `coredisc`($n, \{flag = 0\}$). A *fundamental discriminant* is an integer of the form $t \equiv 1 \pmod{4}$ or $4t \equiv 8, 12 \pmod{16}$, with t squarefree (i.e. 1 or the discriminant of a quadratic number field). Given a nonzero integer n , this routine returns the (unique) fundamental discriminant d such that $n = df^2$, f a positive rational number. If $flag$ is nonzero, returns the two-element row vector $[d, f]$. If n is congruent to 0 or 1 modulo 4, f is an integer, and a half-integer otherwise.

By convention, `coredisc(0, 1)` returns $[0, 1]$.

Note that `quaddisc(n)` returns the same value as `coredisc(n)`, and also works with rational inputs $n \in \mathbf{Q}^*$.

The library syntax is `GEN coredisc0(GEN n, long flag)`. Also available are `GEN coredisc(GEN n)` ($flag = 0$) and `GEN coredisc2(GEN n)` ($flag = 1$)

3.8.23 `dirdiv`(x, y). x and y being vectors of perhaps different lengths but with $y[1] \neq 0$ considered as Dirichlet series, computes the quotient of x by y , again as a vector.

The library syntax is `GEN dirdiv(GEN x, GEN y)`.

3.8.24 direuler($p = a, b, \text{expr}, \{c\}$). Computes the Dirichlet series attached to the Euler product of expression expr as p ranges through the primes from a to b . expr must be a polynomial or rational function in another variable than p (say X) and $\text{expr}(X)$ is understood as the local factor $\text{expr}(p^{-s})$.

The series is output as a vector of coefficients. If c is omitted, output the first b coefficients of the series; otherwise, output the first c coefficients. The following command computes the **sigma** function, attached to $\zeta(s)\zeta(s-1)$:

```
? direuler(p=2, 10, 1/((1-X)*(1-p*X)))
%1 = [1, 3, 4, 7, 6, 12, 8, 15, 13, 18]

? direuler(p=2, 10, 1/((1-X)*(1-p*X)), 5) \\ fewer terms
%2 = [1, 3, 4, 7, 6]
```

Setting $c < b$ is useless (the same effect would be achieved by setting $b = c$). If $c > b$, the computed coefficients are “missing” Euler factors:

```
? direuler(p=2, 10, 1/((1-X)*(1-p*X)), 15) \\ more terms, no longer = sigma !
%3 = [1, 3, 4, 7, 6, 12, 8, 15, 13, 18, 0, 28, 0, 24, 24]
```

The library syntax is **direuler**(void *E, GEN (*eval)(void*,GEN), GEN a, GEN b)

3.8.25 dirmul(x, y). x and y being vectors of perhaps different lengths representing the Dirichlet series $\sum_n x_n n^{-s}$ and $\sum_n y_n n^{-s}$, computes the product of x by y , again as a vector.

```
? dirmul(vector(10,n,1), vector(10,n,moebius(n)))
%1 = [1, 0, 0, 0, 0, 0, 0, 0, 0, 0]
```

The product length is the minimum of $\#x*v(y)$ and $\#y*v(x)$, where $v(x)$ is the index of the first nonzero coefficient.

```
? dirmul([0,1], [0,1]);
%2 = [0, 0, 0, 1]
```

The library syntax is **GEN dirmul**(GEN x, GEN y).

3.8.26 dirpowerssum($N, x, \{f\}, \{both = 0\}$). For positive integer N and complex number x , return the sum $f(1)1^x + f(2)2^x + \dots + f(N)N^x$, where f is a completely multiplicative function. If f is omitted, return $1^x + \dots + N^x$. When $N \leq 0$, the function returns 0. If **both** is set, return the pair for arguments (x, f) and $(-1-x, \overline{f})$. If **both=2**, assume in addition that f is real-valued (which is true when f is omitted, i.e. represents the constant function $f(n) = 1$).

Caveat. when `both` is set, the present implementation assumes that $|f(n)|$ is either 0 or 1, which is the case for Dirichlet characters.

A vector-valued multiplicative function f is allowed, in which case the above conditions must be met componentwise and the vector length must be constant.

Unlike variants using `dirpowers(N,x)`, this function uses $O(\sqrt{N})$ memory instead of $O(N)$. And it is faster for large N . The return value is usually a floating point number, but it will be exact if the result is an integer. On the other hand, rational numbers are converted to floating point approximations, since they are likely to blow up for large N .

```
? dirpowers(5, 2)
%1 = [1, 4, 9, 16, 25]
? vecsum(%)
%2 = 55
? dirpowersssum(5, 2)
%3 = 55
? dirpowersssum(5, -2)
%4 = 1.4636111111111111111111111111111111111111111111111111111
? \p200
? s = 1/2 + I * sqrt(3); N = 10^7;
? dirpowersssum(N, s);
time = 11,425 ms.
? vecsum(dirpowers(N, s))
time = 19,365 ms.
? dirpowersssum(N, s, n->kronecker(-23,n))
time = 10,981 ms.
```

The `dirpowersssum` commands work with default stack size, the `dirpowers` one requires a stacksize of at least 5GB.

The library syntax is `dirpowersssumfun(ulong N, GEN x, void *E, GEN (*f)(void*, ulong, long), long prec)`. When $f = \text{NULL}$, one may use `GEN dirpowersssum(ulong N, GEN x, long prec)`.

3.8.27 divisors($x, \{flag = 0\}$). Creates a row vector whose components are the divisors of x . The factorization of x (as output by `factor`) can be used instead. If $flag = 1$, return pairs $[d, \text{factor}(d)]$.

By definition, these divisors are the products of the irreducible factors of n , as produced by `factor(n)`, raised to appropriate powers (no negative exponent may occur in the factorization). If n is an integer, they are the positive divisors, in increasing order.

```
? divisors(12)
%1 = [1, 2, 3, 4, 6, 12]
? divisors(12, 1) \ include their factorization
%2 = [[1, matrix(0,2)], [2, Mat([2, 1])], [3, Mat([3, 1])],
      [4, Mat([2, 2])], [6, [2, 1; 3, 1]], [12, [2, 2; 3, 1]]]
? divisors(x^4 + 2*x^3 + x^2) \ also works for polynomials
%3 = [1, x, x^2, x + 1, x^2 + x, x^3 + x^2, x^2 + 2*x + 1,
      x^3 + 2*x^2 + x, x^4 + 2*x^3 + x^2]
```

This function requires a lot of memory if x has many divisors. The following idiom runs through all divisors using very little memory, in no particular order this time:

```
F = factor(x); P = F[,1]; E = F[,2];
forvec(e = vectorv(#E,i,[0,E[i]]), d = factorback(P,e); ...)
```

If the factorization of d is also desired, then $[P, e]$ almost provides it but not quite: e may contain 0 exponents, which are not allowed in factorizations. These must be sieved out as in:

```
? tofact(P,E) = matreduce(Mat([P,E]));
? tofact([2,3,5,7]~, [4,0,2,0]~)
%4 =
[2 4]
[5 2]
```

We can then run the above loop with `tofact(P,e)` instead of, or together with, `factorback`.

The library syntax is `GEN divisors0(GEN x, long flag)`. The functions `GEN divisors(GEN N)` ($flag = 0$) and `GEN divisors_factored(GEN N)` ($flag = 1$) are also available.

3.8.28 divisorslenstra(N, r, s). Given three integers $N > s > r \geq 0$ such that $(r, s) = 1$ and $s^3 > N$, find all divisors d of N such that $d \equiv r \pmod{s}$. There are at most 11 such divisors (Lenstra).

```
? N = 245784; r = 19; s = 65 ;
? divisorslenstra(N, r, s)
%2 = [19, 84, 539, 1254, 3724, 245784]
? [ d | d <- divisors(N), d % s == r ]
%3 = [19, 84, 539, 1254, 3724, 245784]
```

When the preconditions are not met, the result is undefined:

```
? N = 4484075232; r = 7; s = 1303; s^3 > N
%4 = 0
? divisorslenstra(N, r, s)
? [ d | d <- divisors(N), d % s == r ]
%6 = [7, 2613, 9128, 19552, 264516, 3407352, 344928864]
```

(Divisors were missing but $s^3 < N$.)

The library syntax is `GEN divisorslenstra(GEN N, GEN r, GEN s)`.

3.8.29 eulerphi(x). Euler's ϕ (totient) function of the integer $|x|$, in other words $|(\mathbf{Z}/x\mathbf{Z})^*|$.

```
? eulerphi(40)
%1 = 16
```

According to this definition we let $\phi(0) := 2$, since $\mathbf{Z}^* = \{-1, 1\}$; this is consistent with `znstar(0)`: we have `znstar(n).no = eulerphi(n)` for all $n \in \mathbf{Z}$.

The library syntax is `GEN eulerphi(GEN x)`.

3.8.30 factor($x, \{D\}$). Factor x over domain D ; if D is omitted, it is determined from x . For instance, if x is an integer, it is factored in $\mathbf{Z}$, if it is a polynomial with rational coefficients, it is factored in $\mathbf{Q}[x]$, etc., see below for details. The result is a two-column matrix: the first contains the irreducibles dividing x (rational or Gaussian primes, irreducible polynomials), and the second the exponents. By convention, 0 is factored as 0^1 .

$x \in \mathbf{Q}$. See `factorint` for the algorithms used. The factorization includes the unit -1 when $x < 0$ and all other factors are positive; a denominator is factored with negative exponents. The factors are sorted in increasing order.

```
? factor(-7/106)
%1 =
[-1  1]
[ 2 -1]
[ 7  1]
[53 -1]
```

By convention, 1 is factored as `matrix(0,2)` (the empty factorization, printed as `[]`).

Large rational “primes” $> 2^{64}$ in the factorization are in fact *pseudoprimes* (see `ispseudo-prime`), a priori not rigorously proven primes. Use `isprime` to prove primality of these factors, as in

```
? fa = factor(2^2^7 + 1)
%2 =
[59649589127497217 1]
[5704689200685129054721 1]
? isprime( fa[,1] )
%3 = [1, 1]~  \\ both entries are proven primes
```

Another possibility is to globally set the default `factor_proven`, which will perform a rigorous primality proof for each pseudoprime factor but will slow down PARI.

A `t_INT` argument D can be added, meaning that we only trial divide by all primes $p < D$ and the `addprimes` entries, then skip all expensive factorization methods. The limit D must be nonnegative. In this case, one entry in the factorization may be a composite number: all factors less than D^2 and primes from the `addprimes` table are actual primes. But (at most) one entry may not verify this criterion, and it may be prime or composite: it is only known to be coprime to all other entries and not a pure power.

```
? factor(2^2^7 + 1, 10^5)
%4 =
[340282366920938463463374607431768211457 1]
```

Deprecated feature. Setting $D = 0$ is the same as setting it to `factorlimit + 1`.

This routine uses trial division and perfect power tests, and should not be used for huge values of D (at most 10^9 , say): `factorint(, 1 + 8)` will in general be faster. The latter does not guarantee that all small prime factors are found, but it also finds larger factors and in a more efficient way.

```
? F = (2^2^7 + 1) * 1009 * (10^5+3); factor(F, 10^5)  \\ fast, incomplete
time = 0 ms.
%5 =
[1009 1]
[34029257539194609161727850866999116450334371 1]
? factor(F, 10^9)  \\ slow
time = 3,260 ms.
%6 =
[1009 1]
[100003 1]
[340282366920938463463374607431768211457 1]
? factorint(F, 1+8)  \\ much faster and all small primes were found
time = 8 ms.
%7 =
[1009 1]
[100003 1]
[340282366920938463463374607431768211457 1]
? factor(F)  \\ complete factorization
time = 60 ms.
%8 =
[1009 1]
[100003 1]
[59649589127497217 1]
[5704689200685129054721 1]
```

$x \in \mathbf{Q}(i)$. The factorization is performed with Gaussian primes in $\mathbf{Z}[i]$ and includes Gaussian units in $\{\pm 1, \pm i\}$; factors are sorted by increasing norm. Except for a possible leading unit, the Gaussian factors are normalized: rational factors are positive and irrational factors have positive imaginary part.

Unless `factor_proven` is set, large factors are actually pseudoprimes, not proven primes; a rational factor is prime if less than 2^{64} and an irrational one if its norm is less than 2^{64} .

```
? factor(5*I)
%9 =
[ 2 + I 1]
[1 + 2*I 1]
```

One can force the factorization of a rational number by setting the domain $D = I$:

```
? factor(-5, I)
```

```
%10 =
[      I 1]
[  2 + I 1]
[1 + 2*I 1]
? factorback(%)
%11 = -5
```

Univariate polynomials and rational functions. PARI can factor univariate polynomials in $K[t]$. The following base fields K are currently supported: $\mathbf{Q}$, $\mathbf{R}$, $\mathbf{C}$, $\mathbf{Q}_p$, finite fields and number fields. See `factormod` and `factorff` for the algorithms used over finite fields and `nffactor` for the algorithms over number fields. The irreducible factors are sorted by increasing degree and normalized: they are monic except when $K = \mathbf{Q}$ where they are primitive in $\mathbf{Z}[t]$.

The content is *not* included in the factorization, in particular `factorback` will in general recover the original x only up to multiplication by an element of K^* : when $K \neq \mathbf{Q}$, this scalar is `pollead(x)` (since irreducible factors are monic); and when $K = \mathbf{Q}$ you can either ask for the $\mathbf{Q}$ -content explicitly or use `factorback`:

```
? P = t^2 + 5*t/2 + 1; F = factor(P)
%12 =
[t + 2 1]
[2*t + 1 1]
? content(P, 1) \\ Q-content
%13 = 1/2
? pollead(factorback(F)) / pollead(P)
%14 = 2
```

You can specify K using the optional “domain” argument D as follows

- $K = \mathbf{Q}$: D a rational number (`t_INT` or `t_FRAC`),
- $K = \mathbf{Z}/p\mathbf{Z}$ with p prime : D a `t_INTMOD` modulo p ; factoring modulo a composite number is not supported.
- $K = \mathbf{F}_q$: D a `t_FFELT` encoding the finite field; you can also use a `t_POLMOD` of `t_INTMOD` modulo a prime p but this is usually less convenient;
- $K = \mathbf{Q}[X]/(T)$ a number field : D a `t_POLMOD` modulo T ,
- $K = \mathbf{Q}(i)$ (alternate syntax for special case): $D = I$,
- $K = \mathbf{Q}(w)$ a quadratic number field (alternate syntax for special case): D a `t_QUAD`,
- $K = \mathbf{R}$: D a real number (`t_REAL`); truncate the factorization at accuracy `precision(D)`. If x is inexact and `precision(x)` is less than `precision(D)`, then the precision of x is used instead.
- $K = \mathbf{C}$: D a complex number with a `t_REAL` component, e.g. `I * 1.`; truncate the factorization as for $K = \mathbf{R}$,
- $K = \mathbf{Q}_p$: D a `t_PADIC`; truncate the factorization at p -adic accuracy `padicprec(D)`, possibly less if x is inexact with insufficient p -adic accuracy;

```
? T = x^2+1;
? factor(T, 1); \\ over Q
```

```

? factor(T, Mod(1,3))           \\ over F_3
? factor(T, ffgen(ffinit(3,2,'t))^0) \\ over F_{3^2}
? factor(T, Mod(Mod(1,3), t^2+t+2)) \\ over F_{3^2}, again
? factor(T, 0(3^6))             \\ over Q_3, precision 6
? factor(T, 1.)                 \\ over R, current precision
? factor(T, I*1.)               \\ over C
? factor(T, Mod(1, y^3-2))      \\ over Q(2^{1/3})

```

In most cases, it is possible and simpler to call a specialized variant rather than use the above scheme:

```

? factormod(T, 3)                \\ over F_3
? factormod(T, [t^2+t+2, 3])     \\ over F_{3^2}
? factormod(T, ffgen([3,2], 't)) \\ over F_{3^2}
? factorpadic(T, 3,6)           \\ over Q_3, precision 6
? nffactor(y^3-2, T)            \\ over Q(2^{1/3})
? polroots(T)                   \\ over C
? polrootsreal(T)               \\ over R (real polynomial)

```

It is also possible to let the routine use the smallest field containing all coefficients, taking into account quotient structures induced by `t_INTMODs` and `t_POLMODs` (e.g. if a coefficient in $\mathbf{Z}/n\mathbf{Z}$ is known, all rational numbers encountered are first mapped to $\mathbf{Z}/n\mathbf{Z}$; different moduli will produce an error):

```

? T = x^2+1;
? factor(T);                     \\ over Q
? factor(T*Mod(1,3))            \\ over F_3
? factor(T*ffgen(ffinit(3,2,'t))^0) \\ over F_{3^2}
? factor(T*Mod(Mod(1,3), t^2+t+2)) \\ over F_{3^2}, again
? factor(T*(1 + 0(3^6)))        \\ over Q_3, precision 6
? factor(T*1.)                  \\ over R, current precision
? factor(T*(1.+0.*I))           \\ over C
? factor(T*Mod(1, y^3-2))       \\ over Q(2^{1/3})

```

Multiplying by a suitable field element equal to $1 \in K$ in this way is error-prone and is not recommended. Factoring existing polynomials with obvious fields of coefficients is fine, the domain argument `D` should be used instead ad hoc conversions.

Note on inexact polynomials. Polynomials with inexact coefficients (e.g. floating point or p -adic numbers) are first rounded to an exact representation, then factored to (potentially) infinite accuracy and we return a truncated approximation of that virtual factorization. To avoid pitfalls, we advise to only factor *exact* polynomials:

```

? factor(x^2-1+0(2^2)) \\ rounded to x^2 + 3, irreducible in Q_2
%1 =
[(1 + 0(2^2))*x^2 + 0(2^2)*x + (1 + 2 + 0(2^2)) 1]
? factor(x^2-1+0(2^3)) \\ rounded to x^2 + 7, reducible !
%2 =
[ (1 + 0(2^3))*x + (1 + 2 + 0(2^3)) 1]
[(1 + 0(2^3))*x + (1 + 2^2 + 0(2^3)) 1]
? factor(x^2-1, 0(2^2)) \\ no ambiguity now

```

```
%3 =
[      (1 + 0(2^2))*x + (1 + 0(2^2)) 1]
[(1 + 0(2^2))*x + (1 + 2 + 0(2^2)) 1]
```

Note about inseparable polynomials. Polynomials with inexact coefficients are considered to be squarefree: indeed, there exist a squarefree polynomial arbitrarily close to the input, and they cannot be distinguished at the input accuracy. This means that irreducible factors are repeated according to their apparent multiplicity. On the contrary, using a specialized function such as `factorpadic` with an *exact* rational input yields the correct multiplicity when the (now exact) input is not separable. Compare:

```
? factor(z^2 + 0(5^2))
%1 =
[(1 + 0(5^2))*z + 0(5^2) 1]
[(1 + 0(5^2))*z + 0(5^2) 1]
? factor(z^2, 0(5^2))
%2 =
[1 + 0(5^2))*z + 0(5^2) 2]
```

Multivariate polynomials and rational functions. PARI recursively factors *multivariate* polynomials in $K[t_1, \dots, t_d]$ for the same fields K as above and the argument D is used in the same way to specify K . The irreducible factors are sorted by their main variable (least priority first) then by increasing degree.

```
? factor(x^2 + y^2, Mod(1,5))
%1 =
[      x + Mod(2, 5)*y 1]
[Mod(1, 5)*x + Mod(3, 5)*y 1]
? factor(x^2 + y^2, 0(5^2))
%2 =
[ (1 + 0(5^2))*x + (0(5^2)*y^2 + (2 + 5 + 0(5^2))*y + 0(5^2)) 1]
[(1 + 0(5^2))*x + (0(5^2)*y^2 + (3 + 3*5 + 0(5^2))*y + 0(5^2)) 1]
? lift(%)
%3 =
[ x + 7*y 1]
[x + 18*y 1]
```

Note that the implementation does not really support inexact real fields (**R** or **C**) and usually misses factors even if the input is exact:

```
? factor(x^2 + y^2, I) \\ over Q(i)
%4 =
[x - I*y 1]
[x + I*y 1]
? factor(x^2 + y^2, I*1.) \\ over C
%5 =
[x^2 + y^2 1]
```

The library syntax is `GEN factor0(GEN x, GEN D = NULL)`.

`GEN factor(GEN x) GEN boundfact(GEN x, ulong lim).`

3.8.31 factorback($f, \{e\}$). Gives back the factored object corresponding to a factorization. The integer 1 corresponds to the empty factorization.

If e is present, e and f must be vectors of the same length (e being integral), and the corresponding factorization is the product of the $f[i]^{e[i]}$.

If not, and f is vector, it is understood as in the preceding case with e a vector of 1s: we return the product of the $f[i]$. Finally, f can be a regular factorization, as produced with any **factor** command. A few examples:

```
? factor(12)
%1 =
[2 2]
[3 1]
? factorback(%)
%2 = 12
? factorback([2,3], [2,1])    \\ 2^2 * 3^1
%3 = 12
? factorback([5,2,3])
%4 = 30
```

The library syntax is `GEN factorback2(GEN f, GEN e = NULL)`. Also available is `GEN factorback(GEN f)` (case $e = \text{NULL}$).

3.8.32 factorcantor(x, p). This function is obsolete, use `factormod`.

The library syntax is `GEN factmod(GEN x, GEN p)`.

3.8.33 factorff($x, \{p\}, \{a\}$). Obsolete, kept for backward compatibility: use `factormod`.

The library syntax is `GEN factorff(GEN x, GEN p = NULL, GEN a = NULL)`.

3.8.34 factorial(x). Factorial of x . The expression $x!$ gives a result which is an integer, while `factorial( $x$ )` gives a real number.

The library syntax is `GEN mpfactr(long x, long prec)`. `GEN mpfact(long x)` returns $x!$ as a `t_INT`.

3.8.35 factorint($x, \{flag = 0\}$). Factors the integer n into a product of pseudoprimes (see `ispseudoprime`), using a combination of the Shanks SQUFOF and Pollard Rho method (with modifications due to Brent), Lenstra's ECM (with modifications by Montgomery), and MPQS (the latter adapted from the LiDIA code with the kind permission of the LiDIA maintainers), as well as a search for pure powers. The output is a two-column matrix as for **factor**: the first column contains the "prime" divisors of n , the second one contains the (positive) exponents.

By convention 0 is factored as 0^1 , and 1 as the empty factorization; also the divisors are by default not proven primes if they are larger than 2^{64} , they only failed the BPSW compositeness test (see `ispseudoprime`). Use `isprime` on the result if you want to guarantee primality or set the `factor_proven` default to 1. Entries of the private prime tables (see `addprimes`) are also included as is.

This gives direct access to the integer factoring engine called by most arithmetical functions. *flag* is optional; its binary digits mean 1: avoid MPQS, 2: skip first stage ECM (we may still fall back to it later), 4: avoid Rho and SQUFOF, 8: don't run final ECM (as a result, a huge composite may be declared to be prime). Note that a (strong) probabilistic primality test is used; thus composites might not be detected, although no example is known.

You are invited to play with the flag settings and watch the internals at work by using `gp`'s `debug` default parameter (level 3 shows just the outline, 4 turns on time keeping, 5 and above show an increasing amount of internal details).

The library syntax is `GEN factorint(GEN x, long flag)`.

3.8.36 factormod(*f*, {*D*}, {*flag* = 0}). Factors the polynomial *f* over the finite field defined by the domain *D* as follows:

- *D* = *p* a prime: factor over $\mathbf{F}_p$;
- *D* = [*T*, *p*] for a prime *p* and *T*(*y*) an irreducible polynomial over $\mathbf{F}_p$: factor over $\mathbf{F}_p[y]/(T)$ (as usual the main variable of *T* must have lower priority than the main variable of *f*);
- *D* a `t_FFELT`: factor over the attached field;
- *D* omitted: factor over the field of definition of *f*, which must be a finite field.

The coefficients of *f* must be operation-compatible with the corresponding finite field. The result is a two-column matrix, the first column being the irreducible polynomials dividing *f*, and the second the exponents. By convention, the 0 polynomial factors as 0^1 ; a nonzero constant polynomial has empty factorization, a 0×2 matrix. The irreducible factors are ordered by increasing degree and the result is canonical: it will not change across multiple calls or sessions.

```
? factormod(x^2 + 1, 3)  \\ over F_3
%1 =
[Mod(1, 3)*x^2 + Mod(1, 3) 1]
? liftall( factormod(x^2 + 1, [t^2+1, 3]) ) \\ over F_9
%2 =
[ x + t 1]
[x + 2*t 1]
\\ same, now letting GP choose a model
? T = ffinit(3,2,'t)
%3 = Mod(1, 3)*t^2 + Mod(1, 3)*t + Mod(2, 3)
? liftall( factormod(x^2 + 1, [T, 3]) )
%4 =  \\ t is a root of T!
[ x + (t + 2) 1]
[x + (2*t + 1) 1]
? t = ffgen(t^2+Mod(1,3)); factormod(x^2 + t^0) \\ same using t_FFELT
%5 =
[ x + t 1]
[x + 2*t 1]
? factormod(x^2+Mod(1,3))
%6 =
[Mod(1, 3)*x^2 + Mod(1, 3) 1]
? liftall( factormod(x^2 + Mod(Mod(1,3), y^2+1)) )
```

```
%7 =
[ x + y 1]
[x + 2*y 1]
```

If *flag* is nonzero, outputs only the *degrees* of the irreducible polynomials (for example to compute an *L*-function). By convention, a constant polynomial (including the 0 polynomial) has empty factorization. The degrees appear in increasing order but need not correspond to the ordering with *flag* = 0 when multiplicities are present.

```
? f = x^3 + 2*x^2 + x + 2;
? factormod(f, 5) \\ (x+2)^2 * (x+3)
%1 =
[Mod(1, 5)*x + Mod(2, 5) 2]
[Mod(1, 5)*x + Mod(3, 5) 1]
? factormod(f, 5, 1) \\ (deg 1) * (deg 1)^2
%2 =
[1 1]
[1 2]
```

The library syntax is GEN factormod0(GEN f, GEN D = NULL, long flag).

3.8.37 factormodDDF(*f*, {*D*}). Distinct-degree factorization of the squarefree polynomial *f* over the finite field defined by the domain *D* as follows:

- *D* = *p* a prime: factor over $\mathbf{F}_p$;
- *D* = [*T*, *p*] for a prime *p* and *T* an irreducible polynomial over $\mathbf{F}_p$: factor over $\mathbf{F}_p[x]/(T)$;
- *D* a t_FFELT: factor over the attached field;
- *D* omitted: factor over the field of definition of *f*, which must be a finite field.

If *f* is not squarefree, the result is undefined. The coefficients of *f* must be operation-compatible with the corresponding finite field. The result is a two-column matrix:

- the first column contains monic (squarefree, pairwise coprime) polynomials dividing *f*, all of whose irreducible factors have the same degree *d*;
- the second column contains the degrees of the irreducible factors.

The factorization is ordered by increasing degree *d* of irreducible factors, and the result is obviously canonical. This function is somewhat faster than full factorization.

```
? f = (x^2 + 1) * (x^2-1);
? factormodSQF(f,3) \\ squarefree over F_3
%2 =
[Mod(1, 3)*x^4 + Mod(2, 3) 1]
? factormodDDF(f, 3)
%3 =
[Mod(1, 3)*x^2 + Mod(2, 3) 1] \\ two degree 1 factors
[Mod(1, 3)*x^2 + Mod(1, 3) 2] \\ irred of degree 2
? for(i=1,10^5,factormodDDF(f,3))
time = 424 ms.
```

```

? for(i=1,10^5,factormod(f,3)) \\ full factorization is a little slower
time = 464 ms.
? liftall( factormodDDF(x^2 + 1, [3, t^2+1]) ) \\ over F_9
%6 =
[x^2 + 1 1] \\ product of two degree 1 factors
? t = ffgen(t^2+Mod(1,3)); factormodDDF(x^2 + t^0) \\ same using t_FFELT
%7 =
[x^2 + 1 1]
? factormodDDF(x^2-Mod(1,3))
%8 =
[Mod(1, 3)*x^2 + Mod(2, 3) 1]

```

The library syntax is `GEN factormodDDF(GEN f, GEN D = NULL)`.

3.8.38 factormodSQF($f, \{D\}$). Squarefree factorization of the polynomial f over the finite field defined by the domain D as follows:

- $D = p$ a prime: factor over $\mathbf{F}_p$;
- $D = [T, p]$ for a prime p and T an irreducible polynomial over $\mathbf{F}_p$: factor over $\mathbf{F}_p[x]/(T)$;
- D a `t_FFELT`: factor over the attached field;
- D omitted: factor over the field of definition of f , which must be a finite field.

The coefficients of f must be operation-compatible with the corresponding finite field. The result is a two-column matrix:

- the first column contains monic squarefree pairwise coprime polynomials dividing f ;
- the second column contains the power to which the polynomial in column 1 divides f ;

This is somewhat faster than full factorization. The factors are ordered by increasing exponent and the result is obviously canonical.

```

? f = (x^2 + 1)^3 * (x^2-1)^2;
? factormodSQF(f, 3) \\ over F_3
%1 =
[Mod(1, 3)*x^2 + Mod(2, 3) 2]
[Mod(1, 3)*x^2 + Mod(1, 3) 3]
? for(i=1,10^5,factormodSQF(f,3))
time = 192 ms.
? for(i=1,10^5,factormod(f,3)) \\ full factorization is slower
time = 409 ms.
? liftall( factormodSQF((x^2 + 1)^3, [3, t^2+1]) ) \\ over F_9
%4 =
[x^2 + 1 3]
? t = ffgen(t^2+Mod(1,3)); factormodSQF((x^2 + t^0)^3) \\ same using t_FFELT
%5 =
[x^2 + 1 3]
? factormodSQF(x^8 + x^7 + x^6 + x^2 + x + Mod(1,2))

```

```
%6 =
[
      Mod(1, 2)*x + Mod(1, 2) 2]
[Mod(1, 2)*x^2 + Mod(1, 2)*x + Mod(1, 2) 3]
```

The library syntax is `GEN factormodSQF(GEN f, GEN D = NULL)`.

3.8.39 factormodcyclo($n, p, \{single = 0\}, \{v = 'x\}$). Factors n -th cyclotomic polynomial $\Phi_n(x)$ mod p , where p is a prime number not dividing n . Much faster than `factormod(polcyclo(n), p)`; the irreducible factors should be identical and given in the same order. If *single* is set, return a single irreducible factor; else (default) return all the irreducible factors. Note that repeated calls of this function with the *single* flag set may return different results because the algorithm is probabilistic. Algorithms used are as follows.

Let $F = \mathbf{Q}(\zeta_n)$. Let K be the splitting field of p in F and e the conductor of K . Then $\Phi_n(x)$ and $\Phi_e(x)$ have the same number of irreducible factors mod p and there is a simple algorithm constructing irreducible factors of $\Phi_n(x)$ from irreducible factors of $\Phi_e(x)$. So we may assume n is equal to the conductor of K . Let d be the order of p in $(\mathbf{Z}/n\mathbf{Z})^\times$ and $\varphi(n) = df$. Then $\Phi_n(x)$ has f irreducible factors $g_i(x)$ ($1 \leq i \leq f$) of degree d over $\mathbf{F}_p$ or $\mathbf{Z}_p$.

- If d is small, then we factor $g_i(x)$ into d linear factors $g_{ij}(x)$, $1 \leq j \leq d$ in $\mathbf{F}_q[x]$ ($q = p^d$) and construct $G_i(x) = \prod_{j=1}^d g_{ij}(x) \in \mathbf{F}_q[x]$. Then $G_i(x) \in \mathbf{F}_p[x]$ and $g_i(x) = G_i(x)$.

- If f is small, then we work in K , which is a Galois extension of degree f over $\mathbf{Q}$. The Gaussian period $\theta_k = \text{Tr}_{F/K}(\zeta_n^k)$ is a sum of k -th power of roots of $g_i(x)$ and $K = \mathbf{Q}(\theta_1)$.

Now, for each k , there is a polynomial $T_k(x) \in \mathbf{Q}[x]$ satisfying $\theta_k = T_k(\theta_1)$ because all θ_k are in K . Let $T(x) \in \mathbf{Z}[x]$ be the minimal polynomial of θ_1 over $\mathbf{Q}$. We get $\theta_1 \bmod p$ from $T(x)$ and construct $\theta_1, \dots, \theta_d \bmod p$ using $T_k(x)$. Finally we recover $g_i(x)$ from $\theta_1, \dots, \theta_d$ by Newton's formula.

```
? lift(factormodcyclo(15, 11))
%1 = [x^2 + 9*x + 4, x^2 + 4*x + 5, x^2 + 3*x + 9, x^2 + 5*x + 3]
? factormodcyclo(15, 11, 1) \ single
%2 = Mod(1, 11)*x^2 + Mod(5, 11)*x + Mod(3, 11)
? z1 = lift(factormod(polcyclo(12345), 11311)[, 1]);
time = 32,498 ms.
? z2 = factormodcyclo(12345, 11311);
time = 47 ms.
? z1 == z2
%4 = 1
```

The library syntax is `GEN factormodcyclo(long n, GEN p, long single, long v = -1)` where v is a variable number.

3.8.40 ffcompomap(f, g). Let k, l, m be three finite fields and f a (partial) map from l to m and g a (partial) map from k to l , return the (partial) map $f \circ g$ from k to m .

```
a = ffgen([3,5], 'a'); b = ffgen([3,10], 'b'); c = ffgen([3,20], 'c');
m = fembed(a, b); n = fembed(b, c);
rm = ffinvmap(m); rn = ffinvmap(n);
nm = ffcompomap(n, m);
ffmap(n, fmap(m, a)) == fmap(nm, a)
%5 = 1
ffcompomap(rm, rn) == ffinvmap(nm)
%6 = 1
```

The library syntax is GEN ffcompomap(GEN f, GEN g).

3.8.41 fembed(a, b). Given two finite fields elements a and b , return a *map* embedding the definition field of a to the definition field of b . Assume that the latter contains the former.

```
? a = ffgen([3,5], 'a');
? b = ffgen([3,10], 'b');
? m = fembed(a, b);
? A = fmap(m, a);
? minpoly(A) == minpoly(a)
%5 = 1
```

The library syntax is GEN fembed(GEN a, GEN b).

3.8.42 fextend($a, P, \{v\}$). Extend the field K of definition of a by a root of the polynomial $P \in K[X]$ assumed to be irreducible over K . Return $[r, m]$ where r is a root of P in the extension field L and m is a map from K to L , see `fmap`. If v is given, the variable name is used to display the generator of L , else the name of the variable of P is used. A generator of L can be recovered using $b = fgen(r)$. The image of P in $L[X]$ can be recovered using $PL = fmap(m, P)$.

```
? a = ffgen([3,5], 'a');
? P = x^2-a; polisirreducible(P)
%2 = 1
? [r,m] = fextend(a, P, 'b');
? r
%3 = b^9+2*b^8+b^7+2*b^6+b^4+1
? subst(fmap(m, P), x, r)
%4 = 0
? fgen(r)
%5 = b
```

The library syntax is GEN fextend(GEN a, GEN P, long v = -1) where v is a variable number.

3.8.43 fffrobenius($m, \{n = 1\}$). Return the n -th power of the Frobenius map over the field of definition of m .

```
? a = ffggen([3,5], 'a');
? f = fffrobenius(a);
? ffgmap(f,a) == a^3
%3 = 1
? g = fffrobenius(a, 5);
? ffgmap(g,a) == a
%5 = 1
? h = fffrobenius(a, 2);
? h == ffgcompomap(f,f)
%7 = 1
```

The library syntax is GEN fffrobenius(GEN m, long n).

3.8.44 ffggen($k, \{v = 'x\}$). Return a generator for the finite field k as a `t_FFELT`. The field k can be given by

- its order q
- the pair $[p, f]$ where $q = p^f$
- a monic irreducible polynomial with `t_INTMOD` coefficients modulo a prime.
- a `t_FFELT` belonging to k .

If v is given, the variable name is used to display g , else the variable of the polynomial or the `t_FFELT` is used, else x is used. For efficiency, the characteristic is not checked to be prime; similarly if a polynomial is given, we do not check whether it is irreducible.

When only the order is specified, the function uses the polynomial generated by `ffinit` and is deterministic: two calls to the function with the same parameters will always give the same generator.

To obtain a multiplicative generator, call `ffprimroot` on the result (which is randomized). Its minimal polynomial then gives a *primitive* polynomial, which can be used to redefine the finite field so that all subsequent computations use the new primitive polynomial:

```
? g = ffggen(16, 't');
? g.mod \\ recover the underlying polynomial.
%2 = t^4 + t^3 + t^2 + t + 1
? g.pol \\ lift g as a t_POL
%3 = t
? g.p \\ recover the characteristic
%4 = 2
? ffgorder(g) \\ g is not a multiplicative generator
%5 = 5
? a = ffprimroot(g) \\ recover a multiplicative generator
%6 = t^3 + t^2 + t
? ffgorder(a)
%7 = 15
? T = minpoly(a) \\ primitive polynomial
%8 = Mod(1, 2)*x^4 + Mod(1, 2)*x^3 + Mod(1, 2)
```

```
? G = ffgen(T); \\ is now a multiplicative generator
? fforder(G)
%10 = 15
```

The library syntax is `GEN ffgen(GEN k, long v = -1)` where v is a variable number.

To create a generator for a prime finite field, the function `GEN p_to_GEN(GEN p, long v)` returns `ffgen(p,v)^0`.

3.8.45 `ffinit`($p, n, \{v = 'x\}$). Computes a monic polynomial of degree n which is irreducible over $\mathbf{F}_p$, where p is assumed to be prime. This function uses a fast variant of Adleman and Lenstra's algorithm.

It is useful in conjunction with `ffgen`; for instance if $P = \text{ffinit}(3,2)$, you can represent elements in $\mathbf{F}_{3^2}$ in term of $g = \text{ffgen}(P, 't)$. This can be abbreviated as $g = \text{ffgen}(3^2, 't)$, where the defining polynomial P can be later recovered as `g.mod`.

The library syntax is `GEN ffinit(GEN p, long n, long v = -1)` where v is a variable number.

3.8.46 `ffinvmap`(m). m being a map from K to L two finite fields, return the partial map p from L to K such that for all $k \in K$, $p(m(k)) = k$.

```
? a = ffgen([3,5], 'a);
? b = ffgen([3,10], 'b);
? m = ffembed(a, b);
? p = ffinvmap(m);
? u = random(a);
? v = fffmap(m, u);
? fffmap(p, v^2+v+2) == u^2+u+2
%7 = 1
? fffmap(p, b)
%8 = []
```

The library syntax is `GEN ffinvmap(GEN m)`.

3.8.47 `fflog`($x, g, \{o\}$). Discrete logarithm of the finite field element x in base g , i.e. an e in $\mathbf{Z}$ such that $g^e = o$. If present, o represents the multiplicative order of g , see Section 3.8.2; the preferred format for this parameter is `[ord, factor(ord)]`, where `ord` is the order of g . It may be set as a side effect of calling `ffprimroot`. The result is undefined if e does not exist. This function uses

- a combination of generic discrete log algorithms (see `znlog`)
- a cubic sieve index calculus algorithm for large fields of degree at least 5.
- Coppersmith's algorithm for fields of characteristic at most 5.

```
? t = ffgen(ffinit(7,5));
? o = fforder(t)
%2 = 5602 \\ not a primitive root.
? fflog(t^10,t)
%3 = 10
? fflog(t^10,t, o)
%4 = 10
```

```
? g = ffprimroot(t, &o);
? o    \\ order is 16806, bundled with its factorization matrix
%6 = [16806, [2, 1; 3, 1; 2801, 1]]
? fforder(g, o)
%7 = 16806
? fflog(g^10000, g, o)
%8 = 10000
```

The library syntax is GEN fflog(GEN x, GEN g, GEN o = NULL).

3.8.48 fffmap(m, x). Given a (partial) map m between two finite fields, return the image of x by m . The function is applied recursively to the component of vectors, matrices and polynomials. If m is a partial map that is not defined at x , return $[]$.

```
? a = ffgen([3,5], 'a);
? b = ffgen([3,10], 'b);
? m = ffembed(a, b);
? P = x^2+a*x+1;
? Q = fffmap(m,P);
? fffmap(m,poldisc(P)) == poldisc(Q)
%6 = 1
```

The library syntax is GEN fffmap(GEN m, GEN x).

3.8.49 fffmaprel(m, x). Given a (partial) map m between two finite fields, express x as an algebraic element over the codomain of m in a way which is compatible with m . The function is applied recursively to the component of vectors, matrices and polynomials.

```
? a = ffgen([3,5], 'a);
? b = ffgen([3,10], 'b);
? m = ffembed(a, b);
? mi= ffinvmap(m);
? R = fffmaprel(mi,b)
%5 = Mod(b,b^2+(a+1)*b+(a^2+2*a+2))
```

In particular, this function can be used to compute the relative minimal polynomial, norm and trace:

```
? minpoly(R)
%6 = x^2+(a+1)*x+(a^2+2*a+2)
? trace(R)
%7 = 2*a+2
? norm(R)
%8 = a^2+2*a+2
```

The library syntax is GEN fffmaprel(GEN m, GEN x).

3.8.50 ffnbirred($q, n, \{flag = 0\}$). Computes the number of monic irreducible polynomials over $\mathbf{F}_q$ of degree exactly n ($flag = 0$ or omitted) or at most n ($flag = 1$).

The library syntax is GEN ffnbirred0(GEN q, long n, long flag). Also available are GEN ffnbirred(GEN q, long n) (for $flag = 0$) and GEN ffsumnbirred(GEN q, long n) (for $flag = 1$).

3.8.51 `fforder`($x, \{o\}$). Multiplicative order of the finite field element x . If o is present, it represents a multiple of the order of the element, see Section 3.8.2; the preferred format for this parameter is `[N, factor(N)]`, where N is the cardinality of the multiplicative group of the underlying finite field.

```
? t = ffgen(ffinit(nextprime(10^8), 5));
? g = ffprimroot(t, &o);  \\ o will be useful!
? fforder(g^1000000, o)
time = 0 ms.
%5 = 5000001750000245000017150000600250008403
? fforder(g^1000000)
time = 16 ms.  \\ noticeably slower, same result of course
%6 = 5000001750000245000017150000600250008403
```

The library syntax is `GEN fforder(GEN x, GEN o = NULL)`.

3.8.52 `ffprimroot`($x, \{&o\}$). Return a primitive root of the multiplicative group of the definition field of the finite field element x (not necessarily the same as the field generated by x). If present, o is set to a vector `[ord, fa]`, where `ord` is the order of the group and `fa` its factorization `factor(ord)`. This last parameter is useful in `fflog` and `fforder`, see Section 3.8.2.

```
? t = ffgen(ffinit(nextprime(10^7), 5));
? g = ffprimroot(t, &o);
? o[1]
%3 = 1000000950003610006859006516052476098
? o[2]
%4 =
[2 1]
[7 2]
[31 1]
[41 1]
[67 1]
[1523 1]
[10498781 1]
[15992881 1]
[46858913131 1]
? fflog(g^1000000, g, o)
time = 1,312 ms.
%5 = 1000000
```

The library syntax is `GEN ffprimroot(GEN x, GEN *o = NULL)`.

3.8.53 gcd($x, \{y\}$). Creates the greatest common divisor of x and y . If you also need the u and v such that $x * u + y * v = \text{gcd}(x, y)$, use the `gcdext` function. x and y can have rather quite general types, for instance both rational numbers. If y is omitted and x is a vector, returns the gcd of all components of x , i.e. this is equivalent to `content(x)`.

When x and y are both given and one of them is a vector/matrix type, the GCD is again taken recursively on each component, but in a different way. If y is a vector, resp. matrix, then the result has the same type as y , and components equal to `gcd(x, y[i])`, resp. `gcd(x, y[,i])`. Else if x is a vector/matrix the result has the same type as x and an analogous definition. Note that for these types, `gcd` is not commutative.

The algorithm used is a naive Euclid except for the following inputs:

- integers: use modified right-shift binary (“plus-minus” variant).
- univariate polynomials with coefficients in the same number field (in particular rational): use modular gcd algorithm.
- general polynomials: use the subresultant algorithm if coefficient explosion is likely (non modular coefficients).

If u and v are polynomials in the same variable with *inexact* coefficients, their gcd is defined to be scalar, so that

```
? a = x + 0.0; gcd(a,a)
%1 = 1
? b = y*x + 0(y); gcd(b,b)
%2 = y
? c = 4*x + 0(2^3); gcd(c,c)
%3 = 4
```

A good quantitative check to decide whether such a gcd “should be” nontrivial, is to use `polresultant`: a value close to 0 means that a small deformation of the inputs has nontrivial gcd. You may also use `gcdext`, which does try to compute an approximate gcd d and provides u, v to check whether $ux + vy$ is close to d .

The library syntax is `GEN ggcd0(GEN x, GEN y = NULL)`. Also available are `GEN ggcd(GEN x, GEN y)`, if y is not NULL, and `GEN content(GEN x)`, if $y = \text{NULL}$.

3.8.54 gcdext(x, y). Returns $[u, v, d]$ such that d is the gcd of x, y , $x * u + y * v = \text{gcd}(x, y)$, and u and v minimal in a natural sense. The arguments must be integers or polynomials.

```
? [u, v, d] = gcdext(32,102)
%1 = [16, -5, 2]
? d
%2 = 2
? gcdext(x^2-x, x^2+x-2)
%3 = [-1/2, 1/2, x - 1]
```

If x, y are polynomials in the same variable and *inexact* coefficients, then compute u, v, d such that $x * u + y * v = d$, where d approximately divides both x and y ; in particular, we do not obtain `gcd(x,y)` which is *defined* to be a scalar in this case:

```
? a = x + 0.0; gcd(a,a)
%1 = 1
```

```
? gcdext(a,a)
%2 = [0, 1, x + 0.E-28]
? gcdext(x-Pi, 6*x^2-zeta(2))
%3 = [-6*x - 18.8495559, 1, 57.5726923]
```

For inexact inputs, the output is thus not well defined mathematically, but you obtain explicit polynomials to check whether the approximation is close enough for your needs.

The library syntax is GEN `gcdext0(GEN x, GEN y)`.

3.8.55 halfgcd(x, y). Let inputs x and y be both integers, or both polynomials in the same variable. Return a vector $[M, [a, b] \sim]$, where M is an invertible 2×2 matrix such that $M \cdot [x, y] \sim [a, b] \sim$, where b is small. More precisely,

- polynomial case: $\det M$ has degree 0 and we have

$$\deg a \geq \lceil \max(\deg x, \deg y) / 2 \rceil > \deg b.$$

- integer case: $\det M = \pm 1$ and we have

$$a \geq \left\lceil \sqrt{\max(|x|, |y|)} \right\rceil > b.$$

Assuming x and y are nonnegative, then M^{-1} has nonnegative coefficients, and $\det M$ is equal to the sign of both main diagonal terms $M[1, 1]$ and $M[2, 2]$.

The library syntax is GEN `ghalfgcd(GEN x, GEN y)`.

3.8.56 hilbert($x, y, \{p\}$). Hilbert symbol of x and y modulo the prime p , $p = 0$ meaning the place at infinity (the result is undefined if $p \neq 0$ is not prime).

It is possible to omit p , in which case we take $p = 0$ if both x and y are rational, or one of them is a real number. And take $p = q$ if one of x, y is a `t_INTMOD` modulo q or a q -adic. (Incompatible types will raise an error.)

The library syntax is long `hilbert(GEN x, GEN y, GEN p = NULL)`.

3.8.57 isfundamental(D). True (1) if D is equal to 1 or to the discriminant of a quadratic field, false (0) otherwise. D can be input in factored form as for arithmetic functions:

```
? isfundamental(factor(-8))
%1 = 1
\\ count fundamental discriminants up to 10^8
? c = 0; forfactored(d = 1, 10^8, if (isfundamental(d), c++)); c
time = 40,840 ms.
%2 = 30396325
? c = 0; for(d = 1, 10^8, if (isfundamental(d), c++)); c
time = 1min, 33,593 ms. \\ slower !
%3 = 30396325
```

The library syntax is long `isfundamental(GEN D)`.

3.8.58 ispolygonal($x, s, \{&N\}$). True (1) if the integer x is an s -gonal number, false (0) if not. The parameter $s > 2$ must be a `t_INT`. If N is given, set it to n if x is the n -th s -gonal number.

```
? ispolygonal(36, 3, &N)
%1 = 1
? N
```

The library syntax is `long ispolygonal(GEN x, GEN s, GEN *N = NULL)`.

3.8.59 ispower($x, \{k\}, \{&n\}$). If k is given, returns true (1) if x is a k -th power, false (0) if not. What it means to be a k -th power depends on the type of x ; see `issquare` for details.

If k is omitted, only integers and fractions are allowed for x and the function returns the maximal $k \geq 2$ such that $x = n^k$ is a perfect power, or 0 if no such k exist; in particular `ispower(-1)`, `ispower(0)`, and `ispower(1)` all return 0.

If a third argument $&n$ is given and x is indeed a k -th power, sets n to a k -th root of x .

For a `t_FFELT` x , instead of omitting k (which is not allowed for this type), it may be natural to set

```
k = (x.p ^ x.f - 1) / fforder(x)
```

The library syntax is `long ispower(GEN x, GEN k = NULL, GEN *n = NULL)`. Also available is `long gisanypower(GEN x, GEN *pty)` (k omitted).

3.8.60 ispowerful(x). True (1) if x is a powerful integer, false (0) if not; an integer is powerful if and only if its valuation at all primes dividing x is greater than 1.

```
? ispowerful(50)
%1 = 0
? ispowerful(100)
%2 = 1
? ispowerful(5^3*(10^1000+1)^2)
%3 = 1
```

The library syntax is `long ispowerful(GEN x)`.

3.8.61 isprime($x, \{flag = 0\}$). True (1) if x is a prime number, false (0) otherwise. A prime number is a positive integer having exactly two distinct divisors among the natural numbers, namely 1 and itself.

This routine proves or disproves rigorously that a number is prime, which can be very slow when x is indeed a large prime integer. For instance a 1000 digits prime should require 15 to 30 minutes with default algorithms. Use `ispseudoprime` to quickly check for compositeness. Use `primecert` in order to obtain a primality proof instead of a yes/no answer; see also `factor`.

The function accepts vector/matrices arguments, and is then applied componentwise.

If $flag = 0$, use a combination of

- Baillie-Pomerance-Selfridge-Wagstaff compositeness test (see `ispseudoprime`),
- Selfridge “ $p - 1$ ” test if $x - 1$ is smooth enough,
- Adleman-Pomerance-Rumely-Cohen-Lenstra (APRCL) for general medium-sized x (less than 1500 bits),

- Atkin-Morain’s Elliptic Curve Primality Prover (ECPP) for general large x .

If $flag = 1$, use Selfridge-Pocklington-Lehmer “ $p - 1$ ” test; this requires partially factoring various auxilliary integers and is likely to be very slow.

If $flag = 2$, use APRCL only.

If $flag = 3$, use ECPP only.

The library syntax is `GEN gisprime(GEN x, long flag)`.

3.8.62 isprimepower($x, \{&n\}$). If $x = p^k$ is a prime power (p prime, $k > 0$), return k , else return 0. If a second argument $&n$ is given and x is indeed the k -th power of a prime p , sets n to p .

The library syntax is `long isprimepower(GEN x, GEN *n = NULL)`.

3.8.63 ispseudoprime($x, \{flag\}$). True (1) if x is a strong pseudo prime (see below), false (0) otherwise. If this function returns false, x is not prime; if, on the other hand it returns true, it is only highly likely that x is a prime number. Use `isprime` (which is of course much slower) to prove that x is indeed prime. The function accepts vector/matrices arguments, and is then applied componentwise.

If $flag = 0$, checks whether x has no small prime divisors (up to 101 included) and is a Baillie-Pomerance-Selfridge-Wagstaff pseudo prime. Such a pseudo prime passes a Rabin-Miller test for base 2, followed by a Lucas test for the sequence $(P, 1)$, where $P \geq 3$ is the smallest odd integer such that $P^2 - 4$ is not a square mod x . (Technically, we are using an “almost extra strong Lucas test” that checks whether V_n is ± 2 , without computing U_n .)

There are no known composite numbers passing the above test, although it is expected that infinitely many such numbers exist. In particular, all composites $\leq 2^{64}$ are correctly detected (checked using <https://www.cecm.sfu.ca/Pseudoprimes/index-2-to-64.html>).

If $flag > 0$, checks whether x is a strong Miller-Rabin pseudo prime for $flag$ randomly chosen bases (with end-matching to catch square roots of -1).

The library syntax is `GEN gispseudoprime(GEN x, long flag)`.

3.8.64 ispseudoprimepower($x, \{&n\}$). If $x = p^k$ is a pseudo-prime power (p pseudo-prime as per `ispseudoprime`, $k > 0$), return k , else return 0. If a second argument $&n$ is given and x is indeed the k -th power of a prime p , sets n to p .

More precisely, k is always the largest integer such that $x = n^k$ for some integer n and, when $n \leq 2^{64}$ the function returns $k > 0$ if and only if n is indeed prime. When $n > 2^{64}$ is larger than the threshold, the function may return 1 even though n is composite: it only passed an `ispseudoprime(n)` test.

The library syntax is `long ispseudoprimepower(GEN x, GEN *n = NULL)`.

3.8.65 issquare($x, \{&n\}$). True (1) if x is a square, false (0) if not. What “being a square” means depends on the type of x : all `t_COMPLEX` are squares, as well as all nonnegative `t_REAL`; for exact types such as `t_INT`, `t_FRAC` and `t_INTMOD`, squares are numbers of the form s^2 with s in $\mathbf{Z}$, $\mathbf{Q}$ and $\mathbf{Z}/N\mathbf{Z}$ respectively.

```
? issquare(3)          \\ as an integer
%1 = 0
? issquare(3.)         \\ as a real number
%2 = 1
? issquare(Mod(7, 8))  \\ in Z/8Z
%3 = 0
? issquare( 5 + 0(13^4) ) \\ in Q_13
%4 = 0
```

If n is given, a square root of x is put into n .

```
? issquare(4, &n)
%1 = 1
? n
%2 = 2
```

For polynomials, either we detect that the characteristic is 2 (and check directly odd and even-power monomials) or we assume that 2 is invertible and check whether squaring the truncated power series for the square root yields the original input.

For `t_POLMOD` x , we only support `t_POLMOD`s of `t_INTMOD`s encoding finite fields, assuming without checking that the intmod modulus p is prime and that the polmod modulus is irreducible modulo p .

```
? issquare(Mod(Mod(2,3), x^2+1), &n)
%1 = 1
? n
%2 = Mod(Mod(2, 3)*x, Mod(1, 3)*x^2 + Mod(1, 3))
```

The library syntax is `long issquareall(GEN x, GEN *n = NULL)`. Also available is `long issquare(GEN x)`. Deprecated GP-specific functions `GEN gissquare(GEN x)` and `GEN gissquareall(GEN x, GEN *pt)` return `gen_0` and `gen_1` instead of a boolean value.

3.8.66 issquarefree(x). True (1) if x is squarefree, false (0) if not. Here x can be an integer or a polynomial with coefficients in an integral domain.

```
? issquarefree(12)
%1 = 0
? issquarefree(6)
%2 = 1
? issquarefree(x^3+x^2)
%3 = 0
? issquarefree(Mod(1,4)*(x^2+x+1))    \\ Z/4Z is not a domain !
***   at top-level: issquarefree(Mod(1,4)*(x^2+x+1))
***   ^-----
*** issquarefree: impossible inverse in Fp_inv: Mod(2, 4).
```

A polynomial is declared squarefree if $\gcd(x, x')$ is 1. In particular a nonzero polynomial with inexact coefficients is considered to be squarefree. Note that this may be inconsistent with `factor`,

which first rounds the input to some exact approximation before factoring in the appropriate domain; this is correct when the input is not close to an inseparable polynomial (the resultant of x and x' is not close to 0).

An integer can be input in factored form as in arithmetic functions.

```
? issquarefree(factor(6))
%1 = 1
\\ count squarefree integers up to 10^8
? c = 0; for(d = 1, 10^8, if (issquarefree(d), c++)); c
time = 3min, 2,590 ms.
%2 = 60792694
? c = 0; forfactored(d = 1, 10^8, if (issquarefree(d), c++)); c
time = 45,348 ms. \\ faster !
%3 = 60792694
```

The library syntax is `long issquarefree(GEN x)`.

3.8.67 istotient($x, \{&N\}$). True (1) if $x = \phi(n)$ for some integer n , false (0) if not.

```
? istotient(14)
%1 = 0
? istotient(100)
%2 = 0
```

If N is given, set $N = n$ as well.

```
? istotient(4, &n)
%1 = 1
? n
%2 = 10
```

The library syntax is `long istotient(GEN x, GEN *N = NULL)`.

3.8.68 kronecker(x, y). Kronecker symbol $(x|y)$, where x and y must be of type integer. By definition, this is the extension of Legendre symbol to $\mathbf{Z} \times \mathbf{Z}$ by total multiplicativity in both arguments with the following special rules for $y = 0, -1$ or 2 :

- $(x|0) = 1$ if $|x| = 1$ and 0 otherwise.
- $(x|-1) = 1$ if $x \geq 0$ and -1 otherwise.
- $(x|2) = 0$ if x is even and 1 if $x = 1, -1 \bmod 8$ and -1 if $x = 3, -3 \bmod 8$.

The library syntax is `long kronecker(GEN x, GEN y)`.

3.8.69 lcm($x, \{y\}$). Least common multiple of x and y , i.e. such that $\text{lcm}(x, y) * \text{gcd}(x, y) = x * y$, up to units. If y is omitted and x is a vector, returns the lcm of all components of x . For integer arguments, return the nonnegative lcm.

When x and y are both given and one of them is a vector/matrix type, the LCM is again taken recursively on each component, but in a different way. If y is a vector, resp. matrix, then the result has the same type as y , and components equal to $\text{lcm}(x, y[i])$, resp. $\text{lcm}(x, y[,i])$. Else if x is a vector/matrix the result has the same type as x and an analogous definition. Note that for these types, lcm is not commutative.

Note that $\text{lcm}(v)$ is quite different from

```
l = v[1]; for (i = 1, #v, l = lcm(l, v[i]))
```

Indeed, $\text{lcm}(v)$ is a scalar, but l may not be (if one of the $v[i]$ is a vector/matrix). The computation uses a divide-conquer tree and should be much more efficient, especially when using the GMP multiprecision kernel (and more subquadratic algorithms become available):

```
? v = vector(10^5, i, random);
? lcm(v);
time = 546 ms.
? l = v[1]; for (i = 1, #v, l = lcm(l, v[i]))
time = 4,561 ms.
```

The library syntax is GEN glcm0(GEN x, GEN y = NULL).

3.8.70 logint($x, b, \{&z\}$). Return the largest non-negative integer e so that $b^e \leq x$, where $b > 1$ is an integer and $x \geq 1$ is a real number. If the parameter z is present, set it to b^e .

```
? logint(1000, 2)
%1 = 9
? 2^9
%2 = 512
? logint(1000, 2, &z)
%3 = 9
? z
%4 = 512
? logint(Pi^2, 2, &z)
%5 = 3
? z
%6 = 8
```

The number of digits used to write x in base b is $1 + \text{logint}(x, b)$:

```
? #digits(1000!, 10)
%5 = 2568
? logint(1000!, 10)
%6 = 2567
```

This function may conveniently replace

```
floor( log(x) / log(b) )
```

which may not give the correct answer since PARI does not guarantee exact rounding.

The library syntax is long logint0(GEN x, GEN b, GEN *z = NULL).

3.8.71 moebius(x). Moebius μ -function of $|x|$; x must be a nonzero integer.

The library syntax is `long moebius(GEN x)`.

3.8.72 nextprime(x). Finds the smallest pseudoprime (see `ispseudoprime`) greater than or equal to x . x can be of any real type. Note that if x is a pseudoprime, this function returns x and not the smallest pseudoprime strictly larger than x . To rigorously prove that the result is prime, use `isprime`.

```
? nextprime(2)
%1 = 2
? nextprime(Pi)
%2 = 5
? nextprime(-10)
%3 = 2 \\ primes are positive
```

Despite the name, please note that the function is not guaranteed to return a prime number, although no counter-example is known at present. The return value *is* a guaranteed prime if $x \leq 2^{64}$. To rigorously prove that the result is prime in all cases, use `isprime`.

The library syntax is `GEN nextprime(GEN x)`.

3.8.73 numdiv(x). Number of divisors of $|x|$. x must be of type integer.

The library syntax is `GEN numdiv(GEN x)`.

3.8.74 omega(x). Number of distinct prime divisors of $|x|$. x must be of type integer.

```
? factor(392)
%1 =
[2 3]
[7 2]
? omega(392)
%2 = 2; \\ without multiplicity
? bigomega(392)
%3 = 5; \\ = 3+2, with multiplicity
```

The library syntax is `long omega(GEN x)`.

3.8.75 precprime(x). Finds the largest pseudoprime (see `ispseudoprime`) less than or equal to x ; the input x can be of any real type. Returns 0 if $x \leq 1$. Note that if x is a prime, this function returns x and not the largest prime strictly smaller than x .

```
? precprime(2)
%1 = 2
? precprime(Pi)
%2 = 3
? precprime(-10)
%3 = 0 \\ primes are positive
```

The function name comes from *preceding prime*. Despite the name, please note that the function is not guaranteed to return a prime number (although no counter-example is known at present); the return value *is* a guaranteed prime if $x \leq 2^{64}$. To rigorously prove that the result is prime in all cases, use `isprime`.

The library syntax is `GEN precprime(GEN x)`.

3.8.76 `prime(n)`. The n^{th} prime number

```
? prime(10^9)
%1 = 22801763489
```

Uses checkpointing and a naive $O(n)$ algorithm. Will need about 30 minutes for n up to 10^{11} ; make sure to start gp with `primelimit` at least $\sqrt{p_n}$, e.g. the value $\sqrt{n \log(n \log n)}$ is guaranteed to be sufficient.

The library syntax is `GEN prime(long n)`.

3.8.77 `primecert(N, {flag = 0}, {partial = 0})`. If N is a prime, return a PARI Primality Certificate for the prime N , as described below. Otherwise, return 0. A Primality Certificate c can be checked using `primecertisvalid(c)`.

If $flag = 0$ (default), return an ECPP certificate (Atkin-Morain)

If $flag = 0$ and $partial > 0$, return a (potentially) partial ECPP certificate.

A PARI ECPP Primality Certificate for the prime N is either a prime integer $N < 2^{64}$ or a vector $\mathbf{C}$ of length ℓ whose i th component $\mathbf{C}[i]$ is a vector $[N_i, t_i, s_i, a_i, P_i]$ of length 5 where $N_1 = N$. It is said to be *valid* if for each $i = 1, \dots, \ell$, all of the following conditions are satisfied

- N_i is a positive integer
- t_i is an integer such that $t_i^2 < 4N_i$
- s_i is a positive integer which divides m_i where $m_i = N_i + 1 - t_i$
- If we set $q_i = \frac{m_i}{s_i}$, then
 - $q_i > (N_i^{1/4} + 1)^2$
 - $q_i = N_{i+1}$ if $1 \leq i < \ell$
 - $q_\ell \leq 2^{64}$ is prime

- a_i is an integer

• $\mathbf{P}[i]$ is a vector of length 2 representing the affine point $P_i = (x_i, y_i)$ on the elliptic curve $E : y^2 = x^3 + a_i x + b_i$ modulo N_i where $b_i = y_i^2 - x_i^3 - a_i x_i$ satisfying the following:

- $m_i P_i = \infty$
- $s_i P_i \neq \infty$

Using the following theorem, the data in the vector $\mathbf{C}$ allows to recursively certify the primality of N (and all the q_i) under the single assumption that q_ℓ be prime.

Theorem. If N is an integer and there exist positive integers m, q and a point P on the elliptic curve $E : y^2 = x^3 + ax + b$ defined modulo N such that $q > (N^{1/4} + 1)^2$, q is a prime divisor of m , $mP = \infty$ and $\frac{m}{q}P \neq \infty$, then N is prime.

A partial certificate is identical except that the condition $q_\ell \leq 2^{64}$ is replaced by $q_\ell \leq 2^{partial}$. Such partial certificate C can be extended to a full certificate by calling $C = \text{primecert}(C)$, or to a longer partial certificate by calling $C = \text{primecert}(C, b)$ with $b < partial$.

```
? primecert(10^35 + 69)
%1 = [[10000000000000000000000000000000069, 5468679110354
52074, 2963504668391148, 0, [60737979324046450274283740674
208692, 24368673584839493121227731392450025]], [3374383076
4501150277, -11610830419, 734208843, 0, [26740412374402652
72 4, 6367191119818901665]], [45959444779, 299597, 2331, 0
, [18022351516, 9326882 51]]]
? primecert(nextprime(2^64))
%2 = [[18446744073709551629, -8423788454, 160388, 1, [1059
8342506117936052, 2225259013356795550]]]
? primecert(6)
%3 = 0
? primecert(41)
%4 = 41
? N = 2^2000+841;
? Cp1 = primecert(N,,1500); \\ partial certificate
time = 16,018 ms.
? Cp2 = primecert(Cp1,,1000); \\ (longer) partial certificate
time = 5,890 ms.
? C = primecert(Cp2); \\ full certificate for N
time = 1,777 ms.
? primecertisvalid(C)
%9 = 1
? primecert(N);
time = 23.625 ms.
```

As the last command shows, attempting a succession of partial certificates should be about as fast as a direct computation.

If $flag = 1$ (very slow), return an $N - 1$ certificate (Pocklington Lehmer)

A PARI $N - 1$ Primality Certificate for the prime N is either a prime integer $N < 2^{64}$ or a pair $[N, C]$, where C is a vector with ℓ elements which are either a single integer $p_i < 2^{64}$ or a triple $[p_i, a_i, C_i]$ with $p_i > 2^{64}$ satisfying the following properties:

- p_i is a prime divisor of $N - 1$;
- a_i is an integer such that $a_i^{N-1} \equiv 1 \pmod{N}$ and $a_i^{(N-1)/p_i} - 1$ is coprime with N ;
- C_i is an $N - 1$ Primality Certificate for p_i
- The product F of the $p_i^{v_{p_i}(N-1)}$ is strictly larger than $N^{1/3}$. Provided that all p_i are indeed primes, this implies that any divisor of N is congruent to 1 modulo F .

3.8.81 primes(n). Creates a row vector whose components are the first n prime numbers. (Returns the empty vector for $n \leq 0$.) A `t_VEC` $n = [a, b]$ is also allowed, in which case the primes in $[a, b]$ are returned

```
? primes(10)      \\ the first 10 primes
%1 = [2, 3, 5, 7, 11, 13, 17, 19, 23, 29]
? primes([0,29])  \\ the primes up to 29
%2 = [2, 3, 5, 7, 11, 13, 17, 19, 23, 29]
? primes([15,30])
%3 = [17, 19, 23, 29]
```

The library syntax is `GEN primes0(GEN n)`.

3.8.82 qfbclassno($D, \{flag = 0\}$). Ordinary class number of the quadratic order of discriminant D , for “small” values of D .

- if $D > 0$ or $flag = 1$, use a $O(|D|^{1/2})$ algorithm (compute $L(1, \chi_D)$ with the approximate functional equation). This is slower than `quadclassunit` as soon as $|D| \approx 10^2$ or so and is not meant to be used for large D .

- if $D < 0$ and $flag = 0$ (or omitted), use a $O(|D|^{1/4})$ algorithm (Shanks’s baby-step/giant-step method). It should be faster than `quadclassunit` for small values of D , say $|D| < 10^{18}$.

Important warning. In the latter case, this function only implements part of Shanks’s method (which allows to speed it up considerably). It gives unconditionally correct results for $|D| < 2 \cdot 10^{10}$, but may give incorrect results for larger values if the class group has many cyclic factors. We thus recommend to double-check results using the function `quadclassunit`, which is about 2 to 3 times slower in the range $|D| \in [10^{10}, 10^{18}]$, assuming GRH. We currently have no counter-examples but they should exist: we would appreciate a bug report if you find one.

Warning. Contrary to what its name implies, this routine does not compute the number of classes of binary primitive forms of discriminant D , which is equal to the *narrow* class number. The two notions are the same when $D < 0$ or the fundamental unit ε has negative norm; when $D > 0$ and $N\varepsilon > 0$, the number of classes of forms is twice the ordinary class number. This is a problem which we cannot fix for backward compatibility reasons. Use the following routine if you are only interested in the number of classes of forms:

```
? QFBclassno(D) = qfbclassno(D) * if (D > 0 && quadunitnorm(D) > 0, 2, 1)
? QFBclassno(136)
%1 = 4
? qfbclassno(136)
%2 = 2
? quadunitnorm(136)
%3 = 1
? bnfarrow(bnfinit(x^2 - 136)).cyc
%4 = [4] \\ narrow class group is cyclic ~ Z/4Z
```

Note that the use of `bnfarrow` above is only valid because 136 is a fundamental discriminant: that function is asymptotically faster (and returns the group structure, not only its order) but only supports *maximal* orders. Here are a few more examples:

```
? qfbclassno(400000028) \\ D > 0: slow
time = 3,140 ms.
```

```

%1 = 1
? quadclassunit(400000028).no
time = 20 ms. \\ much faster, assume GRH
%2 = 1
? qfbclassno(-400000028) \\ D < 0: fast enough
time = 0 ms.
%3 = 7253
? quadclassunit(-400000028).no
time = 0 ms.
%4 = 7253

```

See also `qfbhclassno`.

The library syntax is GEN `qfbclassno0(GEN D, long flag)`.

3.8.83 `qfbcomp`(x, y). composition of the binary quadratic forms x and y , with reduction of the result.

```

? x=Qfb(2,3,-10);y=Qfb(5,3,-4);
? qfbcomp(x,y)
%2 = Qfb(-2, 9, 1)
? qfbcomp(x,y)==qfbred(qfbcompraw(x,y))
%3 = 1

```

The library syntax is GEN `qfbcomp(GEN x, GEN y)`.

3.8.84 `qfbcompraw`(x, y). composition of the binary quadratic forms x and y , without reduction of the result. This is useful e.g. to compute a generating element of an ideal. The result is undefined if x and y do not have the same discriminant.

```

? x=Qfb(2,3,-10);y=Qfb(5,3,-4);
? qfbcompraw(x,y)
%2 = Qfb(10, 3, -2)
? x=Qfb(2,3,-10);y=Qfb(1,-1,1);
? qfbcompraw(x,y)
*** at top-level: qfbcompraw(x,y)
*** ^-----
*** qfbcompraw: inconsistent qfbcompraw t_QFB , t_QFB.

```

The library syntax is GEN `qfbcompraw(GEN x, GEN y)`.

3.8.85 `qfbcornacchia`(d, n). Solves the equation $x^2 + dy^2 = n$ in integers x and y , where $d > 0$ and n is prime. Returns the empty vector `[]` when no solution exists. It is also allowed to try $n = 4$ times a prime but the answer is then guaranteed only if d is 3 mod 4; more precisely if $d \not\equiv 3 \pmod{4}$, the algorithm may fail to find a non-primitive solution.

This function is a special case of `qfbsolve` applied to the principal form in the imaginary quadratic order of discriminant $-4d$ (returning the solution with non-negative x and y). As its name implies, `qfbcornacchia` uses Cornacchia's algorithm and runs in time quasi-linear in $\log n$ (using `halfgcd`); in practical ranges, `qfbcornacchia` should be about twice faster than `qfbsolve` unless we indicate to the latter that its second argument is prime (see below).

```

? qfbcornacchia(1, 113)

```

```

%1 = [8, 7]
? qfbsolve(Qfb(1,0,1), 113)
%2 = [8, 7]
? qfbcornacchia(1, 4*113) \\ misses the non-primitive solution 2*[8,7]
%3 = []
? qfbcornacchia(1, 4*109) \\ finds a non-primitive solution
%4 = [20, 6]
? p = 122838793181521; isprime(p)
%5 = 1
? qfbcornacchia(24, p)
%6 = [10547339, 694995]
? Q = Qfb(1,0,24); qfbsolve(Q,p)
%7 = [10547339, 694995]
? for (i=1, 10^5, qfbsolve(Q, p))
time = 345 ms.
? for (i=1, 10^5, qfbcornacchia(24,p)) \\ faster
time = 251 ms.
? for (i=1, 10^5, qfbsolve(Q, Mat([p,1]))) \\ just as fast
time = 251 ms.

```

We used `Mat([p,1])` to indicate that p^1 was the integer factorization of p , i.e., that p is prime. Without it, `qfbsolve` attempts to factor p and wastes a little time.

The library syntax is `GEN qfbcornacchia(GEN d, GEN n)`.

3.8.86 qfbhclassno(x). Hurwitz class number of x , when x is nonnegative and congruent to 0 or 3 modulo 4, and 0 for other values. For $x > 5 \cdot 10^5$, we assume the GRH, and use `quadclassunit` with default parameters.

```

? qfbhclassno(1) \\ not 0 or 3 mod 4
%1 = 0
? qfbhclassno(3)
%2 = 1/3
? qfbhclassno(4)
%3 = 1/2
? qfbhclassno(23)
%4 = 3

```

The library syntax is `GEN hclassno(GEN x)`.

3.8.87 qfbnucomp(x, y, L). composition of the primitive positive definite binary quadratic forms x and y (type `t_QFB`) using the NUCOMP and NUDUPL algorithms of Shanks, à la Atkin. L is any positive constant, but for optimal speed, one should take $L = |D/4|^{1/4}$, i.e. `sqrtnint(abs(D)>>2,4)`, where D is the common discriminant of x and y . When x and y do not have the same discriminant, the result is undefined.

The current implementation is slower than the generic routine for small D , and becomes faster when D has about 45 bits.

The library syntax is `GEN nucomp(GEN x, GEN y, GEN L)`. Also available is `GEN nudupl(GEN x, GEN L)` when $x = y$.

3.8.88 qfbnupow($x, n, \{L\}$). n -th power of the primitive positive definite binary quadratic form x using Shanks's NUCOMP and NUDUPL algorithms; if set, L should be equal to `sqrtnint(abs(D)>>2,4)`, where $D < 0$ is the discriminant of x .

The current implementation is slower than the generic routine for small discriminant D , and becomes faster for $D \approx 2^{45}$.

The library syntax is `GEN nupow(GEN x, GEN n, GEN L = NULL)`.

3.8.89 qfbpow(x, n). n -th power of the binary quadratic form x , computed with reduction (i.e. using `qfbcomp`).

The library syntax is `GEN qfbpow(GEN x, GEN n)`.

3.8.90 qfbpowraw(x, n). n -th power of the binary quadratic form x , computed without doing any reduction (i.e. using `qfbcompaw`). Here n must be nonnegative and $n < 2^{31}$.

The library syntax is `GEN qfbpowraw(GEN x, long n)`.

3.8.91 qfbprimeform(x, p). Prime binary quadratic form of discriminant x whose first coefficient is p , where $|p|$ is a prime number. By abuse of notation, $p = \pm 1$ is also valid and returns the unit form. Returns an error if x is not a quadratic residue mod p , or if $x < 0$ and $p < 0$. (Negative definite `t_QFB` are not implemented.)

The library syntax is `GEN primeform(GEN x, GEN p)`.

3.8.92 qfbred($x, \{flag = 0\}, \{isd\}, \{sd\}$). Reduces the binary quadratic form x (updating Shanks's distance function d if $x = [q, d]$ is an extended *indefinite* form). If $flag$ is 1, the function performs a single reduction step, and a complete reduction otherwise.

The arguments isd , sd , if present, supply the values of $\left\lfloor \sqrt{D} \right\rfloor$, and $\sqrt{D}$ respectively, where D is the discriminant (this is not checked). If $d < 0$ these values are useless.

The library syntax is `GEN qfbred0(GEN x, long flag, GEN isd = NULL, GEN sd = NULL)`. Also available is `GEN qfbred(GEN x)` ($flag$ is 0, isd and sd are NULL).

3.8.93 qfbredsl2($x, \{isD\}$). Reduction of the (real or imaginary) binary quadratic form x , returns $[y, g]$ where y is reduced and g in $SL(2, \mathbf{Z})$ is such that $g \cdot x = y$; isD , if present, must be equal to `sqrtnint(D)`, where $D > 0$ is the discriminant of x .

The action of g on x can be computed using `qfeval(x, g)`

```
? q1 = Qfb(33947, -39899, 11650);
? [q2, U] = qfbredsl2(q1)
%2 = [Qfb(749, 2207, -1712), [-1, 3; -2, 5]]
? qfeval(q1, U)
%3 = Qfb(749, 2207, -1712)
```

The library syntax is `GEN qfbredsl2(GEN x, GEN isD = NULL)`.

3.8.94 qfbsolve($Q, n, \{flag = 0\}$). Solve the equation $Q(x, y) = n$ in coprime integers x and y (primitive solutions), where Q is a binary quadratic form and n an integer, up to the action of the special orthogonal group $G = SO(Q, \mathbf{Z})$, which is isomorphic to the group of units of positive norm of the quadratic order of discriminant $D = \text{disc}Q$. If $D > 0$, G is infinite. If $D < -4$, G is of order 2, if $D = -3$, G is of order 6 and if $D = -4$, G is of order 4.

Binary digits of *flag* mean: 1: return all solutions if set, else a single solution; return 1 if a single solution is wanted (bit unset) but none exist. 2: also include imprimitive solutions.

When *flag* = 2 (return a single solution, possibly imprimitive), the algorithm returns a solution with minimal content; in particular, a primitive solution exists if and only if one is returned.

The integer n can also be given by its factorization matrix $fa = \text{factor}(n)$ or by the pair $[n, fa]$.

```
? qfbsolve(Qfb(1,0,2), 603) \\ a single primitive solution
%1 = [5, 17]

? qfbsolve(Qfb(1,0,2), 603, 1) \\ all primitive solutions
%2 = [[5, 17], [-19, -11], [19, -11], [5, -17]]

? qfbsolve(Qfb(1,0,2), 603, 2) \\ a single, possibly imprimitive solution
%3 = [5, 17] \\ actually primitive

? qfbsolve(Qfb(1,0,2), 603, 3) \\ all solutions
%4 = [[5, 17], [-19, -11], [19, -11], [5, -17], [-21, 9], [-21, -9]]

? N = 2^128+1; F = factor(N);
? qfbsolve(Qfb(1,0,1), [N,F], 1)
%3 = [[-16382350221535464479, 8479443857936402504],
      [18446744073709551616, -1], [-18446744073709551616, -1],
      [16382350221535464479, 8479443857936402504]]
```

For fixed Q , assuming the factorisation of n is given, the algorithm runs in probabilistic polynomial time in $\log p$, where p is the largest prime divisor of n , through the computation of square roots of D modulo $4p$). The dependency on Q is more complicated: polynomial time in $\log |D|$ if Q is imaginary, but exponential time if Q is real (through the computation of a full cycle of reduced forms). In the latter case, note that **bnfisprincipal** provides a solution in heuristic subexponential time assuming the GRH.

The library syntax is **GEN qfbsolve(GEN Q, GEN n, long flag)**.

3.8.95 quadclassunit($D, \{flag = 0\}, \{tech = []\}$). Buchmann-McCurley's sub-exponential algorithm for computing the class group of a quadratic order of discriminant D . By default, the results are conditional on the GRH.

This function should be used instead of **qfbclassno** or **quadregulator** when $D < -10^{25}$, $D > 10^{10}$, or when the *structure* is wanted. It is a special case of **bnfinit**, which is slower, but more robust.

The result is a vector v whose components should be accessed using member functions:

- $v.\text{no}$: the class number
- $v.\text{cyc}$: a vector giving the structure of the class group as a product of cyclic groups;
- $v.\text{gen}$: a vector giving generators of those cyclic groups (as binary quadratic forms).

- ***v.reg***: the regulator, computed to an accuracy which is the maximum of an internal accuracy determined by the program and the current default (note that once the regulator is known to a small accuracy it is trivial to compute it to very high accuracy, see the tutorial).

- ***v.normfu*** (for positive D only) return the norm of the fundamental unit, either 1 or -1 . Note that a result of -1 is unconditional and no longer depends on the GRH.

The *flag* is obsolete and should be left alone. In older versions, it supposedly computed the narrow class group when $D > 0$, but this did not work at all; use the general function **bnfnarrow**.

Optional parameter *tech* is a row vector of the form $[c_1, c_2]$, where $c_1 \leq c_2$ are nonnegative real numbers which control the execution time and the stack size, see 3.13.7. The parameter is used as a threshold to balance the relation finding phase against the final linear algebra. Increasing the default c_1 means that relations are easier to find, but more relations are needed and the linear algebra will be harder. The default value for c_1 is 0 and means that it is taken equal to c_2 . The parameter c_2 is mostly obsolete and should not be changed, but we still document it for completeness: we compute a tentative class group by generators and relations using a factorbase of prime ideals $\leq c_1(\log |D|)^2$, then prove that ideals of norm $\leq c_2(\log |D|)^2$ do not generate a larger group. By default an optimal c_2 is chosen, so that the result is provably correct under the GRH — a result of Grenié and Molteni states that $c_2 = 23/6 \approx 3.83$ is fine (and even $c_2 = 15/4 \approx 3.75$ for large $|D| > 2.41E8$). But it is possible to improve on this algorithmically. You may provide a smaller c_2 , it will be ignored (we use the provably correct one); you may provide a larger c_2 than the default value, which results in longer computing times for equally correct outputs (under GRH).

The library syntax is **GEN quadclassunit0(GEN D, long flag, GEN tech = NULL, long prec)**. If you really need to experiment with the *tech* parameter, it will be more convenient to use **GEN Buchquad(GEN D, double c1, double c2, long prec)**.

3.8.96 quaddisc(x). Discriminant of the étale algebra $\mathbf{Q}(\sqrt{x})$, where $x \in \mathbf{Q}^*$. This is the same as **coredisc(d)** where d is the integer squarefree part of x , so $x = df^2$ with $f \in \mathbf{Q}^*$ and $d \in \mathbf{Z}$. This returns 0 for $x = 0$, 1 for x square and the discriminant of the quadratic field $\mathbf{Q}(\sqrt{x})$ otherwise.

```
? quaddisc(7)
%1 = 28
? quaddisc(-7)
%2 = -7
```

The library syntax is **GEN quaddisc(GEN x)**.

3.8.97 quadgen($D, \{v = 'w\}$). Creates the quadratic number $\omega = (a + \sqrt{D})/2$ where $a = 0$ if $D \equiv 0 \pmod{4}$, $a = 1$ if $D \equiv 1 \pmod{4}$, so that $(1, \omega)$ is an integral basis for the quadratic order of discriminant D . D must be an integer congruent to 0 or 1 modulo 4, which is not a square. If v is given, the variable name is used to display g else 'w' is used.

```
? w = quadgen(5, 'w); w^2 - w - 1
%1 = 0
? w = quadgen(0, 'w)
*** at top-level: w=quadgen(0)
***      ^-----
*** quadgen: domain error in quadpoly: issquare(disc) = 1
```

The library syntax is **GEN quadgen0(GEN D, long v = -1)** where v is a variable number.

When v does not matter, the function **GEN quadgen(GEN D)** is also available.

3.8.98 quadhilbert(D). Relative equation defining the Hilbert class field of the quadratic field of discriminant D .

If $D < 0$, uses complex multiplication (Schertz’s variant).

If $D > 0$ Stark units are used and (in rare cases) a vector of extensions may be returned whose compositum is the requested class field. See **bnrstark** for details.

The library syntax is **GEN quadhilbert(GEN D, long prec)**.

3.8.99 quadpoly($D, \{v = 'x\}$). Creates the “canonical” quadratic polynomial (in the variable v) corresponding to the discriminant D , i.e. the minimal polynomial of **quadgen(D)**. D must be an integer congruent to 0 or 1 modulo 4, which is not a square.

```
? quadpoly(5,'y)
%1 = y^2 - y - 1
? quadpoly(0,'y)
*** at top-level: quadpoly(0,'y)
***                               ^-----
*** quadpoly: domain error in quadpoly: issquare(disc) = 1
```

The library syntax is **GEN quadpoly0(GEN D, long v = -1)** where v is a variable number.

3.8.100 quadray(D, f). Relative equation for the ray class field of conductor f for the quadratic field of discriminant D using analytic methods. A **bnf** for $x^2 - D$ is also accepted in place of D .

For $D < 0$, uses the σ function and Schertz’s method.

For $D > 0$, uses Stark’s conjecture, and a vector of relative equations may be returned. See **bnrstark** for more details.

The library syntax is **GEN quadray(GEN D, GEN f, long prec)**.

3.8.101 quadregulator(D). Regulator of the quadratic order of positive discriminant D in time $\tilde{O}(D^{1/2})$ using the continued fraction algorithm. Raise an error if D is not a discriminant (fundamental or not) or if D is a square. The function **quadclassunit** is asymptotically faster (and also in practice for $D > 10^{10}$ or so) but depends on the GRH.

The library syntax is **GEN quadregulator(GEN D, long prec)**.

3.8.102 quadunit($D, \{v = 'w\}$). A fundamental unit u of the real quadratic order of discriminant D . The integer D must be congruent to 0 or 1 modulo 4 and not a square; the result is a quadratic number (see Section 3.8.97). If D is not a fundamental discriminant, the algorithm is wasteful: if $D = df^2$ with d fundamental, it will be faster to compute **quadunit(d)** then raise it to the power **quadunitindex(d, f)**; or keep it in factored form.

If v is given, the variable name is used to display u else ‘w’ is used. The algorithm computes the continued fraction of $(1 + \sqrt{D})/2$ or $\sqrt{D}/2$ (see GTM 138, algorithm 5.7.2). Although the continued fraction length is only $O(\sqrt{D})$, the function still runs in time $\tilde{O}(D)$, in part because the output size is not polynomially bounded in terms of $\log D$. See **bnfinit** and **bnfunits** for a better alternative for large D , running in time subexponential in $\log D$ and returning the fundamental units in compact form (as a short list of S -units of size $O(\log D)^3$ raised to possibly large exponents).

The library syntax is **GEN quadunit0(GEN D, long v = -1)** where v is a variable number.

When v does not matter, the function **GEN quadunit(GEN D)** is also available.

Important remark. Assuming GRH, using `bnfinit` is *much* faster, running in time subexponential in $\log D$ (instead of exponential for `quadunitnorm`). We give examples for the maximal order:

```
? GRHunitnorm(bnf) = vecprod(bnfsignunit(bnf)[,1])
? bnf = bnfinit(x^2 - 17345, 1); GRHunitnorm(bnf)
%2 = -1
? bnf = bnfinit(x^2 - nextprime(2^60), 1); GRHunitnorm(bnf)
time = 119 ms.
%3 = -1
? quadunitnorm(nextprime(2^60))
time = 24,086 ms.
%4 = -1
```

Note that if the result is -1 , it is unconditional because (if GRH is false) it could happen that our tentative fundamental unit in *bnf* is actually a power u^k of the true fundamental unit, but we would still have $\text{Norm}(u) = -1$ (and k odd). We can also remove the GRH assumption when the result is 1 with a little more work:

```
? v = bnfunits(bnf)[1][1] \\ a unit in factored form
? v[,2] %= 2;
? nfeltissquare(bnf, nffactorback(bnf, v))
%7 = 0
```

Under GRH, we know that v is the fundamental unit, but as above it can be a power u^k of the true fundamental unit u . But the final two lines prove that v is not a square, hence k is odd and $\text{Norm}(u)$ must also be 1. We modified the factorization matrix giving v by reducing all exponents modulo 2: this allows to compute `nffactorback` even when the factorization involves huge exponents. And of course the new v is a square if and only if the original one was.

The library syntax is `long quadunitnorm(GEN D)`.

3.8.105 ramanujantau($n, \{ell = 12\}$). Compute the value of Ramanujan's tau function at an individual n , assuming the truth of the GRH (to compute quickly class numbers of imaginary quadratic fields using `quadclassunit`). If `ell` is 16, 18, 20, 22, or 26, same for the newform of level 1 and corresponding weight. Otherwise, compute the coefficient of the trace form at n . The complexity is in $\tilde{O}(n^{1/2})$ using $O(\log n)$ space.

If all values up to N are required, then

$$\sum \tau(n)q^n = q \prod_{n \geq 1} (1 - q^n)^{24}$$

and more generally, setting $u = \ell - 13$ and $C = 2/\zeta(-u)$ for $\ell > 12$,

$$\sum \tau_\ell(n)q^n = q \prod_{n \geq 1} (1 - q^n)^{24} \left(1 + C \sum_{n \geq 1} n^u q^n / (1 - q^n) \right)$$

produces them in time $\tilde{O}(N)$, against $\tilde{O}(N^{3/2})$ for individual calls to `ramanujantau`; of course the space complexity then becomes $\tilde{O}(N)$. For other values of ℓ , `mfcoefs(mftraceform([1,ell]),N)` is much faster.

```
? tauvec(N) = Vec(q*eta(q + O(q^N))^24);
```

```
? N = 10^4; v = tauvec(N);
time = 26 ms.
? ramanujantau(N)
%3 = -482606811957501440000
? w = vector(N, n, ramanujantau(n)); \\ much slower !
time = 13,190 ms.
? v == w
%4 = 1
```

The library syntax is GEN ramanujantau(GEN n, long ell).

3.8.106 randomprime($\{N = 2^{31}\}, \{q\}$). Returns a strong pseudo prime (see `ispseudoprime`) in $[2, N - 1]$. A `t_VEC` $N = [a, b]$ is also allowed, with $a \leq b$ in which case a pseudo prime $a \leq p \leq b$ is returned; if no prime exists in the interval, the function will run into an infinite loop. If the upper bound is less than 2^{64} the pseudo prime returned is a proven prime.

```
? randomprime(100)
%1 = 71
? randomprime([3,100])
%2 = 61
? randomprime([1,1])
*** at top-level: randomprime([1,1])
*** ^-----
*** randomprime: domain error in randomprime:
*** floor(b) - max(ceil(a),2) < 0
? randomprime([24,28]) \\ infinite loop
```

If the optional parameter q is an integer, return a prime congruent to $1 \bmod q$; if q is an `intmod`, return a prime in the given congruence class. If the class contains no prime in the given interval, the function will raise an exception if the class is not invertible, else run into an infinite loop

```
? randomprime(100, 4) \\ 1 mod 4
%1 = 71
? randomprime(100, 4)
%2 = 13
? randomprime([10,100], Mod(2,5))
%3 = 47
? randomprime(100, Mod(0,2)) \\ silly but works
%4 = 2
? randomprime([3,100], Mod(0,2)) \\ not invertible
*** at top-level: randomprime([3,100],Mod(0,2))
*** ^-----
*** randomprime: elements not coprime in randomprime:
0
2
? randomprime(100, 97) \\ infinite loop
```

The library syntax is GEN randomprime0(GEN N = NULL, GEN q = NULL). Also available is GEN randomprime(GEN N = NULL).

3.8.107 removeprimes($\{x = []\}$). Removes the primes listed in x from the prime number table. In particular `removeprimes(addprimes())` empties the extra prime table. x can also be a single integer. List the current extra primes if x is omitted.

The library syntax is `GEN removeprimes(GEN x = NULL)`.

3.8.108 sigma($x, \{k = 1\}$). Sum of the k^{th} powers of the positive divisors of $|x|$. x and k must be of type integer.

The library syntax is `GEN sumdivk(GEN x, long k)`. Also available is `GEN sumdiv(GEN n)`, for $k = 1$.

3.8.109 sqrtint($x, \{&r\}$). Returns the integer square root of x , i.e. the largest integer y such that $y^2 \leq x$, where x a nonnegative real number. If r is present, set it to the remainder $r = x - y^2$, which satisfies $0 \leq r < 2y + 1$. Further, when x is an integer, r is an integer satisfying $0 \leq r \leq 2y$.

```
? x = 120938191237; sqrtint(x)
%1 = 347761
? sqrt(x)
%2 = 347761.68741970412747602130964414095216
? y = sqrtint(x, &r); r
%3 = 478116
? x - y^2
%4 = 478116
? sqrtint(9/4, &r) \\ not 3/2 !
%5 = 1
? r
%6 = 5/4
```

The library syntax is `GEN sqrtint0(GEN x, GEN *r = NULL)`. Also available is `GEN sqrtint(GEN a)`.

3.8.110 sqrtnint(x, n). Returns the integer n -th root of x , i.e. the largest integer y such that $y^n \leq x$, where x is a nonnegative real number.

```
? N = 120938191237; sqrtnint(N, 5)
%1 = 164
? N^(1/5)
%2 = 164.63140849829660842958614676939677391
? sqrtnint(Pi^2, 3)
%3 = 2
```

The special case $n = 2$ is `sqrtint`

The library syntax is `GEN sqrtnint(GEN x, long n)`.

3.8.111 sumdedekind(h, k). Returns the Dedekind sum attached to the integers h and k , corresponding to a fast implementation of

$$s(h, k) = \sum_{n=1}^{k-1} \left(\frac{n}{k} \right) \left(\frac{h \cdot n}{k} - \frac{1}{2} \right)$$

The library syntax is `GEN sumdedekind(GEN h, GEN k)`.

3.8.112 sumdigits($n, \{B = 10\}$). Sum of digits in the integer n , when written in base B .

```
? sumdigits(123456789)
%1 = 45
? sumdigits(123456789, 2)
%2 = 16
? sumdigits(123456789, -2)
%3 = 15
```

Note that the sum of bits in n is also returned by `hammingweight`. This function is much faster than `vecsum(digits(n,B))` when B is 10 or a power of 2, and only slightly faster in other cases.

The library syntax is `GEN sumdigits0(GEN n, GEN B = NULL)`. Also available is `GEN sumdigits(GEN n)`, for $B = 10$.

3.8.113 znchar(D). Given a datum D describing a group $(\mathbf{Z}/N\mathbf{Z})^*$ and a Dirichlet character χ , return the pair `[G, chi]`, where `G` is `znstar(N, 1)` and `chi` is a GP character.

The following possibilities for D are supported

- a nonzero `t_INT` congruent to 0,1 modulo 4, return the real character modulo D given by the Kronecker symbol $(D/.);$
- a `t_INTMOD Mod(m, N)`, return the Conrey character modulo N of index m (see `znconreylog`).
- a modular form space as per `mfini([N,k,χ])` or a modular form for such a space, return the underlying Dirichlet character χ (which may be defined modulo a divisor of N but need not be primitive).

In the remaining cases, `G` is initialized by `znstar(N, 1)`.

- a pair `[G, chi]`, where `chi` is a standard GP Dirichlet character $c = (c_j)$ on `G` (generic character `t_VEC` or Conrey characters `t_COL` or `t_INT`); given generators $G = \oplus (\mathbf{Z}/d_j\mathbf{Z})g_j$, $\chi(g_j) = e(c_j/d_j)$.

- a pair `[G, chin]`, where `chin` is a *normalized* representation $[n, \tilde{c}]$ of the Dirichlet character c ; $\chi(g_j) = e(\tilde{c}_j/n)$ where n is minimal (order of χ).

```
? [G,chi] = znchar(-3);
? G.cyc
%2 = [2]
? chareval(G, chi, 2)
%3 = 1/2
? kronecker(-3,2)
%4 = -1
? znchartokronecker(G,chi)
%5 = -3
? mf = mfini([28, 5/2, Mod(2,7)]); [f] = mfbasis(mf);
? [G,chi] = znchar(mf); [G.mod, chi]
%7 = [7, [2]~]
? [G,chi] = znchar(f); chi
%8 = [28, [0, 2]~]
```

The library syntax is `GEN znchar(GEN D)`.

3.8.114 zncharconductor(G, chi). Let G be attached to $(\mathbf{Z}/q\mathbf{Z})^*$ (as per $G = \text{znstar}(q, 1)$) and chi be a Dirichlet character on $(\mathbf{Z}/q\mathbf{Z})^*$ (see Section 3.8.3 or ??character). Return the conductor of chi :

```
? G = znstar(126000, 1);
? zncharconductor(G,11)  \\ primitive
%2 = 126000
? zncharconductor(G,1)   \\ trivial character, not primitive!
%3 = 1
? zncharconductor(G,1009) \\ character mod 5^3
%4 = 125
```

The library syntax is GEN zncharconductor(GEN G, GEN chi).

3.8.115 znchardecompose(G, chi, Q). Let $N = \prod_p p^{e_p}$ and a Dirichlet character χ , we have a decomposition $\chi = \prod_p \chi_p$ into character modulo N where the conductor of χ_p divides p^{e_p} ; it equals p^{e_p} for all p if and only if χ is primitive.

Given a *znstar* G describing a group $(\mathbf{Z}/N\mathbf{Z})^*$, a Dirichlet character chi and an integer Q , return $\prod_{p|(Q,N)} \chi_p$. For instance, if $Q = p$ is a prime divisor of N , the function returns χ_p (as a character modulo N), given as a Conrey character (t_COL).

```
? G = znstar(40, 1);
? G.cyc
%2 = [4, 2, 2]
? chi = [2, 1, 1];
? chi2 = znchardecompose(G, chi, 2)
%4 = [1, 1, 0]~
? chi5 = znchardecompose(G, chi, 5)
%5 = [0, 0, 2]~
? znchardecompose(G, chi, 3)
%6 = [0, 0, 0]~
? c = charmul(G, chi2, chi5)
%7 = [1, 1, 2]~  \\ t_COL: in terms of Conrey generators !
? znconreychar(G,c)
%8 = [2, 1, 1]  \\ t_VEC: in terms of SNF generators
```

The library syntax is GEN znchardecompose(GEN G, GEN chi, GEN Q).

3.8.116 znchargauss($G, \text{chi}, \{a = 1\}$). Given a Dirichlet character χ on $G = (\mathbf{Z}/N\mathbf{Z})^*$ (see znchar), return the complex Gauss sum

$$g(\chi, a) = \sum_{n=1}^N \chi(n) e(an/N)$$

```
? [G,chi] = znchar(-3); \\ quadratic Gauss sum: I*sqrt(3)
? znchargauss(G,chi)
%2 = 1.7320508075688772935274463415058723670*I
? [G,chi] = znchar(5);
? znchargauss(G,chi)  \\ sqrt(5)
```

```
%2 = 2.2360679774997896964091736687312762354
? G = znstar(300,1); chi = [1,1,12]~;
? znchargauss(G,chi) / sqrt(300) - exp(2*I*Pi*11/25) \\ = 0
%4 = 2.350988701644575016 E-38 + 1.4693679385278593850 E-39*I
? lfuntheta([G,chi], 1) \\ = 0
%5 = -5.79[...] E-39 - 2.71[...] E-40*I
```

The library syntax is GEN znchargauss(GEN G, GEN chi, GEN a = NULL, long bitprec)

3.8.117 zncharinduce(G, χ, N). Let G be attached to $(\mathbf{Z}/q\mathbf{Z})^*$ (as per $G = \text{znstar}(q,1)$) and let χ be a Dirichlet character on $(\mathbf{Z}/q\mathbf{Z})^*$, given by

- a `t_VEC`: a standard character on `bid.gen`,
- a `t_INT` or a `t_COL`: a Conrey index in $(\mathbf{Z}/q\mathbf{Z})^*$ or its Conrey logarithm; see Section 3.8.3 or `??character`.

Let N be a multiple of q , return the character modulo N extending χ . As usual for arithmetic functions, the new modulus N can be given as a `t_INT`, via a factorization matrix or a pair `[N, factor(N)]`, or by `znstar(N,1)`.

```
? G = znstar(4, 1);
? chi = znconreylog(G,1); \\ trivial character mod 4
? zncharinduce(G, chi, 80) \\ now mod 80
%3 = [0, 0, 0]~
? zncharinduce(G, 1, 80) \\ same using directly Conrey label
%4 = [0, 0, 0]~
? G2 = znstar(80, 1);
? zncharinduce(G, 1, G2) \\ same
%4 = [0, 0, 0]~
? chi = zncharinduce(G, 3, G2) \\ extend the nontrivial character mod 4
%5 = [1, 0, 0]~
? [G0,chi0] = znchartoprimitive(G2, chi);
? G0.mod
%7 = 4
? chi0
%8 = [1]~
```

Here is a larger example:

```
? G = znstar(126000, 1);
? label = 1009;
? chi = znconreylog(G, label)
%3 = [0, 0, 0, 14, 0]~
? [G0,chi0] = znchartoprimitive(G, label); \\ works also with 'chi'
? G0.mod
%5 = 125
? chi0 \\ primitive character mod 5^3 attached to chi
%6 = [14]~
? G0 = znstar(N0, 1);
? zncharinduce(G0, chi0, G) \\ induce back
```

```
%8 = [0, 0, 0, 14, 0]~
? znconreyexp(G, %)
%9 = 1009
```

The library syntax is `GEN zncharinduce(GEN G, GEN chi, GEN N)`.

3.8.118 `zncharisodd( $G, chi$ )`. Let G be attached to $(\mathbf{Z}/N\mathbf{Z})^*$ (as per $G = \text{znstar}(N, 1)$) and let chi be a Dirichlet character on $(\mathbf{Z}/N\mathbf{Z})^*$, given by

- a `t_VEC`: a standard character on $G.\text{gen}$,
- a `t_INT` or a `t_COL`: a Conrey index in $(\mathbf{Z}/q\mathbf{Z})^*$ or its Conrey logarithm; see Section 3.8.3 or `??character`.

Return 1 if and only if $chi(-1) = -1$ and 0 otherwise.

```
? G = znstar(8, 1);
? zncharisodd(G, 1)  \\ trivial character
%2 = 0
? zncharisodd(G, 3)
%3 = 1
? chareval(G, 3, -1)
%4 = 1/2
```

The library syntax is `long zncharisodd(GEN G, GEN chi)`.

3.8.119 `znchartokronecker( $G, chi, \{flag = 0\}$ )`. Let G be attached to $(\mathbf{Z}/N\mathbf{Z})^*$ (as per $G = \text{znstar}(N, 1)$) and let chi be a Dirichlet character on $(\mathbf{Z}/N\mathbf{Z})^*$, given by

- a `t_VEC`: a standard character on bid.gen ,
- a `t_INT` or a `t_COL`: a Conrey index in $(\mathbf{Z}/q\mathbf{Z})^*$ or its Conrey logarithm; see Section 3.8.3 or `??character`.

If $flag = 0$, return the discriminant D if chi is real equal to the Kronecker symbol $(D/.)$ and 0 otherwise. The discriminant D is fundamental if and only if chi is primitive.

If $flag = 1$, return the fundamental discriminant attached to the corresponding primitive character.

```
? G = znstar(8,1); CHARS = [1,3,5,7]; \\ Conrey labels
? apply(t->znchartokronecker(G,t), CHARS)
%2 = [4, -8, 8, -4]
? apply(t->znchartokronecker(G,t,1), CHARS)
%3 = [1, -8, 8, -4]
```

The library syntax is `GEN znchartokronecker(GEN G, GEN chi, long flag)`.

3.8.120 znchartoprimitive(G, chi). Let G be attached to $(\mathbf{Z}/q\mathbf{Z})^*$ (as per $G = \text{znstar}(q, 1)$) and chi be a Dirichlet character on $(\mathbf{Z}/q\mathbf{Z})^*$, of conductor $q_0 \mid q$. Return $[G_0, \text{chi0}]$, where chi0 is the primitive character attached to chi and G_0 is $\text{znstar}(q_0, 1)$; the character chi0 is returned as a Conrey logarithm unless chi is primitive, in which case chi0 is identical to chi .

```
? G = znstar(126000, 1);
? [G0,chi0] = znchartoprimitive(G,11)
? G0.mod
%3 = 126000
? chi0
%4 = 11
? [G0,chi0] = znchartoprimitive(G,1);\ \ trivial character, not primitive!
? G0.mod
%6 = 1
? chi0
%7 = []~
? [G0,chi0] = znchartoprimitive(G,1009)
? G0.mod
%4 = 125
? chi0
%5 = [14]~
```

Note that `znconreyconductor` is more efficient since it can return χ_0 and its conductor q_0 without needing to initialize G_0 . The price to pay is a more cryptic format and the need to initialize G_0 later, but that needs to be done only once for all characters with conductor q_0 .

The library syntax is `GEN znchartoprimitive(GEN G, GEN chi)`.

3.8.121 znconreychar(G, m). Given a *znstar* G attached to $(\mathbf{Z}/q\mathbf{Z})^*$ (as per $G = \text{znstar}(q, 1)$), this function returns the Dirichlet character attached to $m \in (\mathbf{Z}/q\mathbf{Z})^*$ via Conrey's logarithm, which establishes a “canonical” bijection between $(\mathbf{Z}/q\mathbf{Z})^*$ and its dual.

Let $q = \prod_p p^{e_p}$ be the factorization of q into distinct primes. For all odd p with $e_p > 0$, let g_p be the element in $(\mathbf{Z}/q\mathbf{Z})^*$ which is

- congruent to 1 mod q/p^{e_p} ,
- congruent mod p^{e_p} to the smallest positive integer that generates $(\mathbf{Z}/p^2\mathbf{Z})^*$.

For $p = 2$, we let g_4 (if $2^{e_2} \geq 4$) and g_8 (if furthermore $2^{e_2} \geq 8$) be the elements in $(\mathbf{Z}/q\mathbf{Z})^*$ which are

- congruent to 1 mod $q/2^{e_2}$,
- $g_4 = -1 \bmod 2^{e_2}$,
- $g_8 = 5 \bmod 2^{e_2}$.

Then the g_p (and the extra g_4 and g_8 if $2^{e_2} \geq 2$) are independent generators of $(\mathbf{Z}/q\mathbf{Z})^*$, i.e. every m in $(\mathbf{Z}/q\mathbf{Z})^*$ can be written uniquely as $\prod_p g_p^{m_p}$, where m_p is defined modulo the order o_p of g_p and $p \in S_q$, the set of prime divisors of q together with 4 if $4 \mid q$ and 8 if $8 \mid q$. Note that the g_p are in general *not* SNF generators as produced by `znstar` whenever $\omega(q) \geq 2$, although their number is the same. They however allow to handle the finite abelian group $(\mathbf{Z}/q\mathbf{Z})^*$ in a fast and elegant way. (Which unfortunately does not generalize to ray class groups or Hecke characters.)

The Conrey logarithm of m is the vector $(m_p)_{p \in S_q}$, obtained via `znconreylog`. The Conrey character $\chi_q(m, \cdot)$ attached to $m \bmod q$ maps each $g_p, p \in S_q$ to $e(m_p/o_p)$, where $e(x) = \exp(2i\pi x)$. This function returns the Conrey character expressed in the standard PARI way in terms of the SNF generators `G.gen`.

```
? G = znstar(8,1);
? G.cyc
%2 = [2, 2]  \\ Z/2 x Z/2
? G.gen
%3 = [7, 3]
? znconreychar(G,1)  \\ 1 is always the trivial character
%4 = [0, 0]
? znconreychar(G,2)  \\ 2 is not coprime to 8 !!!
***   at top-level: znconreychar(G,2)
***                               ^-----
*** znconreychar: elements not coprime in Zideallog:
    2
    8
***   Break loop: type 'break' to go back to GP prompt
break>
? znconreychar(G,3)
%5 = [0, 1]
? znconreychar(G,5)
%6 = [1, 1]
? znconreychar(G,7)
%7 = [1, 0]
```

We indeed get all 4 characters of $(\mathbf{Z}/8\mathbf{Z})^*$.

For convenience, we allow to input the *Conrey logarithm* of m instead of m :

```
? G = znstar(55, 1);
? znconreychar(G,7)
%2 = [7, 0]
? znconreychar(G, znconreylog(G,7))
%3 = [7, 0]
```

The library syntax is `GEN znconreychar(GEN G, GEN m)`.

3.8.122 znconreyconductor($G, \text{chi}, \{\&\text{chi0}\}$). Let G be attached to $(\mathbf{Z}/q\mathbf{Z})^*$ (as per `G = znstar(q, 1)`) and chi be a Dirichlet character on $(\mathbf{Z}/q\mathbf{Z})^*$, given by

- a `t_VEC`: a standard character on `bid.gen`,
- a `t_INT` or a `t_COL`: a Conrey index in $(\mathbf{Z}/q\mathbf{Z})^*$ or its Conrey logarithm; see Section 3.8.3 or `??character`.

Return the conductor of chi , as the `t_INT` `bid.mod` if chi is primitive, and as a pair $[N, \text{faN}]$ (with `faN` the factorization of N) otherwise.

If `chi0` is present, set it to the Conrey logarithm of the attached primitive character.

```
? G = znstar(126000, 1);
? znconreyconductor(G,11)  \\ primitive
```

```

%2 = 126000
? znconreyconductor(G,1)    \\ trivial character, not primitive!
%3 = [1, matrix(0,2)]
? N0 = znconreyconductor(G,1009, &chi0) \\ character mod 5^3
%4 = [125, Mat([5, 3])]
? chi0
%5 = [14]~
? G0 = znstar(N0, 1);      \\ format [N,factor(N)] accepted
? znconreyexp(G0, chi0)
%7 = 9
? znconreyconductor(G0, chi0) \\ now primitive, as expected
%8 = 125

```

The group G_0 is not computed as part of `znconreyconductor` because it needs to be computed only once per conductor, not once per character.

The library syntax is `GEN znconreyconductor(GEN G, GEN chi, GEN *chi0 = NULL)`.

3.8.123 `znconreyexp(G, chi)`. Given a *znstar* G attached to $(\mathbf{Z}/q\mathbf{Z})^*$ (as per $G = \text{znstar}(q, 1)$), this function returns the Conrey exponential of the character chi : it returns the integer $m \in (\mathbf{Z}/q\mathbf{Z})^*$ such that $\text{znconreylog}(G, m)$ is chi .

The character chi is given either as a

- `t_VEC`: in terms of the generators `G.gen`;
- `t_COL`: a Conrey logarithm.

```

? G = znstar(126000, 1)
? znconreylog(G,1)
%2 = [0, 0, 0, 0, 0]~
? znconreyexp(G,%)
%3 = 1
? G.cyc \\ SNF generators
%4 = [300, 12, 2, 2, 2]
? chi = [100, 1, 0, 1, 0]; \\ some random character on SNF generators
? znconreylog(G, chi) \\ in terms of Conrey generators
%6 = [0, 3, 3, 0, 2]~
? znconreyexp(G, %) \\ apply to a Conrey log
%7 = 18251
? znconreyexp(G, chi) \\ ... or a char on SNF generators
%8 = 18251
? znconreychar(G,%)
%9 = [100, 1, 0, 1, 0]

```

The library syntax is `GEN znconreyexp(GEN G, GEN chi)`.

3.8.124 znconreylog(G, m). Given a *znstar* attached to $(\mathbf{Z}/q\mathbf{Z})^*$ (as per $G = \text{znstar}(q, 1)$), this function returns the Conrey logarithm of $m \in (\mathbf{Z}/q\mathbf{Z})^*$.

Let $q = \prod_p p^{e_p}$ be the factorization of q into distinct primes, where we assume $e_2 = 0$ or $e_2 \geq 2$. (If $e_2 = 1$, we can ignore 2 from the factorization, as if we replaced q by $q/2$, since $(\mathbf{Z}/q\mathbf{Z})^* \sim (\mathbf{Z}/(q/2)\mathbf{Z})^*$.)

For all odd p with $e_p > 0$, let g_p be the element in $(\mathbf{Z}/q\mathbf{Z})^*$ which is

- congruent to 1 mod q/p^{e_p} ,
- congruent mod p^{e_p} to the smallest positive integer that generates $(\mathbf{Z}/p^2\mathbf{Z})^*$.

For $p = 2$, we let g_4 (if $2^{e_2} \geq 4$) and g_8 (if furthermore $2^{e_2} \geq 8$) be the elements in $(\mathbf{Z}/q\mathbf{Z})^*$ which are

- congruent to 1 mod $q/2^{e_2}$,
- $g_4 = -1 \bmod 2^{e_2}$,
- $g_8 = 5 \bmod 2^{e_2}$.

Then the g_p (and the extra g_4 and g_8 if $2^{e_2} \geq 2$) are independent generators of $\mathbf{Z}/q\mathbf{Z}^*$, i.e. every m in $(\mathbf{Z}/q\mathbf{Z})^*$ can be written uniquely as $\prod_p g_p^{m_p}$, where m_p is defined modulo the order o_p of g_p and $p \in S_q$, the set of prime divisors of q together with 4 if $4 \mid q$ and 8 if $8 \mid q$. Note that the g_p are in general *not* SNF generators as produced by **znstar** whenever $\omega(q) \geq 2$, although their number is the same. They however allow to handle the finite abelian group $(\mathbf{Z}/q\mathbf{Z})^*$ in a fast and elegant way. (Which unfortunately does not generalize to ray class groups or Hecke characters.)

The Conrey logarithm of m is the vector $(m_p)_{p \in S_q}$. The inverse function **znconreyexp** recovers the Conrey label m from a character.

```
? G = znstar(126000, 1);
? znconreylog(G,1)
%2 = [0, 0, 0, 0, 0]~
? znconreyexp(G, %)
%3 = 1
? znconreylog(G,2)  \\ 2 is not coprime to modulus !!!
***   at top-level: znconreylog(G,2)
***               ^-----
*** znconreylog: elements not coprime in Zideallog:
    2
    126000
***   Break loop: type 'break' to go back to GP prompt
break>
? znconreylog(G,11)  \\ wrt. Conrey generators
%4 = [0, 3, 1, 76, 4]~
? log11 = ideallog(,11,G)  \\ wrt. SNF generators
%5 = [178, 3, -75, 1, 0]~
```

For convenience, we allow to input the ordinary discrete log of m , **ideallog**(m, bid), which allows to convert discrete logs from **bid.gen** generators to Conrey generators.

```
? znconreylog(G, log11)
%7 = [0, 3, 1, 76, 4]~
```

We also allow a character (`t_VEC`) on `bid.gen` and return its representation on the Conrey generators.

```
? G.cyc
%8 = [300, 12, 2, 2, 2]
? chi = [10,1,0,1,1];
? znconreylog(G, chi)
%10 = [1, 3, 3, 10, 2]~
? n = znconreyexp(G, chi)
%11 = 84149
? znconreychar(G, n)
%12 = [10, 1, 0, 1, 1]
```

The library syntax is `GEN znconreylog(GEN G, GEN m)`.

3.8.125 zncoppersmith($P, N, X, \{B = N\}$). Coppersmith's algorithm. N being an integer and $P \in \mathbf{Z}[t]$, finds in polynomial time in $\log(N)$ and $d = \deg(P)$ all integers x with $|x| \leq X$ such that

$$\gcd(N, P(x)) \geq B.$$

This is a famous application of the LLL algorithm meant to help in the factorization of N . Notice that P may be reduced modulo $N\mathbf{Z}[t]$ without affecting the situation. The parameter X must not be too large: assume for now that the leading coefficient of P is coprime to N , then we must have

$$d \log X \log N < \log^2 B,$$

i.e., $X < N^{1/d}$ when $B = N$. Let now P_0 be the gcd of the leading coefficient of P and N . In applications to factorization, we should have $P_0 = 1$; otherwise, either $P_0 = N$ and we can reduce the degree of P , or P_0 is a non trivial factor of N . For completeness, we nevertheless document the exact conditions that X must satisfy in this case: let $p := \log_N P_0$, $b := \log_N B$, $x := \log_N X$, then

- either $p \geq d/(2d-1)$ is large and we must have $xd < 2b-1$;
 - or $p < d/(2d-1)$ and we must have both $p < b < 1-p+p/d$ and $x(d+p(1-2d)) < (b-p)^2$.
- Note that this reduces to $xd < b^2$ when $p = 0$, i.e., the condition described above.

Some x larger than X may be returned if you are very lucky. The routine runs in polynomial time in $\log N$ and d but the smaller B , or the larger X , the slower. The strength of Coppersmith method is the ability to find roots modulo a general *composite* N : if N is a prime or a prime power, `polrootsmod` or `polrootspadic` will be much faster.

We shall now present two simple applications. The first one is finding nontrivial factors of N , given some partial information on the factors; in that case B must obviously be smaller than the largest nontrivial divisor of N .

```
setrand(1); \\ to make the example reproducible
[a,b] = [10^30, 10^31]; D = 20;
p = randomprime([a,b]);
q = randomprime([a,b]); N = p*q;
\\ assume we know 0) p | N; 1) p in [a,b]; 2) the last D digits of p
p0 = p % 10^D;
? L = zncoppersmith(10^D*x + p0, N, b \ 10^D, a)
time = 1ms.
```

```
%6 = [738281386540]
? gcd(L[1] * 10^D + p0, N) == p
%7 = 1
```

and we recovered p , faster than by trying all possibilities $x < 10^{11}$.

The second application is an attack on RSA with low exponent, when the message x is short and the padding P is known to the attacker. We use the same RSA modulus N as in the first example:

```
setrand(1);
P = random(N);    \\ known padding
e = 3;            \\ small public encryption exponent
X = floor(N^0.3); \\ N^(1/e - epsilon)
x0 = random(X);   \\ unknown short message
C = lift( (Mod(x0,N) + P)^e ); \\ known ciphertext, with padding P
zncoppersmith((P + x)^3 - C, N, X)

\\ result in 244ms.
%14 = [2679982004001230401]
? %[1] == x0
%15 = 1
```

We guessed an integer of the order of 10^{18} , almost instantly.

The library syntax is `GEN zncoppersmith(GEN P, GEN N, GEN X, GEN B = NULL)`.

3.8.126 znlog($x, g, \{o\}$). This function allows two distinct modes of operation depending on g :

- if g is the output of `znstar` (with initialization), we compute the discrete logarithm of x with respect to the generators contained in the structure. See `ideallog` for details.
- else g is an explicit element in $(\mathbf{Z}/N\mathbf{Z})^*$, we compute the discrete logarithm of x in $(\mathbf{Z}/N\mathbf{Z})^*$ in base g . The rest of this entry describes the latter possibility.

The result is `[]` when x is not a power of g , though the function may also enter an infinite loop in this case.

If present, o represents the multiplicative order of g , see Section 3.8.2; the preferred format for this parameter is `[ord, factor(ord)]`, where `ord` is the order of g . This provides a definite speedup when the discrete log problem is simple:

```
? p = nextprime(10^4); g = znprimroot(p); o = [p-1, factor(p-1)];
? for(i=1,10^4, znlog(i, g, o))
time = 163 ms.
? for(i=1,10^4, znlog(i, g))
time = 200 ms. \\ a little slower
```

The result is undefined if g is not invertible mod N or if the supplied order is incorrect.

This function uses

- a combination of generic discrete log algorithms (see below).
- in $(\mathbf{Z}/N\mathbf{Z})^*$ when N is prime: a quadratic sieve index calculus method, suitable for $N < 10^{60}$, say, is used for large prime divisors of the order.

The generic discrete log algorithms are:

- Pohlig-Hellman algorithm, to reduce to groups of prime order q , where $q|p-1$ and p is an odd prime divisor of N ,
- Shanks baby-step/giant-step ($q < 2^{32}$ is small),
- Pollard rho method ($q > 2^{32}$).

The latter two algorithms require $O(\sqrt{q})$ operations in the group on average, hence will not be able to treat cases where $q > 10^{30}$, say. In addition, Pollard rho is not able to handle the case where there are no solutions: it will enter an infinite loop.

```
? g = znprimroot(101)
%1 = Mod(2,101)
? znlog(5, g)
%2 = 24
? g^24
%3 = Mod(5, 101)

? G = znprimroot(2 * 101^10)
%4 = Mod(110462212541120451003, 220924425082240902002)
? znlog(5, G)
%5 = 76210072736547066624
? G^% == 5
%6 = 1
? N = 2^4*3^2*5^3*7^4*11; g = Mod(13, N); znlog(g^110, g)
%7 = 110
? znlog(6, Mod(2,3))  \\ no solution
%8 = []
```

For convenience, g is also allowed to be a p -adic number:

```
? g = 3+0(5^10); znlog(2, g)
%1 = 1015243
? g^%
%2 = 2 + 0(5^10)
```

The library syntax is `GEN znlog0(GEN x, GEN g, GEN o = NULL)`. The function `GEN znlog(GEN x, GEN g, GEN o)` is also available

3.8.127 znorder($x, \{o\}$). x must be an integer mod n , and the result is the order of x in the multiplicative group $(\mathbf{Z}/n\mathbf{Z})^*$. Returns an error if x is not invertible. The parameter o , if present, represents a nonzero multiple of the order of x , see Section 3.8.2; the preferred format for this parameter is `[ord, factor(ord)]`, where `ord = eulerphi(n)` is the cardinality of the group.

The library syntax is `GEN znorder(GEN x, GEN o = NULL)`.

3.8.128 znprimroot(n). Returns a primitive root (generator) of $(\mathbf{Z}/n\mathbf{Z})^*$, whenever this latter group is cyclic ($n = 4$ or $n = 2p^k$ or $n = p^k$, where p is an odd prime and $k \geq 0$). If the group is not cyclic, the function will raise an exception. If n is a prime power, then the smallest positive primitive root is returned. This may not be true for $n = 2p^k$, p odd.

Note that this function requires factoring $p - 1$ for p as above, in order to determine the exact order of elements in $(\mathbf{Z}/n\mathbf{Z})^*$: this is likely to be costly if p is large.

The library syntax is `GEN znprimroot(GEN n)`.

3.8.129 znstar($n, \{flag = 0\}$). Gives the structure of the multiplicative group $(\mathbf{Z}/n\mathbf{Z})^*$. The output G depends on the value of $flag$:

- $flag = 0$ (default), an abelian group structure $[h, d, g]$, where $h = \phi(n)$ is the order (`G.no`), d (`G.cyc`) is a k -component row-vector d of integers d_i such that $d_i > 1$, $d_i \mid d_{i-1}$ for $i \geq 2$ and

$$(\mathbf{Z}/n\mathbf{Z})^* \simeq \prod_{i=1}^k (\mathbf{Z}/d_i\mathbf{Z}),$$

and g (`G.gen`) is a k -component row vector giving generators of the image of the cyclic groups $\mathbf{Z}/d_i\mathbf{Z}$.

- $flag = 1$ the result is a `bid` structure; this allows computing discrete logarithms using `znlog` (also in the noncyclic case!).

```
? G = znstar(40)
%1 = [16, [4, 2, 2], [Mod(17, 40), Mod(21, 40), Mod(11, 40)]]
? G.no    \ \ eulerphi(40)
%2 = 16
? G.cyc   \ \ cycle structure
%3 = [4, 2, 2]
? G.gen   \ \ generators for the cyclic components
%4 = [Mod(17, 40), Mod(21, 40), Mod(11, 40)]
? apply(znorder, G.gen)
%5 = [4, 2, 2]
```

For user convenience, we define `znstar(0)` as `[2, [2], [-1]]`, corresponding to $\mathbf{Z}^*$, but $flag = 1$ is not implemented in this trivial case.

The library syntax is `GEN znstar0(GEN n, long flag)`.

3.8.130 znsubgroupgenerators($H, \{flag = 0\}$). Finds a minimal set of generators for the subgroup of $(\mathbf{Z}/f\mathbf{Z})^*$ given by a vector (or vectorsmall) H of length f : for $1 \leq a \leq f$, $H[a]$ is 1 or 0 according as $a \in H_F$ or $a \notin H_F$. In most PARI functions, subgroups of an abelian group are given as HNF left-divisors of a diagonal matrix, representing the discrete logarithms of the subgroup generators in terms of a fixed generators for the group cyclic components. The present function allows to convert an enumeration of the subgroup elements to this representation as follows:

```
? G = znstar(f, 1);
? v = znsubgroupgenerators(H);
? subHNF(G, v) = mathnfmodid(Mat([znlog(h, G) | h<-v]), G.cyc);
```

The function `subHNF` can be applied to any elements of $(\mathbf{Z}/f\mathbf{Z})^*$, yielding the subgroup they generate, but using `znsubgroupgenerators` first allows to reduce the number of discrete logarithms to be computed.

For example, if $H = \{1, 4, 11, 14\} \subset (\mathbf{Z}/15\mathbf{Z})^\times$, then we have

```
? f = 15; H = vector(f); H[1]=H[4]=H[11]=H[14] = 1;
? v = znsubgroupgenerators(H)
%2 = [4, 11]
? G = znstar(f, 1); G.cyc
%3 = [4, 2]
? subHNF(G, v)
%4 =
[2 0]
[0 1]
? subHNF(G, [1,4,11,14])
%5 =
[2 0]
[0 1]
```

This function is mostly useful when f is large and H has small index: if H has few elements, one may just use `subHNF` directly on the elements of H . For instance, let $K = \mathbf{Q}(\zeta_p, \sqrt{m}) \subset L = \mathbf{Q}(\zeta_f)$, where p is a prime, $\sqrt{m}$ is a quadratic number and f is the conductor of the abelian extension $K/\mathbf{Q}$. The following GP script creates H as the Galois group of L/K , as a subgroup of $(\mathbf{Z}/f\mathbf{Z})^*$:

```
HK(m, p, flag = 0)=
{ my(d = quaddisc(m), f = lcm(d, p), H);
  H = vectorsmall(f, a, a % p == 1 && kronecker(d,a) > 0);
  [f, znsubgroupgenerators(H,flag)];
}
? [f, v] = HK(36322, 5)
time = 193 ms.
%1 = [726440, [41, 61, 111, 131]]
? G = znstar(f,1); G.cyc
%2 = [1260, 12, 2, 2, 2, 2]
? A = subHNF(G, v)
%3 =
[2 0 1 1 0 1]
[0 4 0 0 0 2]
[0 0 1 0 0 0]
[0 0 0 1 0 0]
[0 0 0 0 1 0]
[0 0 0 0 0 1]
\\ Double check
? p = 5; d = quaddisc(36322);
? w = select(a->a % p == 1 && kronecker(d,a) > 0, [1..f]); #w
time = 133 ms.
%5 = 30240 \\ w enumerates the elements of H
```

```
? subHNF(G, w) == A \\ same result, about twice slower
time = 242 ms.
%6 = 1
```

This shows that $K = \mathbf{Q}(\sqrt{36322}, \zeta_5)$ is contained in $\mathbf{Q}(\zeta_{726440})$ and $H = \langle 41, 61, 111, 131 \rangle$. Note that $H = \langle 41 \rangle \langle 61 \rangle \langle 111 \rangle \langle 131 \rangle$ is not a direct product. If $flag = 1$, then the function finds generators which decompose H to direct factors:

```
? HK(36322, 5, 1)
%3 = [726440, [41, 31261, 324611, 506221]]
```

This time $H = \langle 41 \rangle \times \langle 31261 \rangle \times \langle 324611 \rangle \times \langle 506221 \rangle$.

The library syntax is `GEN znsubgroupgenerators(GEN H, long flag)`.

3.9 Polynomials and power series.

We group here all functions which are specific to polynomials or power series. Many other functions which can be applied on these objects are described in the other sections. Also, some of the functions described here can be applied to other types.

3.9.1 $O(p^e)$. If p is an integer greater than 2, returns a p -adic 0 of precision e . In all other cases, returns a power series zero with precision given by ev , where v is the X -adic valuation of p with respect to its main variable.

The library syntax is `GEN ggrando()`. `GEN zeropadic(GEN p, long e)` for a p -adic and `GEN zerozer(long v, long e)` for a power series zero in variable v .

3.9.2 $\text{bezoutres}(A, B, \{v\})$. Deprecated alias for `polresultantext`

The library syntax is `GEN polresultantext0(GEN A, GEN B, long v = -1)` where v is a variable number.

3.9.3 $\text{deriv}(x, \{v\})$. Derivative of x with respect to the main variable if v is omitted, and with respect to v otherwise. The derivative of a scalar type is zero, and the derivative of a vector or matrix is done componentwise. One can use x' as a shortcut if the derivative is with respect to the main variable of x ; and also use x'' , etc., for multiple derivatives although `derivn` is often preferable.

By definition, the main variable of a `t_POLMOD` is the main variable among the coefficients from its two polynomial components (representative and modulus); in other words, assuming a `polmod` represents an element of $R[X]/(T(X))$, the variable X is a mute variable and the derivative is taken with respect to the main variable used in the base ring R .

```
? f = (x/y)^5;
? deriv(f)
%2 = 5/y^5*x^4
? f'
%3 = 5/y^5*x^4
? deriv(f, 'x) \\ same since 'x is the main variable
%4 = 5/y^5*x^4
? deriv(f, 'y)
%5 = -5/y^6*x^5
```

This function also operates on closures, in which case the variable must be omitted. It returns a closure performing a numerical differentiation as per `derivnum`:

```
? f(x) = x^2;  
? g = deriv(f)  
? g(1)  
%3 = 2.000000000000000000000000000000000000  
? f(x) = sin(exp(x));  
? deriv(f)(0)  
%5 = 0.54030230586813971740093660744297660373  
? cos(1)  
%6 = 0.54030230586813971740093660744297660373
```

The library syntax is `GEN deriv(GEN x, long v = -1)` where `v` is a variable number.

3.9.4 derivn($x, n, \{v\}$). n -th derivative of x with respect to the main variable if v is omitted, and with respect to v otherwise; the integer n must be nonnegative. The derivative of a scalar type is zero, and the derivative of a vector or matrix is done componentwise. One can use x' , x'' , etc., as a shortcut if the derivative is with respect to the main variable of x .

By definition, the main variable of a `t_POLMOD` is the main variable among the coefficients from its two polynomial components (representative and modulus); in other words, assuming a `polmod` represents an element of $R[X]/(T(X))$, the variable X is a mute variable and the derivative is taken with respect to the main variable used in the base ring R .

```
? f = (x/y)^5;
? derivn(f, 2)
%2 = 20/y^5*x^3
? f''
%3 = 20/y^5*x^3
? derivn(f, 2, 'x') \\ same since 'x' is the main variable
%4 = 20/y^5*x^3
? derivn(f, 2, 'y')
%5 = 30/y^7*x^5
```

This function also operates on closures, in which case the variable must be omitted. It returns a closure performing a numerical differentiation as per `derivnum`:

```
? f(x) = x^10;  
? g = derivn(f, 5)  
? g(1)  
%3 = 30240.0000000000000000000000000000000000  
  
? derivn(zeta, 2)(0)  
%4 = -2.0063564559085848512101000267299604382  
?  
? zeta''(0)  
%5 = -2.0063564559085848512101000267299604382
```

The library syntax is `GEN derivn(GEN x, long n, long v = -1)` where `v` is a variable number.

3.9.5 diffop($x, v, d, \{n = 1\}$). Let v be a vector of variables, and d a vector of the same length, return the image of x by the n -power (1 if n is not given) of the differential operator D that assumes the value $d[i]$ on the variable $v[i]$. The value of D on a scalar type is zero, and D applies componentwise to a vector or matrix. When applied to a `t_POLMOD`, if no value is provided for the variable of the modulus, such value is derived using the implicit function theorem.

Examples. This function can be used to differentiate formal expressions: if $E = \exp(X^2)$ then we have $E' = 2 * X * E$. We derivate $X * \exp(X^2)$ as follows:

```
? diffop(E*X, [X,E], [1,2*X*E])
%1 = (2*X^2 + 1)*E
```

Let Sin and Cos be two function such that $\text{Sin}^2 + \text{Cos}^2 = 1$ and $\text{Cos}' = -\text{Sin}$. We can differentiate Sin/Cos as follows, PARI inferring the value of Sin' from the equation:

```
? diffop(Mod('Sin/'Cos, 'Sin^2+'Cos^2-1), ['Cos], [-'Sin])
%1 = Mod(1/Cos^2, Sin^2 + (Cos^2 - 1))
```

Compute the Bell polynomials (both complete and partial) via the Faa di Bruno formula:

```
Bell(k,n=-1)=
{ my(x, v, dv, var = i->eval(Str("X",i)));
  v = vector(k, i, if (i==1, 'E, var(i-1)));
  dv = vector(k, i, if (i==1, 'X*var(1)*'E, var(i)));
  x = diffop('E,v,dv,k) / 'E;
  if (n < 0, subst(x,'X,1), polcoef(x,n,'X));
}
```

The library syntax is `GEN diffop0(GEN x, GEN v, GEN d, long n)`.

For $n = 1$, the function `GEN diffop(GEN x, GEN v, GEN d)` is also available.

3.9.6 eval(x). Replaces in x the formal variables by the values that have been assigned to them after the creation of x . This is mainly useful in GP, and not in library mode. Do not confuse this with substitution (see `subst`).

If x is a character string, `eval( $x$ )` executes x as a GP command, as if directly input from the keyboard, and returns its output.

```
? x1 = "one"; x2 = "two";
? n = 1; eval(Str("x", n))
%2 = "one"
? f = "exp"; v = 1;
? eval(Str(f, "(", v, ")"))
%4 = 2.7182818284590452353602874713526624978
```

Note that the first construct could be implemented in a simpler way by using a vector $\mathbf{x} = ["one", "two"]$; $\mathbf{x}[n]$, and the second by using a closure $\mathbf{f} = \mathbf{exp}$; $\mathbf{f}(v)$. The final example is more interesting:

```
? genmat(u,v) = matrix(u,v,i,j, eval( Str("x",i,j) ));
? genmat(2,3)  \\ generic 2 x 3 matrix
%2 =
[x11 x12 x13]
```

```
[x21 x22 x23]
```

A syntax error in the evaluation expression raises an `e_SYNTAX` exception, which can be trapped as usual:

```
? 1a
***   syntax error, unexpected variable name, expecting $end or `;`: 1a
***                                     ^_
? E(expr) =
{
    iferr(eval(expr),
          e, print("syntax error"),
          errname(e) == "e_SYNTAX");
}
? E("1+1")
%1 = 2
? E("1a")
syntax error
```

The library syntax is `geval(GEN x)`.

3.9.7 factorpadic(*pol*, *p*, *r*). *p*-adic factorization of the polynomial *pol* to precision *r*, the result being a two-column matrix as in `factor`. Note that this is not the same as a factorization over $\mathbf{Z}/p^r\mathbf{Z}$ (polynomials over that ring do not form a unique factorization domain, anyway), but approximations in $\mathbf{Q}/p^r\mathbf{Z}$ of the true factorization in $\mathbf{Q}_p[X]$.

```
? factorpadic(x^2 + 9, 3,5)
%1 =
[(1 + 0(3^5))*x^2 + 0(3^5)*x + (3^2 + 0(3^5)) 1]
? factorpadic(x^2 + 1, 5,3)
%2 =
[ (1 + 0(5^3))*x + (2 + 5 + 2*5^2 + 0(5^3)) 1]
[(1 + 0(5^3))*x + (3 + 3*5 + 2*5^2 + 0(5^3)) 1]
```

The factors are normalized so that their leading coefficient is a power of *p*. The method used is a modified version of the round 4 algorithm of Zassenhaus.

If *pol* has inexact `t_PADIC` coefficients, this is not always well-defined; in this case, the polynomial is first made integral by dividing out the *p*-adic content, then lifted to $\mathbf{Z}$ using `truncate` coefficientwise. Hence we actually factor exactly a polynomial which is only *p*-adically close to the input. To avoid pitfalls, we advise to only factor polynomials with exact rational coefficients.

The library syntax is `factorpadic(GEN f, GEN p, long r)`. The function `factorpadic0` is deprecated, provided for backward compatibility.

3.9.8 `fft(w, P)`. Let $w = [1, z, \dots, z^{N-1}]$ from some primitive N -roots of unity z where N is a power of 2, and P be a polynomial $< N$, return the unnormalized discrete Fourier transform of P , $\{P(w[i]), 1 \leq i \leq N\}$. Also allow P to be a vector $[p_0, \dots, p_n]$ representing the polynomial $\sum_i p_i X^i$. Composing `fft` and `fftinv` returns N times the original input coefficients.

```
? w = rootsof1(4); fft(w, x^3+x+1)
%1 = [3, 1, -1, 1]
? fftinv(w, %)
%2 = [4, 4, 0, 4]
? Polrev(%) / 4
%3 = x^3 + x + 1
? w = powers(znprimroot(5),3); fft(w, x^3+x+1)
%4 = [Mod(3,5),Mod(1,5),Mod(4,5),Mod(1,5)]
? fftinv(w, %)
%5 = [Mod(4,5),Mod(4,5),Mod(0,5),Mod(4,5)]
```

The library syntax is `GEN FFT(GEN w, GEN P)`.

3.9.9 `fftinv(w, P)`. Let $w = [1, z, \dots, z^{N-1}]$ from some primitive N -roots of unity z where N is a power of 2, and P be a polynomial $< N$, return the unnormalized discrete Fourier transform of P , $\{P(1/w[i]), 1 \leq i \leq N\}$. Also allow P to be a vector $[p_0, \dots, p_n]$ representing the polynomial $\sum_i p_i X^i$. Composing `fft` and `fftinv` returns N times the original input coefficients.

```
? w = rootsof1(4); fft(w, x^3+x+1)
%1 = [3, 1, -1, 1]
? fftinv(w, %)
%2 = [4, 4, 0, 4]
? Polrev(%) / 4
%3 = x^3 + x + 1

? N = 512; w = rootsof1(N); T = random(1000 * x^(N-1));
? U = fft(w, T);
time = 3 ms.
? V = vector(N, i, subst(T, 'x, w[i]));
time = 65 ms.
? exponent(V - U)
%7 = -97
? round(Polrev(fftinv(w,U) / N)) == T
%8 = 1
```

The library syntax is `GEN FFTinv(GEN w, GEN P)`.

3.9.10 `intformal(x, {v})`. formal integration of x with respect to the variable v (wrt. the main variable if v is omitted). Since PARI cannot represent logarithmic or arctangent terms, any such term in the result will yield an error:

```
? intformal(x^2)
%1 = 1/3*x^3
? intformal(x^2, y)
%2 = y*x^2
? intformal(1/x)
*** at top-level: intformal(1/x)
*** ^-----
*** intformal: domain error in intformal: residue(series, pole) != 0
```

The argument x can be of any type. When x is a rational function, we assume that the base ring is an integral domain of characteristic zero.

By definition, the main variable of a `t_POLMOD` is the main variable among the coefficients from its two polynomial components (representative and modulus); in other words, assuming a `polmod` represents an element of $R[X]/(T(X))$, the variable X is a main variable and the integral is taken with respect to the main variable used in the base ring R . In particular, it is meaningless to integrate with respect to the main variable of `x.mod`:

```
? intformal(Mod(1,x^2+1), 'x)
*** intformal: incorrect priority in intformal: variable x = x
```

The library syntax is `GEN integ(GEN x, long v = -1)` where v is a variable number.

3.9.11 `padicappr(pol, a)`. Vector of p -adic roots of the polynomial pol congruent to the p -adic number a modulo p , and with the same p -adic precision as a . The number a can be an ordinary p -adic number (type `t_PADIC`, i.e. an element of $\mathbf{Z}_p$) or can be an integral element of a finite *unramified* extension $\mathbf{Q}_p[X]/(T)$ of $\mathbf{Q}_p$, given as a `t_POLMOD` `Mod(A, T)` at least one of whose coefficients is a `t_PADIC` and T irreducible modulo p . In this case, the result is the vector of roots belonging to the same extension of $\mathbf{Q}_p$ as a . The polynomial pol should have exact coefficients; if not, its coefficients are first rounded to $\mathbf{Q}$ or $\mathbf{Q}[X]/(T)$ and this is the polynomial whose roots we consider.

The library syntax is `GEN padicappr(GEN pol, GEN a)`. Also available is `GEN Zp_appr(GEN f, GEN a)` when a is a `t_PADIC`.

3.9.12 `padicfields(p, N, {flag = 0})`. Returns a vector of polynomials generating all the extensions of degree N of the field $\mathbf{Q}_p$ of p -adic rational numbers; N is allowed to be a 2-component vector $[n, d]$, in which case we return the extensions of degree n and discriminant p^d .

The list is minimal in the sense that two different polynomials generate nonisomorphic extensions; in particular, the number of polynomials is the number of classes of nonisomorphic extensions. If P is a polynomial in this list, α is any root of P and $K = \mathbf{Q}_p(\alpha)$, then α is the sum of a uniformizer and a (lift of a) generator of the residue field of K ; in particular, the powers of α generate the ring of p -adic integers of K .

If $flag = 1$, replace each polynomial P by a vector $[P, e, f, d, c]$ where e is the ramification index, f the residual degree, d the valuation of the discriminant, and c the number of conjugate fields. If $flag = 2$, only return the *number* of extensions in a fixed algebraic closure (Krasner's formula), which is much faster.

The library syntax is `GEN padicfields0(GEN p, GEN N, long flag)`. Also available is `GEN padicfields(GEN p, long n, long d, long flag)`, which computes extensions of $\mathbf{Q}_p$ of degree n and discriminant p^d .

3.9.13 polchebyshev($n, \{flag = 1\}, \{a = 'x\}$). Returns the n^{th} Chebyshev polynomial of the first kind T_n ($flag = 1$) or the second kind U_n ($flag = 2$), evaluated at a ('x' by default). Both series of polynomials satisfy the 3-term relation

$$P_{n+1} = 2xP_n - P_{n-1},$$

and are determined by the initial conditions $U_0 = T_0 = 1$, $T_1 = x$, $U_1 = 2x$. In fact $T'_n = nU_{n-1}$ and, for all complex numbers z , we have $T_n(\cos z) = \cos(nz)$ and $U_{n-1}(\cos z) = \sin(nz)/\sin z$. If $n \geq 0$, then these polynomials have degree n . For $n < 0$, T_n is equal to T_{-n} and U_n is equal to $-U_{-2-n}$. In particular, $U_{-1} = 0$.

The library syntax is `GEN polchebyshev_eval(long n, long flag, GEN a = NULL)`. Also available are `GEN polchebyshev(long n, long flag, long v)`, `GEN polchebyshev1(long n, long v)` and `GEN polchebyshev2(long n, long v)` for T_n and U_n respectively.

3.9.14 polclass($D, \{inv = 0\}, \{x = 'x\}$). Return a polynomial in $\mathbf{Z}[x]$ generating the Hilbert class field for the imaginary quadratic discriminant D . If inv is 0 (the default), use the modular j -function and return the classical Hilbert polynomial, otherwise use a class invariant. The following invariants correspond to the different values of inv , where f denotes Weber's function `weber`, and $w_{p,q}$ the double eta quotient given by $w_{p,q} = \frac{\eta(x/p)\eta(x/q)}{\eta(x)\eta(x/pq)}$

The invariants $w_{p,q}$ are not allowed unless they satisfy the following technical conditions ensuring they do generate the Hilbert class field and not a strict subfield:

- if $p \neq q$, we need them both noninert, prime to the conductor of $\mathbf{Z}[\sqrt{D}]$. Let P, Q be prime ideals above p and q ; if both are unramified, we further require that $P^{\pm 1}Q^{\pm 1}$ be all distinct in the class group of $\mathbf{Z}[\sqrt{D}]$; if both are ramified, we require that $PQ \neq 1$ in the class group.

- if $p = q$, we want it split and prime to the conductor and the prime ideal above it must have order $\neq 1, 2, 4$ in the class group.

Invariants are allowed under the additional conditions on D listed below.

- 0 : j
- 1 : f , $D = 1 \bmod 8$ and $D = 1, 2 \bmod 3$;
- 2 : f^2 , $D = 1 \bmod 8$ and $D = 1, 2 \bmod 3$;
- 3 : f^3 , $D = 1 \bmod 8$;
- 4 : f^4 , $D = 1 \bmod 8$ and $D = 1, 2 \bmod 3$;
- 5 : $\gamma_2 = j^{1/3}$, $D = 1, 2 \bmod 3$;
- 6 : $w_{2,3}$, $D = 1 \bmod 8$ and $D = 1, 2 \bmod 3$;
- 8 : f^8 , $D = 1 \bmod 8$ and $D = 1, 2 \bmod 3$;
- 9 : $w_{3,3}$, $D = 1 \bmod 2$ and $D = 1, 2 \bmod 3$;
- 10: $w_{2,5}$, $D \neq 60 \bmod 80$ and $D = 1, 2 \bmod 3$;

- 14: $w_{2,7}$, $D = 1 \bmod 8$;
- 15: $w_{3,5}$, $D = 1, 2 \bmod 3$;
- 21: $w_{3,7}$, $D = 1 \bmod 2$ and 21 does not divide D
- 23: $w_{2,3}^2$, $D = 1, 2 \bmod 3$;
- 24: $w_{2,5}^2$, $D = 1, 2 \bmod 3$;
- 26: $w_{2,13}$, $D \neq 156 \bmod 208$;
- 27: $w_{2,7}^2$, $D \neq 28 \bmod 112$;
- 28: $w_{3,3}^2$, $D = 1, 2 \bmod 3$;
- 35: $w_{5,7}$, $D = 1, 2 \bmod 3$;
- 39: $w_{3,13}$, $D = 1 \bmod 2$ and $D = 1, 2 \bmod 3$;

The algorithm for computing the polynomial does not use the floating point approach, which would evaluate a precise modular function in a precise complex argument. Instead, it relies on a faster Chinese remainder based approach modulo small primes, in which the class invariant is only defined algebraically by the modular polynomial relating the modular function to j . So in fact, any of the several roots of the modular polynomial may actually be the class invariant, and more precise assertions cannot be made.

For instance, while `polclass(D)` returns the minimal polynomial of $j(\tau)$ with τ (any) quadratic integer for the discriminant D , the polynomial returned by `polclass(D, 5)` can be the minimal polynomial of any of $\gamma_2(\tau)$, $\zeta_3\gamma_2(\tau)$ or $\zeta_3^2\gamma_2(\tau)$, the three roots of the modular polynomial $j = \gamma_2^3$, in which j has been specialised to $j(\tau)$.

The modular polynomial is given by $j = \frac{(f^{24}-16)^3}{f^{24}}$ for Weber's function f .

For the double eta quotients of level $N = pq$, all functions are covered such that the modular curve $X_0^+(N)$, the function field of which is generated by the functions invariant under $\Gamma^0(N)$ and the Fricke–Atkin–Lehner involution, is of genus 0 with function field generated by (a power of) the double eta quotient w . This ensures that the full Hilbert class field (and not a proper subfield) is generated by class invariants from these double eta quotients. Then the modular polynomial is of degree 2 in j , and of degree $\psi(N) = (p+1)(q+1)$ in w .

```
? polclass(-163)
%1 = x + 262537412640768000
? polclass(-51, , 'z)
%2 = z^2 + 5541101568*z + 6262062317568
? polclass(-151,1)
x^7 - x^6 + x^5 + 3*x^3 - x^2 + 3*x + 1
```

The library syntax is `GEN polclass(GEN D, long inv, long x = -1)` where `x` is a variable number.

3.9.15 polcoef($x, n, \{v\}$). Coefficient of degree n of the polynomial x , with respect to the main variable if v is omitted, with respect to v otherwise. If n is greater than the degree, the result is zero.

Naturally applies to scalars (polynomial of degree 0), as well as to rational functions whose denominator is a monomial. It also applies to power series: if n is less than the valuation, the result is zero. If it is greater than the largest significant degree, then an error message is issued.

The library syntax is `GEN polcoef(GEN x, long n, long v = -1)` where v is a variable number.

3.9.16 polcoeff($x, n, \{v\}$). Deprecated alias for `polcoef`.

The library syntax is `GEN polcoeff(GEN x, long n, long v = -1)` where v is a variable number.

3.9.17 polcyclo($n, \{a = 'x\}$). n -th cyclotomic polynomial, evaluated at a (' x ' by default). The integer n must be positive.

Algorithm used: reduce to the case where n is squarefree; to compute the cyclotomic polynomial, use $\Phi_{np}(x) = \Phi_n(x^p)/\Phi(x)$; to compute it evaluated, use $\Phi_n(x) = \prod_{d|n} (x^d - 1)^{\mu(n/d)}$. In the evaluated case, the algorithm assumes that $a^d - 1$ is either 0 or invertible, for all $d \mid n$. If this is not the case (the base ring has zero divisors), use `subst(polcyclo(n), x, a)`.

The library syntax is `GEN polcyclo_eval(long n, GEN a = NULL)`. The variant `GEN polcyclo(long n, long v)` returns the n -th cyclotomic polynomial in variable v .

3.9.18 polcyclofactors(f). Returns a vector of polynomials, whose product is the product of distinct cyclotomic polynomials dividing f .

```
? f = x^10+5*x^8-x^7+8*x^6-4*x^5+8*x^4-3*x^3+7*x^2+3;
? v = polcyclofactors(f)
%2 = [x^2 + 1, x^2 + x + 1, x^4 - x^3 + x^2 - x + 1]
? apply(poliscycloprod, v)
%3 = [1, 1, 1]
? apply(poliscyclo, v)
%4 = [4, 3, 10]
```

In general, the polynomials are products of cyclotomic polynomials and not themselves irreducible:

```
? g = x^8+2*x^7+6*x^6+9*x^5+12*x^4+11*x^3+10*x^2+6*x+3;
? polcyclofactors(g)
%2 = [x^6 + 2*x^5 + 3*x^4 + 3*x^3 + 3*x^2 + 2*x + 1]
? factor(%[1])
%3 =
[          x^2 + x + 1 1]
[x^4 + x^3 + x^2 + x + 1 1]
```

The library syntax is `GEN polcyclofactors(GEN f)`.

3.9.19 poldegree($x, \{v\}$). Degree of the polynomial x in the main variable if v is omitted, in the variable v otherwise.

The degree of 0 is $-\infty$. The degree of a nonzero scalar is 0. Finally, when x is a nonzero polynomial or rational function, returns the ordinary degree of x . Raise an error otherwise.

The library syntax is GEN gppoldegree(GEN x, long v = -1) where v is a variable number. Also available is long poldegree(GEN x, long v), which returns $-\text{LONG_MAX}$ if $x = 0$ and the degree as a long integer.

3.9.20 poldisc($pol, \{v\}$). Discriminant of the polynomial pol in the main variable if v is omitted, in v otherwise. Uses a modular algorithm over $\mathbf{Z}$ or $\mathbf{Q}$, and the subresultant algorithm otherwise.

```
? T = x^4 + 2*x+1;
? poldisc(T)
%2 = -176
? poldisc(T^2)
%3 = 0
```

For convenience, the function also applies to types $\mathbf{t_QUAD}$ and $\mathbf{t_QFB}$:

```
? z = 3*quadgen(8) + 4;
? poldisc(z)
%2 = 8
? q = Qfb(1,2,3);
? poldisc(q)
%4 = -8
```

The library syntax is GEN poldisc0(GEN pol, long v = -1) where v is a variable number.

3.9.21 poldiscfactors($T, \{flag = 0\}$). Given a polynomial T with integer coefficients, return $[D, faD]$ where D is the discriminant of T and faD is a cheap partial factorization of $|D|$: entries in its first column are coprime and not perfect powers but need not be primes. The factors are obtained by a combination of trial division, testing for perfect powers, factorizations in coprimes, and computing Euclidean remainder sequences for (T, T') modulo composite factors d of D (which is likely to produce 0-divisors in $\mathbf{Z}/d\mathbf{Z}$). If $flag$ is 1, finish the factorization using **factorint**.

```
? T = x^3 - 6021021*x^2 + 12072210077769*x - 8092423140177664432;
? [D,faD] = poldiscfactors(T); print(faD); D
[3, 3; 7, 2; 373, 2; 500009, 2; 24639061, 2]
%2 = -27937108625866859018515540967767467

? T = x^3 + 9*x^2 + 27*x - 125014250689643346789780229390526092263790263725;
? [D,faD] = poldiscfactors(T); print(faD)
[2, 6; 3, 3; 125007125141751093502187, 4]
? [D,faD] = poldiscfactors(T, 1); print(faD)
[2, 6; 3, 3; 500009, 12; 1000003, 4]
```

The library syntax is GEN poldiscfactors(GEN T, long flag).

3.9.22 poldiscreduced(f). Reduced discriminant vector of the (integral, monic) polynomial f . This is the vector of elementary divisors of $\mathbf{Z}[\alpha]/f'(\alpha)\mathbf{Z}[\alpha]$, where α is a root of the polynomial f . The components of the result are all positive, and their product is equal to the absolute value of the discriminant of f .

The library syntax is GEN `reduceddiscsmith(GEN f)`.

3.9.23 polfromroots($a, \{v = x\}$). Returns the monic polynomial in variable v whose roots are the components of the vector a with multiplicities, that is $\prod_i (x - a_i)$.

```
? polfromroots([1,2,3])
%1 = x^3 - 6*x^2 + 11*x - 6
? polfromroots([z, -z], 'y)
%2 = y^2 - z^2
```

The library syntax is GEN `polfromroots(GEN a, long v = -1)` where v is a variable number.

3.9.24 polgraeffe(f). Returns the Graeffe transform g of f , such that $g(x^2) = f(x)f(-x)$.

The library syntax is GEN `polgraeffe(GEN f)`.

3.9.25 polhensellift(A, B, p, e). Given a prime p , an integral polynomial A whose leading coefficient is a p -unit, a vector B of integral polynomials that are monic and pairwise relatively prime modulo p , and whose product is congruent to $A/\text{lc}(A)$ modulo p , lift the elements of B to polynomials whose product is congruent to A modulo p^e .

More generally, if T is an integral polynomial irreducible mod p , and B is a factorization of A over the finite field $\mathbf{F}_p[t]/(T)$, you can lift it to $\mathbf{Z}_p[t]/(T, p^e)$ by replacing the p argument with $[p, T]$:

```
? { T = t^3 - 2; p = 7; A = x^2 + t + 1;
  B = [x + (3*t^2 + t + 1), x + (4*t^2 + 6*t + 6)];
  r = polhensellift(A, B, [p, T], 6) }
%1 = [x + (20191*t^2 + 50604*t + 75783), x + (97458*t^2 + 67045*t + 41866)]
? liftall( r[1] * r[2] * Mod(Mod(1,p^6),T) )
%2 = x^2 + (t + 1)
```

The library syntax is GEN `polhensellift(GEN A, GEN B, GEN p, long e)`.

3.9.26 polhermite($n, \{a = 'x\}, \{flag = 0\}$). n^{th} Hermite polynomial H_n evaluated at a (' x ' by default), i.e.

$$H_n(x) = (-1)^n e^{x^2} \frac{d^n}{dx^n} e^{-x^2}.$$

If $flag$ is nonzero and $n > 0$, return $[H_{n-1}(a), H_n(a)]$.

```
? polhermite(5)
%1 = 32*x^5 - 160*x^3 + 120*x
? polhermite(5, -2) \\ H_5(-2)
%2 = 16
? polhermite(5, 1)
%3 = [16*x^4 - 48*x^2 + 12, 32*x^5 - 160*x^3 + 120*x]
? polhermite(5, -2, 1)
%4 = [76, 16]
```

The library syntax is `GEN polhermite_eval0(long n, GEN a = NULL, long flag)`. The variant `GEN polhermite(long n, long v)` returns the n -th Hermite polynomial in variable v . To obtain $H_n(a)$, use `GEN polhermite_eval(long n, GEN a)`.

3.9.27 polinterpolate($X, \{Y\}, \{t = 'x\}, \{\&e\}$). Given the data vectors X and Y of the same length n (X containing the x -coordinates, and Y the corresponding y -coordinates), this function finds the interpolating polynomial P of minimal degree passing through these points and evaluates it at t . If Y is omitted, the polynomial P interpolates the $(i, X[i])$.

```
? v = [1, 2, 4, 8, 11, 13];
? P = polinterpolate(v) \\ formal interpolation
%1 = 7/120*x^5 - 25/24*x^4 + 163/24*x^3 - 467/24*x^2 + 513/20*x - 11
? [ subst(P,'x,a) | a <- [1..6] ]
%2 = [1, 2, 4, 8, 11, 13]
? polinterpolate(v,, 10) \\ evaluate at 10
%3 = 508
? subst(P, x, 10)
%4 = 508

? P = polinterpolate([1,2,4], [9,8,7])
%5 = 1/6*x^2 - 3/2*x + 31/3
? [subst(P, 'x, a) | a <- [1,2,4]]
%6 = [9, 8, 7]
? P = polinterpolate([1,2,4], [9,8,7], 0)
%7 = 31/3
```

If the goal is to extrapolate a function at a unique point, it is more efficient to use the t argument rather than interpolate formally then evaluate:

```
? x0 = 1.5;
? v = vector(20, i, random([-10,10]));
? for(i=1,10^3, subst(polinterpolate(v),'x, x0))
time = 352 ms.
? for(i=1,10^3, polinterpolate(v,,x0))
time = 111 ms.

? v = vector(40, i, random([-10,10]));
? for(i=1,10^3, subst(polinterpolate(v), 'x, x0))
time = 3,035 ms.
? for(i=1,10^3, polinterpolate(v,, x0))
time = 436 ms.
```

The threshold depends on the base field. Over small prime finite fields, interpolating formally first is more efficient

```
? bench(p, N, T = 10^3) =
{ my (v = vector(N, i, random(Mod(0,p))));
  my (x0 = Mod(3, p), t1, t2);
  gettime();
  for(i=1, T, subst(polinterpolate(v), 'x, x0));
  t1 = gettime();
  for(i=1, T, polinterpolate(v,, x0));
  t2 = gettime(); [t1, t2];
```

```

    }
? p = 101;
? bench(p, 4, 10^4) \\ both methods are equivalent
%3 = [39, 40]
? bench(p, 40) \\ with 40 points formal is much faster
%4 = [45, 355]

```

As the cardinality increases, formal interpolation requires more points to become interesting:

```

? p = nextprime(2^128);
? bench(p, 4) \\ formal is slower
%3 = [16, 9]
? bench(p, 10) \\ formal has become faster
%4 = [61, 70]
? bench(p, 100) \\ formal is much faster
%5 = [1682, 9081]

? p = nextprime(10^500);
? bench(p, 4) \\ formal is slower
%7 = [72, 354]
? bench(p, 20) \\ formal is still slower
%8 = [1287, 962]
? bench(p, 40) \\ formal has become faster
%9 = [3717, 4227]
? bench(p, 100) \\ faster but relatively less impressive
%10 = [16237, 32335]

```

If t is a complex numeric value and e is present, e will contain an error estimate on the returned value. More precisely, let P be the interpolation polynomial on the given n points; there exist a subset of $n - 1$ points and Q the attached interpolation polynomial such that $e = \text{exponent}(P(t) - Q(t))$ (Neville's algorithm).

```

? f(x) = 1 / (1 + 25*x^2);
? x0 = 975/1000;
? test(X) =
  { my (P, e);
    P = polinterpolate(X, [f(x) | x <- X], x0, &e);
    [ exponent(P - f(x0)), e ];
  }
\\ equidistant nodes vs. Chebyshev nodes
? test( [-10..10] / 10 )
%4 = [6, 5]
? test( polrootsreal(polchebyshev(21)) )
%5 = [-15, -10]
? test( [-100..100] / 100 )
%7 = [93, 97] \\ P(x0) is way different from f(x0)
? test( polrootsreal(polchebyshev(201)) )
%8 = [-60, -55]

```

This is an example of Runge's phenomenon: increasing the number of equidistant nodes makes extrapolation much worse. Note that the error estimate is not a guaranteed upper bound (cf %4), but is reasonably tight in practice.

Numerical stability. The interpolation is performed in a numerically stable way using $\prod_{j \neq i} (X[i] - X[j])$ instead of $Q'(X[i])$ with $Q = \prod_i (x - X[i])$. Centering the interpolation points $X[i]$ around 0, thereby reconstructing $P(x - m)$, for a suitable m will further reduce the numerical error.

The library syntax is `GEN polint(GEN X, GEN Y = NULL, GEN t = NULL, GEN *e = NULL)`

3.9.28 polisclass(P). P being a monic irreducible polynomial with integer coefficients, return 0 if P is not a class polynomial for the j -invariant, otherwise return the discriminant $D < 0$ such that $P = \text{polclass}(D)$.

```
? polisclass(polclass(-47))
%1 = -47
? polisclass(x^5+x+1)
%2 = 0
? apply(polisclass,factor(poldisc(polmodular(5)))[,1])
%3 = [-16,-4,-3,-11,-19,-64,-36,-24,-51,-91,-99,-96,-84]~
```

The library syntax is `long polisclass(GEN P)`.

3.9.29 poliscyclo(f). Returns 0 if f is not a cyclotomic polynomial, and $n > 0$ if $f = \Phi_n$, the n -th cyclotomic polynomial.

```
? poliscyclo(x^4-x^2+1)
%1 = 12
? poliscyclo(12)
%2 = x^4 - x^2 + 1
? poliscyclo(x^4-x^2-1)
%3 = 0
```

The library syntax is `long poliscyclo(GEN f)`.

3.9.30 poliscycloprod(f). Returns 1 if f is a product of cyclotomic polynomial, and 0 otherwise.

```
? f = x^6+x^5-x^3+x+1;
? poliscycloprod(f)
%2 = 1
? factor(f)
%3 =
[ x^2 + x + 1 1]
[ x^4 - x^2 + 1 1]
? [ poliscyclo(T) | T <- %[,1] ]
%4 = [3, 12]
? polycyclo(3) * polycyclo(12)
%5 = x^6 + x^5 - x^3 + x + 1
```

The library syntax is `long poliscycloprod(GEN f)`.

3.9.31 polisirreducible(*pol*). *pol* being a polynomial (univariate in the present version 2.17.4), returns 1 if *pol* is nonconstant and irreducible, 0 otherwise. Irreducibility is checked over the smallest base field over which *pol* seems to be defined.

The library syntax is `long polisirreducible(GEN pol)`.

3.9.32 pollaguerre(*n*, {*a* = 0}, {*b* = 'x}, {*flag* = 0}). n^{th} Laguerre polynomial $L_n^{(a)}$ of degree *n* and parameter *a* evaluated at *b* ('x by default), i.e.

$$L_n^{(a)}(x) = \frac{x^{-a}e^x}{n!} \frac{d^n}{dx^n}(e^{-x}x^{n+a}).$$

If *flag* is 1, return $[L_{n-1}^{(a)}(b), L_n^{(a)}(b)]$.

The library syntax is `GEN pollaguerre_eval0(long n, GEN a = NULL, GEN b = NULL, long flag)`. To obtain the *n*-th Laguerre polynomial in variable *v*, use `GEN pollaguerre(long n, GEN a, GEN b, long v)`. To obtain $L_n^{(a)}(b)$, use `GEN pollaguerre_eval(long n, GEN a, GEN b)`.

3.9.33 pollead(*x*, {*v*}). Leading coefficient of the polynomial or power series *x*. This is computed with respect to the main variable of *x* if *v* is omitted, with respect to the variable *v* otherwise.

The library syntax is `GEN pollead(GEN x, long v = -1)` where *v* is a variable number.

3.9.34 pollegendre(*n*, {*a* = 'x}, {*flag* = 0}). n^{th} Legendre polynomial P_n evaluated at *a* ('x by default), where

$$P_n(x) = \frac{1}{2^n n!} \frac{d^n}{dx^n}(x^2 - 1)^n.$$

If *flag* is 1, return $[P_{n-1}(a), P_n(a)]$.

The library syntax is `GEN pollegendre_eval0(long n, GEN a = NULL, long flag)`. To obtain the *n*-th Legendre polynomial P_n in variable *v*, use `GEN pollegendre(long n, long v)`. To obtain $P_n(a)$, use `GEN pollegendre_eval(long n, GEN a)`.

3.9.35 polmodular(*L*, {*inv* = 0}, {*x* = 'x}, {*y* = 'y}, {*derivs* = 0}). Return the modular polynomial of prime level *L* in variables *x* and *y* for the modular function specified by *inv*. If *inv* is 0 (the default), use the modular *j* function, if *inv* is 1 use the Weber-*f* function, and if *inv* is 5 use $\gamma_2 = \sqrt{[3]}j$. See `polclass` for the full list of invariants. If *x* is given as `Mod(j, p)` or an element *j* of a finite field (as a `t_FFELT`), then return the modular polynomial of level *L* evaluated at *j*. If *j* is from a finite field and *derivs* is nonzero, then return a triple where the last two elements are the first and second derivatives of the modular polynomial evaluated at *j*.

```
? polmodular(3)
%1 = x^4 + (-y^3 + 2232*y^2 - 1069956*y + 36864000)*x^3 + ...
? polmodular(7, 1, , 'J)
%2 = x^8 - J^7*x^7 + 7*J^4*x^4 - 8*J*x + J^8
? polmodular(7, 5, 7*ffgen(19)^0, 'j)
%3 = j^8 + 4*j^7 + 4*j^6 + 8*j^5 + j^4 + 12*j^2 + 18*j + 18
? polmodular(7, 5, Mod(7,19), 'j)
%4 = Mod(1, 19)*j^8 + Mod(4, 19)*j^7 + Mod(4, 19)*j^6 + ...
? u = ffgen(5)^0; T = polmodular(3,0,, 'j)*u;
```

```

? polmodular(3, 0, u, 'j, 1)
%6 = [j^4 + 3*j^2 + 4*j + 1, 3*j^2 + 2*j + 4, 3*j^3 + 4*j^2 + 4*j + 2]
? subst(T,x,u)
%7 = j^4 + 3*j^2 + 4*j + 1
? subst(T',x,u)
%8 = 3*j^2 + 2*j + 4
? subst(T'',x,u)
%9 = 3*j^3 + 4*j^2 + 4*j + 2

```

The library syntax is GEN `polmodular(long L, long inv, GEN x = NULL, long y = -1, long derivs)` where `y` is a variable number.

3.9.36 polrecip(*pol*). Reciprocal polynomial of *pol* with respect to its main variable, i.e. the coefficients of the result are in reverse order; *pol* must be a polynomial.

```

? polrecip(x^2 + 2*x + 3)
%1 = 3*x^2 + 2*x + 1
? polrecip(2*x + y)
%2 = y*x + 2

```

The library syntax is GEN `polrecip(GEN pol)`.

3.9.37 polresultant(*x*, *y*, {*v*}, {*flag* = 0}). Resultant of the two polynomials *x* and *y* with exact entries, with respect to the main variables of *x* and *y* if *v* is omitted, with respect to the variable *v* otherwise. The algorithm assumes the base ring is a domain. If you also need the *u* and *v* such that $x * u + y * v = \text{Res}(x, y)$, use the `polresultanttext` function.

If *flag* = 0 (default), uses the algorithm best suited to the inputs, either the subresultant algorithm (Lazard/Ducos variant, generic case), a modular algorithm (inputs in $\mathbf{Q}[X]$) or Sylvester's matrix (inexact inputs).

If *flag* = 1, uses the determinant of Sylvester's matrix instead; this should always be slower than the default.

If *x* or *y* are multivariate with a huge *polynomial* content, it is advisable to remove it before calling this function. Compare:

```

? a = polcyclo(7) * ((t+1)/(t+2))^100;
? b = polcyclo(11) * ((t+2)/(t+3))^100;
? polresultant(a,b);
time = 3,833 ms.
? ca = content(a); cb = content(b); \
  polresultant(a/ca, b/cb) * ca^poldegree(b) * cb^poldegree(a); \\ instantaneous

```

The function only removes rational denominators and does not compute automatically the content because it is generically small and potentially *very* expensive (e.g. in multivariate contexts). The choice is yours, depending on your application.

The library syntax is GEN `polresultant0(GEN x, GEN y, long v = -1, long flag)` where *v* is a variable number.

3.9.38 polresultanttext($A, B, \{v\}$). Finds polynomials U and V such that $A*U+B*V = R$, where R is the resultant of U and V with respect to the main variables of A and B if v is omitted, and with respect to v otherwise. Returns the row vector $[U, V, R]$. The algorithm used (subresultant) assumes that the base ring is a domain.

```
? A = x*y; B = (x+y)^2;
? [U,V,R] = polresultanttext(A, B)
%2 = [-y*x - 2*y^2, y^2, y^4]
? A*U + B*V
%3 = y^4
? [U,V,R] = polresultanttext(A, B, y)
%4 = [-2*x^2 - y*x, x^2, x^4]
? A*U+B*V
%5 = x^4
```

The library syntax is `GEN polresultanttext0(GEN A, GEN B, long v = -1)` where v is a variable number. Also available is `GEN polresultanttext(GEN x, GEN y)`.

3.9.39 polroots(T). Complex roots of the polynomial T , given as a column vector where each root is repeated according to its multiplicity and given as floating point complex numbers at the current `realprecision`:

```
? polroots(x^2)
%1 = [0.E-38 + 0.E-38*I, 0.E-38 + 0.E-38*I]~
? polroots(x^3+1)
%2 = [-1.00... + 0.E-38*I, 0.50... - 0.866...*I, 0.50... + 0.866...*I]~
```

The algorithm used is a modification of Schönhage's root-finding algorithm, due to and originally implemented by Gourdon. It runs in polynomial time in $\deg(T)$ and the precision. If furthermore T has rational coefficients, roots are guaranteed to the required relative accuracy. If the input polynomial T is exact, then the ordering of the roots does not depend on the precision: they are ordered by increasing $|\Im z|$, then by increasing $\Re z$; in case of tie (conjugates), the root with negative imaginary part comes first.

The library syntax is `GEN roots(GEN T, long prec)`.

3.9.40 polrootsbound($T, \{\tau = 0.01\}$). Return a sharp upper bound B for the modulus of the largest complex root of the polynomial T with complex coefficients with relative error τ . More precisely, we have $|z| \leq B$ for all roots and there exist one root such that $|z_0| \geq B \exp(-2\tau)$. Much faster than either `polroots` or `polrootsreal`.

```
? T=polttchebi(500);
? vecmax(abs(polroots(T)))
time = 5,706 ms.
%2 = 0.99999506520185816611184481744870013191
? vecmax(abs(polrootsreal(T)))
time = 1,972 ms.
%3 = 0.99999506520185816611184481744870013191
? polrootsbound(T)
time = 217 ms.
%4 = 1.0098792554165905155
? polrootsbound(T, log(2)/2) \\ allow a factor 2, much faster
```

```

time = 51 ms.
%5 = 1.4065759938190154354
? polrootsbound(T, 1e-4)
time = 504 ms.
%6 = 1.0000920717983847741
? polrootsbound(T, 1e-6)
time = 810 ms.
%7 = 0.9999960628901692905
? polrootsbound(T, 1e-10)
time = 1,351 ms.
%8 = 0.9999950652993869760

```

The library syntax is GEN polrootsbound(GEN T, GEN tau = NULL).

3.9.41 polrootsff($x, \{p\}, \{a\}$). Obsolete, kept for backward compatibility: use polrootsmod.

The library syntax is GEN polrootsff(GEN x, GEN p = NULL, GEN a = NULL).

3.9.42 polrootsmod($f, \{D\}$). Vector of roots of the polynomial f over the finite field defined by the domain D as follows:

- $D = p$ a prime: factor over $\mathbf{F}_p$;
- $D = [T, p]$ for a prime p and $T(y)$ an irreducible polynomial over $\mathbf{F}_p$: factor over $\mathbf{F}_p[y]/(T)$ (as usual the main variable of T must have lower priority than the main variable of f);
- D a `t_FFELT`: factor over the attached field;
- D omitted: factor over the field of definition of f , which must be a finite field.

Multiple roots are *not* repeated.

```

? polrootsmod(x^2-1,2)
%1 = [Mod(1, 2)]~
? polrootsmod(x^2+1,3)
%2 = []~
? polrootsmod(x^2+1, [y^2+1,3])
%3 = [Mod(Mod(1, 3)*y, Mod(1, 3)*y^2 + Mod(1, 3)),
      Mod(Mod(2, 3)*y, Mod(1, 3)*y^2 + Mod(1, 3)) ]~
? polrootsmod(x^2 + Mod(1,3))
%4 = []~
? liftall( polrootsmod(x^2 + Mod(Mod(1,3),y^2+1)) )
%5 = [y, 2*y]~
? t = ffgen(y^2+Mod(1,3)); polrootsmod(x^2 + t^0)
%6 = [y, 2*y]~

```

The library syntax is GEN polrootsmod(GEN f, GEN D = NULL).

3.9.43 polrootspadic(f, p, r). Vector of p -adic roots of the polynomial pol , given to p -adic precision r ; the integer p is assumed to be a prime. Multiple roots are *not* repeated. Note that this is not the same as the roots in $\mathbf{Z}/p^r\mathbf{Z}$, rather it gives approximations in $\mathbf{Z}/p^r\mathbf{Z}$ of the true roots living in $\mathbf{Q}_p$:

```
? polrootspadic(x^3 - x^2 + 64, 2, 4)
%1 = [2^3 + 0(2^4), 2^3 + 0(2^4), 1 + 0(2^4)]~
? polrootspadic(x^3 - x^2 + 64, 2, 5)
%2 = [2^3 + 0(2^5), 2^3 + 2^4 + 0(2^5), 1 + 0(2^5)]~
```

As the second commands show, the first two roots *are* distinct in $\mathbf{Q}_p$, even though they are equal modulo 2^4 .

More generally, if T is an integral polynomial irreducible mod p and f has coefficients in $\mathbf{Q}[t]/(T)$, the argument p may be replaced by the vector $[T, p]$; we then return the roots of f in the unramified extension $\mathbf{Q}_p[t]/(T)$.

```
? polrootspadic(x^3 - x^2 + 64*y, [y^2+y+1, 2], 5)
%3 = [Mod((2^3 + 0(2^5))*y + (2^3 + 0(2^5)), y^2 + y + 1),
      Mod((2^3 + 2^4 + 0(2^5))*y + (2^3 + 2^4 + 0(2^5)), y^2 + y + 1),
      Mod(1 + 0(2^5), y^2 + y + 1)]~
```

If pol has inexact `t_PADIC` coefficients, this need not well-defined; in this case, the polynomial is first made integral by dividing out the p -adic content, then lifted to $\mathbf{Z}$ using `truncate` coefficientwise. Hence the roots given are approximations of the roots of an exact polynomial which is p -adically close to the input. To avoid pitfalls, we advise to only factor polynomials with exact rational coefficients.

The library syntax is `GEN polrootspadic(GEN f, GEN p, long r)`.

3.9.44 polrootsreal($T, \{ab\}$). Real roots of the polynomial T with real coefficients, multiple roots being included according to their multiplicity. If the polynomial does not have rational coefficients, it is first rescaled and rounded. The roots are given to a relative accuracy of `realprecision`. If argument ab is present, it must be a vector $[a, b]$ with two components (of type `t_INT`, `t_FRAC` or `t_INFINITY`) and we restrict to roots belonging to that closed interval.

```
? \p9
? polrootsreal(x^2-2)
%1 = [-1.41421356, 1.41421356]~
? polrootsreal(x^2-2, [1,+oo])
%2 = [1.41421356]~
? polrootsreal(x^2-2, [2,3])
%3 = []~
? polrootsreal((x-1)*(x-2), [2,3])
%4 = [2.00000000]~
```

The algorithm used is a modification of Uspensky's method (relying on Descartes's rule of sign), following Rouillier and Zimmerman's article "Efficient isolation of a polynomial real roots" (<https://hal.inria.fr/inria-00072518/>). Barring bugs, it is guaranteed to converge and to give the roots to the required accuracy.

Remark. If the polynomial T is of the form $Q(x^h)$ for some $h \geq 2$ and ab is omitted, the routine will apply the algorithm to Q (restricting to nonnegative roots when h is even), then take h -th roots. On the other hand, if you want to specify ab , you should apply the routine to Q yourself and a suitable interval $[a', b']$ using approximate h -th roots adapted to your problem: the function will not perform this change of variables if ab is present.

The library syntax is `GEN realroots(GEN T, GEN ab = NULL, long prec)`.

3.9.45 polsturm($T, \{ab\}$). Number of distinct real roots of the real polynomial T . If the argument ab is present, it must be a vector $[a, b]$ with two real components (of type `t_INT`, `t_REAL`, `t_FRAC` or `t_INFINITY`) and we count roots belonging to that closed interval.

If possible, you should stick to exact inputs, that is avoid `t_REALs` in T and the bounds a, b : the result is then guaranteed and we use a fast algorithm (Uspensky's method, relying on Descartes's rule of sign, see `polrootsreal`). Otherwise, the polynomial is rescaled and rounded first and the result may be wrong due to that initial error. If only a or b is inexact, on the other hand, the interval is first thickened using rational endpoints and the result remains guaranteed unless there exist a root *very* close to a nonrational endpoint (which may be missed or unduly included).

```
? T = (x-1)*(x-2)*(x-3);
? polsturm(T)
%2 = 3
? polsturm(T, [-oo,2])
%3 = 2
? polsturm(T, [1/2,+oo])
%4 = 3
? polsturm(T, [1, Pi])  \\ Pi inexact: not recommended !
%5 = 3
? polsturm(T*1., [0, 4])  \\ T*1. inexact: not recommended !
%6 = 3
? polsturm(T^2, [0, 4])  \\ not squarefree: roots are not repeated!
%7 = 3
```

The library syntax is `long RgX_sturmpart(GEN T, GEN ab)` or `long sturm(GEN T)` (for the case $ab = \text{NULL}$). The function `long sturmpart(GEN T, GEN a, GEN b)` is obsolete and deprecated.

3.9.46 polsubcyclo($n, d, \{v = 'x\}$). Gives polynomials (in variable v) defining the (Abelian) subextensions of degree d of the cyclotomic field $\mathbf{Q}(\zeta_n)$, where $d \mid \phi(n)$.

If there is exactly one such extension the output is a polynomial, else it is a vector of polynomials, possibly empty. To get a vector in all cases, use `concat([], polsubcyclo(n,d))`.

Each such polynomial is the minimal polynomial for a Gaussian period $\text{Tr}_{\mathbf{Q}(\zeta_f)/L}(\zeta_f)$, where L is the degree d subextension of $\mathbf{Q}(\zeta_n)$ and $f \mid n$ is its conductor. In Galois-theoretic terms, $L = \mathbf{Q}(\zeta_n)^H$, where H runs through all index d subgroups of $(\mathbf{Z}/n\mathbf{Z})^*$.

The function `galoissubcyclo` allows to specify exactly which sub-Abelian extension should be computed by giving H .

Complexity. Ignoring logarithmic factors, `polsubcyclo` runs in time $O(n)$. The function `polsubcyclofast` returns different, less canonical, polynomials but runs in time $O(d^4)$, again ignoring logarithmic factors; thus it can handle much larger values of n .

The library syntax is `GEN polsubcyclo(long n, long d, long v = -1)` where v is a variable number.

3.9.47 polsubcyclofast($n, d, \{s = 0\}, \{exact = 0\}$). If $1 \leq d \leq 6$ or a prime, finds an equation for the subfields of $\mathbf{Q}(\zeta_n)$ with Galois group C_d ; the special value $d = -4$ provides the subfields with group $V_4 = C_2 \times C_2$. Contrary to `polsubcyclo`, the output is always a (possibly empty) vector of polynomials. If $s = 0$ (default) all signatures, otherwise $s = 1$ (resp., -1) for totally real (resp., totally complex). Set `exact = 1` for subfields of conductor n .

The argument n can be given as in arithmetic functions: as an integer, as a factorization matrix, or (preferred) as a pair `[N, factor(N)]`.

Comparison with polsubcyclo. First `polsubcyclofast` does not usually return Gaussian periods, but ad hoc polynomials which do generate the same field. Roughly speaking (ignoring logarithmic factors), the complexity of `polsubcyclo` is independent of d and the complexity of `polsubcyclofast` is independent of n . Ignoring logarithmic factors, `polsubcyclo` runs in time $O(n)$ and `polsubcyclofast` in time $O(d^4)$. So the latter is *much* faster than `polsubcyclo` if n is large, but gets slower as d increases and becomes unusable for $d \geq 40$ or so.

```
? polsubcyclo(10^7+19,7);
time = 1,852 ms.
? polsubcyclofast(10^7+19,7);
time = 15 ms.

? polsubcyclo(10^17+21,5); \\ won't finish
*** polsubcyclo: user interrupt after 2h
? polsubcyclofast(10^17+21,5);
time = 3 ms.

? polsubcyclofast(10^17+3,7);
time = 26 ms.

? polsubcyclo(10^6+117,13);
time = 193 ms.
? polsubcyclofast(10^6+117,13);
time = 50 ms.

? polsubcyclofast(10^6+199,19);
time = 202 ms.
? polsubcyclo(10^6+199,19); \\ about as fast
time = 3191ms.

? polsubcyclo(10^7+271,19);
time = 2,067 ms.
? polsubcyclofast(10^7+271,19);
time = 201 ms.
```

The library syntax is `GEN polsubcyclofast(GEN n, long d, long s, long exact)`.

3.9.48 polysylvestermatrix(x, y). Forms the Sylvester matrix corresponding to the two polynomials x and y , where the coefficients of the polynomials are put in the columns of the matrix (which is the natural direction for solving equations afterwards). The use of this matrix can be essential when dealing with polynomials with inexact entries, since polynomial Euclidean division doesn't make much sense in this case.

The library syntax is GEN `sylvestermatrix`(GEN x , GEN y).

3.9.49 polysym(x, n). Creates the column vector of the symmetric powers of the roots of the polynomial x up to power n , using Newton's formula.

The library syntax is GEN `polysym`(GEN x , long n).

3.9.50 poltchebi($n, \{v = 'x\}$). Deprecated alias for `polchebyshev`

The library syntax is GEN `polchebyshev1`(long n , long $v = -1$) where v is a variable number.

3.9.51 polteichmuller(T, p, r). Given $T \in \mathbf{F}_p[X]$ return the polynomial $P \in \mathbf{Z}_p[X]$ whose roots (resp. leading coefficient) are the Teichmuller lifts of the roots (resp. leading coefficient) of T , to p -adic precision r . If T is monic, P is the reduction modulo p^r of the unique monic polynomial congruent to T modulo p such that $P(X^p) = 0 \pmod{P(X), p^r}$.

```
? T = ffinit(3, 3, 't)
%1 = Mod(1,3)*t^3 + Mod(1,3)*t^2 + Mod(1,3)*t + Mod(2,3)
? P = polteichmuller(T,3,5)
%2 = t^3 + 166*t^2 + 52*t + 242
? subst(P, t, t^3) % (P*Mod(1,3^5))
%3 = Mod(0, 243)
? [algdep(a+0(3^5),2) | a <- Vec(P)]
%4 = [x - 1, 5*x^2 + 1, x^2 + 4*x + 4, x + 1]
```

When T is monic and irreducible mod p , this provides a model $\mathbf{Q}_p[X]/(P)$ of the unramified extension $\mathbf{Q}_p[X]/(T)$ where the Frobenius has the simple form $X \bmod P \mapsto X^p \bmod P$.

The library syntax is GEN `polteichmuller`(GEN T , ulong p , long r).

3.9.52 poltomonic($T, \{\&L\}$). Let $T \in \mathbf{Q}[x]$ be a nonzero polynomial; returns U monic in $\mathbf{Z}[x]$ such that $U(x) = CT(x/L)$ for some $C, L \in \mathbf{Q}$. If the pointer argument $\&L$ is present, set L to L .

```
? poltomonic(9*x^2 - 1/2)
%1 = x^2 - 2
? U = poltomonic(9*x^2 - 1/2, &L)
%2 = x^2 - 2
? L
%3 = 6
? U / subst(9*x^2 - 1/2, x, x/L)
%4 = 4
```

This function does not compute discriminants or maximal orders and runs with complexity almost linear in the input size. If T is already monic with integer coefficient, `poltomonic` may still transform it if $\mathbf{Z}[x]/(T)$ is contained in a trivial subring of the maximal order, generated by Lx :

```
? poltomonic(x^2 + 4, &L)
```

```
%5 = x^2 + 1
? L
%6 = 1/2
```

If T is irreducible, the functions **polredabs** (exponential time) and **polredbest** (polynomial time) also find a monic integral generating polynomial for the number field $\mathbf{Q}[x]/(T)$, with explicit guarantees on its size, but are orders of magnitude slower.

The library syntax is **GEN poltmonic(GEN T, GEN *L = NULL)**.

3.9.53 polzagier(n, m). Creates Zagier's polynomial $P_n^{(m)}$ used in the functions **sumalt** and **sumpos** (with *flag* = 1), see “Convergence acceleration of alternating series”, Cohen et al., *Experiment. Math.*, vol. 9, 2000, pp. 3–12.

If $m < 0$ or $m \geq n$, $P_n^{(m)} = 0$. We have $P_n := P_n^{(0)}$ is $T_n(2x - 1)$, where T_n is the Legendre polynomial of the second kind. For $n > m > 0$, $P_n^{(m)}$ is the m -th difference with step 2 of the sequence $n^{m+1}P_n$; in this case, it satisfies

$$2P_n^{(m)}(\sin^2 t) = \frac{d^{m+1}}{dt^{m+1}}(\sin(2t)^m \sin(2(n-m)t)).$$

The library syntax is **GEN polzag(long n, long m)**.

3.9.54 seralgdep(s, p, r). finds a linear relation between powers $(1, s, \dots, s^p)$ of the series s , with polynomial coefficients of degree $\leq r$. In case no relation is found, return 0.

```
? s = 1 + 10*y - 46*y^2 + 460*y^3 - 5658*y^4 + 77740*y^5 + 0(y^6);
? seralgdep(s, 2, 2)
%2 = -x^2 + (8*y^2 + 20*y + 1)
? subst(%, x, s)
%3 = 0(y^6)
? seralgdep(s, 1, 3)
%4 = (-77*y^2 - 20*y - 1)*x + (310*y^3 + 231*y^2 + 30*y + 1)
? seralgdep(s, 1, 2)
%5 = 0
```

The series main variable must not be x , so as to be able to express the result as a polynomial in x .

The library syntax is **GEN seralgdep(GEN s, long p, long r)**.

3.9.55 serconvol(x, y). Convolution (or Hadamard product) of the two power series x and y ; in other words if $x = \sum a_k * X^k$ and $y = \sum b_k * X^k$ then **serconvol**(x, y) = $\sum a_k * b_k * X^k$.

The library syntax is **GEN convol(GEN x, GEN y)**.

3.9.56 serdiffdep(s, p, r). Find a linear relation between the derivatives ($s, s', \dots, s^p$) of the series s and 1, with polynomial coefficients of degree $\leq r$. In case no relation is found, return 0, otherwise return $[E, P]$ such that $E(d)(S) = P$ where d is the standard derivation.

```
? S = sum(i=0, 50, binomial(3*i,i)*T^i) + 0(T^51);
? serdiffdep(S, 3, 3)
%2 = [(27*T^2 - 4*T)*x^2 + (54*T - 2)*x + 6, 0]
? (27*T^2 - 4*T)*S'' + (54*T - 2)*S' + 6*S
%3 = 0(T^50)

? S = exp(T^2) + T^2;
? serdiffdep(S, 3, 3)
%5 = [x-2*T, -2*T^3+2*T]
? S'-2*T*S
%6 = 2*T-2*T^3+0(T^17)
```

The series main variable must not be x , so as to be able to express the result as a polynomial in x .

The library syntax is GEN `serdiffdep(GEN s, long p, long r)`.

3.9.57 serlaplace(x). x must be a power series with nonnegative exponents or a polynomial. If $x = \sum (a_k/k!) * X^k$ then the result is $\sum a_k * X^k$.

The library syntax is GEN `laplace(GEN x)`.

3.9.58 serreverse(s). Reverse power series of s , i.e. the series t such that $t(s) = x$; s must be a power series whose valuation is exactly equal to one.

```
? \ps 8
? t = serreverse(tan(x))
%2 = x - 1/3*x^3 + 1/5*x^5 - 1/7*x^7 + 0(x^8)
? tan(t)
%3 = x + 0(x^8)
```

The library syntax is GEN `serreverse(GEN s)`.

3.9.59 subst(x, y, z). Replace the simple variable y by the argument z in the “polynomial” expression x . If z is a vector, return the vector of the evaluated expressions `subst(x, y, z[i])`.

Every type is allowed for x , but if it is not a genuine polynomial (or power series, or rational function), the substitution will be done as if the scalar components were polynomials of degree zero. In particular, beware that:

```
? subst(1, x, [1,2; 3,4])
%1 =
[1 0]
[0 1]

? subst(1, x, Mat([0,1]))
*** at top-level: subst(1,x,Mat([0,1]))
*** ^-----
*** subst: forbidden substitution by a non square matrix.
```

If x is a power series, z must be either a polynomial, a power series, or a rational function. If x is a vector, matrix or list, the substitution is applied to each individual entry.

Use the function **substvec** to replace several variables at once, or the function **substpol** to replace a polynomial expression.

The library syntax is **GEN gsubst**(**GEN x**, **long y**, **GEN z**) where **y** is a variable number.

3.9.60 substpol(x, y, z). Replace the “variable” y by the argument z in the “polynomial” expression x . Every type is allowed for x , but the same behavior as **subst** above apply.

The difference with **subst** is that y is allowed to be any polynomial here. The substitution is done moding out all components of x (recursively) by $y - t$, where t is a new free variable of lowest priority. Then substituting t by z in the resulting expression. For instance

```
? substpol(x^4 + x^2 + 1, x^2, y)
%1 = y^2 + y + 1
? substpol(x^4 + x^2 + 1, x^3, y)
%2 = x^2 + y*x + 1
? substpol(x^4 + x^2 + 1, (x+1)^2, y)
%3 = (-4*y - 6)*x + (y^2 + 3*y - 3)
```

The library syntax is **GEN gsubstpol**(**GEN x**, **GEN y**, **GEN z**). Further, **GEN gdeflate**(**GEN T**, **long v**, **long d**) attempts to write $T(x)$ in the form $t(x^d)$, where $x = \text{pol_x}(v)$, and returns **NULL** if the substitution fails (for instance in the example **%2** above).

3.9.61 substvec(x, v, w). v being a vector of monomials of degree 1 (variables), w a vector of expressions of the same length, replace in the expression x all occurrences of v_i by w_i . The substitutions are done simultaneously; more precisely, the v_i are first replaced by new variables in x , then these are replaced by the w_i :

```
? substvec([x,y], [x,y], [y,x])
%1 = [y, x]
? substvec([x,y], [x,y], [y,x+y])
%2 = [y, x + y]      \\ not [y, 2*y]
```

As in **subst**, variables may be replaced by a vector of values, in which case the cartesian product is returned:

```
? substvec([x,y], [x,y], [[1,2], 3])
%3 = [[1, 3], [2, 3]]
? substvec([x,y], [x,y], [[1,2], [3,4]])
%4 = [[1, 3], [2, 3], [1, 4], [2, 4]]
```

The library syntax is **GEN gsubstvec**(**GEN x**, **GEN v**, **GEN w**).

3.9.62 sumformal($f, \{v\}$). formal sum of the polynomial expression f with respect to the main variable if v is omitted, with respect to the variable v otherwise; it is assumed that the base ring has characteristic zero. In other words, considering f as a polynomial function in the variable v , returns F , a polynomial in v vanishing at 0, such that $F(b) - F(a) = \text{sum}_{v=a+1}^b f(v)$:

```
? sumformal(n)  \\ 1 + ... + n
%1 = 1/2*n^2 + 1/2*n
? f(n) = n^3+n^2+1;
? F = sumformal(f(n))  \\ f(1) + ... + f(n)
%3 = 1/4*n^4 + 5/6*n^3 + 3/4*n^2 + 7/6*n
? sum(n = 1, 2000, f(n)) == subst(F, n, 2000)
%4 = 1
? sum(n = 1001, 2000, f(n)) == subst(F, n, 2000) - subst(F, n, 1000)
%5 = 1
? sumformal(x^2 + x*y + y^2, y)
%6 = y*x^2 + (1/2*y^2 + 1/2*y)*x + (1/3*y^3 + 1/2*y^2 + 1/6*y)
? x^2 * y + x * sumformal(y) + sumformal(y^2) == %
%7 = 1
```

The library syntax is GEN `sumformal(GEN f, long v = -1)` where v is a variable number.

3.9.63 taylor($x, t, \{d = \text{seriesprecision}\}$). Taylor expansion around 0 of x with respect to the simple variable t . x can be of any reasonable type, for example a rational function. Contrary to `Ser`, which takes the valuation into account, this function adds $O(t^d)$ to all components of x .

```
? taylor(x/(1+y), y, 5)
%1 = (y^4 - y^3 + y^2 - y + 1)*x + O(y^5)
? Ser(x/(1+y), y, 5)
***      at top-level: Ser(x/(1+y),y,5)
***      ^-----
*** Ser: main variable must have higher priority in gtoSer.
```

The library syntax is GEN `tayl(GEN x, long t, long precd1)` where t is a variable number.

3.9.64 thue($tnf, a, \{sol\}$). Returns all solutions of the equation $P(x, y) = a$ in integers x and y , where tnf was created with `thueinit( $P$ )`. If present, sol must contain the solutions of $\text{Norm}(x) = a$ modulo units of positive norm in the number field defined by P (as computed by `bnfisintnorm`). If there are infinitely many solutions, an error is issued.

It is allowed to input directly the polynomial P instead of a tnf , in which case, the function first performs `thueinit( $P, 0$ )`. This is very wasteful if more than one value of a is required.

If tnf was computed without assuming GRH (flag 1 in `thueinit`), then the result is unconditional. Otherwise, it depends in principle of the truth of the GRH, but may still be unconditionally correct in some favorable cases. The result is conditional on the GRH if $a \neq \pm 1$ and P has a single irreducible rational factor, whose attached tentative class number h and regulator R (as computed assuming the GRH) satisfy

- $h > 1$,
- $R/0.2 > 1.5$.

Here's how to solve the Thue equation $x^{13} - 5y^{13} = -4$:

```
? tnf = thueinit(x^13 - 5);
? thue(tnf, -4)
%1 = [[1, 1]]
```

In this case, one checks that `bnfinit(x^13 - 5).no` is 1. Hence, the only solution is $(x, y) = (1, 1)$ and the result is unconditional. On the other hand:

```
? P = x^3-2*x^2+3*x-17; tnf = thueinit(P);
? thue(tnf, -15)
%2 = [[1, 1]] \\ a priori conditional on the GRH.
? K = bnfinit(P); K.no
%3 = 3
? K.reg
%4 = 2.8682185139262873674706034475498755834
```

This time the result is conditional. All results computed using this particular *tnf* are likewise conditional, *except* for a right-hand side of ± 1 . The above result is in fact correct, so we did not just disprove the GRH:

```
? tnf = thueinit(x^3-2*x^2+3*x-17, 1 /*unconditional*/);
? thue(tnf, -15)
%4 = [[1, 1]]
```

Note that reducible or nonmonic polynomials are allowed:

```
? tnf = thueinit((2*x+1)^5 * (4*x^3-2*x^2+3*x-17), 1);
? thue(tnf, 128)
%2 = [[-1, 0], [1, 0]]
```

Reducible polynomials are in fact much easier to handle.

Note. When P is irreducible without a real root, the default strategy is to use brute force enumeration in time $|a|^{1/\deg P}$ and avoid computing a tough *bnf* attached to P , see `thueinit`. Besides reusing a quantity you might need for other purposes, the default argument *sol* can also be used to use a different strategy and prove that there are no solutions; of course you need to compute a *bnf* on you own to obtain *sol*. If there *are* solutions this won't help unless P is quadratic, since the enumeration will be performed in any case.

The library syntax is `GEN thue(GEN tnf, GEN a, GEN sol = NULL)`.

3.9.65 thueinit($P, \{flag = 0\}$). Initializes the *tnf* corresponding to P , a nonconstant univariate polynomial with integer coefficients. The result is meant to be used in conjunction with `thue` to solve Thue equations $P(X/Y)Y^{\deg P} = a$, where a is an integer. Accordingly, P must either have at least two distinct irreducible factors over $\mathbf{Q}$, or have one irreducible factor T with degree > 2 or two conjugate complex roots: under these (necessary and sufficient) conditions, the equation has finitely many integer solutions.

```
? S = thueinit(t^2+1);
? thue(S, 5)
%2 = [[-2, -1], [-2, 1], [-1, -2], [-1, 2], [1, -2], [1, 2], [2, -1], [2, 1]]
? S = thueinit(t+1);
*** at top-level: thueinit(t+1)
*** ^-----
*** thueinit: domain error in thueinit: P = t + 1
```

The hardest case is when $\deg P > 2$ and P is irreducible with at least one real root. The routine then uses Bilu-Hanrot's algorithm.

If *flag* is nonzero, certify results unconditionally. Otherwise, assume GRH, this being much faster of course. In the latter case, the result may still be unconditionally correct, see **thue**. For instance in most cases where P is reducible (not a pure power of an irreducible), *or* conditional computed class groups are trivial *or* the right hand side is ± 1 , then results are unconditional.

Note. The general philosophy is to disprove the existence of large solutions then to enumerate bounded solutions naively. The implementation will overflow when there exist huge solutions and the equation has degree > 2 (the quadratic imaginary case is special, since we can stick to **bnfisintnorm**, there are no fundamental units):

```
? thue(t^3+2, 10^30)
***   at top-level: L=thue(t^3+2,10^30)
***               ^-----
*** thue: overflow in thue (SmallSols): y <= 80665203789619036028928.
? thue(x^2+2, 10^30)  \\ quadratic case much easier
%1 = [[-10000000000000000, 0], [10000000000000000, 0]]
```

Note. It is sometimes possible to circumvent the above, and in any case obtain an important speed-up, if you can write $P = Q(x^d)$ for some $d > 1$ and Q still satisfying the **thueinit** hypotheses. You can then solve the equation attached to Q then eliminate all solutions (x, y) such that either x or y is not a d -th power.

```
? thue(x^4+1, 10^40); \\ stopped after 10 hours
? filter(L,d) =
    my(x,y); [[x,y] | v<-L, ispower(v[1],d,&x)&&ispower(v[2],d,&y)];
? L = thue(x^2+1, 10^40);
? filter(L, 2)
%4 = [[0, 100000000000], [100000000000, 0]]
```

The last 2 commands use less than 20ms.

Note. When P is irreducible without a real root, the equation can be solved unconditionnally in time $|a|^{1/\deg P}$. When this latter quantity is huge and the equation has no solutions, this fact may still be ascertained via arithmetic conditions but this now implies solving norm equations, computing a *bnf* and possibly assuming the GRH. When there is no real root, the code does not compute a *bnf* (with certification if *flag* = 1) if it expects this to be an “easy” computation (because the result would only be used for huge values of a). See **thue** for a way to compute an expensive *bnf* on your own and still get a result where this default cheap strategy fails.

The library syntax is **GEN thueinit(GEN P, long flag, long prec)**.

3.10 Vectors, matrices, linear algebra and sets.

Note that most linear algebra functions operating on subspaces defined by generating sets (such as `mathnf`, `qflll`, etc.) take matrices as arguments. As usual, the generating vectors are taken to be the *columns* of the given matrix.

Since PARI does not have a strong typing system, scalars live in unspecified commutative base rings. It is very difficult to write robust linear algebra routines in such a general setting. We thus assume that the base ring is a domain and work over its field of fractions. If the base ring is *not* a domain, one gets an error as soon as a nonzero pivot turns out to be noninvertible. Some functions, e.g. `mathnf` or `mathnfmod`, specifically assume that the base ring is $\mathbf{Z}$.

3.10.1 `algdep`($z, k, \{flag = 0\}$). z being real/complex, or p -adic, finds a polynomial (in the variable 'x') of degree at most k , with integer coefficients, having z as approximate root. Note that the polynomial which is obtained is not necessarily the “correct” one. In fact it is not even guaranteed to be irreducible. One can check the closeness either by a polynomial evaluation (use `subst`), or by computing the roots of the polynomial given by `algdep` (use `polroots` or `polrootspadic`).

Internally, `linddep`($[1, z, \dots, z^k], flag$) is used. A nonzero value of *flag* may improve on the default behavior if the input number is known to a *huge* accuracy, and you suspect the last bits are incorrect: if $flag > 0$ the computation is done with an accuracy of *flag* decimal digits; to get meaningful results, the parameter *flag* should be smaller than the number of correct decimal digits in the input. But default values are usually sufficient, so try without *flag* first:

```
? \p200
? z = 2^(1/6)+3^(1/5);
? algdep(z, 30);          \\ right in 63ms
? algdep(z, 30, 100);    \\ wrong in 39ms
? algdep(z, 30, 170);    \\ right in 61ms
? algdep(z, 30, 200);    \\ wrong in 146ms
? \p250
? z = 2^(1/6)+3^(1/5);    \\ recompute to new, higher, accuracy !
? algdep(z, 30);          \\ right in 68ms
? algdep(z, 30, 200);    \\ right in 68ms
? \p500
? algdep(2^(1/6)+3^(1/5), 30); \\ right in 138ms
? \p1000
? algdep(2^(1/6)+3^(1/5), 30); \\ right in 276s
```

The changes in `realprecision` only affect the quality of the initial approximation to $2^{1/6} + 3^{1/5}$, `algdep` itself uses exact operations. The size of its operands depend on the accuracy of the input of course: a more accurate input means slower operations.

Proceeding by increments of 5 digits of accuracy, `algdep` with default *flag* produces its first correct result at 195 digits, and from then on a steady stream of correct results:

```
\\ assume T contains the correct result, for comparison
forstep(d=100, 250, 5, \
  localprec(d);          \
  print(d, " ", algdep(2^(1/6)+3^(1/5),30) == T))
```

This example is the test case studied in a 2000 paper by Borwein and Lisonek: Applications of integer relation algorithms, *Discrete Math.*, **217**, p. 65–82. The version of PARI tested there was

1.39, which succeeded reliably from precision 265 on, in about 1000 as much time as the current version (on slower hardware of course).

Note that this function does not work if z is a power series. The function `seralgdep` can be used in this case to find linear relations with polynomial coefficients between powers of z .

The library syntax is `GEN algdep0(GEN z, long k, long flag)`. Also available is `GEN algdep(GEN z, long k) (flag = 0)`.

3.10.2 bestapprnf($V, T, \{rootT\}$). T being an integral polynomial and V being a scalar, vector, or matrix with complex coefficients, return a reasonable approximation of V with polmods modulo T . T can also be any number field structure, in which case the minimal polynomial attached to the structure ($T.pol$) is used. The $rootT$ argument, if present, must be an element of `polroots( $T$ )` (or $T.pol$), i.e. a complex root of T fixing an embedding of $\mathbf{Q}[x]/(T)$ into $\mathbf{C}$.

```
? bestapprnf(sqrt(5), polcyclo(5))
%1 = Mod(-2*x^3 - 2*x^2 - 1, x^4 + x^3 + x^2 + x + 1)
? bestapprnf(sqrt(5), polcyclo(5), exp(4*I*Pi/5))
%2 = Mod(2*x^3 + 2*x^2 + 1, x^4 + x^3 + x^2 + x + 1)
```

When the output has huge rational coefficients, try to increase the working `realbitprecision`: if the answer does not stabilize, consider that the reconstruction failed. Beware that if T is not Galois over $\mathbf{Q}$, some embeddings may not allow to reconstruct V :

```
? T = x^3-2; vT = polroots(T); z = 3*2^(1/3)+1;
? bestapprnf(z, T, vT[1])
%2 = Mod(3*x + 1, x^3 - 2)
? bestapprnf(z, T, vT[2])
%3 = 4213714286230872/186454048314072  \\ close to 3*2^(1/3) + 1
```

The library syntax is `GEN bestapprnf(GEN V, GEN T, GEN rootT = NULL, long prec)`.

3.10.3 charpoly($A, \{v = 'x\}, \{flag = 5\}$). characteristic polynomial of A with respect to the variable v , i.e. determinant of $v * I - A$ if A is a square matrix.

```
? charpoly([1,2;3,4]);
%1 = x^2 - 5*x - 2
? charpoly([1,2;3,4],, 't)
%2 = t^2 - 5*t - 2
```

If A is not a square matrix, the function returns the characteristic polynomial of the map “multiplication by A ” if A is a scalar:

```
? charpoly(Mod(x+2, x^3-2))
%1 = x^3 - 6*x^2 + 12*x - 10
? charpoly(I)
%2 = x^2 + 1
? charpoly(quadgen(5))
%3 = x^2 - x - 1
? charpoly(ffgen(ffinit(2,4)))
%4 = Mod(1, 2)*x^4 + Mod(1, 2)*x^3 + Mod(1, 2)*x^2 + Mod(1, 2)*x + Mod(1, 2)
```

The value of *flag* is only significant for matrices, and we advise to stick to the default value. Let n be the dimension of A .

If *flag* = 0, same method (Le Verrier's) as for computing the adjoint matrix, i.e. using the traces of the powers of A . Assumes that $n!$ is invertible; uses $O(n^4)$ scalar operations.

If *flag* = 1, uses Lagrange interpolation which is usually the slowest method. Assumes that $n!$ is invertible; uses $O(n^4)$ scalar operations.

If *flag* = 2, uses the Hessenberg form. Assumes that the base ring is a field. Uses $O(n^3)$ scalar operations, but suffers from coefficient explosion unless the base field is finite or $\mathbf{R}$.

If *flag* = 3, uses Berkowitz's division free algorithm, valid over any ring (commutative, with unit). Uses $O(n^4)$ scalar operations.

If *flag* = 4, x must be integral. Uses a modular algorithm: Hessenberg form for various small primes, then Chinese remainders.

If *flag* = 5 (default), uses the "best" method given x . This means we use Berkowitz unless the base ring is $\mathbf{Z}$ (use *flag* = 4) or a field where coefficient explosion does not occur, e.g. a finite field or the reals (use *flag* = 2).

The library syntax is `GEN charpoly0(GEN A, long v = -1, long flag)` where v is a variable number. Also available are `GEN charpoly(GEN x, long v)` (*flag* = 5), `GEN caract(GEN A, long v)` (*flag* = 1), `GEN carhess(GEN A, long v)` (*flag* = 2), `GEN carberkowitz(GEN A, long v)` (*flag* = 3) and `GEN caradj(GEN A, long v, GEN *pt)`. In this last case, if *pt* is not NULL, **pt* receives the address of the adjoint matrix of A (see `matadjoint`), so both can be obtained at once.

3.10.4 concat($x, \{y\}$). Concatenation of x and y . If x or y is not a vector or matrix, it is considered as a one-dimensional vector. All types are allowed for x and y , but the sizes must be compatible. Note that matrices are concatenated horizontally, i.e. the number of rows stays the same. Using transpositions, one can concatenate them vertically, but it is often simpler to use `matconcat`.

```
? x = matid(2); y = 2*matid(2);
? concat(x,y)
%2 =
[1 0 2 0]
[0 1 0 2]
? concat(x~,y~)~
%3 =
[1 0]
[0 1]
[2 0]
[0 2]
? matconcat([x;y])
%4 =
[1 0]
[0 1]
[2 0]
[0 2]
```

To concatenate vectors sideways (i.e. to obtain a two-row or two-column matrix), use `Mat` instead, or `matconcat`:

```
? x = [1,2];
? y = [3,4];
? concat(x,y)
%3 = [1, 2, 3, 4]

? Mat([x,y]~)
%4 =
[1 2]
[3 4]
? matconcat([x;y])
%5 =
[1 2]
[3 4]
```

Concatenating a row vector to a matrix having the same number of columns will add the row to the matrix (top row if the vector is x , i.e. comes first, and bottom row otherwise).

The empty matrix `[]` is considered to have a number of rows compatible with any operation, in particular concatenation. (Note that this is *not* the case for empty vectors `[]` or `[]~`.)

If y is omitted, x has to be a row vector or a list, in which case its elements are concatenated, from left to right, using the above rules.

```
? concat([1,2], [3,4])
%1 = [1, 2, 3, 4]
? a = [[1,2]~, [3,4]~]; concat(a)
%2 =
[1 3]
[2 4]

? concat([1,2; 3,4], [5,6]~)
%3 =
[1 2 5]
[3 4 6]
? concat([%, [7,8]~, [1,2,3,4]])
%5 =
[1 2 5 7]
[3 4 6 8]
[1 2 3 4]
```

The library syntax is `GEN gconcat(GEN x, GEN y = NULL)`. `GEN gconcat1(GEN x)` is a shortcut for `gconcat(x,NULL)`.

3.10.5 dirpowers(n, x). For nonnegative n and complex number x , return the vector with n components $[1^x, 2^x, \dots, n^x]$.

```
? dirpowers(5, 2)
%1 = [1, 4, 9, 16, 25]
? dirpowers(5, 1/2)
%2 = [1, 1.414..., 1.732..., 2.000..., 2.236...]
```

When $n \leq 0$, the function returns the empty vector `[]`.

The library syntax is `GEN dirpowers(long n, GEN x, long prec)`.

3.10.6 forqfvec(v, q, b, expr). q being a square and symmetric integral matrix representing a positive definite quadratic form, evaluate `expr` for all pairs of nonzero vectors $(-v, v)$ such that $q(v) \leq b$. The formal variable v runs through representatives of all such pairs in turn.

```
? forqfvec(v, [3,2;2,3], 3, print(v))
[0, 1]~
[1, 0]~
[-1, 1]~
```

The library syntax is `void forqfvec0(GEN v, GEN q = NULL, GEN b)`. The following functions are also available: `void forqfvec(void *E, long (*fun)(void *, GEN, GEN, double), GEN q, GEN b)`: Evaluate `fun(E,U,v,m)` on all v such that $q(Uv) < b$, where U is a `t_MAT`, v is a `t_VECSMALL` and $m = q(v)$ is a C double. The function `fun` must return 0, unless `forqfvec` should stop, in which case, it should return 1.

`void forqfvec1(void *E, long (*fun)(void *, GEN), GEN q, GEN b)`: Evaluate `fun(E,v)` on all v such that $q(v) < b$, where v is a `t_COL`. The function `fun` must return 0, unless `forqfvec` should stop, in which case, it should return 1.

3.10.7 lindep($v, \{flag = 0\}$). finds a small nontrivial integral linear combination between components of v . If none can be found return an empty vector.

If v is a vector with real/complex entries we use a floating point (variable precision) LLL algorithm. If $flag = 0$ the accuracy is chosen internally using a crude heuristic. If $flag > 0$ the computation is done with an accuracy of $flag$ decimal digits. To get meaningful results in the latter case, the parameter $flag$ should be smaller than the number of correct decimal digits in the input.

```
? lindep([sqrt(2), sqrt(3), sqrt(2)+sqrt(3)])
%1 = [-1, -1, 1]~
```

If v is p -adic, $flag$ is ignored and the algorithm LLL-reduces a suitable (dual) lattice.

```
? lindep([1, 2 + 3 + 3^2 + 3^3 + 3^4 + 0(3^5)])
%2 = [1, -2]~
```

If v is a matrix (or a vector of column vectors, or a vector of row vectors), $flag$ is ignored and the function returns a non trivial kernel vector if one exists, else an empty vector.

```
? lindep([1,2,3;4,5,6;7,8,9])
%3 = [1, -2, 1]~
? lindep([[1,0], [2,0]])
%4 = [2, -1]~
? lindep([[1,0], [0,1]])
```

```
%5 = []~
```

If v contains polynomials or power series over some base field, finds a linear relation with coefficients in the field.

```
? lindep([x*y, x^2 + y, x^2*y + x*y^2, 1])
%4 = [y, y, -1, -y^2]~
```

For better control, it is preferable to use `t_POL` rather than `t_SER` in the input, otherwise one gets a linear combination which is t -adically small, but not necessarily 0. Indeed, power series are first converted to the minimal absolute accuracy occurring among the entries of v (which can cause some coefficients to be ignored), then truncated to polynomials:

```
? v = [t^2+0(t^4), 1+0(t^2)]; L=lindep(v)
%1 = [1, 0]~
? v*L
%2 = t^2+0(t^4)  \\ small but not 0
```

The library syntax is `GEN lindep0(GEN v, long flag)`.

3.10.8 matadjoint($M, \{flag = 0\}$). adjoint matrix of M , i.e. a matrix N of cofactors of M , satisfying $M * N = \det(M) * \text{Id}$. M must be a (not necessarily invertible) square matrix of dimension n . If $flag$ is 0 or omitted, we try to use Leverrier-Faddeev's algorithm, which assumes that $n!$ invertible. If it fails or $flag = 1$, computes $T = \text{charpoly}(M)$ independently first and returns $(-1)^{n-1}(T(x) - T(0))/x$ evaluated at M .

```
? a = [1,2,3;3,4,5;6,7,8] * Mod(1,4);
? matadjoint(a)
%2 =
[Mod(1, 4) Mod(1, 4) Mod(2, 4)]
[Mod(2, 4) Mod(2, 4) Mod(0, 4)]
[Mod(1, 4) Mod(1, 4) Mod(2, 4)]
```

Both algorithms use $O(n^4)$ operations in the base ring. Over a field, they are usually slower than computing the characteristic polynomial or the inverse of M directly.

The library syntax is `GEN matadjoint0(GEN M, long flag)`. Also available are `GEN adj(GEN x)` ($flag = 0$) and `GEN adjsafe(GEN x)` ($flag = 1$).

3.10.9 matcompanion(x). The left companion matrix to the nonzero polynomial x .

The library syntax is `GEN matcompanion(GEN x)`.

3.10.10 matconcat(v). Returns a `t_MAT` built from the entries of v , which may be a `t_VEC` (concatenate horizontally), a `t_COL` (concatenate vertically), or a `t_MAT` (concatenate vertically each column, and concatenate vertically the resulting matrices). The entries of v are always considered as matrices: they can themselves be `t_VEC` (seen as a row matrix), a `t_COL` seen as a column matrix), a `t_MAT`, or a scalar (seen as an 1×1 matrix).

```
? A=[1,2;3,4]; B=[5,6]~; C=[7,8]; D=9;
? matconcat([A, B]) \\ horizontal
%1 =
[1 2 5]
[3 4 6]
? matconcat([A, C]~) \\ vertical
%2 =
[1 2]
[3 4]
[7 8]
? matconcat([A, B; C, D]) \\ block matrix
%3 =
[1 2 5]
[3 4 6]
[7 8 9]
```

If the dimensions of the entries to concatenate do not match up, the above rules are extended as follows:

- each entry $v_{i,j}$ of v has a natural length and height: 1×1 for a scalar, $1 \times n$ for a `t_VEC` of length n , $n \times 1$ for a `t_COL`, $m \times n$ for an $m \times n$ `t_MAT`
- let H_i be the maximum over j of the lengths of the $v_{i,j}$, let L_j be the maximum over i of the heights of the $v_{i,j}$. The dimensions of the (i,j) -th block in the concatenated matrix are $H_i \times L_j$.
- a scalar $s = v_{i,j}$ is considered as s times an identity matrix of the block dimension $\min(H_i, L_j)$
- blocks are extended by 0 columns on the right and 0 rows at the bottom, as needed.

```
? matconcat([1, [2,3]~, [4,5,6]~]) \\ horizontal
%4 =
[1 2 4]
[0 3 5]
[0 0 6]
? matconcat([1, [2,3], [4,5,6]]~) \\ vertical
%5 =
[1 0 0]
[2 3 0]
[4 5 6]
? matconcat([B, C; A, D]) \\ block matrix
%6 =
[5 0 7 8]
[6 0 0 0]
```

```

[1 2 9 0]
[3 4 0 9]
? U=[1,2;3,4]; V=[1,2,3;4,5,6;7,8,9];
? matconcat(matdiagonal([U, V])) \\ block diagonal
%7 =
[1 2 0 0 0]
[3 4 0 0 0]
[0 0 1 2 3]
[0 0 4 5 6]
[0 0 7 8 9]

```

The library syntax is `GEN matconcat(GEN v)`.

3.10.11 matdet($x, \{flag = 0\}$). Determinant of the square matrix x .

If $flag = 0$, uses an appropriate algorithm depending on the coefficients:

- integer entries: modular method due to Dixon, Pernet and Stein.
- real or p -adic entries: classical Gaussian elimination using maximal pivot.
- intmod entries: classical Gaussian elimination using first nonzero pivot.
- other cases: Gauss-Bareiss.

If $flag = 1$, uses classical Gaussian elimination with appropriate pivoting strategy (maximal pivot for real or p -adic coefficients). This is usually worse than the default.

The library syntax is `GEN det0(GEN x, long flag)`. Also available are `GEN det(GEN x)` ($flag = 0$), `GEN det2(GEN x)` ($flag = 1$) and `GEN ZM_det(GEN x)` for integer entries.

3.10.12 matdetint(B). Let B be an $m \times n$ matrix with integer coefficients. The *determinant* D of the lattice generated by the columns of B is the square root of $\det(B^T B)$ if B has maximal rank m , and 0 otherwise.

This function uses the Gauss-Bareiss algorithm to compute a positive *multiple* of D . When B is square, the function actually returns $D = |\det B|$.

This function is useful in conjunction with `mathnfmod`, which needs to know such a multiple. If the rank is maximal but the matrix is nonsquare, you can obtain D exactly using

```
matdet( mathnfmod(B, matdetint(B)) )
```

Note that as soon as one of the dimensions gets large (m or n is larger than 20, say), it will often be much faster to use `mathnf(B, 1)` or `mathnf(B, 4)` directly.

The library syntax is `GEN detint(GEN B)`.

3.10.13 matdetmod(x, d). Given a matrix x with `t_INT` entries and d an arbitrary positive integer, return the determinant of x modulo d .

```
? A = [4,2,3; 4,5,6; 7,8,9]
? matdetmod(A,27)
%2 = 9
```

Note that using the generic function `matdet` on a matrix with `t_INTMOD` entries uses Gaussian reduction and will fail in general when the modulus is not prime.

```
? matdet(A * Mod(1,27))
*** at top-level: matdet(A*Mod(1,27))
***      ^-----
*** matdet: impossible inverse in Fl_inv: Mod(3, 27).
```

The library syntax is `GEN matdetmod(GEN x, GEN d)`.

3.10.14 matdiagonal(x). x being a vector, creates the diagonal matrix whose diagonal entries are those of x .

```
? matdiagonal([1,2,3]);
%1 =
[1 0 0]
[0 2 0]
[0 0 3]
```

Block diagonal matrices are easily created using `matconcat`:

```
? U=[1,2;3,4]; V=[1,2,3;4,5,6;7,8,9];
? matconcat(matdiagonal([U, V]))
%1 =
[1 2 0 0 0]
[3 4 0 0 0]
[0 0 1 2 3]
[0 0 4 5 6]
[0 0 7 8 9]
```

The library syntax is `GEN diagonal(GEN x)`.

3.10.15 mateigen($x, \{flag = 0\}$). Returns the (complex) eigenvectors of x as columns of a matrix. If $flag = 1$, return $[L, H]$, where L contains the eigenvalues and H the corresponding eigenvectors; multiple eigenvalues are repeated according to the eigenspace dimension (which may be less than the eigenvalue multiplicity in the characteristic polynomial).

This function first computes the characteristic polynomial of x and approximates its complex roots (λ_i), then tries to compute the eigenspaces as kernels of the $x - \lambda_i$. This algorithm is ill-conditioned and is likely to miss kernel vectors if some roots of the characteristic polynomial are close, in particular if it has multiple roots.

```
? A = [13,2; 10,14]; mateigen(A)
%1 =
[-1/2 2/5]
```

```

[ 1 1]
? [L,H] = mateigen(A, 1);
? L
%3 = [9, 18]
? H
%4 =
[-1/2 2/5]
[ 1 1]
? A * H == H * matdiagonal(L)
%5 = 1

```

For symmetric matrices, use `qfjacobi` instead; for Hermitian matrices, compute

```

A = real(x);
B = imag(x);
y = matconcat([A, -B; B, A]);

```

and apply `qfjacobi` to y .

The library syntax is `GEN mateigen(GEN x, long flag, long prec)`. Also available is `GEN eigen(GEN x, long prec)` ($flag = 0$)

3.10.16 matfrobenius($M, \{flag\}, \{v = 'x\}$). Returns the Frobenius form of the square matrix M . If $flag = 1$, returns only the elementary divisors as a vector of polynomials in the variable v . If $flag = 2$, returns a two-components vector $[F, B]$ where F is the Frobenius form and B is the basis change so that $M = B^{-1}FB$.

The library syntax is `GEN matfrobenius(GEN M, long flag, long v = -1)` where v is a variable number.

3.10.17 mathess(x). Returns a matrix similar to the square matrix x , which is in upper Hessenberg form (zero entries below the first subdiagonal).

The library syntax is `GEN hess(GEN x)`.

3.10.18 mathilbert(n). Creates the Hilbert matrix of order $n \geq 0$, i.e. the square matrix H whose coefficient $H[i, j]$ is $1/(i + j - 1)$. This matrix is ill-conditioned but its inverse has integer entries.

The library syntax is `GEN mathilbert(long n)`.

3.10.19 mathnf($M, \{flag = 0\}$). Let R be a Euclidean ring, equal to $\mathbf{Z}$ or to $K[X]$ for some field K . If M is a (not necessarily square) matrix with entries in R , this routine finds the *upper triangular* Hermite normal form of M . If the rank of M is equal to its number of rows, this is a square matrix. In general, the columns of the result form a basis of the R -module spanned by the columns of M .

The values of $flag$ are:

- 0 (default): only return the Hermite normal form H
- 1 (complete output): return $[H, U]$, where H is the Hermite normal form of M , and U is a transformation matrix such that $MU = [0|H]$. The matrix U belongs to $GL(R)$. When M has a large kernel, the entries of U are in general huge.

For these two values, we use a naive algorithm, which behaves well in small dimension only. Larger values correspond to different algorithms, are restricted to *integer* matrices, and all output the unimodular matrix U . From now on all matrices have integral entries.

- $flag = 4$, returns $[H, U]$ as in “complete output” above, using a variant of LLL reduction along the way. The matrix U is provably small in the L_2 sense, and often close to optimal; but the reduction is in general slow, although provably polynomial-time.

If $flag = 5$, uses Batut’s algorithm and output $[H, U, P]$, such that H and U are as before and P is a permutation of the rows such that P applied to MU gives H . This is in general faster than $flag = 4$ but the matrix U is usually worse; it is heuristically smaller than with the default algorithm.

When the matrix is dense and the dimension is large (bigger than 100, say), $flag = 4$ will be fastest. When M has maximal rank, then

```
H = mathnfmod(M, matdetint(M))
```

will be even faster. You can then recover U as $M^{-1}H$.

```
? M = matrix(3,4,i,j,random([-5,5]))
%1 =
[ 0 2  3  0]
[-5 3 -5 -5]
[ 4 3 -5  4]
? [H,U] = mathnf(M, 1);
? U
%3 =
[-1 0 -1 0]
[ 0 5  3 2]
[ 0 3  1 1]
[ 1 0  0 0]
? H
%5 =
[19 9 7]
[ 0 9 1]
[ 0 0 1]
? M*U
%6 =
[0 19 9 7]
[0  0 9 1]
[0  0 0 1]
```

For convenience, M is allowed to be a `t_VEC`, which is then automatically converted to a `t_MAT`, as per the `Mat` function. For instance to solve the generalized extended gcd problem, one may use

```
? v = [116085838, 181081878, 314252913, 10346840];
? [H,U] = mathnf(v, 1);
? U
```

```
%2 =
[ 103 -603    15  -88]
[-146   13 -1208  352]
[   58  220   678 -167]
[-362 -144   381 -101]
? v*U
%3 = [0, 0, 0, 1]
```

This also allows to input a matrix as a `t_VEC` of `t_COLS` of the same length (which `Mat` would concatenate to the `t_MAT` having those columns):

```
? v = [[1,0,4]~, [3,3,4]~, [0,-4,-5]~]; mathnf(v)
%1 =
[47 32 12]
[ 0  1  0]
[ 0  0  1]
```

The library syntax is `GEN mathnf0(GEN M, long flag)`. Also available are `GEN hnf(GEN M)` ($flag = 0$) and `GEN hnfall(GEN M)` ($flag = 1$). To reduce *huge* relation matrices (sparse with small entries, say dimension 400 or more), you can use the pair `hnfspec` / `hnfadd`. Since this is quite technical and the calling interface may change, they are not documented yet. Look at the code in `basemath/hnf_snf.c`.

3.10.20 `mathnfmod`(x, d). If x is a (not necessarily square) matrix of maximal rank with integer entries, and d is a multiple of the (nonzero) determinant of the lattice spanned by the columns of x , finds the *upper triangular* Hermite normal form of x .

If the rank of x is equal to its number of rows, the result is a square matrix. In general, the columns of the result form a basis of the lattice spanned by the columns of x . Even when d is known, this is in general slower than `mathnf` but uses much less memory.

The library syntax is `GEN hnfmod(GEN x, GEN d)`.

3.10.21 `mathnfmodid`(x, d). Outputs the (upper triangular) Hermite normal form of x concatenated with the diagonal matrix with diagonal d . Assumes that x has integer entries. Variant: if d is an integer instead of a vector, concatenate d times the identity matrix.

```
? m=[0,7;-1,0;-1,-1]
%1 =
[ 0  7]
[-1  0]
[-1 -1]
? mathnfmodid(m, [6,2,2])
%2 =
[2 1 1]
[0 1 0]
[0 0 1]
? mathnfmodid(m, 10)
%3 =
```

```
[10 7 3]
[ 0 1 0]
[ 0 0 1]
```

The library syntax is `GEN hnfmodid(GEN x, GEN d)`.

3.10.22 mathouseholder(Q, v). applies a sequence Q of Householder transforms, as returned by `matqr( $M, 1$ )` to the vector or matrix v .

```
? m = [2,1; 3,2]; \\ some random matrix
? [Q,R] = matqr(m);
? Q
%3 =
[-0.554... -0.832...]
[-0.832... 0.554...]
? R
%4 =
[-3.605... -2.218...]
[0          0.277...]
? v = [1, 2]~; \\ some random vector
? Q * v
%6 = [-2.218..., 0.277...]~
? [q,r] = matqr(m, 1);
? exponent(r - R) \\ r is the same as R
%8 = -128
? q \\ but q has a different structure
%9 = [[0.0494..., [5.605..., 3]]]
? mathouseholder(q, v) \\ applied to v
%10 = [-2.218..., 0.277...]~
```

The point of the Householder structure is that it efficiently represents the linear operator $v \mapsto Qv$ in a more stable way than expanding the matrix Q :

```
? m = mathilbert(20); v = vectorv(20,i,i^2+1);
? [Q,R] = matqr(m);
? [q,r] = matqr(m, 1);
? \p100
? [q2,r2] = matqr(m, 1); \\ recompute at higher accuracy
? exponent(R - r)
%5 = -127
? exponent(R - r2)
%6 = -127
? exponent(mathouseholder(q,v) - mathouseholder(q2,v))
%7 = -119
? exponent(Q*v - mathouseholder(q2,v))
%8 = 9
```

We see that R is OK with or without a flag to `matqr` but that multiplying by Q is considerably less precise than applying the sequence of Householder transforms encoded by q .

The library syntax is `GEN mathouseholder(GEN Q, GEN v)`.

3.10.23 `matid`(n). Creates the $n \times n$ identity matrix.

The library syntax is `GEN matid(long n)`.

3.10.24 `matimage`($x, \{flag = 0\}$). Gives a basis for the image of the matrix x as columns of a matrix. A priori the matrix can have entries of any type. If $flag = 0$, use standard Gauss pivot. If $flag = 1$, use `mat supplement` (much slower: keep the default flag!).

The library syntax is `GEN matimage0(GEN x, long flag)`. Also available is `GEN image(GEN x)` ($flag = 0$).

3.10.25 `matimagecompl`(x). Gives the vector of the column indices which are not extracted by the function `matimage`, as a permutation (`t_VECSMALL`). Hence the number of components of `matimagecompl(x)` plus the number of columns of `matimage(x)` is equal to the number of columns of the matrix x .

The library syntax is `GEN imagecompl(GEN x)`.

3.10.26 `matimagemod`($x, d, \&U$). Gives a Howell basis (unique representation for submodules of $(\mathbf{Z}/d\mathbf{Z})^n$) for the image of the matrix x modulo d as columns of a matrix H . The matrix x must have `t_INT` entries, and d can be an arbitrary positive integer. If U is present, set it to a matrix such that $AU = H$.

```
? A = [2,1;0,2];
? matimagemod(A,6,&U)
%2 =
[1 0]
[0 2]
? U
%3 =
[5 1]
[3 4]
? (A*U)%6
%4 =
[1 0]
[0 2]
```

Caveat. In general the number of columns of the Howell form is not the minimal number of generators of the submodule. Example:

```
? matimagemod([1;2],4)
%5 =
[2 1]
[0 2]
```

Caveat 2. In general the matrix U is not invertible, even if A and H have the same size. Example:

```
? matimagemod([4,1;0,4],8,&U)
%6 =
[2 1]
[0 4]
? U
%7 =
[0 0]
[2 1]
```

The library syntax is GEN `matimagemod(GEN x, GEN d, GEN *U = NULL)`.

3.10.27 matindexrank(M). M being a matrix of rank r , returns a vector with two `t_VECSMALL` components y and z of length r giving a list of rows and columns respectively (starting from 1) such that the extracted matrix obtained from these two vectors using `vecextract( $M, y, z$ )` is invertible. The vectors y and z are sorted in increasing order.

The library syntax is GEN `indexrank(GEN M)`.

3.10.28 matintersect(x, y). x and y being two matrices with the same number of rows, finds a basis of the vector space equal to the intersection of the spaces spanned by the columns of x and y respectively. For efficiency, the columns of x (resp. y) should be independent.

The faster function `idealintersect` can be used to intersect fractional ideals (projective $\mathbf{Z}_K$ modules of rank 1); the slower but more general function `nfhnf` can be used to intersect general $\mathbf{Z}_K$ -modules.

The library syntax is GEN `intersect(GEN x, GEN y)`.

3.10.29 matinverseimage(x, y). Given a matrix x and a column vector or matrix y , returns a preimage z of y by x if one exists (i.e such that $xz = y$), an empty vector or matrix otherwise. The complete inverse image is $z + \text{Ker}x$, where a basis of the kernel of x may be obtained by `matker`.

```
? M = [1,2;2,4];
? matinverseimage(M, [1,2]~)
%2 = [1, 0]~
? matinverseimage(M, [3,4]~)
%3 = []~      \\ no solution
? matinverseimage(M, [1,3,6;2,6,12])
%4 =
[1 3 6]
[0 0 0]
? matinverseimage(M, [1,2;3,4])
%5 = [;]      \\ no solution
? K = matker(M)
%6 =
[-2]
[1]
```

The library syntax is GEN `inverseimage(GEN x, GEN y)`.

3.10.30 `matinvmod`(x, d). Computes a left inverse of the matrix x modulo d . The matrix x must have `t_INT` entries, and d can be an arbitrary positive integer.

```
? A = [3,1,2;1,2,1;3,1,1];
? U = matinvmod(A,6)
%2 =
[1 1 3]
[2 3 5]
[1 0 5]
? (U*A)%6
%3 =
[1 0 0]
[0 1 0]
[0 0 1]
? matinvmod(A,5)
*** at top-level: matinvmod(A,5)
*** ^-----
*** matinvmod: impossible inverse in gen_inv: 0.
```

The library syntax is `GEN matinvmod(GEN x, GEN d)`.

3.10.31 `matisdiagonal`(x). Returns true (1) if x is a diagonal matrix, false (0) if not.

The library syntax is `int isdiagonal(GEN x)`.

3.10.32 `matker`($x, \{flag = 0\}$). Gives a basis for the kernel of the matrix x as columns of a matrix. The matrix can have entries of any type, provided they are compatible with the generic arithmetic operations (+, $\times$ and /).

If x is known to have integral entries, set $flag = 1$.

The library syntax is `GEN matker0(GEN x, long flag)`. Also available are `GEN ker(GEN x)` ($flag = 0$), `GEN ZM_ker(GEN x)` ($flag = 1$).

3.10.33 `matkerint`($x, \{flag = 0\}$). Gives an LLL-reduced $\mathbf{Z}$ -basis for the lattice equal to the kernel of the matrix x with rational entries. $flag$ is deprecated, kept for backward compatibility. The function `matsolvemod` allows to solve more general linear systems over $\mathbf{Z}$.

The library syntax is `GEN matkerint0(GEN x, long flag)`. Use directly `GEN kerint(GEN x)` if x is known to have integer entries, and `Q_primpart` first otherwise.

3.10.34 matkernelmod($x, d, \&im$). Gives a Howell basis (unique representation for submodules of $(\mathbf{Z}/d\mathbf{Z})^n$, cf. **matimagemod**) for the kernel of the matrix x modulo d as columns of a matrix. The matrix x must have `t_INT` entries, and d can be an arbitrary positive integer. If im is present, set it to a basis of the image of x (which is computed on the way).

```
? A = [1,2,3;5,1,4]
%1 =
[1 2 3]
[5 1 4]
? K = matkernelmod(A,6)
%2 =
[2 1]
[2 1]
[0 3]
? (A*K)%6
%3 =
[0 0]
[0 0]
```

The library syntax is `GEN matkernelmod(GEN x, GEN d, GEN *im = NULL)`.

3.10.35 matmuldiagonal(x, d). Product of the matrix x by the diagonal matrix whose diagonal entries are those of the vector d . Equivalent to, but much faster than $x * \text{matdiagonal}(d)$.

The library syntax is `GEN matmuldiagonal(GEN x, GEN d)`.

3.10.36 matmultodiagonal(x, y). Product of the matrices x and y assuming that the result is a diagonal matrix. Much faster than $x * y$ in that case. The result is undefined if $x * y$ is not diagonal.

The library syntax is `GEN matmultodiagonal(GEN x, GEN y)`.

3.10.37 matpascal($n, \{q\}$). Creates as a matrix the lower triangular Pascal triangle of order $x + 1$ (i.e. with binomial coefficients up to x). If q is given, compute the q -Pascal triangle (i.e. using q -binomial coefficients).

The library syntax is `GEN matqpascal(long n, GEN q = NULL)`. Also available is `GEN matpascal(GEN x)`.

3.10.38 matpermanent(x). Permanent of the square matrix x using Ryser's formula in Gray code order.

```
? n = 20; m = matrix(n,n,i,j, i!=j);
? matpermanent(m)
%2 = 895014631192902121
? n! * sum(i=0,n, (-1)^i/i!)
%3 = 895014631192902121
```

This function runs in time $O(2^n n)$ for a matrix of size n and is not implemented for n large.

The library syntax is `GEN matpermanent(GEN x)`.

3.10.39 matqr($M, \{flag = 0\}$). Returns $[Q, R]$, the QR-decomposition of the square invertible matrix M with real entries: Q is orthogonal and R upper triangular. If $flag = 1$, the orthogonal matrix is returned as a sequence of Householder transforms: applying such a sequence is stabler and faster than multiplication by the corresponding Q matrix. More precisely, if

```
[Q,R] = matqr(M);
[q,r] = matqr(M, 1);
```

then $r = R$ and `mathouseholder(q, M)` is (close to) R ; furthermore

```
mathouseholder(q, matid(#M)) == Q~
```

the inverse of Q . This function raises an error if the precision is too low or x is singular.

The library syntax is `GEN matqr(GEN M, long flag, long prec)`.

3.10.40 matrank(x). Rank of the matrix x .

The library syntax is `long rank(GEN x)`.

3.10.41 matreduce(m). Let m be a factorization matrix, i.e., a 2-column matrix whose columns contains arbitrary “generators” and integer “exponents” respectively. Returns the canonical form of m : the first column is sorted with unique elements and the second one contains the merged “exponents” (exponents of identical entries in the first column of m are added, rows attached to 0 exponents are deleted). The generators are sorted with respect to the universal `cmp` routine; in particular, this function is the identity on true integer factorization matrices, but not on other factorizations (in products of polynomials or maximal ideals, say). It is idempotent.

For convenience, this function also allows a vector m , which is handled as a factorization with all exponents equal to 1, as in `factorback`.

```
? A=[x,2;y,4]; B=[x,-2; y,3; 3, 4]; C=matconcat([A,B]~)
%1 =
[x  2]
[y  4]
[x -2]
[y  3]
[3  4]

? matreduce(C)
%2 =
[3 4]
[y 7]

? matreduce([x,x,y,x,z,x,y]) \\ vector argument
%3 =
[x 4]
[y 2]
[z 1]
```

The following one-line functions will list elements occurring exactly once (resp. more than once) in the vector or list v :

```

unique(v) = [ x[1] | x <- matreduce(v)~, x[2] == 1 ];
duplicates(v) = [ x[1] | x <- matreduce(v)~, x[2] > 1 ];
? v = [0,1,2,3,1,2];
? unique(v)
%2 = [0, 3]
? duplicates(v)
%3 = [1, 2]

```

The library syntax is `GEN matreduce(GEN m)`.

3.10.42 matrix($m, \{n = m\}, \{X\}, \{Y\}, \{expr = 0\}$). Creation of the $m \times n$ matrix whose coefficients are given by the expression *expr*. There are two formal parameters in *expr*, the first one (*X*) corresponding to the rows, the second (*Y*) to the columns, and *X* goes from 1 to *m*, *Y* goes from 1 to *n*. If one of the last 3 parameters is omitted, fill the matrix with zeroes. If *n* is omitted, return a square $m \times m$ matrix.

3.10.43 matrixqz($A, \{p = 0\}$). *A* being an $m \times n$ matrix in $M_{m,n}(\mathbf{Q})$, let $\text{Im}_{\mathbf{Q}}A$ (resp. $\text{Im}_{\mathbf{Z}}A$) the $\mathbf{Q}$ -vector space (resp. the $\mathbf{Z}$ -module) spanned by the columns of *A*. This function has varying behavior depending on the sign of *p*:

If $p \geq 0$, *A* is assumed to have maximal rank $n \leq m$. The function returns a matrix $B \in M_{m,n}(\mathbf{Z})$, with $\text{Im}_{\mathbf{Q}}B = \text{Im}_{\mathbf{Q}}A$, such that the GCD of all its $n \times n$ minors is coprime to *p*; in particular, if $p = 0$ (default), this GCD is 1.

If $p = -1$, returns a basis of the lattice $\mathbf{Z}^m \cap \text{Im}_{\mathbf{Z}}A$.

If $p = -2$, returns a basis of the lattice $\mathbf{Z}^m \cap \text{Im}_{\mathbf{Q}}A$.

Caveat. ($p = -1$ or -2) For efficiency reason, we do not compute the HNF of the resulting basis.

```

? minors(x) = vector(#x[,1], i, matdet(x[~i,]));
? A = [3,1/7; 5,3/7; 7,5/7]; minors(A)
%1 = [4/7, 8/7, 4/7]  \\ determinants of all 2x2 minors
? B = matrixqz(A)
%2 =
[3 1]
[5 2]
[7 3]
? minors(%)
%3 = [1, 2, 1]  \\ B integral with coprime minors
? matrixqz(A,-1)
%4 =
[3 1]
[5 3]
[7 5]
? matrixqz(A,-2)
%5 =
[3 1]
[5 2]

```

[7 3]

The library syntax is `GEN matrixqz0(GEN A, GEN p = NULL)`.

3.10.44 `matsize`(x). x being a vector or matrix, returns a row vector with two components, the first being the number of rows (1 for a row vector), the second the number of columns (1 for a column vector).

The library syntax is `GEN matsize(GEN x)`.

3.10.45 `matsnf`($X, \{flag = 0\}$). If X is a (singular or nonsingular) matrix outputs the vector of elementary divisors of X , i.e. the diagonal of the Smith normal form of X , normalized so that $d_n \mid d_{n-1} \mid \dots \mid d_1$. X must have integer or polynomial entries; in the latter case, X must be a square matrix.

The binary digits of *flag* mean:

1 (complete output): if set, outputs $[U, V, D]$, where U and V are two unimodular matrices such that UXV is the diagonal matrix D . Otherwise output only the diagonal of D . If X is not a square matrix, then D will be a square diagonal matrix padded with zeros on the left or the top.

4 (cleanup): if set, cleans up the output. This means that elementary divisors equal to 1 will be deleted, i.e. outputs a shortened vector D' instead of D . If complete output was required, returns $[U', V', D']$ so that $U'XV' = D'$ holds. If this flag is set, X is allowed to be of the form ‘vector of elementary divisors’ or $[U, V, D]$ as would normally be output with the cleanup flag unset.

If v is an output from `matsnf` and p is a power of an irreducible element, then `snfrank(v, p)` returns the p -rank of the attached module.

```
? X = [27,0; 0,3; 1,1; 0,0]; matsnf(X)
%1 = [0, 0, 3, 1]
? [U,V,D] = v = matsnf(X, 1); U*X*V == D
%2
? U
%3 =
[0 0 0 1]
[1 9 -27 0]
[0 1 0 0]
[0 0 1 0]
? V
%4 =
[-1 1]
[ 1 0]
? snfrank(v, 3)
%5 = 3
```

Continuing the same example after cleanup:

```
? [U,V,D] = v = matsnf(X, 1+4); U*X*V == D
%6 = 1
? D
```

```

%7 =
[0]
[0]
[3]
? snfrank(v, 3)
%8 = 3
? snfrank(v, 2)
%9 = 2

```

The library syntax is `GEN matsnf0(GEN X, long flag)`.

3.10.46 matsolve(M, B). Let M be a left-invertible matrix and B a column vector such that there exists a solution X to the system of linear equations $MX = B$; return the (unique) solution X . This has the same effect as, but is faster, than $M^{-1} * B$. Uses Dixon p -adic lifting method if M and B are integral and Gaussian elimination otherwise. When there is no solution, the function returns an X such that $MX - B$ is nonzero although it has at least $\#M$ zero entries:

```

? M = [1,2;3,4;5,6];
? B = [4,6,8]~; X = matsolve(M, B)
%2 = [-2, 3]~
? M*X == B
%3 = 1
? B = [1,2,4]~; X = matsolve(M, [1,2,4]~)
%4 = [0, 1/2]~
? M*X - B
%5 = [0, 0, -1]~

```

Raises an exception if M is not left-invertible, even if there is a solution:

```

? M = [1,1;1,1]; matsolve(M, [1,1]~)
*** at top-level: matsolve(M,[1,1]~)
***      ^-----
*** matsolve: impossible inverse in gauss: [1, 1; 1, 1].

```

The function also works when B is a matrix and we return the unique matrix solution X provided it exists. Again, if there is no solution, the function returns an X such that $MX - B$ is nonzero although it has at least $\#M$ zero rows.

The library syntax is `GEN gauss(GEN M, GEN B)`.

3.10.47 matsolvemod($M, D, B, \{flag = 0\}$). M being any integral matrix, D a column vector of nonnegative integer moduli, and B an integral column vector, gives an integer solution to the system of congruences $\sum_i m_{i,j} x_j \equiv b_i \pmod{d_i}$ if one exists, otherwise returns the integer zero. Note that we explicitly allow $d_i = 0$ corresponding to an equality in $\mathbf{Z}$. Shorthand notation: B (resp. D) can be given as a single integer, in which case all the b_i (resp. d_i) above are taken to be equal to B (resp. D). Again, $D = 0$ solves the linear system of equations over $\mathbf{Z}$.

```

? M = [1,2;3,4];
? matsolvemod(M, [3,4]~, [1,2]~)
%2 = [10, 0]~
? matsolvemod(M, 3, 1) \\ M X = [1,1]~ over F_3

```

```
%3 = [2, 1]~
? matsolvemod(M, [3,0]~, [1,2]~) \\ x + 2y = 1 (mod 3), 3x + 4y = 2 (in Z)
%4 = [6, -4]~
? matsolvemod(M, 0, [1,2]~) \\ no solution in Z for x + 2y = 1, 3x + 4y = 2
```

If $flag = 1$, all solutions are returned in the form of a two-component row vector $[x, u]$, where x is an integer solution to the system of congruences and u is a matrix whose columns give a basis of the homogeneous system (so that all solutions can be obtained by adding x to any linear combination of columns of u). If no solution exists, returns zero.

The library syntax is GEN `matsolvemod`(GEN M, GEN D, GEN B, long flag). Also available are GEN `gaussmodulo`(GEN M, GEN D, GEN B) ($flag = 0$) and GEN `gaussmodulo2`(GEN M, GEN D, GEN B) ($flag = 1$).

3.10.48 matsupplement(x). Assuming that the columns of the matrix x are linearly independent (if they are not, an error message is issued), finds a square invertible matrix whose first columns are the columns of x , i.e. supplement the columns of x to a basis of the whole space.

```
? matsupplement([1;2])
%1 =
[1 0]
[2 1]
```

Raises an error if x has 0 columns, since (due to a long standing design bug), the dimension of the ambient space (the number of rows) is unknown in this case:

```
? matsupplement(matrix(2,0))
*** at top-level: matsupplement(matrix
*** ^-----
*** matsupplement: sorry, suppl [empty matrix] is not yet implemented.
```

The library syntax is GEN `suppl`(GEN x).

3.10.49 mattranspose(x). Transpose of x (also $x^~$). This has an effect only on vectors and matrices.

The library syntax is GEN `gtrans`(GEN x).

3.10.50 minpoly($A, \{v = 'x\}$). minimal polynomial of A with respect to the variable v , i.e. the monic polynomial P of minimal degree (in the variable v) such that $P(A) = 0$.

The library syntax is GEN `minpoly`(GEN A, long v = -1) where v is a variable number.

3.10.51 norml2(x). Square of the L^2 -norm of x . More precisely, if x is a scalar, `norml2( $x$ )` is defined to be the square of the complex modulus of x (real `t_QUAD`s are not supported). If x is a polynomial, a (row or column) vector or a matrix, `norml2( $x$ )` is defined recursively as $\sum_i \text{norml2}(x_i)$, where (x_i) run through the components of x . In particular, this yields the usual $\sum_i |x_i|^2$ (resp. $\sum_{i,j} |x_{i,j}|^2$) if x is a polynomial or vector (resp. matrix) with complex components.

```
? norml2( [ 1, 2, 3 ] )      \\ vector
%1 = 14
? norml2( [ 1, 2; 3, 4 ] )    \\ matrix
%2 = 30
? norml2( 2*I + x )
%3 = 5
? norml2( [ [1,2], [3,4], 5, 6 ] )  \\ recursively defined
%4 = 91
```

The library syntax is GEN `gnorml2(GEN x)`.

3.10.52 normlp($x, \{p = oo\}$). L^p -norm of x ; sup norm if p is omitted or `+oo`. More precisely, if x is a scalar, `normlp( $x, p$ )` is defined to be `abs( $x$ )`. If x is a polynomial, a (row or column) vector or a matrix:

- if p is omitted or `+oo`, then `normlp( $x$ )` is defined recursively as $\max_i \text{normlp}(x_i)$, where x_i runs through the components of x . In particular, this yields the usual sup norm if x is a polynomial or vector with complex components.

- otherwise, `normlp( $x, p$ )` is defined recursively as $(\sum_i \text{normlp}^p(x_i, p))^{1/p}$. In particular, this yields the usual $(\sum_i |x_i|^p)^{1/p}$ if x is a polynomial or vector with complex components.

```
? v = [1,-2,3]; normlp(v)      \\ vector
%1 = 3
? normlp(v, +oo)                \\ same, more explicit
%2 = 3
? M = [1,-2;-3,4]; normlp(M)    \\ matrix
%3 = 4
? T = (1+I) + I*x^2; normlp(T)
%4 = 1.4142135623730950488016887242096980786
? normlp([[1,2], [3,4], 5, 6])  \\ recursively defined
%5 = 6

? normlp(v, 1)
%6 = 6
? normlp(M, 1)
%7 = 10
? normlp(T, 1)
%8 = 2.4142135623730950488016887242096980786
```

The library syntax is GEN `gnormlp(GEN x, GEN p = NULL, long prec)`.

3.10.53 powers($x, n, \{x0\}$). For nonnegative n , return the vector with $n + 1$ components $[1, x, \dots, x^n]$ if $x0$ is omitted, and $[x_0, x_0 * x, \dots, x_0 * x^n]$ otherwise.

```
? powers(Mod(3,17), 4)
%1 = [Mod(1, 17), Mod(3, 17), Mod(9, 17), Mod(10, 17), Mod(13, 17)]
? powers(Mat([1,2;3,4]), 3)
%2 = [[1, 0; 0, 1], [1, 2; 3, 4], [7, 10; 15, 22], [37, 54; 81, 118]]
? powers(3, 5, 2)
%3 = [2, 6, 18, 54, 162, 486]
```

When $n < 0$, the function returns the empty vector `[]`.

The library syntax is `GEN gpowers0(GEN x, long n, GEN x0 = NULL)`. Also available is `GEN gpowers(GEN x, long n)` when $x0$ is `NULL`.

3.10.54 qfauto($G, \{fl\}$). G being a square and symmetric matrix with integer entries representing a positive definite quadratic form, outputs the automorphism group of the associate lattice. Since this requires computing the minimal vectors, the computations can become very lengthy as the dimension grows. G can also be given by an `qfisominit` structure. See `qfisominit` for the meaning of fl .

The output is a two-components vector $[o, g]$ where o is the group order and g is the list of generators (as a vector). For each generator H , the equality $G = {}^tHGH$ holds.

The interface of this function is experimental and will likely change in the future.

This function implements an algorithm of Plesken and Souvignier, following Souvignier's implementation.

```
? K = matkerint(Mat(concat([vector(23,i,2*i+1), 51, 145])));
? M = matdiagonal(vector(25,i,if(i==25,-1,1)));
? L24 = K~ * M * K; \\ the Leech lattice
? [o,g] = qfauto(L24); o
%4 = 8315553613086720000
? #g
%5 = 2
```

The library syntax is `GEN qfauto0(GEN G, GEN fl = NULL)`. The function `GEN qfauto(GEN G, GEN fl)` is also available where G is a vector of `zm` matrices.

3.10.55 qfautoexport($qfa, \{flag\}$). qfa being an automorphism group as output by `qfauto`, export the underlying matrix group as a string suitable for (no flags or $flag = 0$) GAP or ($flag = 1$) Magma. The following example computes the size of the matrix group using GAP:

```
? G = qfauto([2,1;1,2])
%1 = [12, [[-1, 0; 0, -1], [0, -1; 1, 1], [1, 1; 0, -1]]]
? s = qfautoexport(G)
%2 = "Group([[ -1, 0], [0, -1]], [[0, -1], [1, 1]], [[1, 1], [0, -1]])"
? extern("echo \"Order(\"s\")\" | gap -q")
%3 = 12
```

The library syntax is `GEN qfautoexport(GEN qfa, long flag)`.

3.10.56 qfbil($x, y, \{q\}$). This function is obsolete, use `qfeval`.

The library syntax is `GEN qfbil(GEN x, GEN y, GEN q = NULL)`.

3.10.57 qfcholesky(q). Given a square symmetric `t_MAT` M , return R such that ${}^tRR = M$, or `[]` if there is no solution.

The library syntax is `GEN qfcholesky(GEN q, long prec)`.

3.10.58 qfcvp($x, t, \{B\}, \{m\}, \{flag = 0\}$). x being a square and symmetric matrix of dimension d representing a positive definite quadratic form, and t a vector of the same dimension d . This function deals with the vectors whose squared distance to t is less than B , enumerated using the Fincke-Pohst algorithm, storing at most m vectors. There is no limit if m is omitted: beware that this may be a huge vector! The vectors are returned in no particular order.

The function searches for the closest vectors to t if B is omitted or ≤ 0 . The behavior is undefined if x is not positive definite (a “precision too low” error is most likely, although more precise error messages are possible). The precise behavior depends on *flag*.

- If *flag* = 0 (default), return $[N, M, V]$, where N is the number of vectors enumerated (possibly larger than m), $M \leq B$ is the maximum squared distance found, and V is a matrix whose columns are found vectors.

- If *flag* = 1, ignore m and return $[M, v]$, where v is a vector at squared distance $M \leq B$. If no vector has distance $\leq B$, return `[]`.

In these two cases, x must have integral *small* entries: more precisely, we definitely must have $d \cdot \|x\|_\infty^2 < 2^{53}$ but even that may not be enough. The implementation uses low precision floating point computations for maximal speed and gives incorrect results when x has large entries. That condition is checked in the code and the routine raises an error if large rounding errors occur.

```
? M = [2,1;1,2]; t = [1/2, -1/2];
? qfcvp(M, t, 0)
%2 = [2, 0.50000000000000000000, [0, 1; 0, -1]]
? qfcvp(M, t, 1.5)
%3 = [4, 1.50000000000000000000, [1, 0, 1, 0; 0, 0, -1, -1]]
```

The library syntax is `GEN qfcvp0(GEN x, GEN t, GEN B = NULL, GEN m = NULL, long flag)`.

3.10.59 qfeval($\{q\}, x, \{y\}$). Evaluate the quadratic form q (given by a symmetric matrix) at the vector x ; if y is present, evaluate the polar form at (x, y) ; if q omitted, use the standard Euclidean scalar product, corresponding to the identity matrix.

Roughly equivalent to `x~* q * y`, but a little faster and more convenient (does not distinguish between column and row vectors):

```
? x = [1,2,3]~; y = [-1,3,1]~; q = [1,2,3;2,2,-1;3,-1,9];
? qfeval(q,x,y)
%2 = 23
? for(i=1,10^6, qfeval(q,x,y))
time = 661ms
? for(i=1,10^6, x~*q*y)
time = 697ms
```

The speedup is noticeable for the quadratic form, compared to $x \sim q * x$, since we save almost half the operations:

```
? for(i=1,10^6, qfeval(q,x))
time = 487ms
```

The special case $q = \text{Id}$ is handled faster if we omit q altogether:

```
? qfeval(,x,y)
%6 = 8
? q = matid(#x);
? for(i=1,10^6, qfeval(q,x,y))
time = 529 ms.
? for(i=1,10^6, qfeval(,x,y))
time = 228 ms.
? for(i=1,10^6, x~*y)
time = 274 ms.
```

We also allow t_MAT s of compatible dimensions for x , and return $x \sim q * x$ in this case as well:

```
? M = [1,2,3;4,5,6;7,8,9]; qfeval(,M) \\ Gram matrix
%5 =
[66 78 90]
[78 93 108]
[90 108 126]
? q = [1,2,3;2,2,-1;3,-1,9];
? for(i=1,10^6, qfeval(q,M))
time = 2,008 ms.
? for(i=1,10^6, M~*q*M)
time = 2,368 ms.
? for(i=1,10^6, qfeval(,M))
time = 1,053 ms.
? for(i=1,10^6, M~*M)
time = 1,171 ms.
```

If q is a t_QFB , it is implicitly converted to the attached symmetric t_MAT . This is done more efficiently than by direct conversion, since we avoid introducing a denominator 2 and rational arithmetic:

```
? q = Qfb(2,3,4); x = [2,3];
? qfeval(q, x)
%2 = 62
? Q = Mat(q)
%3 =
[ 2 3/2]
[3/2 4]
? qfeval(Q, x)
%4 = 62
? for (i=1, 10^6, qfeval(q,x))
time = 758 ms.
```

```
? for (i=1, 10^6, qfeval(Q,x))
time = 1,110 ms.
```

Finally, when x is a `t_MAT` with *integral* coefficients, we allow a `t_QFB` for q and return the binary quadratic form $q \circ M$. Again, the conversion to `t_MAT` is less efficient in this case:

```
? q = Qfb(2,3,4); Q = Mat(q); x = [1,2;3,4];
? qfeval(q, x)
%2 = Qfb(47, 134, 96)
? qfeval(Q,x)
%3 =
[47 67]
[67 96]
? for (i=1, 10^6, qfeval(q,x))
time = 701 ms.
? for (i=1, 10^6, qfeval(Q,x))
time = 1,639 ms.
```

The library syntax is `GEN qfeval0(GEN q = NULL, GEN x, GEN y = NULL)`.

3.10.60 qfgaussred($q, \{flag = 0\}$). decomposition into squares of the quadratic form represented by the symmetric matrix q . If $flag = 0$ (default), the result is a matrix M whose diagonal entries are the coefficients of the squares, and the off-diagonal entries on each line represent the bilinear forms. More precisely, if (a_{ij}) denotes the output, one has

$$q(x) = \sum_i a_{i,i}(x_i + \sum_{j \neq i} a_{i,j}x_j)^2$$

```
? qfgaussred([0,1;1,0])
%1 =
[1/2 1]
[-1 -1/2]
```

This means that $2xy = (1/2)(x+y)^2 - (1/2)(x-y)^2$. Singular matrices are supported, in which case some diagonal coefficients vanish:

```
? qfgaussred([1,1;1,1])
%2 =
[1 1]
[1 0]
```

This means that $x^2 + 2xy + y^2 = (x+y)^2$.

If $flag = 1$, return $[U,V]$ where U is a square matrix and V a vector, such that if $D=\text{matdiagonal}(V)$, $q = {}^tUDU$. More precisely

$$q(x) = \sum_i D_i (\sum_j U_{i,j}x_j)^2$$

and the matrix M is recovered as $M = U + D - 1$.

```
? q = [0,1;1,0];
```

```

? [U,V] = qfgaussred(q,1); D = matdiagonal(V);
? U~*D*U
%5 =
[0 1]
[1 0]
? U+D-1
%6 =
[1/2    1]
[ -1 -1/2]

```

The library syntax is `GEN qfgaussred0(GEN q, long flag)`. See also the functions `GEN qfgaussred(GEN a)` (for `qfgaussred(a,0)`), `GEN qfgaussred2(GEN a)` (for `qfgaussred0(a,1)`). Finally, the function `GEN qfgaussred_positive(GEN q)` assumes that q is positive definite and is a little faster; returns `NULL` if a vector with negative norm occurs (non positive matrix or too many rounding errors).

3.10.61 `qfisom`($G, H, \{fl\}, \{grp\}$). G, H being square and symmetric matrices with integer entries representing positive definite quadratic forms, return an invertible matrix S such that $G = {}^tSH \times S$. This defines an isomorphism between the corresponding lattices. Since this requires computing the minimal vectors, the computations can become very lengthy as the dimension grows. See `qfisominit` for the meaning of fl . If grp is given it must be the automorphism group of H . It will be used to speed up the computation.

G can also be given by an `qfisominit` structure which is preferable if several forms H need to be compared to G .

This function implements an algorithm of Plesken and Souvignier, following Souvignier's implementation.

The library syntax is `GEN qfisom0(GEN G, GEN H, GEN fl = NULL, GEN grp = NULL)`. Also available is `GEN qfisom(GEN G, GEN H, GEN fl, GEN grp)` where G is a vector of `zm`, and H is a `zm`, and grp is either `NULL` or a vector of `zm`.

3.10.62 `qfisominit`($G, \{fl\}, \{m\}$). G being a square and symmetric matrix with integer entries representing a positive definite quadratic form, return an `isom` structure allowing to compute isomorphisms between G and other quadratic forms faster.

The interface of this function is experimental and will likely change in future release.

If present, the optional parameter fl must be a `t_VEC` with two components. It allows to specify the invariants used, which can make the computation faster or slower. The components are

- `fl[1]` Depth of scalar product combination to use.
- `fl[2]` Maximum level of Bacher polynomials to use.

If present, m must be the set of vectors of norm up to the maximal of the diagonal entry of G , either as a matrix or as given by `qfminim`. Otherwise this function computes the minimal vectors so it become very lengthy as the dimension of G grows.

The library syntax is `GEN qfisominit0(GEN G, GEN fl = NULL, GEN m = NULL)`. Also available is `GEN qfisominit(GEN F, GEN fl)` where F is a vector of `zm`.

3.10.63 qfjacobi(A). Apply Jacobi’s eigenvalue algorithm to the real symmetric matrix A . This returns $[L, V]$, where

- L is the vector of (real) eigenvalues of A , sorted in increasing order,
- V is the corresponding orthogonal matrix of eigenvectors of A .

```
? \p19
? A = [1,2;2,1]; mateigen(A)
%1 =
[-1 1]
[ 1 1]
? [L, H] = qfjacobi(A);
? L
%3 = [-1.000000000000000000, 3.000000000000000000]~
? H
%4 =
[ 0.7071067811865475245 0.7071067811865475244]
[-0.7071067811865475244 0.7071067811865475245]
? norml2( (A-L[1])*H[,1] )      \\ approximate eigenvector
%5 = 9.403954806578300064 E-38
? norml2(H*H~ - 1)
%6 = 2.350988701644575016 E-38    \\ close to orthogonal
```

The library syntax is `GEN jacobi(GEN A, long prec)`.

3.10.64 qflll($x, \{flag = 0\}$). LLL algorithm applied to the *columns* of the matrix x . The columns of x may be linearly dependent. The result is by default a unimodular transformation matrix T such that $x \cdot T$ is an LLL-reduced basis of the lattice generated by the column vectors of x . Note that if x is not of maximal rank T will not be square. The LLL parameters are $(0.51, 0.99)$, meaning that the Gram-Schmidt coefficients for the final basis satisfy $|\mu_{i,j}| \leq 0.51$, and the Lovász’s constant is 0.99.

If $flag = 0$ (default), assume that x has either exact (integral or rational) or real floating point entries. The matrix is rescaled, converted to integers and the behavior is then as in $flag = 1$. Computations involving Gram-Schmidt vectors are approximate, with precision varying as needed (Lehmer’s trick, as generalized by Schnorr). Adapted from Nguyen and Stehlé’s algorithm and Stehlé’s code (`fp111-1.3`) as building blocks for the FLATTER (block recursive) algorithm of Heninger and Ryan.

If $flag = 1$, disable use of FLATTER algorithm; use `fp111`. This flag is provided to experiment with the concrete speed-ups allowed by FLATTER, as well as to genuinely disable it on the rare classes of lattices for which it turns out it performs badly: many such classes are detected in the code, which then restricts to stock `fp111`, but new examples may turn up.

If $flag = 2$, x should be an integer matrix whose columns are linearly independent. Returns a partially reduced basis for x , using an unpublished algorithm by Peter Montgomery: a basis is said to be *partially reduced* if $|v_i \pm v_j| \geq |v_i|$ for any two distinct basis vectors v_i, v_j . This is faster than $flag = 1$, esp. when one row is huge compared to the other rows (knapsack-style), and should quickly produce relatively short vectors. The resulting basis is *not* LLL-reduced in general. If LLL reduction is eventually desired, avoid this partial reduction: applying LLL to the partially reduced matrix is significantly *slower* than starting from a knapsack-type lattice.

If $flag = 3$, as $flag = 0$, but the reduction is performed in place: the routine returns $x \cdot T$. This is usually faster for knapsack-type lattices.

If $flag = 4$, as $flag = 0$, returning a vector $[K, T]$ of matrices: the columns of K represent a basis of the integer kernel of x (not LLL-reduced in general) and T is the transformation matrix such that $x \cdot T$ is an LLL-reduced $\mathbf{Z}$ -basis of the image of the matrix x .

If $flag = 5$, case as $flag = 4$, but x may have polynomial coefficients.

If $flag = 8$, same as $flag = 0$, but x may have polynomial coefficients.

```
? \p500
  realprecision = 500 significant digits
? a = 2*cos(2*Pi/97);
? C = 10^450;
? v = powers(a,48); b = round(matconcat([matid(48),C*v]~));
? p = b * qflll(b)[,1]; \\ tiny linear combination of powers of 'a'
  time = 4,470 ms.
? exponent(v * p / C)
%5 = -1418
? p3 = qflll(b,3)[,1]; \\ compute in place, faster
  time = 3,790 ms.
? p3 == p \\ same result
%7 = 1
? p2 = b * qflll(b,2)[,1]; \\ partial reduction: faster, not as good
  time = 343 ms.
? exponent(v * p2 / C)
%9 = -1190
```

The library syntax is `GEN qflll0(GEN x, long flag)`. Also available are `GEN lll(GEN x)` ($flag = 0$), `GEN lllint(GEN x)` ($flag = 1$), and `GEN lllkerim(GEN x)` ($flag = 4$).

3.10.65 qflllgram($G, \{flag = 0\}$). Same as `qflll`, except that the matrix $G = x \sim x$ is the Gram matrix of some lattice vectors x , and not the coordinates of the vectors themselves. In particular, G must now be a square symmetric real matrix, corresponding to a positive quadratic form (not necessarily definite: x needs not have maximal rank). The result is a unimodular transformation matrix T such that $x \cdot T$ is an LLL-reduced basis of the lattice generated by the column vectors of x . See `qflll` for further details about the LLL implementation.

If $flag = 0$ (default), assume that G has either exact (integral or rational) or real floating point entries. The matrix is rescaled, converted to integers. Computations involving Gram-Schmidt vectors are approximate, with precision varying as needed (Lehmer's trick, as generalized by Schnorr). Adapted from Nguyen and Stehlé's algorithm and Stehlé's code (`fp111-1.3`) and FLATTER algorithm for Heninger and Ryan.

If $flag = 1$, disable use of FLATTER algorithm.

$flag = 4$: G has integer entries, gives the kernel and reduced image of x .

$flag = 5$: same as 4, but G may have polynomial coefficients.

The library syntax is `GEN qflllgram0(GEN G, long flag)`. Also available are `GEN lllgram(GEN G)` ($flag = 0$), `GEN lllgramint(GEN G)` ($flag = 1$), and `GEN lllgramkerim(GEN G)` ($flag = 4$).

3.10.66 qfminim($x, \{B\}, \{m\}, \{flag = 0\}$). x being a square and symmetric matrix of dimension d representing a positive definite quadratic form, this function deals with the vectors of x whose norm is less than or equal to B , enumerated using the Fincke-Pohst algorithm, storing at most m pairs of vectors: only one vector is given for each pair $\pm v$. There is no limit if m is omitted: beware that this may be a huge vector! The vectors are returned in no particular order.

The function searches for the minimal nonzero vectors if B is omitted. The behavior is undefined if x is not positive definite (a “precision too low” error is most likely, although more precise error messages are possible). The precise behavior depends on *flag*.

- If *flag* = 0 (default), return $[N, M, V]$, where N is the number of vectors enumerated (an even number, possibly larger than $2m$), $M \leq B$ is the maximum norm found, and V is a matrix whose columns are found vectors.

- If *flag* = 1, ignore m and return $[M, v]$, where v is a nonzero vector of length $M \leq B$. If no nonzero vector has length $\leq B$, return $[]$. If no explicit B is provided, return a vector of smallish norm, namely the vector of smallest length (usually the first one but not always) in an LLL-reduced basis for x .

In these two cases, x must have integral *small* entries: more precisely, we definitely must have $d \cdot \|x\|_\infty^2 < 2^{53}$ but even that may not be enough. The implementation uses low precision floating point computations for maximal speed and gives incorrect results when x has large entries. That condition is checked in the code and the routine raises an error if large rounding errors occur. A more robust, but much slower, implementation is chosen if the following flag is used:

- If *flag* = 2, x can have non integral real entries, but this is also useful when x has large integral entries. Return $[N, M, V]$ as in case *flag* = 0, where M is returned as a floating point number. If x is inexact and B is omitted, the “minimal” vectors in V only have approximately the same norm (up to the internal working accuracy). This version is very robust but still offers no hard and fast guarantee about the result: it involves floating point operations performed at a high floating point precision depending on your input, but done without rigorous tracking of roundoff errors (as would be provided by interval arithmetic for instance). No example is known where the input is exact but the function returns a wrong result.

```
? x = matid(2);
? qfminim(x)  \\ 4 minimal vectors of norm 1: ±[0,1], ±[1,0]
%2 = [4, 1, [0, 1; 1, 0]]
? { x = \\ The Leech lattice
[4, 2, 0, 0, 0,-2, 0, 0, 0, 0, 0, 0, 0, 1,-1, 0, 0, 0, 1, 0,-1, 0, 0, 0,-2;
 2, 4,-2,-2, 0,-2, 0, 0, 0, 0, 0, 0, 0, 0,-1, 0, 0, 0, 0,-1, 0, 1,-1,-1;
 0,-2, 4, 0,-2, 0, 0, 0, 0, 0, 0, 0, 0, 0,-1, 1, 0, 0, 1, 0, 0, 1,-1,-1, 0, 0;
 0,-2, 0, 4, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 1,-1, 0, 0, 0, 1,-1, 0, 1,-1, 1, 0;
 0, 0,-2, 0, 4, 0, 0, 0, 1,-1, 0, 0, 1, 0, 0, 0, 0,-2, 0, 0,-1, 1, 1, 0, 0;
-2, -2,0, 0, 0, 4,-2, 0,-1, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 1, 1, 0,-1, 1, 1;
 0, 0, 0, 0, 0,-2, 4,-2, 0, 0, 0, 0, 0, 1, 0, 0, 0,-1, 0, 0, 0, 1,-1, 0;
 0, 0, 0, 0, 0, 0,-2, 4, 0, 0, 0, 0, 0,-1, 0, 0, 0, 0,-1,-1,-1, 0, 1, 0;
 0, 0, 0, 0, 1,-1, 0, 0, 4, 0,-2, 0, 1, 1, 0,-1, 0, 1, 0, 0, 0, 0, 0, 0, 0;
 0, 0, 0, 0,-1, 0, 0, 0, 0, 4, 0, 0, 1, 1,-1, 1, 0, 0, 0, 1, 0, 0, 1, 0, 0;
 0, 0, 0, 0, 0, 0, 0, 0,-2, 0, 4,-2, 0,-1, 0, 0, 0,-1, 0,-1, 0, 0, 0, 0;
 0, 0, 0, 0, 0, 0, 0, 0, 0, 0,-2, 4,-1, 1, 0, 0,-1, 1, 0, 1, 1, 1,-1, 0;
 1, 0,-1, 1, 1, 0, 0,-1, 1, 1, 0,-1, 4, 0, 0, 1, 0, 1, 1, 0, 1, 0, 1,-1;
-1,-1, 1,-1, 0, 0, 1, 0, 1, 1,-1, 1, 0, 4, 1, 1, 0, 0, 1, 1, 0, 1, 0, 1;
```

```

0, 0, 0, 0, 0, 0, 0, 0, 0, 0, -1, 0, 0, 0, 1, 4, 0, 0, 0, 1, 0, 0, 0, 0, 0;
0, 0, 0, 0, 0, 0, 0, 0, 0, -1, 1, 0, 0, 1, 1, 0, 4, 0, 0, 0, 0, 1, 1, 0, 0;
0, 0, 1, 0, -2, 0, 0, 0, 0, 0, 0, -1, 0, 0, 0, 0, 4, 1, 1, 1, 0, 0, 1, 1;
1, 0, 0, 1, 0, 0, -1, 0, 1, 0, -1, 1, 1, 0, 0, 0, 1, 4, 0, 1, 1, 0, 1, 0;
0, 0, 0, -1, 0, 1, 0, -1, 0, 0, 0, 0, 1, 1, 1, 0, 1, 0, 4, 0, 1, 1, 0, 1;
-1, -1, 1, 0, -1, 1, 0, -1, 0, 1, -1, 1, 0, 1, 0, 0, 1, 1, 0, 4, 0, 0, 1, 1;
0, 0, -1, 1, 1, 0, 0, -1, 0, 0, 0, 1, 1, 0, 0, 1, 0, 1, 1, 0, 4, 1, 0, 1;
0, 1, -1, -1, 1, -1, 1, 0, 0, 0, 0, 1, 0, 1, 0, 1, 0, 0, 1, 0, 1, 4, 0, 1;
0, -1, 0, 1, 0, 1, -1, 1, 0, 1, 0, -1, 1, 0, 0, 0, 1, 1, 0, 1, 0, 0, 4, 1;
-2, -1, 0, 0, 0, 1, 0, 0, 0, 0, 0, 0, -1, 1, 0, 0, 1, 0, 1, 1, 1, 1, 1, 4]; }
? qfminim(x,,0) \\ 0: don't store minimal vectors
time = 121 ms.
%4 = [196560, 4, [;]] \\ 196560 minimal vectors of norm 4
? qfminim(x) \\ store all minimal vectors !
time = 821 ms.
? qfminim(x,,0,2); \\ safe algorithm. Slower and unnecessary here.
time = 5,540 ms.
%6 = [196560, 4.000061035156250000, [;]]
? qfminim(x,,2); \\ safe algorithm; store all minimal vectors
time = 6,602 ms.

```

In this example, storing 0 vectors limits memory use; storing all of them requires a **parisize** about 50MB. All minimal vectors are nevertheless enumerated in both cases of course, which means the speedup is likely to be marginal.

The library syntax is GEN `qfminim0`(GEN `x`, GEN `B` = NULL, GEN `m` = NULL, long `flag`, long `prec`). Also available are GEN `minim`(GEN `x`, GEN `B` = NULL, GEN `m` = NULL) (*flag* = 0), GEN `minim2`(GEN `x`, GEN `B` = NULL, GEN `m` = NULL) (*flag* = 1). GEN `minim_raw`(GEN `x`, GEN `B` = NULL, GEN `m` = NULL) (do not perform LLL reduction on `x` and return NULL on accuracy error). GEN `minim_zm`(GEN `x`, GEN `B` = NULL, GEN `m` = NULL) (*flag* = 0, return vectors as `t_VECSMALL` to save memory)

3.10.67 qfminimize(*G*). Given a square symmetric matrix *G* with rational coefficients, and non-zero determinant, of dimension $n \geq 1$, return [*H*,*U*,*c*] such that $H = c \cdot U^{-1} \cdot G \cdot U$ for some rational *c*, and *H* integral with minimal determinant. The coefficients of *U* are usually nonintegral.

```

? G = matdiagonal([650, -104329, -104329]);
? [H,U,c]=qfminimize(G); H
%2 = [-1,0,0;0,-1,0;0,0,1]
? U
%3 = [0,0,1/5;5/323,-1/323,0;-1/323,-5/323,0]
? c
%4 = 1/26
? c * U~ * G * U
%4 = [-1,0,0;0,-1,0;0,0,1]

```

The library syntax is GEN `qfminimize`(GEN *G*).

3.10.68 qfnorm(*x*, {*q*}). This function is obsolete, use `qfeval`.

The library syntax is GEN `qfnorm`(GEN *x*, GEN *q* = NULL).

3.10.69 qforbits(G, V). Return the orbits of V under the action of the group of linear transformation generated by the set G . It is assumed that G contains minus identity, and only one vector in $\{v, -v\}$ should be given. If G does not stabilize V , the function return 0.

In the example below, we compute representatives and lengths of the orbits of the vectors of norm ≤ 3 under the automorphisms of the lattice $\mathbf{Z}^6$.

```
? Q=matid(6); G=qfauto(Q); V=qfminim(Q,3);
? apply(x->[x[1],#x],qforbits(G,V))
%2 = [[0,0,0,0,0,1]~,6],[0,0,0,0,1,-1]~,30],[0,0,0,1,-1,-1]~,80]]
```

The library syntax is GEN qforbits(GEN G, GEN V).

3.10.70 qfparam($G, sol, \{flag = 0\}$). Coefficients of binary quadratic forms that parametrize the solutions of the ternary quadratic form G , using the particular solution sol . $flag$ is optional and can be 1, 2, or 3, in which case the $flag$ -th form is reduced. The default is $flag = 0$ (no reduction).

```
? G = [1,0,0;0,1,0;0,0,-34];
? M = qfparam(G, qfsolve(G))
%2 =
[ 3 -10 -3]
[-5  -6  5]
[ 1   0  1]
```

Indeed, the solutions can be parametrized as

$$(3x^2 - 10xy - 3y^2)^2 + (-5x^2 - 6xy + 5y^2)^2 - 34(x^2 + y^2)^2 = 0.$$

```
? v = y^2 * M*[1,x/y,(x/y)^2]~
%3 = [3*x^2 - 10*y*x - 3*y^2, -5*x^2 - 6*y*x + 5*y^2, -x^2 - y^2]~
? v~*G*v
%4 = 0
```

The library syntax is GEN qfparam(GEN G, GEN sol, long flag).

3.10.71 qfperfection(G). G being a square and symmetric matrix with integer entries representing a positive definite quadratic form, outputs the perfection rank of the form. That is, gives the rank of the family of the s symmetric matrices $v^t v$, where v runs through the minimal vectors.

A form is perfect if and only if its perfection rank is $d(d+1)/2$ where d is the dimension of G .

The algorithm computes the minimal vectors and its runtime is exponential in d .

The library syntax is GEN qfperfection(GEN G).

3.10.72 qfrep($q, B, \{flag = 0\}$). q being a square and symmetric matrix with integer entries representing a positive definite quadratic form, count the vectors representing successive integers.

- If $flag = 0$, count all vectors. Outputs the vector whose i -th entry, $1 \leq i \leq B$ is half the number of vectors v such that $q(v) = i$.

- If $flag = 1$, count vectors of even norm. Outputs the vector whose i -th entry, $1 \leq i \leq B$ is half the number of vectors such that $q(v) = 2i$.

```
? q = [2, 1; 1, 3];
? qfrep(q, 5)
%2 = Vecsmall([0, 1, 2, 0, 0]) \\ 1 vector of norm 2, 2 of norm 3, etc.
? qfrep(q, 5, 1)
%3 = Vecsmall([1, 0, 0, 1, 0]) \\ 1 vector of norm 2, 0 of norm 4, etc.
```

This routine uses a naive algorithm based on `qfminim`, and will fail if any entry becomes larger than 2^{31} (or 2^{63}).

The library syntax is GEN `qfrep0(GEN q, GEN B, long flag)`.

3.10.73 qfsign(x). Returns $[p, m]$ the signature of the quadratic form represented by the symmetric matrix x . Namely, p (resp. m) is the number of positive (resp. negative) eigenvalues of x . The result is computed using Gaussian reduction.

The library syntax is GEN `qfsign(GEN x)`.

3.10.74 qfsolve(G). Given a square symmetric matrix G of dimension $n \geq 1$, solve over $\mathbf{Q}$ the quadratic equation ${}^tXGX = 0$. The matrix G must have rational coefficients. When G is integral, the argument can also be a vector $[G, F]$ where F is the factorization matrix of the absolute value of the determinant of G .

The solution might be a single nonzero column vector (`t_COL`) or a matrix (whose columns generate a totally isotropic subspace).

If no solution exists, returns an integer, that can be a prime p such that there is no local solution at p , or -1 if there is no real solution, or -2 if $n = 2$ and $-\det G$ is not a square (which implies there is a real solution, but no local solution at some p dividing $\det G$).

```
? G = [1,0,0;0,1,0;0,0,-34];
? qfsolve(G)
%1 = [-3, -5, 1]~
? qfsolve([1,0; 0,2])
%2 = -1 \\ no real solution
? qfsolve([1,0,0;0,3,0; 0,0,-2])
%3 = 3 \\ no solution in Q_3
? qfsolve([1,0; 0,-2])
%4 = -2 \\ no solution, n = 2
```

The library syntax is GEN `qfsolve(GEN G)`.

3.10.75 setbinop($f, X, \{Y\}$). The set whose elements are the $f(x,y)$, where x,y run through X,Y respectively. If Y is omitted, assume that $X = Y$ and that f is symmetric: $f(x,y) = f(y,x)$ for all x,y in X .

```
? X = [1,2,3]; Y = [2,3,4];
? setbinop((x,y)->x+y, X,Y) \\ set X + Y
%2 = [3, 4, 5, 6, 7]
? setbinop((x,y)->x-y, X,Y) \\ set X - Y
%3 = [-3, -2, -1, 0, 1]
? setbinop((x,y)->x+y, X) \\ set 2X = X + X
%2 = [2, 3, 4, 5, 6]
```

The library syntax is `GEN setbinop(GEN f, GEN X, GEN Y = NULL)`.

3.10.76 setdelta(x,y). Symmetric difference of the two sets x and y (see `setisset`). If x or y is not a set, the result is undefined.

```
? a=[1,2,2,3];b=[4,2,3,4];
? setdelta(Set(a), Set(b))
%2 = [1, 4] \\ the symmetric difference of the two sets
? setdelta(a,b)
%3 = [1, 2, 2, 3, 4, 2, 3, 4] \\ undefined result
```

The library syntax is `GEN setdelta(GEN x, GEN y)`.

3.10.77 setintersect(x,y). Intersection of the two sets x and y (see `setisset`). If x or y is not a set, the result is undefined.

The library syntax is `GEN setintersect(GEN x, GEN y)`.

3.10.78 setisset(x). Returns true (1) if x is a set, false (0) if not. In PARI, a set is a row vector whose entries are strictly increasing with respect to a (somewhat arbitrary) universal comparison function. To convert any object into a set (this is most useful for vectors, of course), use the function `Set`.

```
? a = [3, 1, 1, 2];
? setisset(a)
%2 = 0
? Set(a)
%3 = [1, 2, 3]
```

The library syntax is `long setisset(GEN x)`.

3.10.79 setminus(x,y). Difference of the two sets x and y (see `setisset`), i.e. set of elements of x which do not belong to y . If x or y is not a set, the result is undefined.

The library syntax is `GEN setminus(GEN x, GEN y)`.

3.10.80 setsearch($S, x, \{flag = 0\}$). Determines whether x belongs to the set or sorted list S (see `setisset`).

We first describe the default behavior, when $flag$ is zero or omitted. If x belongs to the set S , returns the index j such that $S[j] = x$, otherwise returns 0.

```
? T = [7,2,3,5]; S = Set(T);
? setsearch(S, 2)
%2 = 1
? setsearch(S, 4)      \\ not found
%3 = 0
? setsearch(T, 7)      \\ search in a randomly sorted vector
%4 = 0 \\ WRONG !
```

If S is not a set, we also allow sorted lists with respect to the `cmp` sorting function, without repeated entries, as per `listsort(L,1)`; otherwise the result is undefined.

```
? L = List([1,4,2,3,2]); setsearch(L, 4)
%1 = 0 \\ WRONG !
? listsort(L, 1); L      \\ sort L first
%2 = List([1, 2, 3, 4])
? setsearch(L, 4)
%3 = 4      \\ now correct
```

If $flag$ is nonzero, this function returns the index j where x should be inserted, and 0 if it already belongs to S . This is meant to be used for dynamically growing (sorted) lists, in conjunction with `listinsert`.

```
? L = List([1,5,2,3,2]); listsort(L,1); L
%1 = List([1,2,3,5])
? j = setsearch(L, 4, 1)  \\ 4 should have been inserted at index j
%2 = 4
? listinsert(L, 4, j); L
%3 = List([1, 2, 3, 4, 5])
```

The library syntax is `long setsearch(GEN S, GEN x, long flag)`.

3.10.81 setunion(x, y). Union of the two sets x and y (see `setisset`). If x or y is not a set, the result is undefined.

The library syntax is `GEN setunion(GEN x, GEN y)`.

3.10.82 snfrank($D, \{q = 0\}$). Assuming that D is a Smith normal form (i.e. vector of elementary divisors) for some module and q a power of an irreducible element or 0, returns the minimal number of generators for D/qD . For instance, if $q = p^n$ where p is a prime number, this is the dimension of $(p^{n-1}D)/p^n D$ as an $\mathbf{F}_p$ -vector space. An argument $q = 0$ may be omitted.

```
? snfrank([4,4,2], 2)
%1 = 3
? snfrank([4,4,2], 4)
%2 = 2
? snfrank([4,4,2], 8)
%3 = 0
```

```
? snfrank([4,4,2])  \\ or snfrank([4,4,2], 0)
%4 = 3
```

The function also works for $K[x]$ -modules:

```
? D=matsnf([-x-5,-1,-1,0; 0,x^2+10*x+26,-1,-x-5; 1,-x-5,-x-5,1; -1,0,0,1]);
? snfrank(D, x^2 + 10*x + 27)
%6 = 2
? A=matdiagonal([x-1,x^2+1,x-1,(x^2+1)^2,x,(x-1)^2]); D=matsnf(A);
? snfrank(D,x-1)
%8 = 3
? snfrank(D,(x-1)^2)
%9 = 1
? snfrank(D,(x-1)^3)
%9 = 0
? snfrank(D,x^2+1)
%10 = 2
```

Finally this function supports any output from `matsnf` (e.g., with transformation matrices included, with or without cleanup).

The library syntax is `long snfrank(GEN D, GEN q = NULL)`.

3.10.83 trace(x). This applies to quite general x . If x is not a matrix, it is equal to the sum of x and its conjugate, except for polmods where it is the trace as an algebraic number.

For x a square matrix, it is the ordinary trace. If x is a nonsquare matrix (but not a vector), an error occurs.

The library syntax is `GEN gtrace(GEN x)`.

3.10.84 vecextract($x, y, \{z\}$). Extraction of components of the vector or matrix x according to y . In case x is a matrix, its components are the *columns* of x . The parameter y is a component specifier, which is either an integer, a string describing a range, or a vector.

If y is an integer, it is considered as a mask: the binary bits of y are read from right to left, but correspond to taking the components from left to right. For example, if $y = 13 = (1101)_2$ then the components 1,3 and 4 are extracted.

If y is a vector (`t_VEC`, `t_COL` or `t_VECSMALL`), which must have integer entries, these entries correspond to the component numbers to be extracted, in the order specified.

If y is a string, it can be

- a single (nonzero) index giving a component number (a negative index means we start counting from the end).
- a range of the form " $a..b$ ", where a and b are indexes as above. Any of a and b can be omitted; in this case, we take as default values $a = 1$ and $b = -1$, i.e. the first and last components respectively. We then extract all components in the interval $[a, b]$, in reverse order if $b < a$.

In addition, if the first character in the string is `^`, the complement of the given set of indices is taken.

If z is not omitted, x must be a matrix. y is then the *row* specifier, and z the *column* specifier, where the component specifier is as explained above.

```

? v = [a, b, c, d, e];
? vecextract(v, 5)          \\ mask
%1 = [a, c]
? vecextract(v, [4, 2, 1]) \\ component list
%2 = [d, b, a]
? vecextract(v, "2..4")    \\ interval
%3 = [b, c, d]
? vecextract(v, "-1..-3")  \\ interval + reverse order
%4 = [e, d, c]
? vecextract(v, "^2")      \\ complement
%5 = [a, c, d, e]
? vecextract(matid(3), "2..", "..")
%6 =
[0 1 0]
[0 0 1]

```

The range notations `v[i..j]` and `v[~i]` (for `t_VEC` or `t_COL`) and `M[i..j, k..l]` and friends (for `t_MAT`) implement a subset of the above, in a simpler and *faster* way, hence should be preferred in most common situations. The following features are not implemented in the range notation:

- reverse order,
- omitting either *a* or *b* in *a..b*.

The library syntax is `GEN extract0(GEN x, GEN y, GEN z = NULL)`.

3.10.85 `vecprod(v)`. Return the product of the components of the vector *v*. Return 1 on an empty vector.

```

? vecprod([1,2,3])
%1 = 6
? vecprod([])
%2 = 1

```

The library syntax is `GEN vecprod(GEN v)`.

3.10.86 `vecsearch(v, x, {cmpf})`. Determines whether *x* belongs to the sorted vector or list *v*: return the (positive) index where *x* was found, or 0 if it does not belong to *v*.

If the comparison function `cmpf` is omitted, we assume that *v* is sorted in increasing order, according to the standard comparison function `lex`, thereby restricting the possible types for *x* and the elements of *v* (integers, fractions, reals, and vectors of such). We also transparently allow a `t_VECSMALL` *x* in this case, for the natural ordering of the integers.

If `cmpf` is present, it is understood as a comparison function and we assume that *v* is sorted according to it, see `vecsrt` for how to encode comparison functions.

```

? v = [1,3,4,5,7];
? vecsearch(v, 3)
%2 = 2
? vecsearch(v, 6)
%3 = 0 \\ not in the list
? vecsearch([7,6,5], 5) \\ unsorted vector: result undefined

```

```
%4 = 0
```

Note that if we are sorting with respect to a key which is expensive to compute (e.g. a discriminant), one should rather precompute all keys, sort that vector and search in the vector of keys, rather than searching in the original vector with respect to a comparison function.

By abuse of notation, x is also allowed to be a matrix, seen as a vector of its columns; again by abuse of notation, a `t_VEC` is considered as part of the matrix, if its transpose is one of the matrix columns.

```
? v = vecsort([3,0,2; 1,0,2]) \\ sort matrix columns according to lex order
%1 =
[0 2 3]
[0 2 1]
? vecsearch(v, [3,1]~)
%2 = 3
? vecsearch(v, [3,1]) \\ can search for x or x~
%3 = 3
? vecsearch(v, [1,2])
%4 = 0 \\ not in the list
```

The library syntax is `long vecsearch(GEN v, GEN x, GEN cmpf = NULL)`.

3.10.87 vecsort($x, \{cmpf\}, \{flag = 0\}$). Sorts the vector x in ascending order, using a mergesort method. x must be a list, vector or matrix (seen as a vector of its columns). Note that mergesort is stable, hence the initial ordering of “equal” entries (with respect to the sorting criterion) is not changed.

If `cmpf` is omitted, we use the standard comparison function `lex`, thereby restricting the possible types for the elements of x (integers, fractions or reals and vectors of those). We also transparently allow a `t_VECSMALL` x in this case, for the standard ordering on the integers.

If `cmpf` is present, it is understood as a comparison function and we sort according to it. The following possibilities exist:

- an integer k : sort according to the value of the k -th subcomponents of the components of x .
- a vector: sort lexicographically according to the components listed in the vector. For example, if `cmpf = [2, 1, 3]`, sort with respect to the second component, and when these are equal, with respect to the first, and when these are equal, with respect to the third.
- a comparison function: `t_CLOSURE` with two arguments x and y , and returning a real number which is < 0 , > 0 or $= 0$ if $x < y$, $x > y$ or $x = y$ respectively.
- a key: `t_CLOSURE` with one argument x and returning the value $f(x)$ with respect to which we sort.

```
? vecsort([3,0,2; 1,0,2]) \\ sort columns according to lex order
%1 =
[0 2 3]
[0 2 1]
? vecsort(v, (x,y)->y-x) \\ reverse sort
? vecsort(v, (x,y)->abs(x)-abs(y)) \\ sort by increasing absolute value
```

```
? vecsort(v, abs)  \\ sort by increasing absolute value, using key
? cmpf(x,y) = my(dx = poldisc(x), dy = poldisc(y)); abs(dx) - abs(dy);
? v = [x^2+1, x^3-2, x^4+5*x+1] vecsort(v, cmpf)  \\ comparison function
? vecsort(v, x->abs(poldisc(x)))  \\ key
```

The `abs` and `cmpf` examples show how to use a named function instead of an anonymous function. It is preferable to use a *key* whenever possible rather than include it in the comparison function as above since the key is evaluated $O(n)$ times instead of $O(n \log n)$, where n is the number of entries.

A direct approach is also possible and equivalent to using a sorting key:

```
? T = [abs(poldisc(x)) | x<-v];
? perm = vecsort(T,,1);  \\ indirect sort
? vecextract(v, perm)
```

This also provides the vector T of all keys, which is interesting for instance in later `vecsearch` calls: it is more efficient to sort T ($T = \text{vecextract}(T, \text{perm})$) then search for a key in T rather than to search in v using a comparison function or a key. Note also that `mapisdefined` is often easier to use and faster than `vecsearch`.

The binary digits of *flag* mean:

- 1: indirect sorting of the vector x , i.e. if x is an n -component vector, returns a permutation of $[1, 2, \dots, n]$ which applied to the components of x sorts x in increasing order. For example, `vecextract(x, vecsort(x,,1))` is equivalent to `vecsort(x)`.
- 4: use descending instead of ascending order.
- 8: remove “duplicate” entries with respect to the sorting function (keep the first occurring entry). For example:

```
? vecsort([Pi,Mod(1,2),z], (x,y)->0, 8)  \\ make everything compare equal
%1 = [3.141592653589793238462643383]
? vecsort([[2,3],[0,1],[0,3]], 2, 8)
%2 = [[0, 1], [2, 3]]
```

The library syntax is `GEN vecsort0(GEN x, GEN cmpf = NULL, long flag)`.

3.10.88 `vecsum(v)`. Return the sum of the components of the vector v . Return 0 on an empty vector.

```
? vecsum([1,2,3])
%1 = 6
? vecsum([])
%2 = 0
```

The library syntax is `GEN vecsum(GEN v)`.

3.10.89 vector($n, \{X\}, \{expr = 0\}$). Creates a row vector (type `t_VEC`) with n components whose components are the expression $expr$ evaluated at the integer points between 1 and n . If the last two arguments are omitted, fills the vector with zeroes.

```
? vector(3,i, 5*i)
%1 = [5, 10, 15]
? vector(3)
%2 = [0, 0, 0]
```

The variable X is lexically scoped to each evaluation of $expr$. Any change to X within $expr$ does not affect subsequent evaluations, it still runs 1 to n . A local change allows for example different indexing:

```
vector(10, i, i=i-1; f(i)) \\ i = 0, ..., 9
vector(10, i, i=2*i; f(i)) \\ i = 2, 4, ..., 20
```

This per-element scope for X differs from `for` loop evaluations, as the following example shows:

```
n = 3
v = vector(n); vector(n, i, i++)      ----> [2, 3, 4]
v = vector(n); for (i = 1, n, v[i] = i++) ----> [2, 0, 4]
```

3.10.90 vectorsmall($n, \{X\}, \{expr = 0\}$). Creates a row vector of small integers (type `t_VECSMALL`) with n components whose components are the expression $expr$ evaluated at the integer points between 1 and n .

3.10.91 vectorv($n, \{X\}, \{expr = 0\}$). As `vector`, but returns a column vector (type `t_COL`).

3.11 Transcendental functions.

Since the values of transcendental functions cannot be exactly represented, these functions will always return an inexact object: a real number, a complex number, a p -adic number or a power series. All these objects have a certain finite precision.

As a general rule, which of course in some cases may have exceptions, transcendental functions operate in the following way:

- If the argument is either a real number or an inexact complex number (like `1.0 + I` or `Pi*I` but not `2 - 3*I`), then the computation is done with the precision of the argument. In the example below, we see that changing the precision to 50 digits does not matter, because x only had a precision of 19 digits.

```
? \p 15
  realprecision = 19 significant digits (15 digits displayed)
? x = Pi/4
%1 = 0.785398163397448
? \p 50
  realprecision = 57 significant digits (50 digits displayed)
? sin(x)
%2 = 0.7071067811865475244
```

Note that even if the argument is real, the result may be complex (e.g. `acos(2.0)` or `acosh(0.0)`). See each individual function help for the definition of the branch cuts and choice of principal value.

- If the argument is either an integer, a rational, an exact complex number or a quadratic number, it is first converted to a real or complex number using the current precision, which can be view and manipulated using the defaults `realprecision` (in decimal digits) or `realbitprecision` (in bits). This precision can be changed indifferently

- in decimal digits: use `\p` or `default(realprecision,...)`.
- in bits: use `\pb` or `default(realbitprecision,...)`.

After this conversion, the computation proceeds as above for real or complex arguments.

In library mode, the `realprecision` does not matter; instead the precision is taken from the `prec` parameter which every transcendental function has. As in `gp`, this `prec` is not used when the argument to a function is already inexact. Note that the argument `prec` stands for the length in words of a real number, including codewords. Hence we must have $prec \geq 3$. (Some functions allow a `bitprec` argument instead which allow finer granularity.)

Some accuracies attainable on 32-bit machines cannot be attained on 64-bit machines for parity reasons. For example, an accuracy of 28 decimal digits on 32-bit machines corresponds to `prec` having the value 5, for a mantissa of $3 \times 32 = 96$ bits. But this cannot be attained on 64-bit machines: we can attain either 64 or 128 bits, but values in between.

- If the argument is a `polmod` (representing an algebraic number), then the function is evaluated for every possible complex embedding of that algebraic number. A column vector of results is returned, with one component for each complex embedding. Therefore, the number of components equals the degree of the `t_POLMOD` modulus.

- If the argument is an `intmod` or a p -adic, at present only a few functions like `sqr` (square root), `sqr` (square), `log`, `exp`, powering, `teichmuller` (Teichmüller character) and `agm` (arithmetic-geometric mean) are implemented.

Note that in the case of a 2-adic number, `sqr`(x) may not be identical to $x * x$: for example if $x = 1 + O(2^5)$ and $y = 1 + O(2^5)$ then $x * y = 1 + O(2^5)$ while `sqr`(x) = $1 + O(2^6)$. Here, $x * x$ yields the same result as `sqr`(x) since the two operands are known to be *identical*. The same statement holds true for p -adics raised to the power n , where $v_p(n) > 0$.

Remark. If we wanted to be strictly consistent with the PARI philosophy, we should have $x * y = (4 \bmod 8)$ and `sqr`(x) = $(4 \bmod 32)$ when both x and y are congruent to 2 modulo 4. However, since `intmod` is an exact object, PARI assumes that the modulus must not change, and the result is hence $(0 \bmod 4)$ in both cases. On the other hand, p -adics are not exact objects, hence are treated differently.

- If the argument is a polynomial, a power series or a rational function, it is, if necessary, first converted to a power series using the current series precision, held in the default `seriesprecision`. This precision (the number of significant terms) can be changed using `\ps` or `default(seriesprecision,...)`. Then the Taylor series expansion of the function around $X = 0$ (where X is the main variable) is computed to a number of terms depending on the number of terms of the argument and the function being computed.

Under `gp` this again is transparent to the user. When programming in library mode, however, it is *strongly* advised to perform an explicit conversion to a power series first, as in

```
x = gtoser(x, gvar(x), seriesprec)
```

where the number of significant terms `seriesprec` can be specified explicitly. If you do not do this, a global variable `prec1` is used instead, to convert polynomials and rational functions to a

power series with a reasonable number of terms; tampering with the value of this global variable is *deprecated* and strongly discouraged.

- If the argument is a vector or a matrix, the result is the *componentwise* evaluation of the function. In particular, transcendental functions on square matrices, are not built-in. For this you can use the following technique, which is neither very efficient nor numerical stable, but is often good enough provided we restrict to diagonalizable matrices:

```
mateval(f, M) =
{ my([L, H] = mateigen(M, 1));
  H * matdiagonal(f(L)) * H^(-1);
}
? A = [13,2;10,14];
? a = mateval(sqrt, A) /* approximates  $\sqrt{A}$  */
%2 =
[3.5522847498307933... 0.27614237491539669...]
[1.3807118745769834... 3.69035593728849174...]
? exponent(a^2 - A)
%3 = -123 \\ OK
? b = mateval(exp, A);
? exponent(mateval(log, b) - A)
%5 = -115 \\ tolerable
```

The absolute error depends on the condition number of the base change matrix H and on the largest $|f(\lambda)|$, where λ runs through the eigenvalues. If M is real symmetric, you may use `qfjacobi` instead of `mateigen`.

Of course when the input is not diagonalizable, this function produces junk:

```
? mateval(sqrt, [0,1;0,0])
%6 =      \\ oops ...
[0.E-57 0]
[      0 0]
```

3.11.1 Catalan. Catalan's constant $G = \sum_{n \geq 0} \frac{(-1)^n}{(2n+1)^2} = 0.91596 \dots$. Note that `Catalan` is one of the few reserved names which cannot be used for user variables.

The library syntax is `GEN mpcatalan(long prec)`.

3.11.2 Euler. Euler's constant $\gamma = 0.57721 \dots$. Note that `Euler` is one of the few reserved names which cannot be used for user variables.

The library syntax is `GEN mpeuler(long prec)`.

3.11.3 I. The complex number $\sqrt{-1}$.

The library syntax is `GEN gen_I()`.

3.11.4 Pi. The constant π (3.14159...). Note that `Pi` is one of the few reserved names which cannot be used for user variables.

The library syntax is `GEN mppi(long prec)`.

3.11.5 abs(x). Absolute value of x (modulus if x is complex). Rational functions are not allowed. Contrary to most transcendental functions, an exact argument is *not* converted to a real number before applying **abs** and an exact result is returned if possible.

```
? abs(-1)
%1 = 1
? abs(3/7 + 4/7*I)
%2 = 5/7
? abs(1 + I)
%3 = 1.414213562373095048801688724
```

If x is a polynomial, returns $-x$ if the leading coefficient is real and negative else returns x . For a power series, the constant coefficient is considered instead.

The library syntax is GEN **gabs**(GEN x , long prec).

3.11.6 acos(x). Principal branch of $\cos^{-1}(x) = -i \log(x + i\sqrt{1-x^2})$. In particular, $\Re(\text{acos}(x)) \in [0, \pi]$ and if $x \in \mathbf{R}$ and $|x| > 1$, then $\text{acos}(x)$ is complex. The branch cut is in two pieces: $] -\infty, -1]$, continuous with quadrant II, and $[1, +\infty[$, continuous with quadrant IV. We have $\text{acos}(x) = \pi/2 - \text{asin}(x)$ for all x .

The library syntax is GEN **gacos**(GEN x , long prec).

3.11.7 acosh(x). Principal branch of $\cosh^{-1}(x) = 2 \log(\sqrt{(x+1)/2} + \sqrt{(x-1)/2})$. In particular, $\Re(\text{acosh}(x)) \geq 0$ and $\Im(\text{acosh}(x)) \in] -\pi, \pi]$; if $x \in \mathbf{R}$ and $x < 1$, then $\text{acosh}(x)$ is complex.

The library syntax is GEN **gacosh**(GEN x , long prec).

3.11.8 agm(x, y). Arithmetic-geometric mean of x and y . In the case of complex or negative numbers, the optimal AGM is returned (the largest in absolute value over all choices of the signs of the square roots). p -adic or power series arguments are also allowed. Note that a p -adic agm exists only if x/y is congruent to 1 modulo p (modulo 16 for $p = 2$). x and y cannot both be vectors or matrices.

The library syntax is GEN **agm**(GEN x , GEN y , long prec).

3.11.9 airy(z). Airy $[Ai, Bi]$ functions of argument z .

```
? [A,B] = airy(1);
? A
%2 = 0.13529241631288141552414742351546630617
? B
%3 = 1.2074235949528712594363788170282869954
```

The library syntax is GEN **airy**(GEN z , long prec).

3.11.10 arg(x). Argument of the complex number x , such that $-\pi < \arg(x) \leq \pi$.

The library syntax is GEN **garg**(GEN x , long prec).

3.11.11 asin(x). Principal branch of $\sin^{-1}(x) = -i \log(ix + \sqrt{1-x^2})$. In particular, $\Re(\text{asin}(x)) \in [-\pi/2, \pi/2]$ and if $x \in \mathbf{R}$ and $|x| > 1$ then $\text{asin}(x)$ is complex. The branch cut is in two pieces: $] -\infty, -1]$, continuous with quadrant II, and $[1, +\infty[$ continuous with quadrant IV. The function satisfies $i \text{asin}(x) = \text{asinh}(ix)$.

The library syntax is `GEN gasin(GEN x, long prec)`.

3.11.12 asinh(x). Principal branch of $\sinh^{-1}(x) = \log(x + \sqrt{1+x^2})$. In particular $\Im(\text{asinh}(x)) \in [-\pi/2, \pi/2]$. The branch cut is in two pieces: $] -i\infty, -i]$, continuous with quadrant III and $[+i, +i\infty[$, continuous with quadrant I.

The library syntax is `GEN gasinh(GEN x, long prec)`.

3.11.13 atan(x). Principal branch of $\tan^{-1}(x) = \log((1+ix)/(1-ix))/2i$. In particular the real part of $\text{atan}(x)$ belongs to $] -\pi/2, \pi/2[$. The branch cut is in two pieces: $] -i\infty, -i]$, continuous with quadrant IV, and $]i, +i\infty[$ continuous with quadrant II. The function satisfies $\text{atan}(x) = -i \text{atanh}(ix)$ for all $x \neq \pm i$.

The library syntax is `GEN gatan(GEN x, long prec)`.

3.11.14 atanh(x). Principal branch of $\tanh^{-1}(x) = \log((1+x)/(1-x))/2$. In particular the imaginary part of $\text{atanh}(x)$ belongs to $[-\pi/2, \pi/2]$; if $x \in \mathbf{R}$ and $|x| > 1$ then $\text{atanh}(x)$ is complex.

The library syntax is `GEN gatanh(GEN x, long prec)`.

3.11.15 besselh1(nu, x). H^1 -Bessel function of index nu and argument x .

The library syntax is `GEN hbessel1(GEN nu, GEN x, long prec)`.

3.11.16 besselh2(nu, x). H^2 -Bessel function of index nu and argument x .

The library syntax is `GEN hbessel2(GEN nu, GEN x, long prec)`.

3.11.17 besseli(nu, x). I -Bessel function of index nu and argument x . If x converts to a power series, the initial factor $(x/2)^\nu/\Gamma(\nu+1)$ is omitted (since it cannot be represented in PARI when ν is not integral).

The library syntax is `GEN ibessel(GEN nu, GEN x, long prec)`.

3.11.18 besselj(nu, x). J -Bessel function of index nu and argument x . If x converts to a power series, the initial factor $(x/2)^\nu/\Gamma(\nu+1)$ is omitted (since it cannot be represented in PARI when ν is not integral).

The library syntax is `GEN jbessel(GEN nu, GEN x, long prec)`.

3.11.19 besseljh(n, x). J -Bessel function of half integral index. More precisely, `besseljh( $n, x$ )` computes $J_{n+1/2}(x)$ where n must be of type integer, and x is any element of $\mathbf{C}$. In the present version 2.17.4, this function is not very accurate when x is small.

The library syntax is `GEN jbesselh(GEN n, GEN x, long prec)`.

3.11.20 `besseljzero`($nu, \{k = 1\}$). k -th zero of the J -Bessel function of index nu , close to $\pi(\nu/2 + k - 1/4)$, usually noted $j_{\nu,k}$.

```
? besseljzero(0) \\ first zero of J_0
%1 = 2.4048255576957727686216318793264546431
? besselj(0, %)
%2 = 7.1951595399463653939930598011247182898 E-41
? besseljzero(0, 2) \\ second zero
%3 = 5.5200781102863106495966041128130274252
? besseljzero(I) \\ also works for complex orders, here J_i
%4 = 2.5377... + 1.4753...*I
```

The function uses a Newton iteration due to Temme. If ν is real and nonnegative, the result is guaranteed and the function really returns the k -th positive zero of J_ν . For general ν , the result is not well defined, given by the Newton iteration with $\pi(\nu/2 + k - 1/4)$ as a starting value. (N.B. Using this method for large real ν would give completely different results than the $j_{\nu,k}$ unless k is large enough.)

The library syntax is `GEN besseljzero(GEN nu, long k, long bitprec)`.

3.11.21 `besselk`(nu, x). K -Bessel function of index nu and argument x .

The library syntax is `GEN kbessel(GEN nu, GEN x, long prec)`.

3.11.22 `besseln`(nu, x). Deprecated alias for `bessely`.

The library syntax is `GEN ybessel(GEN nu, GEN x, long prec)`.

3.11.23 `bessely`(nu, x). Y -Bessel function of index nu and argument x .

The library syntax is `GEN ybessel(GEN nu, GEN x, long prec)`.

3.11.24 `besselyzero`($nu, \{k = 1\}$). k -th zero of the Y -Bessel function of index nu , close to $\pi(\nu/2 + k - 3/4)$, usually noted $y_{\nu,k}$.

```
? besselyzero(0) \\ first zero of Y_0
%1 = 0.89357696627916752158488710205833824123
? bessely(0, %)
%2 = 1.8708573650996561952 E-39
? besselyzero(0, 2) \\ second zero
%3 = 3.9576784193148578683756771869174012814
? besselyzero(I) \\ also works for complex orders, here Y_i
%4 = 1.03930... + 1.3266...*I
```

The function uses a Newton iteration due to Temme. If ν is real and nonnegative, the result is guaranteed and the function really returns the k -th positive zero of Y_ν . For general ν , the result is not well defined, given by Newton iteration with $\pi(\nu/2 + k - 3/4)$ as a starting value. (N.B. Using this method for large real ν would give completely different results than the $y_{\nu,k}$ unless k is large enough.)

The library syntax is `GEN besselyzero(GEN nu, long k, long bitprec)`.

3.11.25 cos(x). Cosine of x . Note that, for real x , cosine and sine can be obtained simultaneously as

```
cs(x) = my(z = exp(I*x)); [real(z), imag(z)];
```

and for general complex x as

```
cs2(x) = my(z = exp(I*x), u = 1/z); [(z+u)/2, (z-u)/2];
```

Note that the latter function suffers from catastrophic cancellation when $z^2 \approx \pm 1$.

The library syntax is GEN `gcos(GEN x, long prec)`.

3.11.26 cosh(x). Hyperbolic cosine of x .

The library syntax is GEN `gcosh(GEN x, long prec)`.

3.11.27 cotan(x). Cotangent of x .

The library syntax is GEN `gcotan(GEN x, long prec)`.

3.11.28 cotanh(x). Hyperbolic cotangent of x .

The library syntax is GEN `gcotanh(GEN x, long prec)`.

3.11.29 dilog(x). Principal branch of the dilogarithm of x , i.e. analytic continuation of the power series $\text{Li}_2(x) = \sum_{n \geq 1} x^n/n^2$.

The library syntax is GEN `dilog(GEN x, long prec)`.

3.11.30 eint1($x, \{n\}$). Exponential integral $\int_x^\infty \frac{e^{-t}}{t} dt = \text{incgam}(0, x)$, where the latter expression extends the function definition from real $x > 0$ to all complex $x \neq 0$.

If n is present, we must have $x > 0$; the function returns the n -dimensional vector $[\text{eint1}(x), \dots, \text{eint1}(nx)]$. Contrary to other transcendental functions, and to the default case (n omitted), the values are correct up to a bounded *absolute*, rather than relative, error 10^{-n} , where n is `precision( $x$ )` if x is a `t_REAL` and defaults to `realprecision` otherwise. (In the most important application, to the computation of L -functions via approximate functional equations, those values appear as weights in long sums and small individual relative errors are less useful than controlling the absolute error.) This is faster than repeatedly calling `eint1(i * x)`, but less precise.

The library syntax is GEN `veceint1(GEN x, GEN n = NULL, long prec)`. Also available is GEN `eint1(GEN x, long prec)`.

3.11.31 ellE(k). Complete elliptic integral of the second kind

$$E(k) = \int_0^{\pi/2} (1 - k^2 \sin(t)^2)^{1/2} dt$$

for the complex parameter k using the agm.

In particular, the perimeter of an ellipse of semi-major and semi-minor axes a and b is given by

```
e = sqrt(1 - (b/a)^2); \\ eccentricity
4 * a * ellE(e) \\ perimeter
```

The library syntax is GEN `ellE(GEN k, long prec)`.

3.11.32 ellK(k). Complete elliptic integral of the first kind

$$K(k) = \int_0^{\pi/2} (1 - k^2 \sin(t)^2)^{-1/2} dt$$

for the complex parameter k using the agm.

The library syntax is GEN `ellK(GEN k, long prec)`.

3.11.33 erfc(x). Complementary error function, analytic continuation of $(2/\sqrt{\pi}) \int_x^\infty e^{-t^2} dt = \text{sign}(x) \text{incgam}(1/2, x^2)/\sqrt{\pi}$ for real $x \neq 0$. The latter expression extends the function definition from real x to complex x with positive real part (or zero real part and positive imaginary part). This is extended to the whole complex plane by the functional equation $\text{erfc}(-x) = 2 - \text{erfc}(x)$.

```
? erfc(0)
%1 = 1.0000000000000000000000000000000000000000000000000000000
? erfc(1)
%2 = 0.15729920705028513065877936491739074071
? erfc(1+I)
%3 = -0.31615128169794764488027108024367036903
      - 0.19045346923783468628410886196916244244*I
```

The library syntax is GEN `gerfc(GEN x, long prec)`.

3.11.34 eta(z, {flag = 0}). Variants of Dedekind's η function. If $flag = 0$, return $\prod_{n=1}^\infty (1 - q^n)$, where q depends on x in the following way:

- $q = e^{2i\pi x}$ if x is a *complex number* (which must then have positive imaginary part); notice that the factor $q^{1/24}$ is missing!

- $q = x$ if x is a `t_PADIC`, or can be converted to a *power series* (which must then have positive valuation).

If $flag$ is nonzero, x is converted to a complex number and we return the true η function, $q^{1/24} \prod_{n=1}^\infty (1 - q^n)$, where $q = e^{2i\pi x}$.

The library syntax is GEN `eta0(GEN z, long flag, long prec)`.

Also available is GEN `trueeta(GEN x, long prec)` ($flag = 1$).

3.11.35 exp(x). Exponential of x . p -adic arguments with positive valuation are accepted.

The library syntax is GEN `gexp(GEN x, long prec)`. For a `t_PADIC` x , the function GEN `Qp_exp(GEN x)` is also available.

3.11.36 expm1(x). Return $\exp(x) - 1$, computed in a way that is also accurate when the real part of x is near 0. A naive direct computation would suffer from catastrophic cancellation; PARI's direct computation of $\exp(x)$ alleviates this well known problem at the expense of computing $\exp(x)$ to a higher accuracy when x is small. Using **expm1** is recommended instead:

```
? default(realprecision, 10000); x = 1e-100;
? a = expm1(x);
time = 4 ms.
? b = exp(x)-1;
time = 4 ms.
? default(realprecision, 10040); x = 1e-100;
? c = expm1(x);  \\ reference point
? abs(a-c)/c  \\ relative error in expm1(x)
%7 = 1.4027986153764843997 E-10019
? abs(b-c)/c  \\ relative error in exp(x)-1
%8 = 1.7907031188259675794 E-9919
```

As the example above shows, when x is near 0, `expm1` is more accurate than `exp(x)-1`.

The library syntax is `GEN gexpm1(GEN x, long prec)`.

3.11.37 gamma(s). For s a complex number, evaluates Euler's gamma function, which is the analytic continuation of

$$\Gamma(s) = \int_0^\infty t^{s-1} \exp(-t) dt, \quad \Re(s) > 0.$$

Error if s is a nonpositive integer, where Γ has a (simple) pole.

[illegible]

For s a $\mathfrak{t}$ -PADIC, evaluates the Morita gamma function at s , that is the unique continuous p -adic function on the p -adic integers extending $\Gamma_p(k) = (-1)^k \prod'_{j < k} j$, where the prime means that p does not divide j .

```
? gamma(1/4 + O(5^10))
%1 = 1 + 4*5 + 3*5^4 + 5^6 + 5^7 + 4*5^9 + O(5^10)
? algdep(%1,4)
%2 = x^4 + 4*x^2 + 5
```

The library syntax is GEN `ggamma(GEN s, long prec)`. For a `t_PADIC` x , the function GEN `Qp_gamma(GEN x)` is also available.

3.11.38 gammah(x). Gamma function evaluated at the argument $x + 1/2$.

The library syntax is `GEN ggammah(GEN x, long prec)`.

3.11.39 gammamellininv($G, t, \{m = 0\}$). Returns the value at t of the inverse Mellin transform G initialized by `gammamellinininit`. If the optional parameter m is present, return the m -th derivative $G^{(m)}(t)$.

```
? G = gammamellinininit([0]);
? gammamellininv(G, 2) - 2*exp(-Pi*2^2)
%2 = -4.484155085839414627 E-44
```

The shortcut

```
gammamellininv(A,t,m)
```

for

```
gammamellininv(gammamellinininit(A,m), t)
```

is available.

The library syntax is `GEN gammamellininv(GEN G, GEN t, long m, long bitprec)`.

3.11.40 gammamellinivasymp($A, n, \{m = 0\}$). Return the first n terms of the asymptotic expansion at infinity of the m -th derivative $K^{(m)}(t)$ of the inverse Mellin transform of the function

$$f(s) = \Gamma_{\mathbf{R}}(s + a_1) \dots \Gamma_{\mathbf{R}}(s + a_d),$$

where $\mathbf{A}$ is the vector $[a_1, \dots, a_d]$ and $\Gamma_{\mathbf{R}}(s) = \pi^{-s/2} \Gamma(s/2)$ (Euler's `gamma`). The result is a vector $[M[1] \dots M[n]]$ with $M[1]=1$, such that

$$K^{(m)}(t) = \sqrt{2^{d+1}/dt^{a+m(2/d-1)}} e^{-d\pi t^{2/d}} \sum_{n \geq 0} M[n+1] (\pi t^{2/d})^{-n}$$

with $a = (1 - d + \sum_{1 \leq j \leq d} a_j)/d$. We also allow A to be the output of `gammamellinininit`.

The library syntax is `GEN gammamellinivasymp(GEN A, long precdl, long n)`.

3.11.41 gammamellinininit($A, \{m = 0\}$). Initialize data for the computation by `gammamellininv` of the m -th derivative of the inverse Mellin transform of the function

$$f(s) = \Gamma_{\mathbf{R}}(s + a_1) \dots \Gamma_{\mathbf{R}}(s + a_d)$$

where $\mathbf{A}$ is the vector $[a_1, \dots, a_d]$ and $\Gamma_{\mathbf{R}}(s) = \pi^{-s/2} \Gamma(s/2)$ (Euler's `gamma`). This is the special case of Meijer's G functions used to compute L -values via the approximate functional equation. By extension, A is allowed to be an `Ldata` or an `Linit`, understood as the inverse Mellin transform of the L -function gamma factor.

Caveat. Contrary to the PARI convention, this function guarantees an *absolute* (rather than relative) error bound.

For instance, the inverse Mellin transform of $\Gamma_{\mathbf{R}}(s)$ is $2 \exp(-\pi z^2)$:

```
? G = gammamellininvert([0]);
? gammamellininv(G, 2) - 2*exp(-Pi*2^2)
%2 = -4.484155085839414627 E-44
```

The inverse Mellin transform of $\Gamma_{\mathbf{R}}(s + 1)$ is $2z \exp(-\pi z^2)$, and its second derivative is $4\pi z \exp(-\pi z^2)(2\pi z^2 - 3)$:

```
? G = gammamellininvert([1], 2);
? a(z) = 4*Pi*z*exp(-Pi*z^2)*(2*Pi*z^2-3);
? b(z) = gammamellininv(G,z);
? t(z) = b(z) - a(z);
? t(3/2)
%3 = -1.4693679385278593850 E-39
```

The library syntax is GEN `gammamellininvert(GEN A, long m, long bitprec)`.

3.11.42 hypergeom($\{N\}, \{D\}, z$). General hypergeometric function, where N and D are the vector of parameters in the numerator and denominator respectively, evaluated at the argument z , which may be complex, p -adic or a power series.

This function implements hypergeometric functions

$${}_pF_q((a_i)_{1 \leq i \leq p}, (b_j)_{1 \leq j \leq q}; z) = \sum_{n \geq 0} \frac{\prod_{1 \leq i \leq p} (a_i)_n}{\prod_{1 \leq j \leq q} (b_j)_n} \frac{z^n}{n!},$$

where $(a)_n = a(a+1) \cdots (a+n-1)$ is the rising Pochhammer symbol. For this to make sense, none of the b_j must be a negative or zero integer. The corresponding general GP command is

```
hypergeom([a1,a2,...,ap], [b1,b2,...,bq], z)
```

Whenever $p = 1$ or $q = 1$, a one-element vector can be replaced by the element it contains. Whenever $p = 0$ or $q = 0$, an empty vector can be omitted. For instance `hypergeom(,b,z)` computes ${}_0F_1(;b;z)$.

The non-archimedean cases (z a p -adic or power series) are handled trivially. We now discuss the case of a complex z ; we distinguish three kinds of such functions according to their radius of convergence R :

- $q \geq p$: $R = \infty$.
- $q = p - 1$: $R = 1$. Nonetheless, by integral representations, ${}_pF_q$ can be analytically continued outside the disc of convergence.
- $q \leq p - 2$: $R = 0$. By integral representations, one can make sense of the function in a suitable domain, by analytic continuation.

The list of implemented functions and their domain of validity in our implementation is as follows:

F01: `hypergeom(a,z)` (or `[a]`). This is essentially a Bessel function and computed as such. $R = \infty$.

F10: `hypergeom(a,,z)` This is $(1 - z)^{-a}$.

F11: `hypergeom(a,b,z)` is the Kummer confluent hypergeometric function, computed by summing the series. $R = \infty$

F20: `hypergeom([a,b],,z)`. $R = 0$, computed as

$$\frac{1}{\Gamma(a)} \int_0^\infty t^{a-1} (1-zt)^{-b} e^{-t} dt.$$

F21: `hypergeom([a,b],c,z)` (or `[c]`). $R = 1$, extended by

$$\frac{\Gamma(c)}{\Gamma(b)\Gamma(c-b)} \int_0^1 t^{b-1} (1-t)^{c-b-1} (1-zt)^a dt.$$

This is Gauss's Hypergeometric function, and almost all of the implementation work is done for this function.

F31: `hypergeom([a,b,c],d,z)` (or `[d]`). $R = 0$, computed as

$$\frac{1}{\Gamma(a)} \int_0^\infty t^{a-1} e^{-t} {}_2F_1(b, c; d; tz) dt.$$

F32: $\text{hypergeom}([a,b,c],[d,e],z)$. $R = 1$, extended by

$$\frac{\Gamma(e)}{\Gamma(c)\Gamma(e-c)} \int_0^1 t^{c-1}(1-t)^{e-c-1} {}_2F_1(a, b; d; tz) dt .$$

For other inputs: if $R = \infty$ or $R = 1$ and $|z| < 1 - \varepsilon$ is not too close to the circle of convergence, we simply sum the series.

```
? hypergeom([3,2], 3.4, 0.7) \\ 2F1(3,2; 3.4; 0.7)  
%1 = 7.999999999999999999999999999999999  
? a=5/3; T1=hypergeom([1,1,1],[a,a],1) \\ 3F2(1,1,1; a,a; 1)  
%2 = 3.1958592952314032651578713968927593818  
? T2=hypergeom([2,1,1],[a+1,a+1],1)  
%3 = 1.6752931349345765309211012564734179541  
? T3=hypergeom([2*a-1,1,1],[a+1,a+1],1)  
%4 = 1.9721037126267142061807688820853354440  
? T1 + (a-1)^2/(a^2*(2*a-3)) * (T2-2*(a-1)*T3) \\  
 - gamma(a)^2/((2*a-3)*gamma(2*a-2))  
%5 = -1.880790961315660013 E-37 \\ ~ 0
```

This identity is due to Bercu.

The library syntax is GEN hypergeom(GEN N = NULL, GEN D = NULL, GEN z, long prec)

3.11.43 hyperu(a, b, z). U -confluent hypergeometric function with complex parameters a, b, z . Note that ${}_2F_0(a, b, z) = (-z)^{-a}U(a, a+1-b, -1/z)$,

```
? hyperu(1, 3/2, I)
%1 = 0.23219... - 0.80952...*I
? -I * hypergeom([1, 1+1-3/2], [], -1/I)
%2 = 0.23219... - 0.80952...*I
```

The library syntax is GEN hyperu(GEN a, GEN b, GEN z, long prec).

3.11.44 incgam($s, x, \{g\}$). Incomplete gamma function $\int_x^\infty e^{-t}t^{s-1}dt$, extended by analytic continuation to all complex x, s not both 0. The relative error is bounded in terms of the precision of s (the accuracy of x is ignored when determining the output precision). When g is given, assume that $g = \Gamma(s)$. For small $|x|$, this will speed up the computation.

The library syntax is GEN incgam0(GEN s, GEN x, GEN g = NULL, long prec). Also available is GEN incgam(GEN s, GEN x, long prec).

3.11.45 incgamc(s, x). Complementary incomplete gamma function. The arguments x and s are complex numbers such that s is not a pole of Γ and $|x|/(|s|+1)$ is not much larger than 1 (otherwise the convergence is very slow). The result returned is $\int_0^x e^{-t}t^{s-1}dt$.

The library syntax is GEN incgamc(GEN s, GEN x, long prec).

3.11.46 lambertw($y, \{branch = 0\}$). Lambert W function, solution of the implicit equation $xe^x = y$.

- For real inputs y : If **branch** = 0, principal branch W_0 defined for $y \geq -\exp(-1)$. If **branch** = -1, branch W_{-1} defined for $-\exp(-1) \leq y < 0$.
- For p -adic inputs, p odd: give a solution of $x \exp(x) = y$ if y has positive valuation, of $\log(x) + x = \log(y)$ otherwise.
- For 2-adic inputs: give a solution of $x \exp(x) = y$ if y has valuation > 1 , of $\log(x) + x = \log(y)$ otherwise.

Caveat. Complex values of y are also supported but experimental. The other branches W_k for k not equal to 0 or -1 (set **branch** to k) are also experimental.

For $k \geq 1$, $W_{-1-k}(x) = \overline{W_k(x)}$, and $\Im(W_k(x))$ is close to $(\pi/2)(4k - \text{sign}(x))$.

The library syntax is GEN glambertW(GEN y, long branch, long prec).

3.11.47 lerchphi(z, s, a). Lerch transcendent $\Phi(z, s, a) = \sum_{n \geq 0} z^n (n+a)^{-s}$ and analytically continued, for reasonable values of the arguments.

The library syntax is GEN lerchphi(GEN z, GEN s, GEN a, long prec).

3.11.48 lerchzeta(s, a, lam). Lerch zeta function

$$L(s, a, \lambda) = \sum_{n \geq 0} e^{2\pi i \lambda n} (n+a)^{-s}$$

and analytically continued, for reasonable values of the arguments.

The library syntax is GEN lerchzeta(GEN s, GEN a, GEN lam, long prec).

3.11.49 lngamma(x). Principal branch of the logarithm of the gamma function of x . This function is analytic on the complex plane with nonpositive integers removed, and can have much larger arguments than `gamma` itself.

For x a power series such that $x(0)$ is not a pole of `gamma`, compute the Taylor expansion. (PARI only knows about regular power series and can't include logarithmic terms.)

```
? lngamma(1+x+O(x^2))
%1 = -0.57721566490153286060651209008240243104*x + O(x^2)
? lngamma(x+O(x^2))
*** at top-level: lngamma(x+O(x^2))
*** ^-----
*** lngamma: domain error in lngamma: valuation != 0
? lngamma(-1+x+O(x^2))
*** lngamma: Warning: normalizing a series with 0 leading term.
*** at top-level: lngamma(-1+x+O(x^2))
*** ^-----
*** lngamma: domain error in intformal: residue(series, pole) != 0
```

For x a `t_PADIC`, return the p -adic $\log \Gamma_p$ function, which is the p -adic logarithm of Morita's gamma function for $x \in \mathbf{Z}_p$, and Diamond's function if $|x| > 1$.

```
? lngamma(5+O(5^7))
%2 = 4*5^2 + 4*5^3 + 5^4 + 2*5^5 + O(5^6)
? log(gamma(5+O(5^7)))
%3 = 4*5^2 + 4*5^3 + 5^4 + 2*5^5 + O(5^6)
? lngamma(1/5+O(5^4))
%4 = 4*5^-1 + 4 + 2*5 + 5^2 + 5^3 + O(5^4)
? gamma(1/5+O(5^4))
*** at top-level: gamma(1/5+O(5^4))
*** ^-----
*** gamma: domain error in gamma: v_p(x) < 0
```

The library syntax is `GEN glngamma(GEN x, long prec)`.

3.11.50 log(x). Principal branch of the natural logarithm of $x \in \mathbf{C}^*$, i.e. such that $\Im(\log(x)) \in]-\pi, \pi]$. The branch cut lies along the negative real axis, continuous with quadrant 2, i.e. such that $\lim_{b \rightarrow 0+} \log(a + bi) = \log a$ for $a \in \mathbf{R}^*$. The result is complex (with imaginary part equal to π) if $x \in \mathbf{R}$ and $x < 0$. In general, the algorithm uses the formula

$$\log(x) \approx \frac{\pi}{2\operatorname{agm}(1, 4/s)} - m \log 2,$$

if $s = x2^m$ is large enough. (The result is exact to B bits provided $s > 2^{B/2}$.) At low accuracies, the series expansion near 1 is used.

p -adic arguments are also accepted for x , with the convention that $\log(p) = 0$. Hence in particular $\exp(\log(x))/x$ is not in general equal to 1 but to a $(p-1)$ -th root of unity (or ± 1 if $p = 2$) times a power of p .

The library syntax is `GEN glog(GEN x, long prec)`. For a `t_PADIC` x , the function `GEN Qp_log(GEN x)` is also available.

3.11.51 log1p(x). Return $\log(1+x)$, computed in a way that is also accurate when the real part of x is near 0. This is the reciprocal function of $\expm1(x) = \exp(x) - 1$.

```
? default(realprecision, 10000); x = Pi*1e-100;
? (expm1(log1p(x)) - x) / x
%2 = -7.668242895059371866 E-10019
? (log1p(expm1(x)) - x) / x
%3 = -7.668242895059371866 E-10019
```

When x is small, this function is both faster and more accurate than $\log(1+x)$:

```
? \p38
? x = 1e-20;
? localprec(100); c = log1p(x); \\ reference point
? a = log1p(x); abs((a - c)/c)
%6 = 0.E-38
? b = log(1+x); abs((b - c)/c) \\ slightly less accurate
%7 = 1.5930919111324522770 E-38
? for (i=1,10^5,log1p(x))
time = 81 ms.
? for (i=1,10^5,log(1+x))
time = 100 ms. \\ slower, too
```

The library syntax is GEN `glog1p(GEN x, long prec)`.

3.11.52 polylog(m, x, {flag = 0}). One of the different polylogarithms, depending on *flag*:

If *flag* = 0 or is omitted: m^{th} polylogarithm of x , i.e. analytic continuation of the power series $\text{Li}_m(x) = \sum_{n \geq 1} x^n/n^m$ ($x < 1$). Uses the functional equation linking the values at x and $1/x$ to restrict to the case $|x| \leq 1$, then the power series when $|x|^2 \leq 1/2$, and the power series expansion in $\log(x)$ otherwise.

Using *flag*, computes a modified m^{th} polylogarithm of x . We use Zagier's notations; let $\Re_m$ denote $\Re$ or $\Im$ depending on whether m is odd or even:

If *flag* = 1: compute $\tilde{D}_m(x)$, defined for $|x| \leq 1$ by

$$\Re_m \left(\sum_{k=0}^{m-1} \frac{(-\log |x|)^k}{k!} \text{Li}_{m-k}(x) + \frac{(-\log |x|)^{m-1}}{m!} \log |1-x| \right).$$

If *flag* = 2: compute $D_m(x)$, defined for $|x| \leq 1$ by

$$\Re_m \left(\sum_{k=0}^{m-1} \frac{(-\log |x|)^k}{k!} \text{Li}_{m-k}(x) - \frac{1}{2} \frac{(-\log |x|)^m}{m!} \right).$$

If *flag* = 3: compute $P_m(x)$, defined for $|x| \leq 1$ by

$$\Re_m \left(\sum_{k=0}^{m-1} \frac{2^k B_k}{k!} (\log |x|)^k \text{Li}_{m-k}(x) - \frac{2^{m-1} B_m}{m!} (\log |x|)^m \right).$$

These three functions satisfy the functional equation $f_m(1/x) = (-1)^{m-1} f_m(x)$.

The library syntax is GEN `polylog0(long m, GEN x, long flag, long prec)`. Also available is GEN `gpolylog(long m, GEN x, long prec)` (*flag* = 0).

3.11.53 polylogmult($s, \{z\}, \{t = 0\}$). For s a vector of positive integers and z a vector of complex numbers of the same length, returns the multiple polylogarithm value (MPV)

$$\zeta(s_1, \dots, s_r; z_1, \dots, z_r) = \sum_{n_1 > \dots > n_r > 0} \prod_{1 \leq i \leq r} z_i^{n_i} / n_i^{s_i}.$$

If z is omitted, assume $z = [1, \dots, 1]$, i.e., Multiple Zeta Value. More generally, return Yamamoto's interpolation between ordinary multiple polylogarithms ($t = 0$) and star polylogarithms ($t = 1$), using the condition $n_1 \geq \dots \geq n_r > 0$, evaluated at t .

We must have $|z_1 \cdots z_i| \leq 1$ for all i , and if $s_1 = 1$ we must have $z_1 \neq 1$.

```
? 8*polylogmult([2,1],[-1,1]) - zeta(3)
%1 = 0.E-38
```

Warning. The algorithm used converges when the z_i are ± 1 . It may not converge as some $z_i \neq 1$ becomes too close to 1, even at roots of 1 of moderate order:

```
? polylogmult([2,1], (99+20*I)/101 * [1,1])
*** polylogmult: sorry, polylogmult in this range is not yet implemented.
? polylogmult([2,1], exp(I*Pi/20)* [1,1])
*** polylogmult: sorry, polylogmult in this range is not yet implemented.
```

More precisely, if $y_i := 1/(z_1 \cdots z_i)$ and

$$v := \min_{i < j; y_i \neq 1} |(1 - y_i)y_j| > 1/4$$

then the algorithm computes the value up to a 2^{-b} absolute error in $O(k^2 N)$ operations on floating point numbers of $O(N)$ bits, where $k = \sum_i s_i$ is the weight and $N = b/\log_2(4v)$.

The library syntax is GEN `polylogmult_interpolate`(GEN s , GEN $z = \text{NULL}$, GEN $t = \text{NULL}$, long prec). Also available is GEN `polylogmult`(GEN s , GEN z , long prec) (t is `NULL`).

3.11.54 psi($x, \{der\}$). The ψ -function of x , i.e. the logarithmic derivative $\Gamma'(x)/\Gamma(x)$. If der is set, return the der -th derivative. For s a `t_PADIC`, evaluates the der -th derivative of the Morita ψ function at s .

The library syntax is GEN `gpsi_der`(GEN x , long der , long prec). For a `t_PADIC` x , the function GEN `Qp_psi`(GEN x , long der) is also available. For $der = 0$, GEN `gpsi`(GEN x , long prec) is also available.

3.11.55 rootsof1(N). Return the column vector v of all complex N -th roots of 1, where N is a positive integer. In other words, $v[k] = \exp(2I(k-1)\pi/N)$ for $k = 1, \dots, N$. Rational components (e.g., the roots ± 1 and $\pm I$) are given exactly, not as floating point numbers:

```
? rootsof1(4)
%1 = [1, I, -1, -I]~
? rootsof1(3)
%2 = [1, -1/2 + 0.866025...*I, -1/2 - 0.866025...*I]~
```

The library syntax is GEN `grootsof1`(long N , long prec).

3.11.56 sin(x). Sine of x . Note that, for real x , cosine and sine can be obtained simultaneously as

```
cs(x) = my(z = exp(I*x)); [real(z), imag(z)];
```

and for general complex x as

```
cs2(x) = my(z = exp(I*x), u = 1/z); [(z+u)/2, (z-u)/2];
```

Note that the latter function suffers from catastrophic cancellation when $z^2 \approx \pm 1$.

The library syntax is GEN `gsin(GEN x, long prec)`.

3.11.57 sinc(x). Cardinal sine of x , i.e. $\sin(x)/x$ if $x \neq 0$, 1 otherwise. Note that this function also allows to compute

$$(1 - \cos(x))/x^2 = \text{sinc}(x/2)^2/2$$

accurately near $x = 0$.

The library syntax is GEN `gsinc(GEN x, long prec)`.

3.11.58 sinh(x). Hyperbolic sine of x .

The library syntax is GEN `gsinh(GEN x, long prec)`.

3.11.59 sqr(x). Square of x . This operation is not completely straightforward, i.e. identical to $x*x$, since it can usually be computed more efficiently (roughly one-half of the elementary multiplications can be saved). Also, squaring a 2-adic number increases its precision. For example,

```
? (1 + 0(2^4))^2
%1 = 1 + 0(2^5)
? (1 + 0(2^4)) * (1 + 0(2^4))
%2 = 1 + 0(2^4)
```

Note that this function is also called whenever one multiplies two objects which are known to be *identical*, e.g. they are the value of the same variable, or we are computing a power.

```
? x = (1 + 0(2^4)); x * x
%3 = 1 + 0(2^5)
? (1 + 0(2^4))^4
%4 = 1 + 0(2^6)
```

(note the difference between %2 and %3 above).

The library syntax is GEN `gsqr(GEN x)`.

3.11.60 sqrt(x). Principal branch of the square root of x , defined as $\sqrt{x} = \exp(\log x/2)$. In particular, we have $\text{Arg}(\text{sqrt}(x)) \in]-\pi/2, \pi/2]$, and if $x \in \mathbf{R}$ and $x < 0$, then the result is complex with positive imaginary part.

Intmod a prime p , `t_PADIC` and `t_FFELT` are allowed as arguments. In the first 2 cases (`t_INTMOD`, `t_PADIC`), the square root (if it exists) which is returned is the one whose first p -adic digit is in the interval $[0, p/2]$. For other arguments, the result is undefined.

The library syntax is GEN `gsqrt(GEN x, long prec)`. For a `t_PADIC` x , the function GEN `Qp_sqrt(GEN x)` is also available.

3.11.61 `sqrtn`($x, n, \{&z\}$). Principal branch of the n th root of x , i.e. such that $\text{Arg}(\text{sqrtn}(x)) \in] - \pi/n, \pi/n]$. Intmod a prime and p -adics are allowed as arguments.

If z is present, it is set to a suitable root of unity allowing to recover all the other roots. If it was not possible, z is set to zero. In the case this argument is present and no n th root exist, 0 is returned instead of raising an error.

```
? sqrtn(Mod(2,7), 2)
%1 = Mod(3, 7)
? sqrtn(Mod(2,7), 2, &z); z
%2 = Mod(6, 7)
? sqrtn(Mod(2,7), 3)
*** at top-level: sqrtn(Mod(2,7),3)
***      ^-----
*** sqrtn: nth-root does not exist in gsqrtn.
? sqrtn(Mod(2,7), 3, &z)
%2 = 0
? z
%3 = 0
```

The following script computes all roots in all possible cases:

```
sqrtnall(x,n)=
{ my(V,r,z,r2);
  r = sqrtn(x,n, &z);
  if (!z, error("Impossible case in sqrtn"));
  if (type(x) == "t_INTMOD" || type(x)=="t_PADIC",
    r2 = r*z; n = 1;
    while (r2!=r, r2*=z;n++));
  V = vector(n); V[1] = r;
  for(i=2, n, V[i] = V[i-1]*z);
  V
}
addhelp(sqrtnall,"sqrtnall(x,n):compute the vector of nth-roots of x");
```

The library syntax is GEN `gsqrtn`(GEN x , GEN n , GEN $*z = \text{NULL}$, long prec). If x is a `t_PADIC`, the function GEN `Qp_sqrtn`(GEN x , GEN n , GEN $*z$) is also available.

3.11.62 `tan`(x). Tangent of x .

The library syntax is GEN `gtan`(GEN x , long prec).

3.11.63 `tanh`(x). Hyperbolic tangent of x .

The library syntax is GEN `gtanh`(GEN x , long prec).

3.11.64 teichmuller($x, \{tab\}$). Teichmüller character of the p -adic number x , i.e. the unique $(p-1)$ -th root of unity congruent to $x/p^{v_p(x)}$ modulo p . If x is of the form $[p, n]$, for a prime p and integer n , return the lifts to $\mathbf{Z}$ of the images of $i + O(p^n)$ for $i = 1, \dots, p-1$, i.e. all roots of 1 ordered by residue class modulo p . Such a vector can be fed back to **teichmuller**, as the optional argument **tab**, to speed up later computations.

```
? z = teichmuller(2 + O(101^5))
%1 = 2 + 83*101 + 18*101^2 + 69*101^3 + 62*101^4 + O(101^5)
? z^100
%2 = 1 + O(101^5)
? T = teichmuller([101, 5]);
? teichmuller(2 + O(101^5), T)
%4 = 2 + 83*101 + 18*101^2 + 69*101^3 + 62*101^4 + O(101^5)
```

As a rule of thumb, if more than

$$p / 2(\log_2(p) + \text{hammingweight}(p))$$

values of **teichmuller** are to be computed, then it is worthwhile to initialize:

```
? p = 101; n = 100; T = teichmuller([p,n]); \\ instantaneous
? for(i=1,10^3, vector(p-1, i, teichmuller(i+O(p^n), T)))
time = 60 ms.
? for(i=1,10^3, vector(p-1, i, teichmuller(i+O(p^n))))
time = 1,293 ms.
? 1 + 2*(log(p)/log(2) + hammingweight(p))
%8 = 22.316[...]
```

Here the precomputation induces a speedup by a factor $1293/60 \approx 21.5$.

Caveat. If the accuracy of **tab** (the argument n above) is lower than the precision of x , the *former* is used, i.e. the cached value is not refined to higher accuracy. If the accuracy of **tab** is larger, then the precision of x is used:

```
? Tlow = teichmuller([101, 2]); \\ lower accuracy !
? teichmuller(2 + O(101^5), Tlow)
%10 = 2 + 83*101 + O(101^5) \\ no longer a root of 1
? Thigh = teichmuller([101, 10]); \\ higher accuracy
? teichmuller(2 + O(101^5), Thigh)
%12 = 2 + 83*101 + 18*101^2 + 69*101^3 + 62*101^4 + O(101^5)
```

The library syntax is **GEN teichmuller**(**GEN x**, **GEN tab** = **NULL**).

Also available are the functions **GEN teich**(**GEN x**) (**tab** is **NULL**) as well as **GEN teichmullerinit**(**long p**, **long n**).

3.11.65 theta(q, z). Jacobi sine theta-function

$$\theta_1(z, q) = 2q^{1/4} \sum_{n \geq 0} (-1)^n q^{n(n+1)} \sin((2n+1)z).$$

The library syntax is **GEN theta**(**GEN q**, **GEN z**, **long prec**).

3.11.66 `thetanullk(q,k)`. k -th derivative at $z = 0$ of `theta(q,z)`.

The library syntax is GEN thetanullk(GEN q, long k, long prec).

GEN vecthetanullk(GEN q, long k, long prec) returns the vector of all $\frac{d^i \theta}{dz^i}(q, 0)$ for all odd $i = 1, 3, \dots, 2k-1$. GEN vecthetanullk_tau(GEN tau, long k, long prec) returns vecthetanullk_tau at $q = \exp(2i\pi\tau)$.

3.11.67 `weber(x, {flag = 0})`. One of Weber's three f functions. If $flag = 0$, returns

$$f(x) = \exp(-i\pi/24) \cdot \eta((x+1)/2) / \eta(x) \quad \text{such that} \quad j = (f^{24} - 16)^3 / f^{24},$$

where j is the elliptic j -invariant (see the function `ellj`). If $flag = 1$, returns

$$f_1(x) = \eta(x/2) / \eta(x) \quad \text{such that} \quad j = (f_1^{24} + 16)^3 / f_1^{24}.$$

Finally, if $flag = 2$, returns

$$f_2(x) = \sqrt{2}\eta(2x) / \eta(x) \quad \text{such that} \quad j = (f_2^{24} + 16)^3 / f_2^{24}.$$

Note the identities $f^8 = f_1^8 + f_2^8$ and $ff_1f_2 = \sqrt{2}$.

The library syntax is `GEN weber0(GEN x, long flag, long prec)`. Also available are `GEN weberf(GEN x, long prec)`, `GEN weberf1(GEN x, long prec)` and `GEN weberf2(GEN x, long prec)`.

3.11.68 zeta(s). For $s \neq 1$ a complex number, Riemann's zeta function $\zeta(s) = \sum_{n \geq 1} n^{-s}$, computed using the Euler-Maclaurin summation formula, except when s is of type integer, in which case it is computed using Bernoulli numbers for $s \leq 0$ or $s > 0$ and even, and using modular forms for $s > 0$ and odd. Power series are also allowed:

```
? zeta(2) - Pi^2/6
%1 = 0.E-38
? zeta(1+x+O(x^3))
%2 = 1.0000000000000000000000000000000000000000000000000000000*x^-1 + \
      0.57721566490153286060651209008240243104 + O(x)
```

For $s \neq 1$ a p -adic number, Kubota-Leopoldt zeta function at s , that is the unique continuous p -adic function on the p -adic integers that interpolates the values of $(1 - p^{-k})\zeta(k)$ at negative integers k such that $k \equiv 1 \pmod{p-1}$ (resp. k is odd) if p is odd (resp. $p = 2$). Power series are not allowed in this case.

```
? zeta(-3+0(5^10))  
%1 = 4*5^-1 + 4 + 3*5 + 4*5^3 + 4*5^5 + 4*5^7 + 0(5^9))))  
? (1-5^3) * zeta(-3)  
%2 = -1.0333333333333333333333333333333333333333333333333333333  
? bestappr(%)  
%3 = -31/30  
? zeta(-3+0(5^10)) - (-31/30)  
%4 = 0(5^9)
```

The library syntax is `GEN gzeta(GEN s, long prec)`.

3.11.69 zetahurwitz($s, x, \{\text{der} = 0\}$). Hurwitz zeta function $\zeta(s, x) = \sum_{n \geq 0} (n + x)^{-s}$ and analytically continued, with $s \neq 1$ and x not a negative or zero integer. Note that $\zeta(s, 1) = \zeta(s)$. s can also be a polynomial, rational function, or power series. If **der** is positive, compute the **der**'th derivative with respect to s . Note that the derivative with respect to x is simply $-\zeta(s + 1, x)$.

[illegible]

The derivative can be used to compute Barnes' multiple gamma functions. For instance:

```
? mygamma(z)=exp(zetahurwitz(0,z,1)-zeta'(0));  
/* Alternate way to compute the gamma function */  
? BarnesG(z)=exp(-zetahurwitz(-1,z,1)+(z-1)*lngamma(z)+zeta'(-1));  
/* Barnes G function, satisfying  $G(z+1)=\Gamma(z)G(z)$ : */  
? BarnesG(6)/BarnesG(5)  
% = 24.00000000000000000000000000000000000000000000000000000
```

The library syntax is `GEN zetahurwitz(GEN s, GEN x, long der, long bitprec)`.

3.11.70 zetamult($s, \{t = 0\}$). For s a vector of positive integers such that $s[1] \geq 2$, returns the multiple zeta value (MZV)

$$\zeta(s_1, \dots, s_k) = \sum_{n_1 \geq \dots \geq n_k \geq 0} n_1^{-s_1} \dots n_k^{-s_k}$$

of length k and weight $\sum_i s_i$. More generally, return Yamamoto's t -MZV interpolation evaluated at t : for $t = 0$, this is the ordinary MZV; for $t = 1$, we obtain the MZSV star value, with $\geq$ instead of strict inequalities; and of course, for $t = 'x$ we obtain Yamamoto's one-variable polynomial.

```
? zetamult([2,1]) - zeta(3) \\ Euler's identity
%1 = 0.E-38
? zetamult([2,1], 1) \\ star value
%2 = 2.4041138063191885707994763230228999815
? zetamult([2,1], 'x)
```

```
%3 = 1.20205[...] * x + 1.20205[...]
```

If the bit precision is B , this function runs in time $\tilde{O}(k(B+k)^2)$ if $t = 0$, and $\tilde{O}(kB^3)$ otherwise.

In addition to the above format (**avec**), the function also accepts a binary word format **evect** (each s_i is replaced by s_i bits, all of them 0 but the last one) giving the MZV representation as an iterated integral, and an **index** format (if e is the positive integer attached the **evect** vector of bits, the index is the integer $e + 2^{k-2}$). The function **zetamultconvert** allows to pass from one format to the other; the function **zetamultall** computes simultaneously all MZVs of weight $\sum_{i \leq k} s_i$ up to n .

The library syntax is `GEN zetamult_interpolate(GEN s, GEN t = NULL, long prec)`. Also available is `GEN zetamult(GEN s, long prec)` for $t = 0$.

3.11.71 zetamultall($k, \{flag = 0\}$). List of all multiple zeta values (MZVs) for weight $s_1 + \dots + s_r$ up to k . Binary digits of *flag* mean : 0 = star values if set; 1 = values up to duality if set (see **zetamultdual**, ignored if star values); 2 = values of weight k if set (else all values up to weight k); 3 = return the 2-component vector $[Z, M]$, where M is the vector of the corresponding indices m , i.e., such that **zetamult**($M[i]$) = $Z[i]$. Note that it is necessary to use **zetamultconvert** to have the corresponding **avec** ($s_1, \dots, s_r$).

With the default value *flag* = 0, the function returns a vector with $2^{k-1} - 1$ components whose i -th entry is the MZV of index i (see **zetamult**). If the bit precision is B , this function runs in time $O(2^k k B^2)$ for an output of size $O(2^k B)$.

```
? Z = zetamultall(5); #Z \\ 2^4 - 1 MZVs of weight <= 5
%1 = 15
? Z[10]
%2 = 0.22881039760335375976874614894168879193
? zetamultconvert(10)
%3 = Vecsmall([3, 2]) \\ index 10 corresponds to  $\zeta(3, 2)$ 
? zetamult(%) \\ double check
%4 = 0.22881039760335375976874614894168879193
? zetamult(10) \\ we can use the index directly
%5 = 0.22881039760335375976874614894168879193
```

If we use flag bits 1 and 2, we avoid unnecessary computations and copying, saving a potential factor 4: half the values are in lower weight and computing up to duality save another rough factor 2. Unfortunately, the indexing now no longer corresponds to the new shorter vector of MZVs:

```
? Z = zetamultall(5, 2); #Z \\ up to duality
%6 = 9
? Z = zetamultall(5, 2); #Z \\ only weight 5
%7 = 8
? Z = zetamultall(5, 2 + 4); #Z \\ both
%8 = 4
```

So how to recover the value attached to index 10 ? Flag bit 3 returns the actual indices used:

```
? [Z, M] = zetamultall(5, 2 + 8); M \\ other indices were not included
%9 = Vecsmall([1, 2, 4, 5, 6, 8, 9, 10, 12])
? Z[8] \\ index m = 10 is now in M[8]
%10 = 0.22881039760335375976874614894168879193
```

```
? [Z, M] = zetamultall(5, 2 + 4 + 8); M
%11 = Vecsmall([8, 9, 10, 12])
? Z[3] \\ index m = 10 is now in M[3]
%12 = 0.22881039760335375976874614894168879193
```

The following construction automates the above programmatically, looking up the MZVs of index 10 ($= \zeta(3, 2)$) in all cases, without inspecting the various index sets M visually:

```
? Z[vecsearch(M, 10)] \\ works in all the above settings
%13 = 0.22881039760335375976874614894168879193
```

The library syntax is GEN `zetamultall(long k, long flag, long prec)`.

3.11.72 zetamultconvert($a, \{flag = 1\}$). a being either an `even`, `avec`, or index m , converts into `even` ($flag = 0$), `avec` ($flag = 1$), or index m ($flag = 2$).

```
? zetamultconvert(10)
%1 = Vecsmall([3, 2])
? zetamultconvert(13)
%2 = Vecsmall([2, 2, 1])
? zetamultconvert(10, 0)
%3 = Vecsmall([0, 0, 1, 0, 1])
? zetamultconvert(13, 0)
%4 = Vecsmall([0, 1, 0, 1, 1])
```

The last two lines imply that $[3, 2]$ and $[2, 2, 1]$ are dual (reverse order of bits and swap 0 and 1 in `even` form). Hence they have the same zeta value:

```
? zetamult([3,2])
%5 = 0.22881039760335375976874614894168879193
? zetamult([2,2,1])
%6 = 0.22881039760335375976874614894168879193
```

The library syntax is GEN `zetamultconvert(GEN a, long flag)`.

3.11.73 zetamultdual(s). s being either an `even`, `avec`, or index m , return the dual sequence in `avec` format. The dual of a sequence of length r and weight k has length $k - r$ and weight k . Duality is an involution and zeta values attached to dual sequences are the same:

```
? zetamultdual([4])
%1 = Vecsmall([2, 1, 1])
? zetamultdual(%)
%2 = Vecsmall([4])
? zetamult(%1) - zetamult(%2)
%3 = 0.E-38
```

In `even` form, duality simply reverses the order of bits and swaps 0 and 1:

```
? zetamultconvert([4], 0)
%4 = Vecsmall([0, 0, 0, 1])
? zetamultconvert([2,1,1], 0)
%5 = Vecsmall([0, 1, 1, 1])
```

The library syntax is GEN `zetamultdual(GEN s)`.

3.12 Sums, products, integrals and similar functions.

Although the `gp` calculator is programmable, it is useful to have a number of preprogrammed loops, including sums, products, and a certain number of recursions. Also, a number of functions from numerical analysis like numerical integration and summation of series will be described here.

One of the parameters in these loops must be the control variable, hence a simple variable name. In the descriptions, the letter X will always denote any simple variable name, and represents the formal parameter used in the function. The expression to be summed, integrated, etc. is any legal PARI expression, including of course expressions using loops.

Library mode. Since it is easier to program directly the loops in library mode, these functions are mainly useful for GP programming. On the other hand, numerical routines code a function (to be integrated, summed, etc.) with two parameters named

```
GEN (*eval)(void*,GEN)
void *E;  \\ context: eval(E, x) must evaluate your function at x.
```

see the Libpari manual for details.

Numerical integration. The “double exponential” (DE) univariate integration method is implemented in `intnum` and its variants. Romberg integration is still available under the name `intnumromb`, but superseded. It is possible to compute numerically integrals to thousands of decimal places in reasonable time, as long as the integrand is regular. It is also reasonable to compute numerically integrals in several variables, although more than two becomes lengthy. The integration domain may be noncompact, and the integrand may have reasonable singularities at endpoints. To use `intnum`, you must split the integral into a sum of subintegrals where the function has no singularities except at the endpoints. Polynomials in logarithms are not considered singular, and neglecting these logs, singularities are assumed to be algebraic (asymptotic to $C(x-a)^{-\alpha}$ for some $\alpha > -1$ when x is close to a), or to correspond to simple discontinuities of some (higher) derivative of the function. For instance, the point 0 is a singularity of `abs(x)`.

Assume the bitprecision is b , so we try to achieve an absolute error less than 2^{-b} . DE methods use $O(b \log b)$ function evaluations and should work for both compact and non-compact intervals as long as the integrand is the restriction of an analytic function to a suitable domain and its behaviour at infinity is correctly described. When integrating regular functions on a *compact* interval, away from poles of the integrand, Gauss-Legendre integration (`intnumgauss`) is the best choice, using $O(b)$ function evaluations. To integrate oscillating functions on non-compact interval, the slower but robust `intnumosc` is available, performing Gaussian integration on intervals of length the half-period (or quasi-period) and using Sidi’s mW algorithm to extrapolate their sum. If poles are close to the integration interval, Gaussian integration may run into difficulties and it is then advisable to split the integral using `intnum` to get away from poles, then `intnumosc` for the remainder.

For maximal efficiency, abscissas and integration weights can be precomputed, respectively using `intnuminit` ($O(b^2)$) or `intnumgaussinit` ($O(b^3)$).

Numerical summation.

Many numerical summation methods are available to approximate $\sum_{n \geq n_0} f(n)$ at accuracy 2^{-b} : the overall best choice should be `sumnum`, which uses Euler-MacLaurin (and $O(b \log b)$ function evaluations); initialization time (`sumnuminit`) is $O(b^3)$. Also available are

- Abel-Plana summation (`sumnumap`), also $O(b \log b)$ function evaluations and $O(b^3)$ initialization (`sumnumapinit`) with a larger implied constant;
- Lagrange summation (`sumnumlagrange`) uses $O(b)$ evaluations but more brittle and the asymptotic behaviour of f must be correctly indicated. Initialization (`sumnumlagrangeinit`) can vary from $O(b^2)$ to $O(b^3)$ depending on the asymptotic behaviour.
- Sidi summation (`sumnumsidi`) uses $O(b)$ evaluations and should be more robust than Lagrange summation. No initialization is needed.
- Monien summation (`sumnummonien`) uses $O(b/\log b)$ evaluations but is even more brittle than Lagrange and also has a $O(b^3)$ initialization (`summonieninit`).
- To sum rational functions, use `sumnumrat`.

All the function so far require f to be the restriction to integers of a regular function on the reals, and even on the complex numbers for Monien summation. The following algorithms allow functions defined only on the integers, under assumptions that are hard to verify. They are best used heuristically since they in fact are often valid when those assumptions do not hold, and for instance often yield a result for divergent series (e.g., Borel resummation).

- To sum alternating series, use `sumalt`, which requires $O(b)$ function evaluations.
- To sum functions of a fixed sign, `sumpos` uses van Wijngarten's trick to reduce to an alternating series, for a cost of $O(b \log b)$ function evaluations but beware that f must be evaluated at large integers, of the order of $2^{b/\alpha}$ if we assume that $f(n) = O(1/n^{\alpha+1})$ for some $\alpha > 0$.

3.12.1 `asypnum(expr, {alpha = 1})`. Asymptotic expansion of $expr$, corresponding to a sequence $u(n)$, assuming it has the shape

$$u(n) \approx \sum_{i \geq 0} a_i n^{-i\alpha}$$

with rational coefficients a_i with reasonable height; the algorithm is heuristic and performs repeated calls to `limitnum`, with `alpha` as in `limitnum`. As in `limitnum`, $u(n)$ may be given either by a closure $n \mapsto u(n)$ or as a closure $N \mapsto [u(1), \dots, u(N)]$, the latter being often more efficient.

```
? f(n) = n! / (n^n * exp(-n) * sqrt(n));
? asypnum(f)
%2 = []      \\ failure !
? localprec(57); l = limitnum(f)
%3 = 2.5066282746310005024157652848110452530
? asypnum(n->f(n)/l) \\ normalize
%4 = [1, 1/12, 1/288, -139/51840, -571/2488320, 163879/209018880,
      5246819/75246796800]
```

and we indeed get a few terms of Stirling's expansion. Note that it definitely helps to normalize with a limit computed to higher accuracy (as a rule of thumb, multiply the bit accuracy by 1.612):

```
? l = limitnum(f)
? asympnum(n->f(n) / l) \\ failure again !!!
%6 = []
```

We treat again the example of the Motzkin numbers M_n given in `limitnum`:

```
\\ [M_k, M_{k*2}, ..., M_{k*N}] / (3^n / n^(3/2))
? vM(N, k = 1) =
{ my(q = k*N, V);
  if (q == 1, return ([1/3]));
  V = vector(q); V[1] = V[2] = 1;
  for(n = 2, q - 1,
    V[n+1] = ((2*n + 1)*V[n] + 3*(n - 1)*V[n-1]) / (n + 2));
  f = (n -> 3^n / n^(3/2));
  return (vector(N, n, V[n*k] / f(n*k)));
}
? localprec(100); l = limitnum(n->vM(n,10)); \\ 3/sqrt(12*Pi)
? \p38
? asympnum(n->vM(n,10)/l)
%2 = [1, -3/32, 101/10240, -1617/1638400, 505659/5242880000, ...]
```

If α is not a rational number, loss of accuracy is expected, so it should be precomputed to double accuracy, say:

```
? \p38
? asympnum(n->log(1+1/n^Pi),Pi)
%1 = [0, 1, -1/2, 1/3, -1/4, 1/5]
? localprec(76); a = Pi;
? asympnum(n->log(1+1/n^Pi), a) \\ more terms
%3 = [0, 1, -1/2, 1/3, -1/4, 1/5, -1/6, 1/7, -1/8, 1/9, -1/10, 1/11, -1/12]
? asympnum(n->log(1+1/sqrt(n)),1/2) \\ many more terms
%4 = 49
```

The expression is evaluated for $n = 1, 2, \dots, N$ for an $N = O(B)$ if the current bit accuracy is B . If it is not defined for one of these values, translate or rescale accordingly:

```
? asympnum(n->log(1-1/n)) \\ can't evaluate at n = 1 !
*** at top-level: asympnum(n->log(1-1/n))
*** ^-----
*** in function asympnum: log(1-1/n)
*** ^-----
*** log: domain error in log: argument = 0
? asympnum(n->-log(1-1/(2*n)))
%5 = [0, 1/2, 1/8, 1/24, ...]
? asympnum(n->-log(1-1/(n+1)))
%6 = [0, 1, -1/2, 1/3, -1/4, ...]
```

The library syntax is `asympnum(void *E, GEN (*u)(void *,GEN,long), GEN alpha, long prec)`, where $u(E, n, prec)$ must return either $u(n)$ or $[u(1), \dots, u(n)]$ in precision `prec`. Also available is `GEN asympnum0(GEN u, GEN alpha, long prec)`, where u is a closure as above or a vector of sufficient length.

3.12.2 asympnumraw(*expr*, *N*, {*alpha* = 1}). Return the $N + 1$ first terms of asymptotic expansion of *expr*, corresponding to a sequence $u(n)$, as floating point numbers. Assume that the expansion has the shape

$$u(n) \approx \sum_{i \geq 0} a_i n^{-i\alpha}$$

and return approximation of $[a_0, a_1, \dots, a_N]$. The algorithm is heuristic and performs repeated calls to `limitnum`, with `alpha` as in `limitnum`. As in `limitnum`, $u(n)$ may be given either by a closure $n \mapsto u(n)$ or as a closure $N \mapsto [u(1), \dots, u(N)]$, the latter being often more efficient. This function is related to, but more flexible than, `asympnum`, which requires rational asymptotic expansions.

```
? f(n) = n! / (n^n*exp(-n)*sqrt(n));
? asympnum(f)
%2 = []    \\ failure !
? v = asympnumraw(f, 10);
? v[1] - sqrt(2*Pi)
%4 = 0.E-37
? bestappr(v / v[1], 2^60)
%5 = [1, 1/12, 1/288, -139/51840, -571/2488320, 163879/209018880, ...]
```

and we indeed get a few terms of Stirling's expansion (the first 9 terms are correct). If $u(n)$ has an asymptotic expansion in $n^{-\alpha}$ with α not an integer, the default *alpha* = 1 is inaccurate:

```
? f(n) = (1+1/n^(7/2))^(n^(7/2));
? v1 = asympnumraw(f,10);
? v1[1] - exp(1)
%8 = 4.62... E-12
? v2 = asympnumraw(f,10,7/2);
? v2[1] - exp(1)
%7 0.E-37
```

As in `asympnum`, if `alpha` is not a rational number, loss of accuracy is expected, so it should be precomputed to double accuracy, say.

The library syntax is `asympnumraw(void *E, GEN (*u)(void *,GEN,long), long N, GEN alpha, long prec)`, where `u(E, n, prec)` must return either $u(n)$ or $[u(1), \dots, u(n)]$ in precision `prec`. Also available is `GEN asympnumraw0(GEN u, GEN alpha, long prec)` where *u* is either a closure as above or a vector of sufficient length.

3.12.3 contfraceval(*CF*, *t*, {*lim* = -1}). Given a continued fraction *CF* output by `contfracinit`, evaluate the first *lim* terms of the continued fraction at *t* (all terms if *lim* is negative or omitted; if positive, *lim* must be less than or equal to the length of *CF*).

The library syntax is `GEN contfraceval(GEN CF, GEN t, long lim)`.

3.12.6 intcirc($X = a, R, expr, \{tab\}$). Numerical integration of $(2i\pi)^{-1}expr$ with respect to X on the circle $|X - a| = R$. In other words, when $expr$ is a meromorphic function, sum of the residues in the corresponding disk; tab is as in `intnum`, except that if computed with `intnuminit` it should be with the endpoints `[-1, 1]`.

```
? \p105
? intcirc(s=1, 0.5, zeta(s)) - 1
time = 496 ms.
%1 = 1.2883911040127271720 E-101 + 0.E-118*I
```

The library syntax is `intcirc(void *E, GEN (*eval)(void*,GEN), GEN a,GEN R,GEN tab, long prec)`.

3.12.7 intfuncinit($t = a, b, f, \{m = 0\}$). Initialize tables for use with integral transforms (such as Fourier, Laplace or Mellin transforms) in order to compute

$$\int_a^b f(t)k(t, z) dt$$

for some kernel $k(t, z)$. The endpoints a and b are coded as in `intnum`, f is the function to which the integral transform is to be applied and the nonnegative integer m is as in `intnum`: multiply the number of sampling points roughly by 2^m , hopefully increasing the accuracy. This function is particularly useful when the function f is hard to compute, such as a gamma product.

Limitation. The endpoints a and b must be at infinity, with the same asymptotic behavior. Oscillating types are not supported. This is easily overcome by integrating vectors of functions, see example below.

Examples.

- numerical Fourier transform

$$F(z) = \int_{-\infty}^{+\infty} f(t)e^{-2i\pi zt} dt.$$

First the easy case, assume that f decrease exponentially:

```
f(t) = exp(-t^2);
A = [-oo,1];
B = [+oo,1];
\p200
T = intfuncinit(t = A,B , f(t));
F(z) =
{ my(a = -2*I*Pi*z);
  intnum(t = A,B, exp(a*t), T);
}
? F(1) - sqrt(Pi)*exp(-Pi^2)
%1 = -1.3... E-212
```

Now the harder case, f decrease slowly: we must specify the oscillating behavior. Thus, we cannot precompute usefully since everything depends on the point we evaluate at:

```
f(t) = 1 / (1+ abs(t));
```

```

\p200
\\ Fourier cosine transform
FC(z) =
{ my(a = 2*Pi*z);
  intnum(t = [-oo, a*I], [+oo, a*I], cos(a*t)*f(t));
}
FC(1)

```

• Fourier coefficients: we must integrate over a period, but `intfuncinit` does not support finite endpoints. The solution is to integrate a vector of functions !

```

FourierSin(f, T, k) = \\ first k sine Fourier coeffs
{
  my (w = 2*Pi/T);
  my (v = vector(k+1));
  intnum(t = -T/2, T/2,
    my (z = exp(I*w*t));
    v[1] = z;
    for (j = 2, k, v[j] = v[j-1]*z);
    f(t) * imag(v)) * 2/T;
}
FourierSin(t->sin(2*t), 2*Pi, 10)

```

The same technique can be used instead of `intfuncinit` to integrate $f(t)k(t, z)$ whenever the list of z -values is known beforehand.

Note that the above code includes an unrelated optimization: the $\sin(jwt)$ are computed as imaginary parts of $\exp(ijwt)$ and the latter by successive multiplications.

• numerical Mellin inversion

$$F(z) = (2i\pi)^{-1} \int_{c-i\infty}^{c+i\infty} f(s)z^{-s} ds = (2\pi)^{-1} \int_{-\infty}^{+\infty} f(c+it)e^{-\log z(c+it)} dt.$$

We take $c = 2$ in the program below:

```

f(s) = gamma(s)^3; \\ f(c+it) decrease as exp(-3Pi|t|/2)
c = 2; \\ arbitrary
A = [-oo, 3*Pi/2];
B = [+oo, 3*Pi/2];
T = intfuncinit(t=A, B, f(c + I*t));
F(z) =
{ my (a = -log(z));
  intnum(t=A, B, exp(a*I*t), T)*exp(a*c) / (2*Pi);
}

```

The library syntax is `intfuncinit(void *E, GEN (*eval)(void*, GEN), GEN a, GEN b, long m, long prec)`.

3.12.8 intnum($X = a, b, \text{expr}, \{tab\}$). Numerical integration of expr on $]a, b[$ with respect to X , using the double-exponential method, and thus $O(D \log D)$ evaluation of the integrand in precision D . The integrand may have values belonging to a vector space over the real numbers; in particular, it can be complex-valued or vector-valued. But it is assumed that the function is regular on $]a, b[$. If the endpoints a and b are finite and the function is regular there, the situation is simple:

```
? intnum(x = 0,1, x^2)
%1 = 0.33333333333333333333333333333333
? intnum(x = 0,Pi/2, [cos(x), sin(x)])
%2 = [1.00000000000000000000000000000000, 1.00000000000000000000000000000000]
```

An endpoint equal to $\pm\infty$ is coded as $+\infty$ or $-\infty$, as expected:

```
? intnum(x = 1,+oo, 1/x^2)
%3 = 1.00000000000000000000000000000000
```

In basic usage, it is assumed that the function does not decrease exponentially fast at infinity:

```
? intnum(x=0,+oo, exp(-x))
*** at top-level: intnum(x=0,+oo,exp(-
*** ^-----
*** exp: overflow in expo()).
```

We shall see in a moment how to avoid that last problem, after describing the last *optional* argument tab .

The tab argument. The routine uses weights w_i , which are mostly independent of the function being integrated, evaluated at many sampling points x_i and approximates the integral by $\sum w_i f(x_i)$. If tab is

- a nonnegative integer m , we multiply the number of sampling points by 2^m , hopefully increasing accuracy. Note that the running time increases roughly by a factor 2^m . One may try consecutive values of m until they give the same value up to an accepted error.
- a set of integration tables containing precomputed x_i and w_i as output by `intnuminit`. This is useful if several integrations of the same type are performed (on the same kind of interval and functions, for a given accuracy): we skip a precomputation of $O(D \log D)$ elementary functions in accuracy D , whose running time has the same order of magnitude as the evaluation of the integrand. This is in particular useful for multivariate integrals.

Specifying the behavior at endpoints. This is done as follows. An endpoint a is either given as such (a scalar, real or complex, ∞ or $-\infty$ for $\pm\infty$), or as a two component vector $[a, \alpha]$, to indicate the behavior of the integrand in a neighborhood of a .

If a is finite, the code $[a, \alpha]$ means the function has a singularity of the form $(x - a)^\alpha$, up to logarithms. (If $\alpha \geq 0$, we only assume the function is regular, which is the default assumption.) If a wrong singularity exponent is used, the result will lose decimals:

```
? c = -9/10;
? intnum(x=0, 1, x^c)          \\ assume x^{-9/10} is regular at 0
%1 = 9.9999839078827082322596783301939063944
? intnum(x=[0,c], 1, x^c)      \\ no, it's not
%2 = 10.0000000000000000000000000000000000000000000000000
? intnum(x=[0,c/2], 1, x^c)    \\ using a wrong exponent is bad
```


whis is translated internally to

```
? intnum(t=[0,-1/2],1, 1/sqrt(t))-intnum(t=[0,-1/2],X^2+0(X^4), 1/sqrt(t))
```

For this form the argument *tab* can be used only as an integer, not a table precomputed by `intnuminit`.

We shall now see many examples to get a feeling for what the various parameters achieve. All examples below assume precision is set to 115 decimal digits. We first type

```
? \p 115
```

Apparent singularities. In many cases, apparent singularities can be ignored. For instance, if $f(x) = 1/(\exp(x) - 1) - \exp(-x)/x$, then $\int_0^\infty f(x) dx = \gamma$, Euler's constant **Euler**. But

```
? f(x) = 1/(exp(x)-1) - exp(-x)/x
? intnum(x = 0, [oo,1], f(x)) - Euler
%1 = 0.E-115
```

But close to 0 the function f is computed with an enormous loss of accuracy, and we are in fact lucky that it get multiplied by weights which are sufficiently close to 0 to hide this:

```
? f(1e-200)
%2 = -3.885337784451458142 E84
```

A more robust solution is to define the function differently near special points, e.g. by a Taylor expansion

```
? F = truncate( f(t + 0(t^10)) ); \\ expansion around t = 0
? poldegree(F)
%4 = 7
? g(x) = if (x > 1e-18, f(x), subst(F,t,x)); \\ note that 7 · 18 > 105
? intnum(x = 0, [oo,1], g(x)) - Euler
%2 = 0.E-115
```

It is up to the user to determine constants such as the 10^{-18} and 10 used above.

True singularities. With true singularities the result is worse. For instance

```
? intnum(x = 0, 1, x^(-1/2)) - 2
%1 = -3.5... E-68 \\ only 68 correct decimals
? intnum(x = [0,-1/2], 1, x^(-1/2)) - 2
%2 = 0.E-114 \\ better
```

Oscillating functions.

```
? intnum(x = 0, oo, sin(x) / x) - Pi/2
%1 = 16.19.. \\ nonsense
? intnum(x = 0, [oo,1], sin(x)/x) - Pi/2
%2 = -0.006.. \\ bad
? intnum(x = 0, [oo,-I], sin(x)/x) - Pi/2
%3 = 0.E-115 \\ perfect
? intnum(x = 0, [oo,-I], sin(2*x)/x) - Pi/2 \\ oops, wrong k
%4 = 0.06...
? intnum(x = 0, [oo,-2*I], sin(2*x)/x) - Pi/2
%5 = 0.E-115 \\ perfect

? intnum(x = 0, [oo,-I], sin(x)^3/x) - Pi/4
%6 = -0.0008... \\ bad
? sin(x)^3 - (3*sin(x)-sin(3*x))/4
%7 = 0(x^17)
```

We may use the above linearization and compute two oscillating integrals with endpoints $[oo, -I]$ and $[oo, -3*I]$ respectively, or notice the obvious change of variable, and reduce to the single integral $\frac{1}{2} \int_0^\infty \sin(x)/x \, dx$. We finish with some more complicated examples:

```
? intnum(x = 0, [oo,-I], (1-cos(x))/x^2) - Pi/2
%1 = -0.0003... \\ bad
? intnum(x = 0, 1, (1-cos(x))/x^2) \
+ intnum(x = 1, oo, 1/x^2) - intnum(x = 1, [oo,I], cos(x)/x^2) - Pi/2
%2 = 0.E-115 \\ perfect

? intnum(x = 0, [oo, 1], sin(x)^3*exp(-x)) - 0.3
%3 = -7.34... E-55 \\ bad
? intnum(x = 0, [oo,-I], sin(x)^3*exp(-x)) - 0.3
%4 = 8.9... E-103 \\ better. Try higher m
? tab = intnuminit(0,[oo,-I], 1); \\ double number of sampling points
? intnum(x = 0, oo, sin(x)^3*exp(-x), tab) - 0.3
%6 = 0.E-115 \\ perfect
```

Warning. Like `sumalt`, `intnum` often assigns a reasonable value to diverging integrals. Use these values at your own risk! For example:

```
? intnum(x = 0, [oo, -I], x^2*sin(x))
%1 = -2.0000000000...
```

Note the formula

$$\int_0^\infty \sin(x)x^{-s} \, dx = \cos(\pi s/2)\Gamma(1-s),$$

a priori valid only for $0 < \Re(s) < 2$, but the right hand side provides an analytic continuation which may be evaluated at $s = -2$...

Multivariate integration. Using successive univariate integration with respect to different formal parameters, it is immediate to do naive multivariate integration. But it is important to use a suitable `intnuminit` to precompute data for the *internal* integrations at least!

For example, to compute the double integral on the unit disc $x^2 + y^2 \leq 1$ of the function $x^2 + y^2$, we can write

```
? tab = intnuminit(-1,1);
? intnum(x=-1,1, intnum(y=-sqrt(1-x^2),sqrt(1-x^2), x^2+y^2, tab),tab) - Pi/2
%2 = -7.1... E-115 \\ OK
```

The first *tab* is essential, the second optional. Compare:

```
? tab = intnuminit(-1,1);
time = 4 ms.
? intnum(x=-1,1, intnum(y=-sqrt(1-x^2),sqrt(1-x^2), x^2+y^2));
time = 3,092 ms. \\ slow
? intnum(x=-1,1, intnum(y=-sqrt(1-x^2),sqrt(1-x^2), x^2+y^2, tab), tab);
time = 252 ms. \\ faster
? intnum(x=-1,1, intnum(y=-sqrt(1-x^2),sqrt(1-x^2), x^2+y^2, tab));
time = 261 ms. \\ the internal integral matters most
```

The library syntax is `intnum(void *E, GEN (*eval)(void*,GEN), GEN a,GEN b,GEN tab, long prec)`, where an omitted *tab* is coded as `NULL`.

3.12.9 intnumgauss($X = a, b, expr, \{tab\}$). Numerical integration of *expr* on the compact interval $[a, b]$ with respect to X using Gauss-Legendre quadrature; *tab* is either omitted or precomputed with `intnumgaussinit`. As a convenience, it can be an integer n in which case we call `intnumgaussinit( $n$ )` and use n -point quadrature.

```
? test(n, b = 1) = T=intnumgaussinit(n);\
  intnumgauss(x=-b,b, 1/(1+x^2),T) - 2*atan(b);
? test(0) \\ default
%1 = -9.490148553624725335 E-22
? test(40)
%2 = -6.186629001816965717 E-31
? test(50)
%3 = -1.1754943508222875080 E-38
? test(50, 2) \\ double interval length
%4 = -4.891779568527713636 E-21
? test(90, 2) \\ n must almost be doubled as well!
%5 = -9.403954806578300064 E-38
```

On the other hand, we recommend to split the integral and change variables rather than increasing n too much:

```
? f(x) = 1/(1+x^2);
? b = 100;
? intnumgauss(x=0,1, f(x)) + intnumgauss(x=1,1/b, f(1/x)*(-1/x^2)) - atan(b)
%3 = -1.0579449157400587572 E-37
```

The library syntax is `GEN intnumgauss0(GEN X, GEN b, GEN expr, GEN tab = NULL, long prec)`.

3.12.10 intnumgaussinit($\{n\}$). Initialize tables for n -point Gauss-Legendre integration of a smooth function f on a compact interval $[a, b]$. If n is omitted, make a default choice $n \approx B/4$, where B is `realbitprecision`, suitable for analytic functions on $[-1, 1]$. The error is bounded by

$$\frac{(b-a)^{2n+1}(n!)^4}{(2n+1)!(2n)!} \frac{f^{(2n)}(\xi)}{(2n)!}, \quad a < \xi < b.$$

If r denotes the distance of the nearest pole to the interval $[a, b]$, then this is of the order of $((b-a)/(4r))^{2n}$. In particular, the integral must be subdivided if the interval length $b-a$ becomes close to $4r$. The default choice $n \approx B/4$ makes this quantity of order 2^{-B} when $b-a=r$, as is the case when integrating $1/(1+t)$ on $[0, 1]$ for instance. If the interval length increases, n should be increased as well.

Specifically, the function returns a pair of vectors $[x, w]$, where x contains the nonnegative roots of the n -th Legendre polynomial P_n and w the corresponding Gaussian integration weights $Q_n(x_j)/P'_n(x_j) = 2/((1-x_j^2)P'_n(x_j))^2$ such that

$$\int_{-1}^1 f(t) dt \approx \sum_j w_j f(x_j).$$

```
? T = intnumgaussinit();
? intnumgauss(t=-1,1,exp(t), T) - exp(1)+exp(-1)
%1 = -5.877471754111437540 E-39
? intnumgauss(t=-10,10,exp(t), T) - exp(10)+exp(-10)
%2 = -8.358367809712546836 E-35
? intnumgauss(t=-1,1,1/(1+t^2), T) - Pi/2 \\ b - a = 2r
%3 = -9.490148553624725335 E-22 \\ ... loses half the accuracy

? T = intnumgaussinit(50);
? intnumgauss(t=-1,1,1/(1+t^2), T) - Pi/2
%5 = -1.1754943508222875080 E-38
? intnumgauss(t=-5,5,1/(1+t^2), T) - 2*atan(5)
%6 = -1.2[...]E-8
```

On the other hand, we recommend to split the integral and change variables rather than increasing n too much, see `intnumgauss`.

The library syntax is `GEN intnumgaussinit(long n, long prec)`.

3.12.11 intnuminit($a, b, \{m = 0\}$). Initialize tables for integration from a to b , where a and b are coded as in `intnum`. Only the compactness, the possible existence of singularities, the speed of decrease or the oscillations at infinity are taken into account, and not the values. For instance `intnuminit(-1,1)` is equivalent to `intnuminit(0,Pi)`, and `intnuminit([0,-1/2],oo)` is equivalent to `intnuminit([-1,-1/2], -oo)`; on the other hand, the order matters and `intnuminit([0,-1/2], [1,-1/3])` is *not* equivalent to `intnuminit([0,-1/3], [1,-1/2])` !

If m is present, it must be nonnegative and we multiply the default number of sampling points by 2^m (increasing the running time by a similar factor).

The result is technical and liable to change in the future, but we document it here for completeness. Let $x = \phi(t)$, $t \in]-\infty, \infty[$ be an internally chosen change of variable, achieving double exponential decrease of the integrand at infinity. The integrator `intnum` will compute

$$h \sum_{|n| < N} \phi'(nh) F(\phi(nh))$$

for some integration step h and truncation parameter N . In basic use, let

```
[h, x0, w0, xp, wp, xm, wm] = intnuminit(a,b);
```

- h is the integration step
- $x_0 = \phi(0)$ and $w_0 = \phi'(0)$,
- xp contains the $\phi(nh)$, $0 < n < N$,
- xm contains the $\phi(nh)$, $0 < -n < N$, or is empty.
- wp contains the $\phi'(nh)$, $0 < n < N$,
- wm contains the $\phi'(nh)$, $0 < -n < N$, or is empty.

The arrays xm and wm are left empty when ϕ is an odd function. In complicated situations, `intnuminit` may return up to 3 such arrays, corresponding to a splitting of up to 3 integrals of basic type.

If the functions to be integrated later are of the form $F = f(t)k(t, z)$ for some kernel k (e.g. Fourier, Laplace, Mellin, ...), it is useful to also precompute the values of $f(\phi(nh))$, which is accomplished by `intfuncinit`. The hard part is to determine the behavior of F at endpoints, depending on z .

The library syntax is `GEN intnuminit(GEN a, GEN b, long m, long prec)`.

3.12.12 intnumosc($x = a, expr, H, \{flag = 0\}, \{tab\}$). Numerical integration from a to ∞ of oscillating quasi-periodic function $expr$ of half-period H , meaning that we at least expect the distance between the function's consecutive zeros to be close to H : the sine or cosine functions ($H = \pi$) are paradigmatic examples, but the Bessel J_ν or Y_ν functions ($H = \pi/2$) can also be handled. The integral from a to ∞ is computed by summing the integral between two consecutive multiples of H ; $flag$ determines the summation algorithm used: either 0 (Sidi extrapolation, safe mode), 1 (Sidi extrapolation, unsafe mode), 2 (`sumalt`), 3 (`sumnumlagrange`) or 4 (`sumpos`). For the last two modes (Lagrange and Sumpos), one should input the period $2H$ instead of the half-period H .

The default is $flag = 0$; Sidi summation should be the most robust algorithm; you can try it in unsafe mode when the integrals between two consecutive multiples of H form an alternating series, this should be about twice faster than the default and not lose accuracy. Sumpos should be by far the slowest method, but also very robust and may be able to handle integrals where Sidi fails. Sumalt should be fast but often wrong, especially when the integrals between two consecutive multiples of H do not form an alternating series), and Lagrange should be as fast as Sumalt but more often wrong.

When one of the Sidi modes runs into difficulties, it will return the result to the accuracy believed to be correct (the other modes do not perform extrapolation and do not have this property) :

```

? f(x)=besselj(0,x)^4*log(x+1);
? \pb384
? intnumosc(x = 0, f(x), Pi)
%1 = 0.4549032054850867417 \\ fewer digits than expected !
? bitprecision(%)
%2 = 64
? \g1 \\ increase debug level to see diagnostics
? intnumosc(x = 0, f(x), Pi)
sumsidi: reached accuracy of 23 bits.
%2 = 0.4549032054850867417

```

The algorithm could extrapolate the series to 23 bits of accuracy, then diverged. So only the absolute error is likely to be around 2^{-23} instead of the possible 2^{-64} (or the requested 2^{-384}). We'll come back to this example at the end.

In case of difficulties, you may try to replace the half-(quasi)-period H by a multiple, such as the quasi-period $2H$: since we do not expect alternating behaviour, `sumalt` mode will almost surely be broken, but others may improve, in particular Lagrange or Sumpso.

`tab` is either omitted or precomputed with `intnumgaussinit`; if using Sidi summation in safe mode (`flag = 0`) and precompute `tab`, you should use a precision roughly 50% larger than the target (this is not necessary for any of the other summations).

First an alternating example:

```

? \pb384
\\ Sidi, safe mode
? exponent(intnumosc(x=0,sinc(x),Pi) - Pi/2)
time = 183 ms.
%1 = -383
? exponent(intnumosc(x=0,sinc(x),2*Pi) - Pi/2)
time = 224 ms.
%2 = -383 \\ also works with 2H, a little slower

\\ Sidi, unsafe mode
? exponent(intnumosc(x=0,sinc(x),Pi,1) - Pi/2)
time = 79 ms.
%3 = -383 \\ alternating: unsafe mode is fine and almost twice faster
? exponent(intnumosc(x=0,sinc(x),2*Pi,1) - Pi/2)
time = 86 ms.
%4 = -285 \\ but this time 2H loses accuracy

\\ Sumalt
? exponent(intnumosc(x=0,sinc(x),Pi,2) - Pi/2)
time = 115 ms. \\ sumalt is just as accurate and fast
%5 = -383
? exponent(intnumosc(x=0,sinc(x),2*Pi,2) - Pi/2)
time = 115 ms.
%6 = -10 \\ ...but breaks completely with 2H

\\ Lagrange
? exponent(intnumosc(x=0,sinc(x),Pi,2) - Pi/2)
time = 100 ms. \\ junk

```

```

%7 = 224
? exponent(intnumosc(x=0,sinc(x),2*Pi,2) - Pi/2)
time = 100 ms.
%8 = -238 \\ ...a little better with 2H
\\ Sumpos
? exponent(intnumosc(x=0,sinc(x),Pi,4) - Pi/2)
time = 17,961 ms.
%9 = 7 \\ junk; slow
? exponent(intnumosc(x=0,sinc(x),2*Pi,4) - Pi/2)
time = 19,105 ms.
%10 = -4 \\ still junk

Now a non-alternating one:

? exponent(intnumosc(x=0,sinc(x)^2,Pi) - Pi/2)
time = 277 ms.
%1 = -383 \\ safe mode is still perfect
? exponent(intnumosc(x=0,sinc(x)^2,Pi,1) - Pi/2)
time = 97 ms.
%2 = -284 \\ non-alternating; this time, Sidi's unsafe mode loses accuracy
? exponent(intnumosc(x=0,sinc(x)^2,Pi,2) - Pi/2)
time = 113 ms.
%3 = -10 \\ this time sumalt fails completely
? exponent(intnumosc(x=0,sinc(x)^2,Pi,3) - Pi/2)
time = 103 ms.
%4 = -237 \\ Lagrange loses accuracy (same with 2H = 2*Pi)
? exponent(intnumosc(x=0,sinc(x)^2,Pi,4) - Pi/2)
time = 17,681 ms.
%4 = -381 \\ and Sumpos is good but slow (perfect with 2H)

Exemples of a different flavour:

? exponent(intnumosc(x = 0, besselj(0,x)*sin(3*x), Pi) - 1/sqrt(8))
time = 4,615 ms.
%1 = -385 \\ more expensive but correct
? exponent(intnumosc(x = 0, besselj(0,x)*sin(3*x), Pi, 1) - 1/sqrt(8))
time = 1,424 ms.
%2 = -279 \\ unsafe mode loses some accuracy (other modes return junk)

? S = log(2*Pi)- Euler - 1;
? exponent(intnumosc(t=1, (frac(t)/t)^2, 1/2) - S)
time = 21 ms.
%4 = -6 \\ junk
? exponent(intnumosc(t=1, (frac(t)/t)^2, 1) - S)
time = 66ms.
%5 = -384 \\ perfect with 2H
? exponent(intnumosc(t=1, (frac(t)/t)^2, 1, 1) - S)
time = 20 ms.
%6 = -286 \\ unsafe mode loses accuracy
? exponent(intnumosc(t=1, (frac(t)/t)^2, 1, 3) - S)
time = 30 ms.

```

```
%7 = -236 \\ and so does Lagrange (Sumalt fails)
? exponent(intnumosc(t=1, (frac(t)/t)^2, 1, 4) - S)
time = 2,315 ms.
%8 = -382 \\ Sumpos is perfect but slow
```

Again, Sidi extrapolation behaves well, especially in safe mode, but $2H$ is required here.

If the integrand has singularities close to the interval of integration, it is advisable to split the integral in two: use the more robust `intnum` to handle the singularities, then `intnumosc` for the remainder:

```
? \p38
? f(x) = besselj(0,x)^3 * log(x); \\ mild singularity at 0
? g() = intnumosc(x = 0, f(x), Pi); \\ direct
? h() = intnum(x = 0, Pi, f(x)) + intnumosc(x = Pi, f(x), Pi); \\ split at Pi
? G = g();
time = 293 ms.
? H = h();
time = 320 ms. \\ about as fast
? exponent(G-H)
%6 = -12 \\ at least one of them is junk
? \p77 \\ increase accuracy
? G2=g(); H2=h();
? exponent(G - G2)
%8 = -13 \\ g() is not consistent
? exponent(H - H2)
%9 = -128 \\ not a proof, but h() looks good
```

Finally, here is an exemple where all methods fail, even when splitting the integral, except Sumpos:

```
? \p38
? f(x)=besselj(0,x)^4*log(x+1);
? F = intnumosc(x=0,f(x), Pi, 4)
time = 2,437 ms.
%2 = 0.45489838778971732178155161172638343214
? \p76 \\ double accuracy to check
? exponent(F - intnumosc(x = 0,f(x), Pi, 4))
time = 18,817 ms.
%3 = -122 \\ F was almost perfect
```

The library syntax is `GEN intnumosc0(GEN x, GEN expr, GEN H, long flag, GEN tab = NULL, long prec)`.

3.12.13 intnumromb($X = a, b, \text{expr}, \{\text{flag} = 0\}$). Numerical integration of expr (smooth in $]a, b[$), with respect to X . Suitable for low accuracy; if expr is very regular (e.g. analytic in a large region) and high accuracy is desired, try **intnum** first.

Set $\text{flag} = 0$ (or omit it altogether) when a and b are not too large, the function is smooth, and can be evaluated exactly everywhere on the interval $[a, b]$.

If $\text{flag} = 1$, uses a general driver routine for doing numerical integration, making no particular assumption (slow).

$\text{flag} = 2$ is tailored for being used when a or b are infinite using the change of variable $t = 1/X$. One *must* have $ab > 0$, and in fact if for example $b = +\infty$, then it is preferable to have a as large as possible, at least $a \geq 1$.

If $\text{flag} = 3$, the function is allowed to be undefined at a (but right continuous) or b (left continuous), for example the function $\sin(x)/x$ between $x = 0$ and 1.

The user should not require too much accuracy: **realprecision** about 30 decimal digits (**realbitprecision** about 100 bits) is OK, but not much more. In addition, analytical cleanup of the integral must have been done: there must be no singularities in the interval or at the boundaries. In practice this can be accomplished with a change of variable. Furthermore, for improper integrals, where one or both of the limits of integration are plus or minus infinity, the function must decrease sufficiently rapidly at infinity, which can often be accomplished through integration by parts. Finally, the function to be integrated should not be very small (compared to the current precision) on the entire interval. This can of course be accomplished by just multiplying by an appropriate constant.

Note that infinity can be represented with essentially no loss of accuracy by an appropriate huge number. However beware of real underflow when dealing with rapidly decreasing functions. For example, in order to compute the $\int_0^\infty e^{-x^2} dx$ to 38 decimal digits, then one can set infinity equal to 10 for example, and certainly not to **1e1000**.

The library syntax is **GEN intnumromb(void *E, GEN (*eval)(void*, GEN), GEN a, GEN b, long flag, long bitprec)**, where $\text{eval}(x, E)$ returns the value of the function at x . You may store any additional information required by **eval** in E , or set it to **NULL**.

3.12.14 laurentseries($f, \{M = \text{seriesprecision}\}, \{x = 'x\}$). Expand f as a Laurent series around $x = 0$ to order M . This function computes $f(x + O(x^n))$ until n is large enough: it must be possible to evaluate f on a power series with 0 constant term.

```
? laurentseries(t->sin(t)/(1-cos(t)), 5)
%1 = 2*x^-1 - 1/6*x - 1/360*x^3 - 1/15120*x^5 + 0(x^6)
? laurentseries(log)
*** at top-level: laurentseries(log)
*** ^-----
*** in function laurentseries: log
*** ^---
*** log: domain error in log: series valuation != 0
```

Note that individual Laurent coefficients of order $\leq M$ can be retrieved from $s = \text{laurentseries}(f, M)$ via **polcoef**(**s**,**i**) for any $i \leq M$. The series s may occasionally be more precise than the required $O(x^{M+1})$.

With respect to successive calls to **derivnum**, **laurentseries** is both faster and more precise:

The library syntax is `laurentseries(void *E, GEN (*f)(void*,GEN,long), long M, long long prec)`.

The routine assume that u_n has an asymptotic expansion in $n^{-\alpha}$:

for some a_i . It is purely numerical and heuristic, thus may or may not work on your examples. The expression will be evaluated for $n = 1, 2, \dots, N$ for an $N = O(B)$ at a bit accuracy bounded by $1.612B$.

It is not mandatory to specify α when the u_n have an asymptotic expansion in n^{-1} . However, if the series in n^{-1} is lacunary, specifying α allows faster computation:

355

```
time = 27,271 ms.
%5 = 0.E-1001 \\ still perfect, 4 times faster
```

When u_n has an asymptotic expansion in $n^{-\alpha}$ with α not an integer, leaving α unspecified will bring an inexact limit. Giving a satisfying optional argument improves precision; the program runs faster when the optional argument gives non lacunary series.

```
? \p50
? limitnum(n->(1+1/n^(7/2))^(n^(7/2))) - exp(1)
time = 982 ms.
%6 = 4.13[...] E-12
? limitnum(n->(1+1/n^(7/2))^(n^(7/2)), 1/2) - exp(1)
time = 16,745 ms.
%7 = 0.E-57
? limitnum(n->(1+1/n^(7/2))^(n^(7/2)), 7/2) - exp(1)
time = 105 ms.
%8 = 0.E-57
```

Alternatively, u_n may be given by a closure $N \mapsto [u_1, \dots, u_N]$ which can often be programmed in a more efficient way, for instance when u_{n+1} is a simple function of the preceding terms:

```
? \p2000
? limitnum(n -> 2^(4*n+1)*(n!)^4 / (2*n)! / (2*n+1)! ) - Pi
time = 1,755 ms.
%9 = 0.E-2003
? vu(N) = \\ exploit hypergeometric property
{ my(v = vector(N)); v[1] = 8./3;\
  for (n=2, N, my(q = 4*n^2); v[n] = v[n-1]*q/(q-1));\
  return(v);
}
? limitnum(vu) - Pi \\ much faster
time = 106 ms.
%11 = 0.E-2003
```

All sums and recursions can be handled in the same way. In the above it is essential that u_n be defined as a closure because it must be evaluated at a higher precision than the one expected for the limit. Make sure that the closure does not depend on a global variable which would be computed at a priori fixed accuracy. For instance, precomputing $v1 = 8.0/3$ first and using $v1$ in vu above would be wrong because the resulting vector of values will use the accuracy of $v1$ instead of the ambient accuracy at which `limitnum` will call it.

Alternatively, and more clumsily, u_n may be given by a vector of values: it must be long and precise enough for the extrapolation to make sense. Let B be the current `realbitprecision`, the vector length must be at least $1.102B$ and the values computed with bit accuracy $1.612B$.

```
? limitnum(vector(10,n,(1+1/n)^n))
***
*** limitnum: nonexistent component in limitnum: index < 43
\\ at this accuracy, we must have at least 43 values
? limitnum(vector(43,n,(1+1/n)^n)) - exp(1)
%12 = 0.E-37
? v = vector(43);
```

```

? s = 0; for(i=1,#v, s += 1/i; v[i]= s - log(i));
? limitnum(v) - Euler
%15 = -1.57[...] E-16

? v = vector(43);
\\ ~ 128 bit * 1.612
? localbitprec(207);\
  s = 0; for(i=1,#v, s += 1/i; v[i]= s - log(i));
? limitnum(v) - Euler
%18 = 0.E-38

```

Because of the above problems, the preferred format is thus a closure, given either a single value or the vector of values $[u_1, \dots, u_N]$. The function distinguishes between the two formats by evaluating the closure at $N \neq 1$ and 1 and checking whether it yields vectors of respective length N and 1 or not.

Warning. The expression is evaluated for $n = 1, 2, \dots, N$ for an $N = O(B)$ if the current bit accuracy is B . If it is not defined for one of these values, translate or rescale accordingly:

```

? limitnum(n->log(1-1/n))  \\ can't evaluate at n = 1 !
***   at top-level: limitnum(n->log(1-1/n))
***   ^-----
***   in function limitnum: log(1-1/n)
***   ^-----
*** log: domain error in log: argument = 0
? limitnum(n->-log(1-1/(2*n)))
%19 = -6.11[...] E-58

```

We conclude with a complicated example. Since the function is heuristic, it is advisable to check whether it produces the same limit for $u_n, u_{2n}, \dots, u_{km}$ for a suitable small multiplier k . The following function implements the recursion for the Motzkin numbers M_n which count the number of ways to draw non intersecting chords between n points on a circle:

$$M_n = M_{n-1} + \sum_{i < n-1} M_i M_{n-2-i} = ((n+1)M_{n-1} + (3n-3)M_{n-2})/(n+2).$$

It is known that $M_n^2 \sim \frac{9^{n+1}}{12\pi n^3}$.

```

\\ [M_k, M_{k*2}, ..., M_{k*N}] / (3^n / n^(3/2))
vM(N, k = 1) =
{ my(q = k*N, V);
  if (q == 1, return ([1/3]));
  V = vector(q); V[1] = V[2] = 1;
  for(n = 2, q - 1,
    V[n+1] = ((2*n + 1)*V[n] + 3*(n - 1)*V[n-1]) / (n + 2));
  f = (n -> 3^n / n^(3/2));
  return (vector(N, n, V[n*k] / f(n*k)));
}
? limitnum(vM) - 3/sqrt(12*Pi) \\ complete junk
%1 = 35540390.753542730306762369615276452646
? limitnum(N->vM(N,5)) - 3/sqrt(12*Pi) \\ M_{5n}: better
%2 = 4.130710262178469860 E-25

```

```

? limitnum(N->vM(N,10)) - 3/sqrt(12*Pi) \\ M_{10n}: perfect
%3 = 0.E-38
? \p2000
? limitnum(N->vM(N,10)) - 3/sqrt(12*Pi) \\ also at high accuracy
time = 409 ms.
%4 = 1.1048895470044788191 E-2004

```

In difficult cases such as the above a multiplier of 5 to 10 is usually sufficient. The above example is typical: a good multiplier usually remains sufficient when the requested precision increases!

The library syntax is **limitnum**(void *E, GEN (*u)(void *,GEN,long), GEN alpha, long prec), where **u**(E, n, prec) must return $u(n)$ in precision **prec**. Also available is **GEN limitnum0**(GEN u, GEN alpha, long prec), where u must be a vector of sufficient length as above.

3.12.16 prod($X = a, b, expr, \{x = 1\}$). Product of expression $expr$, initialized at x , the formal parameter X going from a to b . As for **sum**, the main purpose of the initialization parameter x is to force the type of the operations being performed. For example if it is set equal to the integer 1, operations will start being done exactly. If it is set equal to the real 1., they will be done using real numbers having the default precision. If it is set equal to the power series $1 + O(X^k)$ for a certain k , they will be done using power series of precision at most k . These are the three most common initializations.

As an extreme example, compare

```

? prod(i=1, 100, 1 - X^i); \\ this has degree 5050 !!
time = 128 ms.
? prod(i=1, 100, 1 - X^i, 1 + O(X^101))
time = 8 ms.
%2 = 1 - X - X^2 + X^5 + X^7 - X^12 - X^15 + X^22 + X^26 - X^35 - X^40 + \
X^51 + X^57 - X^70 - X^77 + X^92 + X^100 + O(X^101)

```

Of course, in this specific case, it is faster to use **eta**, which is computed using Euler's formula.

```

? prod(i=1, 1000, 1 - X^i, 1 + O(X^1001));
time = 589 ms.
? \ps1000
seriesprecision = 1000 significant terms
? eta(X) - %
time = 8ms.
%4 = O(X^1001)

```

The library syntax is **produit**(GEN a, GEN b, char *expr, GEN x).

3.12.17 prodeuler($p = a, b, expr$). Product of expression $expr$, initialized at 1.0 (i.e. to a floating point number equal to 1 to the current `realprecision`), the formal parameter p ranging over the prime numbers between a and b .

```
? prodeuler(p = 2, 10^4, 1 - p^-2)
%1 = 0.60793306911405513018380499671124428015
? P = 1; forprime(p = 2, 10^4, P *= (1 - p^-2))
? exponent(numerator(P))
%3 = 22953
```

The function returns a floating point number because, as the second expression shows, such products are usually intractably large rational numbers when computed symbolically. If the expression is a rational function, `prodeulerrat` computes the product over all primes:

```
? prodeulerrat(1 - p^-2)
%4 = 0.60792710185402662866327677925836583343
? 6/Pi^2
%3 = 0.60792710185402662866327677925836583343
```

The library syntax is `prodeuler(void *E, GEN (*eval)(void*,GEN), GEN a, GEN b, long prec)`.

3.12.18 prodeulerrat($F, \{s = 1\}, \{a = 2\}$). $\prod_{p \geq a} F(p^s)$, where the product is taken over prime numbers and F is a rational function.

```
? prodeulerrat(1+1/q^3,1)
%1 = 1.1815649490102569125693997341604542605
? zeta(3)/zeta(6)
%2 = 1.1815649490102569125693997341604542606
```

The library syntax is `GEN prodeulerrat(GEN F, GEN s = NULL, long a, long prec)`.

3.12.19 prodinf($X = a, expr, \{flag = 0\}$). infinite product of expression $expr$, the formal parameter X starting at a . The evaluation stops when the relative error of the expression minus 1 is less than the default precision. In particular, divergent products result in infinite loops. The expressions must always evaluate to an element of $\mathbf{C}$.

If $flag = 1$, do the product of the $(1 + expr)$ instead.

The library syntax is `prodinf(void *E, GEN (*eval)(void*,GEN), GEN a, long prec)` ($flag = 0$), or `prodinf1` with the same arguments ($flag = 1$).

3.12.20 prodnumrat(F, a). $\prod_{n \geq a} F(n)$, where $F - 1$ is a rational function of degree less than or equal to -2 .

```
? prodnumrat(1+1/x^2,1)
%1 = 3.6760779103749777206956974920282606665
```

The library syntax is `GEN prodnumrat(GEN F, long a, long prec)`.

3.12.21 solve($X = a, b, expr$). Find a real root of expression $expr$ between a and b . If both a and b are finite, the condition is that $expr(X = a) * expr(X = b) \leq 0$. (You will get an error message `roots must be bracketed in solve` if this does not hold.)

If only one between a and b is finite, say a , then $b = \pm\infty$. The routine will test all $b = a \pm 2^r$, with $r \geq \log_2(|a|)$ until it finds a bracket for the root which satisfies the abovementioned condition.

If both a and b are infinite, the routine will test 0 and all $\pm 2^r$, $r \geq 0$, until it finds a bracket for the root which satisfies the condition.

This routine uses Brent's method and can fail miserably if $expr$ is not defined in the whole of $[a, b]$ (try `solve(x=1, 2, tan(x))`).

The library syntax is `zbrent(void *E, GEN (*eval)(void*, GEN), GEN a, GEN b, long prec)`.

3.12.22 solvestep($X = a, b, step, expr, \{flag = 0\}$). Find zeros of a continuous function in the real interval $[a, b]$ by naive interval splitting. This function is heuristic and may or may not find the intended zeros. Binary digits of $flag$ mean

- 1: return as soon as one zero is found, otherwise return all zeros found;
- 2: refine the splitting until at least one zero is found (may loop indefinitely if there are no zeros);
- 4: do a multiplicative search (we must have $a > 0$ and $step > 1$), otherwise an additive search; $step$ is the multiplicative or additive step.
- 8: refine the splitting until at least one zero is very close to an integer.

```
? solvestep(X=0,10,1,sin(X^2),1)
%1 = 1.7724538509055160272981674833411451828
? solvestep(X=1,12,2,besselj(4,X),4)
%2 = [7.588342434..., 11.064709488...]
```

The library syntax is `solvestep(void *E, GEN (*eval)(void*, GEN), GEN a, GEN b, GEN step, long flag, long prec)`.

3.12.23 sum($X = a, b, expr, \{x = 0\}$). Sum of expression $expr$, initialized at x , the formal parameter going from a to b . As for `prod`, the initialization parameter x may be given to force the type of the operations being performed.

As an extreme example, compare

```
? sum(i=1, 10^4, 1/i); \\ rational number: denominator has 4345 digits.
time = 236 ms.
? sum(i=1, 5000, 1/i, 0.)
time = 8 ms.
%2 = 9.787606036044382264178477904
```

3.12.24 sumalt($X = a, \text{expr}, \{\text{flag} = 0\}$). Numerical summation of the series expr , which should be an alternating series $(-1)^k a_k$, the formal variable X starting at a . Use an algorithm of Cohen, Villegas and Zagier (*Experiment. Math.* **9** (2000), no. 1, 3–12).

If $\text{flag} = 0$, assuming that the a_k are the moments of a positive measure on $[0, 1]$, the relative error is $O(3 + \sqrt{8})^{-n}$ after using a_k for $k \leq n$. If **realprecision** is p , we thus set $n = \log(10)p / \log(3 + \sqrt{8}) \approx 1.3p$; besides the time needed to compute the a_k , $k \leq n$, the algorithm overhead is negligible: time $O(p^2)$ and space $O(p)$.

If $\text{flag} = 1$, use a variant with more complicated polynomials, see **polzagier**. If the a_k are the moments of $w(x)dx$ where w (or only $xw(x^2)$) is a smooth function extending analytically to the whole complex plane, convergence is in $O(14.4^{-n})$. If $xw(x^2)$ extends analytically to a smaller region, we still have exponential convergence, with worse constants. Usually faster when the computation of a_k is expensive. If **realprecision** is p , we thus set $n = \log(10)p / \log(14.4) \approx 0.86p$; besides the time needed to compute the a_k , $k \leq n$, the algorithm overhead is *not* negligible: time $O(p^3)$ and space $O(p^2)$. Thus, even if the analytic conditions for rigorous use are met, this variant is only worthwhile if the a_k are hard to compute, at least $O(p^2)$ individually on average: otherwise we gain a small constant factor (1.5, say) in the number of needed a_k at the expense of a large overhead.

The conditions for rigorous use are hard to check but the routine is best used heuristically: even divergent alternating series can sometimes be summed by this method, as well as series which are not exactly alternating (see for example Section 2.7). It should be used to try and guess the value of an infinite sum. (However, see the example at the end of Section 2.7.1.)

If the series already converges geometrically, **suminf** is often a better choice:

```
? \p38
? sumalt(i = 1, -(-1)^i / i) - log(2)
time = 0 ms.
%1 = 0.E-38
? suminf(i = 1, -(-1)^i / i)  \\ Had to hit Ctrl-C
*** at top-level: suminf(i=1,-(-1)^i/i)
*** ^-----
*** suminf: user interrupt after 10min, 20,100 ms.
? \p1000
? sumalt(i = 1, -(-1)^i / i) - log(2)
time = 90 ms.
%2 = 4.459597722 E-1002
? sumalt(i = 0, (-1)^i / i!) - exp(-1)
time = 670 ms.
%3 = -4.03698781490633483156497361352190615794353338591897830587 E-944
? suminf(i = 0, (-1)^i / i!) - exp(-1)
time = 110 ms.
%4 = -8.39147638 E-1000  \\ faster and more accurate
```

The library syntax is **sumalt**(void *E, GEN (*eval)(void*,GEN),GEN a,long prec). Also available is **sumalt2** with the same arguments ($\text{flag} = 1$).

3.12.25 sumdiv($n, X, expr$). Sum of expression $expr$ over the positive divisors of n . This function is a trivial wrapper essentially equivalent to

```
D = divisors(n);
sum (i = 1, #D, my(X = D[i]); eval(expr))
```

If $expr$ is a multiplicative function, use `sumdivmult`.

3.12.26 sumdivmult($n, d, expr$). Sum of *multiplicative* expression $expr$ over the positive divisors d of n . Assume that $expr$ evaluates to $f(d)$ where f is multiplicative: $f(1) = 1$ and $f(ab) = f(a)f(b)$ for coprime a and b . The library syntax is `sumdivmultexpr(void *E, GEN (*eval)(void*,GEN), GEN d)`

3.12.27 sumeulerrat($F, \{s = 1\}, \{a = 2\}$). $\sum_{p \geq a} F(p^s)$, where the sum is taken over prime numbers and F is a rational function.

```
? sumeulerrat(1/p^2)
%1 = 0.45224742004106549850654336483224793417
? sumeulerrat(1/p, 2)
%2 = 0.45224742004106549850654336483224793417
```

The library syntax is `GEN sumeulerrat(GEN F, GEN s = NULL, long a, long prec)`.

3.12.28 suminf($X = a, expr$). Naive summation of expression $expr$, the formal parameter X going from a to infinity. The evaluation stops when the relative error of the expression is less than the default bit precision for 3 consecutive evaluations. The expressions must evaluate to a complex number.

If the expression tends slowly to 0, like n^{-a} for some $a > 1$, make sure $b = \text{realbitprecision}$ is low: indeed, the algorithm will require $O(2^{b/a})$ function evaluations and we expect only about $b(1 - 1/a)$ correct bits in the answer. If the series is alternating, we can expect b correct bits but the `sumalt` function should be used instead since its complexity is polynomial in b , instead of exponential. More generally, `sumpos` should be used if the terms have a constant sign and `sumnum` if the function is C^∞ .

```
? \pb25
  realbitprecision = 25 significant bits (7 decimal digits displayed)
? exponent(suminf(i = 1, (-1)^i / i) + log(2))
time = 2min, 2,602 ms.
%1 = -29
? \pb45
  realbitprecision = 45 significant bits (13 decimal digits displayed)
? exponent(suminf(i = 1, 1 / i^2) - zeta(2))
time = 2,186 ms.
%2 = -23

\\ alternatives are much faster
? \pb 10000
  realbitprecision = 10000 significant bits (3010 decimal digits displayed)
? exponent(sumalt(i = 1, (-1)^i / i) + log(2))
time = 25 ms.
%3 = -10043
```

```

? \pb 4000
  realbitprecision = 4000 significant bits (1204 decimal digits displayed)))
? exponent(sumpos(i = 1, 1 / i^2) - zeta(2))
time = 22,593 ms.
%4 = -4030

? exponent(sumnum(i = 1, 1 / i^2) - zeta(2))
time = 7,032 ms.
%5 = -4031

\\ but suminf is perfect for geometrically converging series
? exponent(suminf(i = 1, 2^-i) - 1)
time = 25 ms.
%6 = -4003

```

The library syntax is `suminf(void *E, GEN (*eval)(void*,GEN), GEN a, long prec)`.

3.12.29 sumnum($n = a, f, \{tab\}$). Numerical summation of $f(n)$ at high accuracy using Euler-MacLaurin, the variable n taking values from a to $+\infty$, where f is assumed to have positive values and is a C^∞ function; a must be an integer and tab , if given, is the output of `sumnuminit`. The latter precomputes abscissas and weights, speeding up the computation; it also allows to specify the behavior at infinity via `sumnuminit([+oo, asymp])`.

```

? \p500
? z3 = zeta(3);
? sumpos(n = 1, n^-3) - z3
time = 2,332 ms.
%2 = 2.438468843 E-501
? sumnum(n = 1, n^-3) - z3 \\ here slower than sumpos
time = 2,752 ms.
%3 = 0.E-500

```

Complexity. The function f will be evaluated at $O(D \log D)$ real arguments, where $D \approx \text{realprecision} \cdot \log(10)$. The routine is geared towards slowly decreasing functions: if f decreases exponentially fast, then one of `suminf` or `sumpos` should be preferred. If f satisfies the stronger hypotheses required for Monien summation, i.e. if $f(1/z)$ is holomorphic in a complex neighbourhood of $[0, 1]$, then `sumnummonien` will be faster since it only requires $O(D/\log D)$ evaluations:

```

? sumnummonien(n = 1, 1/n^3) - z3
time = 1,985 ms.
%3 = 0.E-500

```

The `tab` argument precomputes technical data not depending on the expression being summed and valid for a given accuracy, speeding up immensely later calls:

```

? tab = sumnuminit();
time = 2,709 ms.
? sumnum(n = 1, 1/n^3, tab) - z3 \\ now much faster than sumpos
time = 40 ms.
%5 = 0.E-500

? tabmon = sumnummonieninit(); \\ Monien summation allows precomputations too
time = 1,781 ms.
? sumnummonien(n = 1, 1/n^3, tabmon) - z3

```

```
time = 2 ms.
%7 = 0.E-500
```

The speedup due to precomputations becomes less impressive when the function f is expensive to evaluate, though:

```
? sumnum(n = 1, lngamma(1+1/n)/n, tab);
time = 14,180 ms.

? sumnummonien(n = 1, lngamma(1+1/n)/n, tabmon); \\ fewer evaluations
time = 717 ms.
```

Behaviour at infinity. By default, `sumnum` assumes that $expr$ decreases slowly at infinity, but at least like $O(n^{-2})$. If the function decreases like n^α for some $-2 < \alpha < -1$, then it must be indicated via

```
tab = sumnuminit([+oo, alpha]); /* alpha < 0 slow decrease */
```

otherwise loss of accuracy is expected. If the functions decreases quickly, like $\exp(-\alpha n)$ for some $\alpha > 0$, then it must be indicated via

```
tab = sumnuminit([+oo, alpha]); /* alpha > 0 exponential decrease */
```

otherwise exponent overflow will occur.

```
? sumnum(n=1,2^-n)
*** at top-level: sumnum(n=1,2^-n)
*** ^----
*** _^_: overflow in expo().
? tab = sumnuminit([+oo,log(2)]); sumnum(n=1,2^-n, tab)
%1 = 1.000[...]
```

As a shortcut, one can also input

```
sumnum(n = [a, asymp], f)
```

instead of

```
tab = sumnuminit(asymp);
sumnum(n = a, f, tab)
```

Further examples.

```
? \p200
? sumnum(n = 1, n^(-2)) - zeta(2) \\ accurate, fast
time = 200 ms.
%1 = -2.376364457868949779 E-212
? sumpos(n = 1, n^(-2)) - zeta(2) \\ even faster
time = 96 ms.
%2 = 0.E-211
? sumpos(n=1,n^(-4/3)) - zeta(4/3) \\ now much slower
time = 13,045 ms.
%3 = -9.980730723049589073 E-210
? sumnum(n=1,n^(-4/3)) - zeta(4/3) \\ fast but inaccurate
time = 365 ms.
%4 = -9.85[...]E-85
? sumnum(n=[1,-4/3],n^(-4/3)) - zeta(4/3) \\ with decrease rate, now accurate
time = 416 ms.
%5 = -4.134874156691972616 E-210
? tab = sumnuminit([+oo,-4/3]);
time = 196 ms.
? sumnum(n=1, n^(-4/3), tab) - zeta(4/3) \\ faster with precomputations
time = 216 ms.
%5 = -4.134874156691972616 E-210
? sumnum(n=1,-log(n)*n^(-4/3), tab) - zeta'(4/3)
time = 321 ms.
%7 = 7.224147951921607329 E-210
```

Note that in the case of slow decrease ($\alpha < 0$), the exact decrease rate must be indicated, while in the case of exponential decrease, a rough value will do. In fact, for exponentially decreasing functions, `sumnum` is given for completeness and comparison purposes only: one of `suminf` or `sumpos` should always be preferred.

```
? sumnum(n=[1, 1], 2^-n) \\ pretend we decrease as exp(-n)
time = 240 ms.
%8 = 1.000[...] \\ perfect
? sumpos(n=1, 2^-n)
%9 = 1.000[...] \\ perfect and instantaneous
```

Beware cancellation. The function $f(n)$ is evaluated for huge values of n , so beware of cancellation in the evaluation:

```
? f(n) = 2 - 1/n - 2*n*log(1+1/n); \\ result is O(1/n^2)
? z = -2 + log(2*Pi) - Euler;
? sumnummonien(n=1, f(n)) - z
time = 149 ms.
%12 = 0.E-212 \\ perfect
? sumnum(n=1, f(n)) - z
time = 116 ms.
%13 = -948.216[...] \\ junk
```

As `sumnum(n=1, print(n))` shows, we evaluate $f(n)$ for $n > 1e233$ and our implementation of f suffers from massive cancellation since we are summing two terms of the order of $O(1)$ for a result in $O(1/n^2)$. You can either rewrite your sum so that individual terms are evaluated without cancellation or locally replace $f(n)$ by an accurate asymptotic expansion:

```
? F = truncate( f(1/x + O(x^30)) );
? sumnum(n=1, if(n > 1e7, subst(F,x,1/n), f(n))) - z
%15 = 1.1 E-212 \\ now perfect
```

The library syntax is `sumnum((void *E, GEN (*eval)(void*, GEN), GEN a, GEN tab, long prec))` where an omitted *tab* is coded as NULL.

3.12.30 sumnumap($n = a, f, \{tab\}$). Numerical summation of $f(n)$ at high accuracy using Abel-Plana, the variable n taking values from a to $+\infty$, where f is holomorphic in the right half-plane $\Re(z) > a$; a must be an integer and *tab*, if given, is the output of `sumnumapinit`. The latter precomputes abscissas and weights, speeding up the computation; it also allows to specify the behavior at infinity via `sumnumapinit([+oo, asymp])`.

```
? \p500
? z3 = zeta(3);
? sumpos(n = 1, n^-3) - z3
time = 2,332 ms.
%2 = 2.438468843 E-501
? sumnumap(n = 1, n^-3) - z3 \\ here slower than sumpos
time = 2,565 ms.
%3 = 0.E-500
```

Complexity. The function f will be evaluated at $O(D \log D)$ real arguments and $O(D)$ complex arguments, where $D \approx \text{realprecision} \cdot \log(10)$. The routine is geared towards slowly decreasing functions: if f decreases exponentially fast, then one of `suminf` or `sumpos` should be preferred. The default algorithm `sumnum` is usually a little *slower* than `sumnumap` but its initialization function `sumnuminit` becomes much faster as `realprecision` increases.

If f satisfies the stronger hypotheses required for Monien summation, i.e. if $f(1/z)$ is holomorphic in a complex neighbourhood of $[0, 1]$, then `sumnummonien` will be faster since it only requires $O(D/\log D)$ evaluations:

```
? sumnummonien(n = 1, 1/n^3) - z3
time = 1,128 ms.
%3 = 0.E-500
```

The `tab` argument precomputes technical data not depending on the expression being summed and valid for a given accuracy, speeding up immensely later calls:

```
? tab = sumnumapinit();
time = 2,567 ms.
? sumnumap(n = 1, 1/n^3, tab) - z3 \\ now much faster than sumpos
time = 39 ms.
%5 = 0.E-500

? tabmon = sumnummonieninit(); \\ Monien summation allows precomputations too
time = 1,125 ms.
? sumnummonien(n = 1, 1/n^3, tabmon) - z3
time = 2 ms.
%7 = 0.E-500
```

The speedup due to precomputations becomes less impressive when the function f is expensive to evaluate, though:

```
? sumnumap(n = 1, lngamma(1+1/n)/n, tab);
time = 10,762 ms.

? sumnummonien(n = 1, lngamma(1+1/n)/n, tabmon); \\ fewer evaluations
time = 205 ms.
```

Behaviour at infinity. By default, `sumnumap` assumes that $expr$ decreases slowly at infinity, but at least like $O(n^{-2})$. If the function decreases like n^α for some $-2 < \alpha < -1$, then it must be indicated via

```
tab = sumnumapinit([+oo, alpha]); /* alpha < 0 slow decrease */
```

otherwise loss of accuracy is expected. If the function decreases quickly, like $\exp(-\alpha n)$ for some $\alpha > 0$, then it must be indicated via

```
tab = sumnumapinit([+oo, alpha]); /* alpha > 0 exponential decrease */
```

otherwise exponent overflow will occur.

```
? sumnumap(n=1,2^-n)
*** at top-level: sumnumap(n=1,2^-n)
*** ^----
*** _^_: overflow in expo().
? tab = sumnumapinit([+oo,log(2)]); sumnumap(n=1,2^-n, tab)
```

```
%1 = 1.000[...]
```

As a shortcut, one can also input

```
sumnumap(n = [a, asymp], f)
```

instead of

```
tab = sumnumapinit(asymp);  
sumnumap(n = a, f, tab)
```

Further examples.

```
? \p200  
? sumnumap(n = 1, n^(-2)) - zeta(2) \\ accurate, fast  
time = 169 ms.  
%1 = -4.752728915737899559 E-212  
? sumpos(n = 1, n^(-2)) - zeta(2) \\ even faster  
time = 79 ms.  
%2 = 0.E-211  
? sumpos(n=1,n^(-4/3)) - zeta(4/3) \\ now much slower  
time = 10,518 ms.  
%3 = -9.980730723049589073 E-210  
? sumnumap(n=1,n^(-4/3)) - zeta(4/3) \\ fast but inaccurate  
time = 309 ms.  
%4 = -2.57[...]  
? sumnumap(n=[1,-4/3],n^(-4/3)) - zeta(4/3) \\ decrease rate: now accurate  
time = 329 ms.  
%6 = -5.418110963941205497 E-210  
? tab = sumnumapinit([+oo,-4/3]);  
time = 160 ms.  
? sumnumap(n=1, n^(-4/3), tab) - zeta(4/3) \\ faster with precomputations  
time = 175 ms.  
%5 = -5.418110963941205497 E-210  
? sumnumap(n=1,-log(n)*n^(-4/3), tab) - zeta'(4/3)  
time = 258 ms.  
%7 = 9.125239518216767153 E-210
```

Note that in the case of slow decrease ($\alpha < 0$), the exact decrease rate must be indicated, while in the case of exponential decrease, a rough value will do. In fact, for exponentially decreasing functions, `sumnumap` is given for completeness and comparison purposes only: one of `suminf` or `sumpos` should always be preferred.

```
? sumnumap(n=[1, 1], 2^-n) \\ pretend we decrease as exp(-n)  
time = 240 ms.  
%8 = 1.000[...]  
? sumpos(n=1, 2^-n)  
%9 = 1.000[...]
```

The library syntax is `sumnumap((void *E, GEN (*eval)(void*,GEN), GEN a, GEN tab, long prec))` where an omitted *tab* is coded as `NULL`.

3.12.31 sumnumapinit($\{asympt\}$). Initialize tables for Abel–Plana summation of a series $\sum f(n)$, where f is holomorphic in a right half-plane. If given, **asympt** is of the form $[+\infty, \alpha]$, as in **intnum** and indicates the decrease rate at infinity of functions to be summed. A positive $\alpha > 0$ encodes an exponential decrease of type $\exp(-\alpha n)$ and a negative $-2 < \alpha < -1$ encodes a slow polynomial decrease of type n^α .

```
? \p200
? sumnumap(n=1, n^-2);
time = 163 ms.
? tab = sumnumapinit();
time = 160 ms.
? sumnumap(n=1, n^-2, tab); \\ faster
time = 7 ms.

? tab = sumnumapinit([+oo, log(2)]); \\ decrease like 2^-n
time = 164 ms.
? sumnumap(n=1, 2^-n, tab) - 1
time = 36 ms.
%5 = 3.0127431466707723218 E-282

? tab = sumnumapinit([+oo, -4/3]); \\ decrease like n^(-4/3)
time = 166 ms.
? sumnumap(n=1, n^(-4/3), tab);
time = 181 ms.
```

The library syntax is `GEN sumnumapinit(GEN asympt = NULL, long prec)`.

3.12.32 sumnuminit($\{asympt\}$). Initialize tables for Euler–MacLaurin delta summation of a series with positive terms. If given, **asympt** is of the form $[+\infty, \alpha]$, as in **intnum** and indicates the decrease rate at infinity of functions to be summed. A positive $\alpha > 0$ encodes an exponential decrease of type $\exp(-\alpha n)$ and a negative $-2 < \alpha < -1$ encodes a slow polynomial decrease of type n^α .

```
? \p200
? sumnum(n=1, n^-2);
time = 200 ms.
? tab = sumnuminit();
time = 188 ms.
? sumnum(n=1, n^-2, tab); \\ faster
time = 8 ms.

? tab = sumnuminit([+oo, log(2)]); \\ decrease like 2^-n
time = 200 ms.
? sumnum(n=1, 2^-n, tab)
time = 44 ms.

? tab = sumnuminit([+oo, -4/3]); \\ decrease like n^(-4/3)
time = 200 ms.
? sumnum(n=1, n^(-4/3), tab);
time = 221 ms.
```

The library syntax is `GEN sumnuminit(GEN asympt = NULL, long prec)`.

3.12.33 sumnumlagrange($n = a, f, \{tab\}$). Numerical summation of $f(n)$ from $n = a$ to $+\infty$ using Lagrange summation; a must be an integer, and the optional argument **tab** is the output of **sumnumlagrangeinit**. By default, the program assumes that the N th remainder has an asymptotic expansion in integral powers of $1/N$. If not, initialize **tab** using **sumnumlagrangeinit(al)**, where the asymptotic expansion of the remainder is integral powers of $1/N^{al}$; al can be equal to 1 (default), $1/2$, $1/3$, or $1/4$, and also equal to 2, but in this latter case it is the N th remainder minus one half of the last summand which has an asymptotic expansion in integral powers of $1/N^2$.

```
? \p1000
? z3 = zeta(3);
? sumpos(n = 1, n^-3) - z3
time = 4,440 ms.
%2 = -2.08[...] E-1001
? sumnumlagrange(n = 1, n^-3) - z3 \\ much faster than sumpos
time = 25 ms.
%3 = 0.E-1001
? tab = sumnumlagrangeinit();
time = 21 ms.
? sumnumlagrange(n = 1, n^-3, tab) - z3
time = 2 ms. /* even faster */
%5 = 0.E-1001
? \p115
? tab = sumnumlagrangeinit([1/3,1/3]);
time = 316 ms.
? sumnumlagrange(n = 1, n^-(7/3), tab) - zeta(7/3)
time = 24 ms.
%7 = 0.E-115
? sumnumlagrange(n = 1, n^(-2/3) - 3*(n^(1/3)-(n-1)^(1/3)), tab) - zeta(2/3)
time = 32 ms.
%8 = 1.0151767349262596893 E-115
```

Complexity. The function f is evaluated at $O(D)$ integer arguments, where $D \approx \text{realprecision} \cdot \log(10)$.

The library syntax is **sumnumlagrange**((void *E, GEN (*eval)(void*, GEN), GEN a, GEN tab, long prec)) where an omitted *tab* is coded as NULL.

3.12.34 sumnumlagrangeinit($\{asympt\}, \{cl\}$). Initialize tables for Lagrange summation of a series. By default, assume that the remainder $R(n) = \sum_{m \geq n} f(m)$ has an asymptotic expansion

$$R(n) = \sum_{m \geq n} f(m) \approx \sum_{i \geq 1} a_i / n^i$$

at infinity. The argument **asympt** allows to specify different expansions:

- a real number β means

$$R(n) = n^{-\beta} \sum_{i \geq 1} a_i / n^i$$

- a **t_CLOSURE** g means

$$R(n) = g(n) \sum_{i \geq 1} a_i / n^i$$

(The preceding case corresponds to $g(n) = n^{-\beta}$.)

• a pair $[\alpha, \beta]$ where β is as above and $\alpha \in \{2, 1, 1/2, 1/3, 1/4\}$. We let $R_2(n) = R(n) - f(n)/2$ and $R_\alpha(n) = R(n)$ for $\alpha \neq 2$. Then

$$R_\alpha(n) = g(n) \sum_{i \geq 1} a_i / n^{i\alpha}$$

Note that the initialization times increase considerable for the α is this list (1/4 being the slowest).

The constant $c1$ is technical and computed by the program, but can be set by the user: the number of interpolation steps will be chosen close to $c1 \cdot B$, where B is the bit accuracy.

```
? \p2000
? sumnumlagrange(n=1, n^-2);
time = 173 ms.
? tab = sumnumlagrangeinit();
time = 172 ms.
? sumnumlagrange(n=1, n^-2, tab);
time = 4 ms.

? \p115
? sumnumlagrange(n=1, n^(-4/3)) - zeta(4/3);
%1 = -0.1093[...] \\ junk: expansion in n^(1/3)
time = 84 ms.
? tab = sumnumlagrangeinit([1/3,0]); \\ alpha = 1/3
time = 336 ms.
? sumnumlagrange(n=1, n^(-4/3), tab) - zeta(4/3)
time = 84 ms.
%3 = 1.0151767349262596893 E-115 \\ now OK

? tab = sumnumlagrangeinit(1/3); \\ alpha = 1, beta = 1/3: much faster
time = 3ms
? sumnumlagrange(n=1, n^(-4/3), tab) - zeta(4/3) \\ ... but wrong
%5 = -0.273825[...] \\ junk !
? tab = sumnumlagrangeinit(-2/3); \\ alpha = 1, beta = -2/3
time = 3ms
? sumnumlagrange(n=1, n^(-4/3), tab) - zeta(4/3)
%6 = 2.030353469852519379 E-115 \\ now OK
```

in The final example with $\zeta(4/3)$, the remainder $R_1(n)$ is of the form $n^{-1/3} \sum_{i \geq 0} a_i / n^i$, i.e. $n^{2/3} \sum_{i \geq 1} a_i / n^i$. This explains the wrong result for $\beta = 1/3$ and the correction with $\beta = -2/3$.

The library syntax is GEN sumnumlagrangeinit(GEN asymp = NULL, GEN c1 = NULL, long prec).

3.12.35 sumnummonien($n = a, f, \{tab\}$). Numerical summation $\sum_{n \geq a} f(n)$ at high accuracy, the variable n taking values from the integer a to $+\infty$ using Monien summation, which assumes that $f(1/z)$ has a complex analytic continuation in a (complex) neighbourhood of the segment $[0, 1]$.

The function f is evaluated at $O(D/\log D)$ real arguments, where $D \approx \text{realprecision} \cdot \log(10)$. By default, assume that $f(n) = O(n^{-2})$ and has a nonzero asymptotic expansion

$$f(n) = \sum_{i \geq 2} a_i n^{-i}$$

at infinity. To handle more complicated behaviors and allow time-saving precomputations (for a given `realprecision`), see `sumnummonieninit`.

The library syntax is `GEN sumnummonien0(GEN n, GEN f, GEN tab = NULL, long prec)`.

3.12.36 sumnummonieninit($\{asympt\}, \{w\}, \{n0 = 1\}$). Initialize tables for Monien summation of a series $\sum_{n \geq n_0} f(n)$ where $f(1/z)$ has a complex analytic continuation in a (complex) neighbourhood of the segment $[0, 1]$.

By default, assume that $f(n) = O(n^{-2})$ and has a nonzero asymptotic expansion

$$f(n) = \sum_{i \geq 2} a_i / n^i$$

at infinity. Note that the sum starts at $i = 2$! The argument `asympt` allows to specify different expansions:

- a real number $\beta > 0$ means

$$f(n) = \sum_{i \geq 1} a_i / n^{i+\beta}$$

(Now the summation starts at 1.)

- a vector $[\alpha, \beta]$ of reals, where we must have $\alpha > 0$ and $\alpha + \beta > 1$ to ensure convergence, means that

$$f(n) = \sum_{i \geq 1} a_i / n^{\alpha i + \beta}$$

Note that `asympt = [1,  $\beta$ ]` is equivalent to `asympt =  $\beta$` .

```
? \p57
? s = sumnum(n = 1, sin(1/sqrt(n)) / n); \\ reference point
? \p38
? sumnummonien(n = 1, sin(1/sqrt(n)) / n) - s
%2 = -0.001[...] \\ completely wrong
? t = sumnummonieninit(1/2); \\ f(n) = sum_i 1 / n^(i+1/2)
? sumnummonien(n = 1, sin(1/sqrt(n)) / n, t) - s
%3 = 0.E-37 \\ now correct
```

(As a matter of fact, in the above summation, the result given by `sumnum` at `\p38` is slightly incorrect, so we had to increase the accuracy to `\p57`.)

The argument w is used to sum expressions of the form

$$\sum_{n \geq n_0} f(n)w(n),$$

for varying f as above, and fixed weight function w , where we further assume that the auxiliary sums

$$g_w(m) = \sum_{n \geq n_0} w(n)/n^{\alpha m + \beta}$$

converge for all $m \geq 1$. Note that for nonnegative integers k , and weight $w(n) = (\log n)^k$, the function $g_w(m) = \zeta^{(k)}(\alpha m + \beta)$ has a simple expression; for general weights, g_w is computed using `sumnum`. The following variants are available

- an integer $k \geq 0$, to code $w(n) = (\log n)^k$;
- a `t_CLOSURE` computing the values $w(n)$, where we assume that $w(n) = O(n^\epsilon)$ for all $\epsilon > 0$;
- a vector $[w, \text{fast}]$, where w is a closure as above and **fast** is a scalar; we assume that $w(n) = O(n^{\text{fast} + \epsilon})$; note that $\mathbf{w} = [w, 0]$ is equivalent to $\mathbf{w} = w$. Note that if w decreases exponentially, `suminf` should be used instead.

The subsequent calls to `sumnummonien` *must* use the same value of n_0 as was used here.

```
? \p300
? sumnummonien(n = 1, n^-2*log(n)) + zeta'(2)
time = 328 ms.
%1 = -1.323[...]E-6 \\ completely wrong, f does not satisfy hypotheses !
? tab = sumnummonieninit(, 1); \\ codes w(n) = log(n)
time = 3,993 ms.
? sumnummonien(n = 1, n^-2, tab) + zeta'(2)
time = 41 ms.
%3 = -5.562684646268003458 E-309 \\ now perfect
? tab = sumnummonieninit(, n->log(n)); \\ generic, slower
time = 9,808 ms.
? sumnummonien(n = 1, n^-2, tab) + zeta'(2)
time = 40 ms.
%5 = -5.562684646268003458 E-309 \\ identical result
```

The library syntax is `GEN sumnummonieninit(GEN asymp = NULL, GEN w = NULL, GEN n0 = NULL, long prec)`.

3.12.37 `sumnumrat`(F, a). $\sum_{n \geq a} F(n)$, where F is a rational function of degree less than or equal to -2 and where poles of F at integers $\geq a$ are omitted from the summation. The argument a must be a `t_INT` or `-oo`.

```
? sumnumrat(1/(x^2+1)^2, 0)
%1 = 1.3068369754229086939178621382829073480
? sumnumrat(1/x^2, -oo) \\ value at x=0 is discarded
%2 = 3.2898681336964528729448303332920503784
? 2*zeta(2)
%3 = 3.2898681336964528729448303332920503784
```

When $\deg F = -1$, we define

$$\sum_{-\infty}^{\infty} F(n) := \sum_{n \geq 0} (F(n) + F(-1-n)) :$$

```
? sumnumrat(1/x, -oo)
%4 = 0.E-38
```

The library syntax is GEN `sumnumrat(GEN F, GEN a, long prec)`.

3.12.38 sumnumsidi($n = a, f, \{safe = 1\}$). Numerical summation of $f(n)$ from $n = a$ to $+\infty$ using Sidi summation; a must be an integer. The optional argument **safe** (set by default to 1) can be set to 0 for a faster but much less robust program; this is likely to lose accuracy when the sum is non-alternating.

```
? \pb3328
? z = zeta(2);
? exponent(sumnumsidi(n = 1, 1/n^2) - z)
time = 1,507 ms.
%2 = -3261 \\ already loses some decimals
? exponent(sumnumsidi(n = 1, 1/n^2, 0) - z)
time = 442 ms. \\ unsafe is much faster
%3 = -2108      \\ ... but very wrong

? l2 = log(2);
? exponent(sumnumsidi(n = 1, (-1)^(n-1)/n) - z)
time = 718 ms.
%5 = -3328 \\ not so slow and perfect
? exponent(sumnumsidi(n = 1, (-1)^(n-1)/n, 0) - z)
time = 504 ms.
%5 = -3328 \\ still perfect in unsafe mode, not so much faster
```

Complexity. If the bitprecision is b , we try to achieve an absolute error less than 2^{-b} . The function f is evaluated at $O(b)$ consecutive integer arguments at bit accuracy $1.56b$ (resp. b) in safe (resp. unsafe) mode.

The library syntax is GEN `sumnumsidi0(GEN n, GEN f, long safe, long prec)`.

3.12.39 sumpos($X = a, expr, \{flag = 0\}$). Numerical summation of the series $expr$, which must be a series of terms having the same sign, the formal variable X starting at a . The algorithm uses Van Wijngaarden's trick for converting such a series into an alternating one, then **sumalt**. For regular functions, the function **sumnum** is in general much faster once the initializations have been made using **sumnuminit**. Contrary to **sumnum**, **sumpos** allows functions defined only at integers:

```
? sumnum(n = 0, 1/n!)
*** at top-level: sumnum(n=1,1/n!)
***                               ^---
*** incorrect type in gtos [integer expected] (t_FRAC).
? sumpos(n = 0, 1/n!) - exp(1)
%2 = -1.0862155548773347717 E-33
```

On the other hand, when the function accepts general real numbers, it is usually advantageous to replace n by $n * 1.0$ in the `sumpos` call in particular when rational functions are involved:

```
? \p500
? sumpos(n = 0, n^7 / (n^9+n+1));
time = 6,108 ms.
? sumpos(n = 0, n *= 1.; n^7 / (n^9+n+1));
time = 2,788 ms.
? sumnumrat(n^7 / (n^9+n+1), 0);
time = 4 ms.
```

In the last example, `sumnumrat` is of course much faster but it only applies to rational functions.

The routine is heuristic and assumes that *expr* is more or less a decreasing function of X . In particular, the result will be completely wrong if *expr* is 0 too often. We do not check either that all terms have the same sign: as `sumalt`, this function should be used to try and guess the value of an infinite sum.

If *flag* = 1, use `sumalt(, 1)` instead of `sumalt(, 0)`, see Section 3.12.24. Requiring more stringent analytic properties for rigorous use, but allowing to compute fewer series terms.

To reach accuracy 10^{-p} , both algorithms require $O(p^2)$ space; furthermore, assuming the terms decrease polynomially (in $O(n^{-C})$), both need to compute $O(p^2)$ terms. The `sumpos(, 1)` variant has a smaller implied constant (roughly 1.5 times smaller). Since the `sumalt(, 1)` overhead is now small compared to the time needed to compute series terms, this last variant should be about 1.5 faster. On the other hand, the achieved accuracy may be much worse: as for `sumalt`, since conditions for rigorous use are hard to check, the routine is best used heuristically.

The library syntax is `sumpos(void *E, GEN (*eval)(void*, GEN), GEN a, long prec)`. Also available is `sumpos2` with the same arguments (*flag* = 1).

3.13 General number fields.

In this section, we describe functions related to general number fields. Functions related to quadratic number fields are found in Section 3.8 (Arithmetic functions).

3.13.1 Number field structures.

Let $K = \mathbf{Q}[X]/(T)$ a number field, $\mathbf{Z}_K$ its ring of integers, $T \in \mathbf{Z}[X]$ is monic. Three basic number field structures can be attached to K in GP:

- *nf* denotes a number field, i.e. a data structure output by `nfinit`. This contains the basic arithmetic data attached to the number field: signature, maximal order (given by a basis `nf.zk`), discriminant, defining polynomial T , etc.
- *bnf* denotes a “Buchmann’s number field”, i.e. a data structure output by `bnfinit`. This contains *nf* and the deeper invariants of the field: units $U(K)$, class group $\text{Cl}(K)$, as well as technical data required to solve the two attached discrete logarithm problems.
- *bnr* denotes a “ray number field”, i.e. a data structure output by `bnrinit`, corresponding to the ray class group structure of the field, for some modulus f . It contains a *bnf*, the modulus f , the ray class group $\text{Cl}_f(K)$ and data attached to the discrete logarithm problem therein.

3.13.2 Algebraic numbers and ideals.

An *algebraic number* belonging to $K = \mathbf{Q}[X]/(T)$ is given as

- a `t_INT`, `t_FRAC` or `t_POL` (implicitly modulo T), or
- a `t_POLMOD` (modulo T), or
- a `t_COL` v of dimension $N = [K : \mathbf{Q}]$, representing the element in terms of the computed integral basis, as `sum(i = 1, N, v[i] * nf.zk[i])`. Note that a `t_VEC` will not be recognized.

An *ideal* is given in any of the following ways:

- an algebraic number in one of the above forms, defining a principal ideal.
- a prime ideal, i.e. a 5-component vector in the format output by `idealprimedec` or `ideal-factor`.
- a `t_MAT`, square and in Hermite Normal Form (or at least upper triangular with nonnegative coefficients), whose columns represent a $\mathbf{Z}$ -basis of the ideal.

One may use `ideallhnf` to convert any ideal to the last (preferred) format.

- an *extended ideal* is a 2-component vector $[I, t]$, where I is an ideal as above and t is an algebraic number, representing the ideal $(t)I$. This is useful whenever `idealred` is involved, implicitly working in the ideal class group, while keeping track of principal ideals. The following multiplicative ideal operations update the principal part: `idealmul`, `idealinv`, `idealsqr`, `idealpow` and `idealred`; e.g. using `idealmul` on $[I, t]$, $[J, u]$, we obtain $[IJ, tu]$. In all other functions, the extended part is silently discarded, e.g. using `idealadd` with the above input produces $I + J$.

The “principal part” t in an extended ideal may be represented in any of the above forms, and *also* as a factorization matrix (in terms of number field elements, not ideals!), possibly the empty factorization matrix `factor(1)` representing 1; the empty matrix `[,]` is also accepted as a synonym for 1. When t is such a factorization matrix, elements stay in factored form, or *famat* for *factorization matrix*, which is a convenient way to avoid coefficient explosion. To recover the conventional expanded form, try `nfactorback`; but many functions already accept *famats* as input, for instance `ideallog`, so expanding huge elements should never be necessary.

3.13.3 Finite abelian groups.

A finite abelian group G in user-readable format is given by its Smith Normal Form as a pair $[h, d]$ or triple $[h, d, g]$. Here h is the cardinality of G , (d_i) is the vector of elementary divisors, and (g_i) is a vector of generators. In short, $G = \oplus_{i \leq n} (\mathbf{Z}/d_i \mathbf{Z}) g_i$, with $d_n \mid \dots \mid d_2 \mid d_1$ and $\prod_i d_i = h$. This information can also be retrieved as $G.\text{no}$, $G.\text{cyc}$ and $G.\text{gen}$.

- a *character* on the abelian group $\oplus (\mathbf{Z}/d_j \mathbf{Z}) g_j$ is given by a row vector $\chi = [a_1, \dots, a_n]$ such that $\chi(\prod_j g_j^{n_j}) = \exp(2\pi i \sum_j a_j n_j / d_j)$.
- given such a structure, a *subgroup* H is input as a square matrix in HNF, whose columns express generators of H on the given generators g_i . Note that the determinant of that matrix is equal to the index $(G : H)$.

3.13.4 Relative extensions.

We now have a look at data structures attached to relative extensions of number fields L/K , and to projective $\mathbf{Z}_K$ -modules. When defining a relative extension L/K , the nf attached to the base field K must be defined by a variable having a lower priority (see Section 2.5.3) than the variable defining the extension. For example, you may use the variable name y to define the base field K , and x to define the relative extension L/K .

Basic definitions.

- rnf denotes a relative number field, i.e. a data structure output by `rnfinit`, attached to the extension L/K . The nf attached to the base field K is `rnf.nf`.

- A *relative matrix* is an $m \times n$ matrix whose entries are elements of K , in any form. Its m columns A_j represent elements in K^n .

- An *ideal list* is a row vector of fractional ideals of the number field nf .

- A *pseudo-matrix* is a 2-component row vector (A, I) where A is a relative $m \times n$ matrix and I an ideal list of length n . If $I = \{\mathfrak{a}_1, \dots, \mathfrak{a}_n\}$ and the columns of A are $(A_1, \dots, A_n)$, this data defines the torsion-free (projective) $\mathbf{Z}_K$ -module $\mathfrak{a}_1 A_1 \oplus \mathfrak{a}_n A_n$.

- An *integral pseudo-matrix* is a 3-component row vector (A, I, J) where $A = (a_{i,j})$ is an $m \times n$ relative matrix and $I = (\mathfrak{b}_1, \dots, \mathfrak{b}_m)$, $J = (\mathfrak{a}_1, \dots, \mathfrak{a}_n)$ are ideal lists, such that $a_{i,j} \in \mathfrak{b}_i \mathfrak{a}_j^{-1}$ for all i, j . This data defines two abstract projective $\mathbf{Z}_K$ -modules $N = \mathfrak{a}_1 \omega_1 \oplus \dots \oplus \mathfrak{a}_n \omega_n$ in K^n , $P = \mathfrak{b}_1 \eta_1 \oplus \dots \oplus \mathfrak{b}_m \eta_m$ in K^m , and a $\mathbf{Z}_K$ -linear map $f : N \rightarrow P$ given by

$$f\left(\sum_j \alpha_j \omega_j\right) = \sum_i \left(a_{i,j} \alpha_j\right) \eta_i.$$

This data defines the $\mathbf{Z}_K$ -module $M = P/f(N)$.

- Any *projective* $\mathbf{Z}_K$ -module M of finite type in K^m can be given by a pseudo matrix (A, I) .

- An arbitrary $\mathbf{Z}_K$ module of finite type in K^m , with nontrivial torsion, is given by an integral pseudo-matrix (A, I, J)

Algebraic numbers in relative extension.

We are given a number field $K = \text{nfinit}(T)$, attached to $K = \mathbf{Q}[Y]/(T)$, $T \in \mathbf{Q}[Y]$, and a relative extension $L = \text{rnfinit}(K, P)$, attached to $L = K[X]/(P)$, $P \in K[X]$. In all contexts (except `rnfelstabstorel` and `rnfeltdown`, see below), an *algebraic number* is given as

- a `t_INT`, `t_FRAC` or `t_POL` in $\mathbf{Q}[Y]$ (implicitly modulo T) or a `t_POL` in $K[X]$ (implicitly modulo P),

- a `t_POLMOD` (modulo T or P), or

- a `t_COL v` of dimension $m = [K : \mathbf{Q}]$, representing the element in terms of the integral basis `K.zk`;

- if an absolute `nf` structure `Labs` was attached to L , via `Labs = nfinit(L)`, then we can also use a `t_COL v` of dimension $[L : \mathbf{Q}]$, representing the element in terms of the computed integral basis `Labs.zk`. Be careful that in the degenerate case $L = K$, then the previous interpretation (with respect to `K.zk`) takes precedence. This is no concern when $K = \mathbf{Q}$ or if $P = X - Y$ (because in that case the primitive polynomial `Labs.pol` defining L of $\mathbf{Q}$ is `nf.pol` and the computation of `nf.zk` is deterministic); but in other cases, the integer bases attached to K and `Labs` may differ.

Special case: `rnfeltabstorel` and `rnfeltdown`. These two functions assume that elements are given in absolute representation (with respect to `Labs.zk` or modulo `Labs.pol` and converts them to relative representation modulo `L.pol`. In these two functions (only), a `t_POL` in X is implicitly understood modulo `Labs.pol` and a `t_COL` of length $[L : \mathbf{Q}]$ refers to the integral basis `Labs.zk` in all cases, including $L = K$.

Pseudo-bases, determinant.

- The pair (A, I) is a *pseudo-basis* of the module it generates if the $\mathfrak{a}_j$ are nonzero, and the A_j are K -linearly independent. We call n the *size* of the pseudo-basis. If A is a relative matrix, the latter condition means it is square with nonzero determinant; we say that it is in Hermite Normal Form (HNF) if it is upper triangular and all the elements of the diagonal are equal to 1.
- For instance, the relative integer basis `rnf.zk` is a pseudo-basis (A, I) of $\mathbf{Z}_L$, where $A = \text{rnf.zk}[1]$ is a vector of elements of L , which are K -linearly independent. Most *rnf* routines return and handle $\mathbf{Z}_K$ -modules contained in L (e.g. $\mathbf{Z}_L$ -ideals) via a pseudo-basis (A', I') , where A' is a relative matrix representing a vector of elements of L in terms of the fixed basis `rnf.zk[1]`
- The *determinant* of a pseudo-basis (A, I) is the ideal equal to the product of the determinant of A by all the ideals of I . The determinant of a pseudo-matrix is the determinant of any pseudo-basis of the module it generates.

3.13.5 Class field theory.

A *modulus*, in the sense of class field theory, is a divisor supported on the real and finite places of K . In PARI terms, this means either an ordinary ideal I as above (no Archimedean component), or a pair $[I, a]$, where a is a vector with r_1 $\{0, 1\}$ -components, corresponding to the infinite part of the divisor. More precisely, the i -th component of a corresponds to the real embedding attached to the i -th real root of `K.roots`. (That ordering is not canonical, but well defined once a defining polynomial for K is chosen.) For instance, $[1, [1, 1]]$ is a modulus for a real quadratic field, allowing ramification at any of the two places at infinity, and nowhere else.

A *bid* or “big ideal” is a structure output by `idealstar` needed to compute in $(\mathbf{Z}_K/I)^*$, where I is a modulus in the above sense. It is a finite abelian group as described above, supplemented by technical data needed to solve discrete log problems.

Finally we explain how to input ray number fields (or *bnr*), using class field theory. These are defined by a triple A, B, C , where the defining set $[A, B, C]$ can have any of the following forms: $[bnr]$, $[bnr, subgroup]$, $[bnr, character]$, $[bnf, mod]$, $[bnf, mod, subgroup]$. The last two forms are kept for backward compatibility, but no longer serve any real purpose (see example below); no newly written function will accept them.

- *bnf* is as output by `bnfinit`, where units are mandatory unless the modulus is trivial; *bnr* is as output by `bnrinit`. This is the ground field K .
- *mod* is a modulus $\mathfrak{f}$, as described above.
- *subgroup* a subgroup of the ray class group modulo $\mathfrak{f}$ of K . As described above, this is input as a square matrix expressing generators of a subgroup of the ray class group `bnr.clgp` on the given generators. We also allow a `t_INT` n for $n \cdot \text{Cl}_{\mathfrak{f}}$.
- *character* is a character χ of the ray class group modulo $\mathfrak{f}$, representing the subgroup $\text{Ker}\chi$.

The corresponding *bnr* is the subfield of the ray class field of K modulo $\mathfrak{f}$, fixed by the given subgroup.

```

? K = bnfinit(y^2+1);
? bnr = bnrinit(K, 13)
? %.clgp
%3 = [36, [12, 3]]
? bnrdisc(bnr); \\ discriminant of the full ray class field
? bnrdisc(bnr, [3,1;0,1]); \\ discriminant of cyclic cubic extension of K
? bnrconductor(bnr, [3,1]); \\ conductor of chi: g1->zeta_12^3, g2->zeta_3

```

We could have written directly

```

? bnrdisc(K, 13);
? bnrdisc(K, 13, [3,1;0,1]);

```

avoiding one `bnrinit`, but this would actually be slower since the `bnrinit` is called internally anyway. And now twice!

3.13.6 General use.

All the functions which are specific to relative extensions, number fields, Buchmann's number fields, Buchmann's number rays, share the prefix `rnf`, `nf`, `bnf`, `bnr` respectively. They take as first argument a number field of that precise type, respectively output by `rnfinit`, `nfinit`, `bnfinit`, and `bnrinit`.

However, and even though it may not be specified in the descriptions of the functions below, it is permissible, if the function expects a *nf*, to use a *bnf* instead, which contains much more information. On the other hand, if the function requires a `bnf`, it will *not* launch `bnfinit` for you, which is a costly operation. Instead, it will give you a specific error message. In short, the types

$$\text{nf} \leq \text{bnf} \leq \text{bnr}$$

are ordered, each function requires a minimal type to work properly, but you may always substitute a larger type.

The data types corresponding to the structures described above are rather complicated. Thus, as we already have seen it with elliptic curves, GP provides “member functions” to retrieve data from these structures (once they have been initialized of course). The relevant types of number fields are indicated between parentheses:

```

bid      (bnr          ) : bid ideal structure.
bnf      (bnr, bnf     ) : Buchmann's number field.
clgp     (bnr, bnf     ) : classgroup. This one admits the following three subclasses:
  cyc          : cyclic decomposition (SNF).
  gen          : generators.
  no           : number of elements.
diff     (bnr, bnf, nf) : the different ideal.
codiff   (bnr, bnf, nf) : the codifferent (inverse of the different in the ideal group).
disc     (bnr, bnf, nf) : discriminant.
fu       (   bnf      ) : fundamental units.
index    (bnr, bnf, nf) : index of the power order in the ring of integers.
mod      (bnr          ) : modulus.
nf       (bnr, bnf, nf) : number field.
pol      (bnr, bnf, nf) : defining polynomial.
r1       (bnr, bnf, nf) : the number of real embeddings.

```

r2 (*bnr*, *bnf*, *nf*) : the number of pairs of complex embeddings.
reg (*bnf*) : regulator.
roots (*bnr*, *bnf*, *nf*) : roots of the polynomial generating the field.
sign (*bnr*, *bnf*, *nf*) : signature [*r1*, *r2*].
t2 (*bnr*, *bnf*, *nf*) : the T_2 matrix (see **nfinit**).
tu (*bnf*) : a generator for the torsion units.
zk (*bnr*, *bnf*, *nf*) : integral basis, i.e. a $\mathbf{Z}$ -basis of the maximal order.
zkst (*bnr*) : structure of $(\mathbf{Z}_K/m)^*$.

The member functions **.codiff**, **.t2** and **.zk** perform a computation and are relatively expensive in large degree: move them out of tight loops and store them in variables.

For instance, assume that $bnf = \mathbf{bnfinit}(pol)$, for some polynomial. Then $bnf.\mathbf{clgp}$ retrieves the class group, and $bnf.\mathbf{clgp.no}$ the class number. If we had set $bnf = \mathbf{nfinit}(pol)$, both would have output an error message. All these functions are completely recursive, thus for instance $bnr.\mathbf{bnf.nf.zk}$ will yield the maximal order of *bnr*, which you could get directly with a simple $bnr.\mathbf{zk}$.

3.13.7 Class group, units, and the GRH.

Some of the functions starting with **bnf** are implementations of the sub-exponential algorithms for finding class and unit groups under GRH, due to Hafner-McCurley, Buchmann and Cohen-Diaz-Olivier. The general call to the functions concerning class groups of general number fields (i.e. excluding **quadclassunit**) involves a polynomial P and a technical vector

$$tech = [c_1, c_2, nrpid],$$

where the parameters are to be understood as follows:

P is the defining polynomial for the number field, which must be in $\mathbf{Z}[X]$, irreducible and monic. In fact, if you supply a nonmonic polynomial at this point, **gp** issues a warning, then *transforms your polynomial* so that it becomes monic. The **nfinit** routine will return a different result in this case: instead of **res**, you get a vector **[res, Mod(a, Q)]**, where $\text{Mod}(\mathbf{a}, \mathbf{Q}) = \text{Mod}(X, P)$ gives the change of variables. In all other routines, the variable change is simply lost.

The *tech* interface is obsolete and you should not tamper with these parameters. Indeed, from version 2.4.0 on,

- the results are always rigorous under GRH (before that version, they relied on a heuristic strengthening, hence the need for overrides).
- the influence of these parameters on execution time and stack size is marginal. They *can* be useful to fine-tune and experiment with the **bnfinit** code, but you will be better off modifying all tuning parameters in the C code (there are many more than just those three). We nevertheless describe it for completeness.

The numbers $c_1 \leq c_2$ are nonnegative real numbers. By default they are chosen so that the result is correct under GRH. For $i = 1, 2$, let $B_i = c_i(\log |d_K|)^2$, and denote by $S(B)$ the set of maximal ideals of K whose norm is less than B . We want $S(B_1)$ to generate $\text{Cl}(K)$ and hope that $S(B_2)$ can be *proven* to generate $\text{Cl}(K)$.

More precisely, $S(B_1)$ is a factorbase used to compute a tentative $\text{Cl}(K)$ by generators and relations. We then check explicitly, using essentially **bnfisprincipal**, that the elements of $S(B_2)$ belong to the span of $S(B_1)$. Under the assumption that $S(B_2)$ generates $\text{Cl}(K)$, we are done.

User-supplied c_i are only used to compute initial guesses for the bounds B_i , and the algorithm increases them until one can *prove* under GRH that $S(B_2)$ generates $\text{Cl}(K)$. A uniform result of Grenié and Molteni says that $c_2 = 4$ is always suitable, but this bound is very pessimistic and a direct algorithm due to Belabas-Diaz-Friedman, improved by Grenié and Molteni, is used to check the condition, assuming GRH. The default values are $c_1 = c_2 = 0$. When c_1 is equal to 0 the algorithm takes it equal to c_2 .

nrpid is the maximal number of small norm relations attached to each ideal in the factor base. Set it to 0 to disable the search for small norm relations. Otherwise, reasonable values are between 4 and 20. The default is 4.

Warning. Make sure you understand the above! By default, most of the **bnf** routines depend on the correctness of the GRH. In particular, any of the class number, class group structure, class group generators, regulator and fundamental units may be wrong, independently of each other. Any result computed from such a **bnf** may be wrong. The only guarantee is that the units given generate a subgroup of finite index in the full unit group. You must use **bnfcertify** to certify the computations unconditionally.

Remarks.

You do not need to supply the technical parameters (under the library you still need to send at least an empty vector, coded as **NULL**). However, should you choose to set some of them, they *must* be given in the requested order. For example, if you want to specify a given value of *nrpid*, you must give some values as well for c_1 and c_2 , and provide a vector $[c_1, c_2, \text{nrpid}]$.

Note also that you can use an *nf* instead of P , which avoids recomputing the integral basis and analogous quantities.

3.13.8 Hecke Grossencharacters.

Hecke Grossencharacters are continuous characters of the idèle class group; they generalize classical Hecke characters on ray class groups obtained through the *bnr* structure.

Let K be a number field, $\mathbf{A}^\times$ its group of idèles. Every Grossencharacter

$$\chi: \mathbf{A}^\times / K^\times \rightarrow \mathbf{C}^\times$$

can be uniquely written $\chi = \chi_0 \| \cdot \|^s$ for some $s \in \mathbf{C}$ and some character χ_0 of the compact group $\mathbf{A}^\times / (K^\times \cdot \mathbf{R}_{>0})$, where $\|a\| = \prod_v |a_v|_v$ is the idèle norm.

Let $\mathfrak{m}$ be a modulus (an integral ideal and a finite set of real places). Let $U(\mathfrak{m})$ be the subgroup of idèles congruent to 1 modulo $\mathfrak{m}$ (units outside $\mathfrak{m}$, positive at real places in $\mathfrak{m}$). The Hecke Grossencharacters defined modulo $\mathfrak{m}$ are the characters of the idèle class group

$$C_K(\mathfrak{m}) = \mathbf{A}^\times / (K^\times \cdot U(\mathfrak{m})),$$

that is, combinations of an archimedean character χ_∞ on the connected component $K_\infty^\times$ and a ray class group character χ_f satisfying a compatibility condition $\chi_\infty(a)\chi_f(a) = 1$ for all units a congruent to 1 modulo $\mathfrak{m}$.

- *gc* denotes a structure allowing to compute with Hecke Grossencharacters.
- **gcharinit**(*bnf*, *mod*) initializes the structure *gc*. The underlying number field and modulus can be accessed using *gc.bnf* and *gc.mod*.

- `gc.cyc` describes the finite abelian group structure of gc , the torsion part corresponding to finite order ray class characters, the exact zeros corresponding to a lattice of infinite order Grossencharacters, and the approximate zero being a placeholder for the complex powers of the idèle norm.

- A Hecke character of modulus $\mathfrak{m}$ is described as a `t_COL` of coordinates corresponding to `gc.cyc`: all the coordinates are integers except the last one, which can be an arbitrary complex number, or omitted instead of 0.

- Hecke Grossencharacters have L -functions and can be given to all `lfun` functions as a 2 components vector `[gc, chi]`, see also Section 3.17.5.

3.13.9 `bnfcertify(bnf, {flag = 0})`. `bnf` being as output by `bnfinit`, checks whether the result is correct, i.e. whether it is possible to remove the assumption of the Generalized Riemann Hypothesis. It is correct if and only if the answer is 1. If it is incorrect, the program may output some error message, or loop indefinitely. You can check its progress by increasing the debug level. The `bnf` structure must contain the fundamental units:

```
? K = bnfinit(x^3+2^2^3+1); bnfcertify(K)
***   at top-level: K=bnfinit(x^3+2^2^3+1);bnfcertify(K)
***                                     ^-----
*** bnfcertify: precision too low in makeunits [use bnfinit(,1)].
? K = bnfinit(x^3+2^2^3+1, 1); \\ include units
? bnfcertify(K)
%3 = 1
```

If `flag` is present, only certify that the class group is a quotient of the one computed in `bnf` (much simpler in general); likewise, the computed units may form a subgroup of the full unit group. In this variant, the units are no longer needed:

```
? K = bnfinit(x^3+2^2^3+1); bnfcertify(K, 1)
%4 = 1
```

The library syntax is `long bnfcertify0(GEN bnf, long flag)`. Also available is `GEN bnfcertify(GEN bnf) (flag = 0)`.

3.13.10 `bnfdecodemodule(nf, m)`. If m is a module as output in the first component of an extension given by `bnrdisc`, outputs the true module.

```
? K = bnfinit(x^2+23); L = bnrdisc(K, 10); s = L[2]
%1 = [[Vecsmall([8]), Vecsmall([1])], [[0, 0, 0]]],
      [[Vecsmall([9]), Vecsmall([1])], [[0, 0, 0]]]]
? bnfdecodemodule(K, s[1][1])
%2 =
[2 0]
[0 1]
? bnfdecodemodule(K, s[2][1])
%3 =
[2 1]
[0 1]
```

The library syntax is `GEN decodemodule(GEN nf, GEN m)`.

3.13.11 bnfinit($P, \{flag = 0\}, \{tech = []\}$). Initializes a **bnf** structure. Used in programs such as **bnfisprincipal**, **bnfisunit** or **bnfnarrow**. By default, the results are conditional on the GRH, see 3.13.7. The result is a 10-component vector *bnf*.

This implements Buchmann's sub-exponential algorithm for computing the class group, the regulator and a system of fundamental units of the general algebraic number field K defined by the irreducible polynomial P with integer coefficients. The meaning of *flag* is as follows:

- *flag* = 0 (default). This is the historical behavior, kept for compatibility reasons and speed. It has severe drawbacks but is likely to be a little faster than the alternative, twice faster say, so only use it if speed is paramount, you obtain a useful speed gain for the fields under consideration, and you are only interested in the field invariants such as the classgroup structure or its regulator. The computations involve exact algebraic numbers which are replaced by floating point embeddings for the sake of speed. If the precision is insufficient, **gp** may not be able to compute fundamental units, nor to solve some discrete logarithm problems. It *may* be possible to increase the precision of the **bnf** structure using **nfnewprec** but this may fail, in particular when fundamental units are large. In short, the resulting **bnf** structure is correct and contains useful information but later function calls to **bnfisprincipal** or **bnrclassfield** may fail.

When *flag* = 1, we keep an exact algebraic version of all floating point data and this allows to guarantee that functions using the structure will always succeed, as well as to compute the fundamental units exactly. The units are computed in compact form, as a product of small S -units, possibly with huge exponents. This flag also allows **bnfisprincipal** to compute generators of principal ideals in factored form as well. Be warned that expanding such products explicitly can take a very long time, but they can easily be mapped to floating point or ℓ -adic embeddings of bounded accuracy, or to $K^*/(K^*)^\ell$, and this is enough for applications. In short, this flag should be used by default, unless you have a very good reason for it, for instance building massive tables of class numbers, and you do not care about units or the effect large units would have on your computation.

tech is a technical vector (empty by default, see 3.13.7). Careful use of this parameter may speed up your computations, but it is mostly obsolete and you should leave it alone.

The components of a *bnf* are technical. In fact: *never access a component directly, always use a proper member function*. However, for the sake of completeness and internal documentation, their description is as follows. We use the notations explained in the book by H. Cohen, *A Course in Computational Algebraic Number Theory*, Graduate Texts in Maths **138**, Springer-Verlag, 1993, Section 6.5, and subsection 6.5.5 in particular.

bnf[1] contains the matrix W , i.e. the matrix in Hermite normal form giving relations for the class group on prime ideal generators $(\mathfrak{p}_i)_{1 \leq i \leq r}$.

bnf[2] contains the matrix B , i.e. the matrix containing the expressions of the prime ideal factorbase in terms of the $\mathfrak{p}_i$. It is an $r \times c$ matrix.

bnf[3] contains the complex logarithmic embeddings of the system of fundamental units which has been found. It is an $(r_1 + r_2) \times (r_1 + r_2 - 1)$ matrix.

bnf[4] contains the matrix M_C'' of Archimedean components of the relations of the matrix $(W|B)$.

bnf[5] contains the prime factor base, i.e. the list of prime ideals used in finding the relations.

bnf[6] contains a dummy 0.

`bnf[7]` or `bnf.nf` is equal to the number field data `nf` as would be given by `bnfinit`.

`bnf[8]` is a vector containing the classgroup `bnf.clgp` as a finite abelian group, the regulator `bnf.reg`, the number of roots of unity and a generator `bnf.tu`, the fundamental units *in expanded form* `bnf.fu`. If the fundamental units were omitted in the `bnf`, `bnf.fu` returns the sentinel value 0. If `flag = 1`, this vector contains also algebraic data corresponding to the fundamental units and to the discrete logarithm problem (see `bnfisprincipal`). In particular, if `flag = 1` we may *only* know the units in factored form: the first call to `bnf.fu` expands them, which may be very costly, then caches the result.

`bnf[9]` is a vector used in `bnfisprincipal` only and obtained as follows. Let $D = UWV$ obtained by applying the Smith normal form algorithm to the matrix W ($= \text{bnf}[1]$) and let U_r be the reduction of U modulo D . The first elements of the factorbase are given (in terms of `bnf.gen`) by the columns of U_r , with Archimedean component g_a ; let also GD_a be the Archimedean components of the generators of the (principal) ideals defined by the `bnf.gen[i]^bnf.cyc[i]`. Then `bnf[9] = [U_r, g_a, GD_a]`, followed by technical exact components which allow to recompute g_a and GD_a to higher accuracy.

`bnf[10]` is by default unused and set equal to 0. This field is used to store further information about the field as it becomes available, which is rarely needed, hence would be too expensive to compute during the initial `bnfinit` call. For instance, the generators of the principal ideals `bnf.gen[i]^bnf.cyc[i]` (during a call to `bnfisprincipal`), or those corresponding to the relations in W and B (when the `bnf` internal precision needs to be increased).

The library syntax is `GEN bnfinit0(GEN P, long flag, GEN tech = NULL, long prec)`

Also available is `GEN Buchall(GEN P, long flag, long prec)`, corresponding to `tech = NULL`, where `flag` is either 0 (default) or `nf_FORCE` (include all data in algebraic form). The function `GEN Buchall_param(GEN P, double c1, double c2, long nrpid, long flag, long prec)` gives direct access to the technical parameters.

3.13.12 `bnfisintnorm(bnf, x, {flag = 0})`. Computes a complete system of solutions (modulo units of positive norm) of the absolute norm equation $\text{Norm}(a) = x$, where a is an integer in `bnf`. If `bnf` has not been certified, the correctness of the result depends on the validity of GRH. If (optional) `flag` is set, allow returning solutions in factored form, which helps a lot when the fundamental units are large (equivalently, when `bnf.reg` is large); having an exact algebraic `bnf` from `bnfinit(,1)` is necessary in this case, else setting the flag will mostly be a no-op.

```
? bnf = bnfinit(x^4-2, 1);
? bnfisintnorm(bnf,7)
%2 = [-x^2 + x - 1, x^2 + x + 1]
? bnfisintnorm(bnf,-7)
%3 = [-x^3 - 1, x^3 + 2*x^2 + 2*x + 1]

? bnf = bnfinit(x^2-2305843005992468481, 1);
? bnfisintnorm(bnf, 2305843008139952128)
\\ stack overflow with 100GB parsize
? bnf.reg \\ fundamental unit is huge
%6 = 14054016.227457155120413774802385952043

? v = bnfisintnorm(bnf, 2305843008139952128, 1); #v
%7 = 31 \\ succeeds instantly
```

```
? s = v[1]; [type(s), matsize(s)]
%8 = ["t_MAT", [165, 2]]    \\ solution 1 is a product of 165 factors
? exponent(s[,2])
%9 = 105
```

The *exponents* have 105 bits, so there is indeed little hope of writing down the solutions in expanded form.

See also `bnfisnorm`.

The library syntax is `GEN bnfisintnorm0(GEN bnf, GEN x, long flag)`. The function `GEN bnfisintnormabs0(GEN bnf, GEN a, long flag)`, where `bnf` is a true *bnf* structure, returns a complete system of solutions modulo units of the absolute norm equation $|\text{Norm}(x)| = |a|$. As fast as `bnfisintnorm`, but solves the two equations $\text{Norm}(x) = \pm a$ simultaneously. The functions `GEN bnfisintnormabs(GEN bnf, GEN a)`, `GEN bnfisintnorm(GEN bnf, GEN a)` correspond to $\text{flag} = 0$.

3.13.13 `bnfisnorm`(*bnf*, *x*, {*flag* = 1}). Tries to tell whether the rational number *x* is the norm of some element *y* in *bnf*. Returns a vector [*a*, *b*] where $x = \text{Norm}(a) * b$. Looks for a solution which is an *S*-unit, with *S* a certain set of prime ideals containing (among others) all primes dividing *x*. If *bnf* is known to be Galois, you may set *flag* = 0 (in this case, *x* is a norm iff *b* = 1). If *flag* is nonzero the program adds to *S* the following prime ideals, depending on the sign of *flag*. If *flag* > 0, the ideals of norm less than *flag*. And if *flag* < 0 the ideals dividing *flag*.

Assuming GRH, the answer is guaranteed (i.e. *x* is a norm iff *b* = 1), if *S* contains all primes less than $4 \log(\text{disc}(\text{Bnf}))^2$, where *Bnf* is the Galois closure of *bnf*.

See also `bnfisintnorm`.

The library syntax is `GEN bnfisnorm(GEN bnf, GEN x, long flag)`.

3.13.14 `bnfisprincipal`(*bnf*, *x*, {*flag* = 1}). *bnf* being the number field data output by `bnfinit`, and *x* being an ideal, this function tests whether the ideal is principal or not. The result is more complete than a simple true/false answer and solves a general discrete logarithm problem. Assume the class group is $\oplus(\mathbf{Z}/d_i\mathbf{Z})g_i$ (where the generators *g_i* and their orders *d_i* are respectively given by `bnf.gen` and `bnf.cyc`). The routine returns a row vector [*e*, *t*], where *e* is a vector of exponents $0 \leq e_i < d_i$, and *t* is a number field element such that

$$x = (t) \prod_i g_i^{e_i}.$$

For *given* *g_i* (i.e. for a given `bnf`), the *e_i* are unique, and *t* is unique modulo units.

In particular, *x* is principal if and only if *e* is the zero vector. Note that the empty vector, which is returned when the class number is 1, is considered to be a zero vector (of dimension 0).

```
? K = bnfinit(y^2+23);
? K.cyc
%2 = [3]
? K.gen
%3 = [[2, 0; 0, 1]]    \\ a prime ideal above 2
? P = idealprimedec(K,3)[1]; \\ a prime ideal above 3
? v = bnfisprincipal(K, P)
```

```

%5 = [[2]~, [3/4, 1/4]~]
? idealmul(K, v[2], idealfactorback(K, K.gen, v[1]))
%6 =
[3 0]
[0 1]
? % == idealhnf(K, P)
%7 = 1

```

The binary digits of *flag* mean:

- 1: If set, outputs $[e, t]$ as explained above, otherwise returns only e , which is easier to compute. The following idiom only tests whether an ideal is principal:

```
is_principal(bnf, x) = !bnfisprincipal(bnf, x, 0);
```

- 2: It may not be possible to recover t , given the initial accuracy to which the **bnf** structure was computed. In that case, a warning is printed and t is set equal to the empty vector $[]~$. If this bit is set, increase the precision and recompute needed quantities until t can be computed. Warning: setting this may induce *lengthy* computations, and the result may be too large to be physically representable in any case. You should consider using *flag* = 4 instead.

- 4: Return t in factored form (compact representation), as a small product of S -units for a small set of finite places S , possibly with huge exponents. This kind of result can be cheaply mapped to $K^*/(K^*)^\ell$ or to $\mathbf{C}$ or $\mathbf{Q}_p$ to bounded accuracy and this is usually enough for applications. Explicitly expanding such a compact representation is possible using **nfactorback** but may be *very* costly. The algorithm is guaranteed to succeed if the **bnf** was computed using **bnfinit**(,1). If not, the algorithm may fail to compute a huge generator in this case (and replace it by $[]~$). This is orders of magnitude faster than *flag* = 2 when the generators are indeed large.

The library syntax is **GEN bnfisprincipal0**(**GEN bnf**, **GEN x**, **long flag**). Instead of the above hardcoded numerical flags, one should rather use an or-ed combination of the symbolic flags **nf_GEN** (include generators, possibly a place holder if too difficult), **nf_GENMAT** (include generators in compact form) and **nf_FORCE** (insist on finding the generators, a no-op if **nf_GENMAT** is included).

3.13.15 bnfissunit(*bnf*, *sfu*, *x*). This function is obsolete, use **bnfisunit**.

The library syntax is **GEN bnfissunit**(**GEN bnf**, **GEN sfu**, **GEN x**).

3.13.16 bnfisunit(*bnf*, *x*, {*U*}). *bnf* being the number field data output by **bnfinit** and *x* being an algebraic number (type integer, rational or polmod), this outputs the decomposition of x on the fundamental units and the roots of unity if x is a unit, the empty vector otherwise. More precisely, if $u_1, \dots, u_r$ are the fundamental units, and ζ is the generator of the group of roots of unity (**bnf.tu**), the output is a vector $[x_1, \dots, x_r, x_{r+1}]$ such that $x = u_1^{x_1} \dots u_r^{x_r} \cdot \zeta^{x_{r+1}}$. The x_i are integers but the last one ($i = r + 1$) is only defined modulo the order w of ζ and is guaranteed to be in $[0, w[$.

Note that *bnf* need not contain the fundamental units explicitly: it may contain the placeholder 0 instead:

```

? setrand(1); bnf = bnfinit(x^2-x-100000);
? bnf.fu
%2 = 0
? u = [119836165644250789990462835950022871665178127611316131167, \

```

```

379554884019013781006303254896369154068336082609238336]~;
? bnfisunit(bnf, u)
%3 = [-1, 0]~

```

The given u is $1/u_1$, where u_1 is the fundamental unit implicitly stored in bnf . In this case, u_1 was not computed and stored in algebraic form since the default accuracy was too low. Re-run the `bnfinit` command at `\g1` or higher to see such diagnostics.

This function allows x to be given in factored form, but it then assumes that x is an actual unit. (Because it is general too costly to check whether this is the case.)

```

? { v = [2, 85; 5, -71; 13, -162; 17, -76; 23, -37; 29, -104; [224, 1]~, -66;
[-86, 1]~, 86; [-241, 1]~, -20; [44, 1]~, 30; [124, 1]~, 11; [125, -1]~, -11;
[-214, 1]~, 33; [-213, -1]~, -33; [189, 1]~, 74; [190, -1]~, 104;
[-168, 1]~, 2; [-167, -1]~, -8]; }
? bnfisunit(bnf,v)
%5 = [1, 0]~

```

Note that v is the fundamental unit of bnf given in compact (factored) form.

If the argument U is present, as output by `bnfunits(bnf, S)`, then the function decomposes x on the S -units generators given in $U[1]$.

```

? bnf = bnfinit(x^4 - x^3 + 4*x^2 + 3*x + 9, 1);
? bnf.sign
%2 = [0, 2]
? S = idealprimedec(bnf,5); #S
%3 = 2
? US = bnfunits(bnf,S);
? g = US[1]; #g \ \ #S = #g, four S-units generators, in factored form
%5 = 4
? g[1]
%6 = [[6, -3, -2, -2]~ 1]
? g[2]
%7 =
[[-1, 1/2, -1/2, -1/2]~ 1]
[ [4, -2, -1, -1]~ 1]
? [nfactorback(bnf, x) | x <- g]
%8 = [[6, -3, -2, -2]~, [-5, 5, 0, 0]~, [-1, 1, -1, 0]~,
[1, -1, 0, 0]~]
? u = [10,-40,24,11]~;
? a = bnfisunit(bnf, u, US)
%9 = [2, 0, 1, 4]~
? nfactorback(bnf, g, a) \ \ prod_i g[i]^a[i] still in factored form
%10 =
[[6, -3, -2, -2]~ 2]
[ [0, 0, -1, -1]~ 1]
[ [2, -1, -1, 0]~ -2]
[ [1, 1, 0, 0]~ 2]
[ [-1, 1, 1, 1]~ -1]

```

```
[ [1, -1, 0, 0]~ 4]
? nffactorback(bnf,%) \\ u = prod_i g[i]^a[i]
%11 = [10, -40, 24, 11]~
```

The library syntax is GEN `bnfisunit0(GEN bnf, GEN x, GEN U = NULL)`. Also available is GEN `bnfisunit(GEN bnf, GEN x)` for $U = \text{NULL}$.

3.13.17 bnfflog(*bnf*, l). Let *bnf* be a *bnf* structure attached to the number field F and let l be a prime number (hereafter denoted ℓ for typographical reasons). Return the logarithmic ℓ -class group $\widetilde{Cl}_F$ of F . This is an abelian group, conjecturally finite (known to be finite if $F/\mathbf{Q}$ is abelian). The function returns if and only if the group is indeed finite (otherwise it would run into an infinite loop). Let $S = \{\mathfrak{p}_1, \dots, \mathfrak{p}_k\}$ be the set of ℓ -adic places (maximal ideals containing ℓ). The function returns $[D, G(\ell), G']$, where

- D is the vector of elementary divisors for $\widetilde{Cl}_F$.
- $G(\ell)$ is the vector of elementary divisors for the (conjecturally finite) abelian group

$$\widetilde{Cl}(\ell) = \{\mathfrak{a} = \sum_{i \leq k} a_i \mathfrak{p}_i : \deg_F \mathfrak{a} = 0\},$$

where the $\mathfrak{p}_i$ are the ℓ -adic places of F ; this is a subgroup of $\widetilde{Cl}$.

- G' is the vector of elementary divisors for the ℓ -Sylow Cl' of the S -class group of F ; the group $\widetilde{Cl}$ maps to Cl' with a simple co-kernel.

The library syntax is GEN `bnfflog(GEN bnf, GEN l)`.

3.13.18 bnfflogdegree(*nf*, A , l). Let *nf* be a *nf* structure attached to a number field F , and let l be a prime number (hereafter denoted ℓ). The ℓ -adified group of idèles of F quotiented by the group of logarithmic units is identified to the ℓ -group of logarithmic divisors $\oplus \mathbf{Z}_\ell[\mathfrak{p}]$, generated by the maximal ideals of F .

The *degree* map $\deg_F$ is additive with values in $\mathbf{Z}_\ell$, defined by $\deg_F \mathfrak{p} = \tilde{f}_{\mathfrak{p}} \deg_\ell p$, where the integer $\tilde{f}_{\mathfrak{p}}$ is as in `bnffloggef` and $\deg_\ell p$ is $\log_\ell p$ for $p \neq \ell$, $\log_\ell(1 + \ell)$ for $p = \ell \neq 2$ and $\log_\ell(1 + 2^2)$ for $p = \ell = 2$.

Let $A = \prod \mathfrak{p}^{n_{\mathfrak{p}}}$ be an ideal and let $\tilde{A} = \sum n_{\mathfrak{p}}[\mathfrak{p}]$ be the attached logarithmic divisor. Return the exponential of the ℓ -adic logarithmic degree $\deg_F A$, which is a natural number.

The library syntax is GEN `bnfflogdegree(GEN nf, GEN A, GEN l)`.

3.13.19 bnffloggef(*nf*, pr). Let *nf* be a *nf* structure attached to a number field F and let pr be a *prid* structure attached to a maximal ideal $\mathfrak{p}/p$. Return $[\tilde{e}(F_{\mathfrak{p}}/\mathbf{Q}_p), \tilde{f}(F_{\mathfrak{p}}/\mathbf{Q}_p)]$ the logarithmic ramification and residue degrees. Let $\mathbf{Q}_p^c/\mathbf{Q}_p$ be the cyclotomic $\mathbf{Z}_p$ -extension, then $\tilde{e} = [F_{\mathfrak{p}} : F_{\mathfrak{p}} \cap \mathbf{Q}_p^c]$ and $\tilde{f} = [F_{\mathfrak{p}} \cap \mathbf{Q}_p^c : \mathbf{Q}_p]$. Note that $\tilde{e}\tilde{f} = e(\mathfrak{p}/p)f(\mathfrak{p}/p)$, where $e(\mathfrak{p}/p)$ and $f(\mathfrak{p}/p)$ denote the usual ramification and residue degrees.

```
? F = nfinit(y^6 - 3*y^5 + 5*y^3 - 3*y + 1);
? bnffloggef(F, idealprimedec(F,2)[1])
%2 = [6, 1]
? bnffloggef(F, idealprimedec(F,5)[1])
%3 = [1, 2]
```

The library syntax is GEN `bnffloggef(GEN nf, GEN pr)`.

3.13.20 bnfarrow(*bnf*). *bnf* being as output by `bnfinit`, computes the narrow class group of *bnf*. The output is a 3-component row vector *v* analogous to the corresponding class group component *bnf.clgp*: the first component is the narrow class number *v.no*, the second component is a vector containing the SNF cyclic components *v.cyc* of the narrow class group, and the third is a vector giving the generators of the corresponding *v.gen* cyclic groups. Note that this function is a special case of `bnrinit`; the *bnf* need not contain fundamental units.

The library syntax is `GEN bnfarrow(GEN bnf)`.

3.13.21 bnfsignunit(*bnf*). *bnf* being as output by `bnfinit`, this computes an $r_1 \times (r_1 + r_2 - 1)$ matrix having ± 1 components, giving the signs of the real embeddings of the fundamental units. The following functions compute generators for the totally positive units:

```
/* exponents of totally positive units generators on K.tu, K.fu */
tpuexpo(K)=
{ my(M, S = bnfsignunit(K), [m,n] = matsize(S));
  \\ m = K.r1, n = r1+r2-1
  S = matrix(m,n, i,j, if (S[i,j] < 0, 1,0));
  S = concat(vectorv(m,i,1), S); \\ add sign(-1)
  M = matkermod(S, 2);
  if (M, mathnfmodid(M, 2), 2*matid(n+1))
}

/* totally positive fundamental units of bnf K */
tpu(K)=
{ my(ex = tpuexpo(K)[,^1]); \\ remove ex[,1], corresponds to 1 or -1
  my(v = concat(K.tu[2], K.fu));
  [ nffactorback(K, v, c) | c <- ex];
}
```

The library syntax is `GEN signunits(GEN bnf)`.

3.13.22 bnfsunit(*bnf*, *S*). Computes the fundamental *S*-units of the number field *bnf* (output by `bnfinit`), where *S* is a list of prime ideals (output by `idealprimedec`). The output is a vector *v* with 6 components.

v[1] gives a minimal system of (integral) generators of the *S*-unit group modulo the unit group.

v[2] contains technical data needed by `bnfissunit`.

v[3] is an obsoleted component, now the empty vector.

v[4] is the *S*-regulator (this is the product of the regulator, the *S*-class number and the natural logarithms of the norms of the ideals in *S*).

v[5] gives the *S*-class group structure, in the usual abelian group format: a vector whose three components give in order the *S*-class number, the cyclic components and the generators.

v[6] is a copy of *S*.

The library syntax is `GEN bnfsunit(GEN bnf, GEN S, long prec)`. Also available is `GEN sunits_mod_units(GEN bnf, GEN S)` which returns only *v*[1].

3.13.23 bnfunits(*bnf*, {*S*}). Return the fundamental units of the number field *bnf* output by *bnfinit*; if *S* is present and is a list of prime ideals, compute fundamental *S*-units instead. The first component of the result contains independent integral *S*-units generators: first nonunits, then r_1+r_2-1 fundamental units, then the torsion unit. The result may be used as an optional argument to *bnfisunit*. The units are given in compact form: no expensive computation is attempted if the *bnf* does not already contain units.

```
? bnf = bnfinit(x^4 - x^3 + 4*x^2 + 3*x + 9, 1);
? bnf.sign    \\ r1 + r2 - 1 = 1
%2 = [0, 2]
? U = bnfunits(bnf); u = U[1];
? #u \\ r1 + r2 = 2 units
%5 = 2;
? u[1] \\ fundamental unit as factorization matrix
%6 =
[[0, 0, -1, -1]~ 1]
[[2, -1, -1, 0]~ -2]
[ [1, 1, 0, 0]~ 2]
[ [-1, 1, 1, 1]~ -1]
? u[2] \\ torsion unit as factorization matrix
%7 =
[[1, -1, 0, 0]~ 1]
? [nffactorback(bnf, z) | z <- u] \\ same units in expanded form
%8 = [[-1, 1, -1, 0]~, [1, -1, 0, 0]~]
```

Now an example involving *S*-units for a nontrivial *S*:

```
? S = idealprimedec(bnf,5); #S
%9 = 2
? US = bnfunits(bnf, S); uS = US[1];
? g = [nffactorback(bnf, z) | z <- uS] \\ now 4 units
%11 = [[6, -3, -2, -2]~, [-5, 5, 0, 0]~, [-1, 1, -1, 0]~, [1, -1, 0, 0]~]
? bnfisunit(bnf, [10, -40, 24, 11]~)
%12 = []~ \\ not a unit
? e = bnfisunit(bnf, [10, -40, 24, 11]~, US)
%13 = [2, 0, 1, 4]~ \\ ...but an S-unit
? nffactorback(bnf, g, e)
%14 = [10, -40, 24, 11]~
? nffactorback(bnf, uS, e) \\ in factored form
%15 =
[[6, -3, -2, -2]~ 2]
[ [0, 0, -1, -1]~ 1]
[ [2, -1, -1, 0]~ -2]
[ [1, 1, 0, 0]~ 2]
[ [-1, 1, 1, 1]~ -1]
[ [1, -1, 0, 0]~ 4]
```

Note that in more complicated cases, any `nffactorback` fully expanding an element in factored form could be *very* expensive. On the other hand, the final example expands a factorization whose components are themselves in factored form, hence the result is a factored form: this is a cheap operation.

The library syntax is `GEN bnfunits(GEN bnf, GEN S = NULL)`.

3.13.24 bnrL1(*bnr*, {*H*}, {*flag* = 0}). Let *bnr* be the number field data output by `bnrinit` and *H* be a square matrix defining a congruence subgroup of the ray class group corresponding to *bnr* (the trivial congruence subgroup if omitted). This function returns, for each character χ of the ray class group which is trivial on *H*, the value at $s = 1$ (or $s = 0$) of the abelian *L*-function attached to χ . For the value at $s = 0$, the function returns in fact for each χ a vector $[r_\chi, c_\chi]$ where

$$L(s, \chi) = c \cdot s^r + O(s^{r+1})$$

near 0.

The argument *flag* is optional, its binary digits mean 1: compute at $s = 0$ if unset or $s = 1$ if set, 2: compute the primitive *L*-function attached to χ if unset or the *L*-function with Euler factors at prime ideals dividing the modulus of *bnr* removed if set (that is $L_S(s, \chi)$, where *S* is the set of infinite places of the number field together with the finite prime ideals dividing the modulus of *bnr*), 3: return also the character if set.

```
K = bnfinit(x^2-229);
bnr = bnrinit(K,1);
bnrL1(bnr)
```

returns the order and the first nonzero term of $L(s, \chi)$ at $s = 0$ where χ runs through the characters of the class group of $K = \mathbf{Q}(\sqrt{229})$. Then

```
bnr2 = bnrinit(K,2);
bnrL1(bnr2,,2)
```

returns the order and the first nonzero terms of $L_S(s, \chi)$ at $s = 0$ where χ runs through the characters of the class group of *K* and *S* is the set of infinite places of *K* together with the finite prime 2. Note that the ray class group modulo 2 is in fact the class group, so `bnrL1(bnr2,0)` returns the same answer as `bnrL1(bnr,0)`.

This function will fail with the message

```
*** bnrL1: overflow in zeta_get_N0 [need too many primes].
```

if the approximate functional equation requires us to sum too many terms (if the discriminant of *K* is too large).

The library syntax is `GEN bnrL1(GEN bnr, GEN H = NULL, long flag, long prec)`.

3.13.25 bnrchar($G, g, \{v\}$). Returns all characters χ on G such that $\chi(g_i) = e(v_i)$, where $e(x) = \exp(2i\pi x)$. G is allowed to be a *bnr* struct (representing a ray class group) or a *znstar* (representing $(\mathbf{Z}/N\mathbf{Z})^*$). If v is omitted, returns all characters that are trivial on the g_i . Else the vectors g and v must have the same length, the g_i must be elements of G , and each v_i is a rational number whose denominator must divide the order of g_i in G .

For convenience, the vector of the g_i can be replaced by a matrix whose columns give their discrete logarithm in G , for instance as given by **bnrisprincipal** if G is a *bnr*; in this particular case, G can be any finite abelian group given by a vector of elementary divisors.

```
? G = bnrinit(bnfinit(x), [160,[1]], 1); /* (Z/160Z)^* */
? G.cyc
%2 = [8, 4, 2]
? g = G.gen;
? bnrchar(G, g, [1/2,0,0])
%4 = [[4, 0, 0]] \\ a unique character
? bnrchar(G, [g[1],g[3]]) \\ all characters trivial on g[1] and g[3]
%5 = [[0, 1, 0], [0, 2, 0], [0, 3, 0], [0, 0, 0]]
? bnrchar(G, [1,0,0;0,1,0;0,0,2])
%6 = [[0, 0, 1], [0, 0, 0]] \\ characters trivial on given subgroup
? G = znstar(75, 1);
? bnrchar(G, [2, 7], [11/20, 1/4])
%8 = [[1, 1]] \\ Dirichlet char: chi(2) = e(11/20), chi(7) = e(1/4)
```

The library syntax is GEN **bnrchar**(GEN G , GEN g , GEN $v = \text{NULL}$).

3.13.26 bnrclassfield(*bnr*, {*subgp*}, {*flag* = 0}). *bnr* being as output by **bnrinit**, returns a relative equation for the class field corresponding to the congruence group defined by (*bnr*, *subgp*) (the full ray class field if *subgp* is omitted). The subgroup can also be a **t_INT** n , meaning $n \cdot \text{Cl}_f$. The function also handles a vector of subgroup, e.g, from **subgrouplist** and returns the vector of individual results in this case.

If *flag* = 0, returns a vector of polynomials such that the compositum of the corresponding fields is the class field; if *flag* = 1 returns a single polynomial; if *flag* = 2 returns a single absolute polynomial.

```
? bnf = bnfinit(y^3+14*y-1); bnf.cyc
%1 = [4, 2]
? pol = bnrclassfield(bnf,,1) \\ Hilbert class field
%2 = x^8 - 2*x^7 + ... + Mod(11*y^2 - 82*y + 116, y^3 + 14*y - 1)
? rnfdisc(bnf,pol)[1]
%3 = 1
? bnr = bnrinit(bnf,3*5*7); bnr.cyc
%4 = [24, 12, 12, 2]
? bnrclassfield(bnr,2) \\ maximal 2-elementary subextension
%5 = [x^2 + (-21*y - 105), x^2 + (-5*y - 25), x^2 + (-y - 5), x^2 + (-y - 1)]
\\ quadratic extensions of maximal conductor
? bnrclassfield(bnr, subgrouplist(bnr,[2]))
%6 = [[x^2 - 105], [x^2 + (-105*y^2 - 1260)], [x^2 + (-105*y - 525)],
      [x^2 + (-105*y - 105)]]
? #bnrclassfield(bnr,subgrouplist(bnr,[2],1)) \\ all quadratic extensions
```

```
%7 = 15
```

When the subgroup contains $n\text{Cl}_f$, where n is fixed, it is advised to directly compute the `bnr` modulo n to avoid expensive discrete logarithms:

```
? bnf = bnfininit(y^2-5); p = 1594287814679644276013;
? bnr = bnrinit(bnf,p); \\ very slow
time = 24,146 ms.
? bnrclassfield(bnr, 2) \\ ... even though the result is trivial
%3 = [x^2 - 1594287814679644276013]
? bnr2 = bnrinit(bnf,p,,2); \\ now fast
time = 1 ms.
? bnrclassfield(bnr2, 2)
%5 = [x^2 - 1594287814679644276013]
```

This will save a lot of time when the modulus contains a maximal ideal whose residue field is large.

The library syntax is `GEN bnrclassfield(GEN bnr, GEN subgp = NULL, long flag, long prec)`.

3.13.27 `bnrclassno`($A, \{B\}, \{C\}$). Let A, B, C define a class field L over a ground field K (of type `[bnr]`, `[bnr, subgroup]`, or `[bnf, modulus]`, or `[bnf, modulus, subgroup]`, Section 3.13.5); this function returns the relative degree $[L : K]$.

In particular if A is a *bnf* (with units), and B a modulus, this function returns the corresponding ray class number modulo B . One can input the attached *bid* (with generators if the subgroup C is non trivial) for B instead of the module itself, saving some time.

This function is faster than `bnrinit` and should be used if only the ray class number is desired. See `bnrclassnolist` if you need ray class numbers for all moduli less than some bound.

The library syntax is `GEN bnrclassno0(GEN A, GEN B = NULL, GEN C = NULL)`. Also available is `GEN bnrclassno(GEN bnf, GEN f)` to compute the ray class number modulo f .

3.13.28 `bnrclassnolist`(*bnf*, *list*). *bnf* being as output by `bnfininit`, and *list* being a list of moduli (with units) as output by `ideallist` or `ideallistarch`, outputs the list of the class numbers of the corresponding ray class groups. To compute a single class number, `bnrclassno` is more efficient.

```
? bnf = bnfininit(x^2 - 2);
? L = ideallist(bnf, 100, 2);
? H = bnrclassnolist(bnf, L);
? H[98]
%4 = [1, 3, 1]
? l = L[1][98]; ids = vector(#l, i, l[i].mod[1])
%5 = [[98, 88; 0, 1], [14, 0; 0, 7], [98, 10; 0, 1]]
```

The weird `l[i].mod[1]`, is the first component of `l[i].mod`, i.e. the finite part of the conductor. (This is cosmetic: since by construction the Archimedean part is trivial, I do not want to see it). This tells us that the ray class groups modulo the ideals of norm 98 (printed as %5) have respectively order 1, 3 and 1. Indeed, we may check directly:

```
? bnrclassno(bnf, ids[2])
%6 = 3
```

The library syntax is `GEN bnrclassnolist(GEN bnf, GEN list)`.

3.13.29 bnrcompositum(A, B). Given two abelian extensions $A = [\mathbf{bnr1}, H1]$ and $B = [\mathbf{bnr2}, H2]$, where $\mathbf{bnr1}$ and $\mathbf{bnr2}$ are two **bnr** structures attached to the same base field, return their compositum as $[\mathbf{bnr}, H]$. The modulus attached to $\mathbf{bnr}$ need not be the conductor of the compositum.

```
? Q = bnfinit(y);
? bnr1 = bnrinit(Q, [7, [1]]); bnr1.cyc
%2 = [6]
? bnr2 = bnrinit(Q, [13, [1]]); bnr2.cyc
%3 = [12]
? H1 = Mat(2); bnrclassfield(bnr1, H1)
%4 = [x^2 + 7]
? H2 = Mat(2); bnrclassfield(bnr2, H2)
%5 = [x^2 - 13]
? [bnr,H] = bnrcompositum([bnr1, H1], [bnr2,H2]);
? bnrclassfield(bnr,H)
%7 = [x^2 - 13, x^2 + 7]
```

The library syntax is `GEN bnrcompositum(GEN A, GEN B)`.

3.13.30 bnrconductor($A, \{B\}, \{C\}, \{flag = 0\}$). Conductor f of the subfield of a ray class field as defined by $[A, B, C]$ (of type $[\mathbf{bnr}]$, $[\mathbf{bnr}, subgroup]$, $[\mathbf{bnf}, modulus]$ or $[\mathbf{bnf}, modulus, subgroup]$, Section 3.13.5)

If $flag = 0$, returns f .

If $flag = 1$, returns $[f, Cl_f, H]$, where Cl_f is the ray class group modulo f , as a finite abelian group; finally H is the subgroup of Cl_f defining the extension.

If $flag = 2$, returns $[f, \mathbf{bnr}(f), H]$, as above except Cl_f is replaced by a **bnr** structure, as output by `bnrinit(f)`, without generators unless the input contained a $\mathbf{bnr}$ with generators.

In place of a subgroup H , this function also accepts a character $\mathbf{chi} = (a_j)$, expressed as usual in terms of the generators $\mathbf{bnr.gen}$: $\chi(g_j) = \exp(2i\pi a_j/d_j)$, where g_j has order $d_j = \mathbf{bnr.cyc}[j]$. In which case, the function returns respectively

If $flag = 0$, the conductor f of $\text{Ker}\chi$.

If $flag = 1$, $[f, Cl_f, \chi_f]$, where χ_f is χ expressed on the minimal ray class group, whose modulus is the conductor.

If $flag = 2$, $[f, \mathbf{bnr}(f), \chi_f]$.

Note. Using this function with $flag \neq 0$ is usually a bad idea and kept for compatibility and convenience only: $flag = 1$ has always been useless, since it is no faster than $flag = 2$ and returns less information; $flag = 2$ is mostly OK with two subtle drawbacks:

- it returns the full bnr attached to the full ray class group, whereas in applications we only need Cl_f modulo N -th powers, where N is any multiple of the exponent of Cl_f/H . Computing directly the conductor, then calling `bnrinit` with optional argument N avoids this problem.
- computing the bnr needs only be done once for each conductor, which is not possible using this function.

For maximal efficiency, the recommended procedure is as follows. Starting from data (character or congruence subgroups) attached to a modulus m , we can first compute the conductors using this function with default $flag = 0$. Then for all data with a common conductor $f \mid m$, compute (once!) the bnr attached to f using `bnrinit` (modulo N -th powers for a suitable N !) and finally map original data to the new bnr using `bnrmap`.

The library syntax is `GEN bnrconductor0(GEN A, GEN B = NULL, GEN C = NULL, long flag)`.

Also available are `GEN bnrconductor(GEN bnr, GEN H, long flag)` and `GEN bnrconductormod(GEN bnr, GEN H, long flag, GEN cycmod)` which returns ray class groups modulo `cycmod`-th powers.

3.13.31 `bnrconductorofchar`(bnr, chi). This function is obsolete, use `bnrconductor`.

The library syntax is `GEN bnrconductorofchar(GEN bnr, GEN chi)`.

3.13.32 `bnrdisc`($A, \{B\}, \{C\}, \{flag = 0\}$). A, B, C defining a class field L over a ground field K (of type `[bnr]`, `[bnr, subgroup]`, `[bnr, character]`, `[bnf, modulus]` or `[bnf, modulus, subgroup]`, Section 3.13.5), outputs data $[N, r_1, D]$ giving the discriminant and signature of L , depending on the binary digits of $flag$:

- 1: if this bit is unset, output absolute data related to $L/\mathbf{Q}$: N is the absolute degree $[L : \mathbf{Q}]$, r_1 the number of real places of L , and D the discriminant of $L/\mathbf{Q}$. Otherwise, output relative data for L/K : N is the relative degree $[L : K]$, r_1 is the number of real places of K unramified in L (so that the number of real places of L is equal to r_1 times N), and D is the relative discriminant ideal of L/K .

- 2: if this bit is set and if the modulus is not the conductor of L , only return 0.

The library syntax is `GEN bnrdisc0(GEN A, GEN B = NULL, GEN C = NULL, long flag)`.

3.13.33 `bnrdisclist`($bnf, bound, \{arch\}$). bnf being as output by `bnfinit` (with units), computes a list of discriminants of Abelian extensions of the number field by increasing modulus norm up to bound $bound$. The ramified Archimedean places are given by $arch$; all possible values are taken if $arch$ is omitted.

The alternative syntax `bnrdisclist(bnf, list)` is supported, where $list$ is as output by `ideal-list` or `ideallistarch` (with units), in which case $arch$ is disregarded.

The output v is a vector, where $v[k]$ is itself a vector w , whose length is the number of ideals of norm k .

- We consider first the case where *arch* was specified. Each component of *w* corresponds to an ideal *m* of norm *k*, and gives invariants attached to the ray class field *L* of *bnf* of conductor $[m, arch]$. Namely, each contains a vector $[m, d, r, D]$ with the following meaning: *m* is the prime ideal factorization of the modulus, $d = [L : \mathbf{Q}]$ is the absolute degree of *L*, *r* is the number of real places of *L*, and *D* is the factorization of its absolute discriminant. We set $d = r = D = 0$ if *m* is not the finite part of a conductor.

- If *arch* was omitted, all $t = 2^{r_1}$ possible values are taken and a component of *w* has the form $[m, [[d_1, r_1, D_1], \dots, [d_t, r_t, D_t]]]$, where *m* is the finite part of the conductor as above, and $[d_i, r_i, D_i]$ are the invariants of the ray class field of conductor $[m, v_i]$, where *v_i* is the *i*-th Archimedean component, ordered by inverse lexicographic order; so $v_1 = [0, \dots, 0]$, $v_2 = [1, 0, \dots, 0]$, etc. Again, we set $d_i = r_i = D_i = 0$ if $[m, v_i]$ is not a conductor.

Finally, each prime ideal $pr = [p, \alpha, e, f, \beta]$ in the prime factorization *m* is coded as the integer $p \cdot n^2 + (f - 1) \cdot n + (j - 1)$, where *n* is the degree of the base field and *j* is such that

```
pr = idealprimedec(nf,p)[j].
```

m can be decoded using `bnfdecodemodule`.

Note that to compute such data for a single field, either `bnrclassno` or `bnrdisc` are (much) more efficient.

The library syntax is `GEN bnrdisc1ist0(GEN bnf, GEN bound, GEN arch = NULL)`.

3.13.34 `bnrgaloisapply`(*bnr*, *mat*, *H*). Apply the automorphism given by its matrix *mat* to the congruence subgroup *H* given as a HNF matrix. The matrix *mat* can be computed with `bnrgaloismatrix`.

The library syntax is `GEN bnrgaloisapply(GEN bnr, GEN mat, GEN H)`.

3.13.35 `bnrgaloismatrix`(*bnr*, *aut*). Return the matrix of the action of the automorphism *aut* of the base field `bnf.nf` on the generators of the ray class field `bnr.gen`. The automorphism *aut* can be given as a polynomial, an algebraic number, or a vector of automorphisms and must stabilize the modulus `bnr.mod`. We also allow a Galois group as output by `galoisinit`, in which case a vector of matrices is returned corresponding to the generators `aut.gen`. Note: This function only makes sense when the ray class field attached to *bnr* is Galois, which is not checked.

The generators `bnr.gen` need not be explicitly computed in the input *bnr*, which saves time: the result is well defined in this case also.

```
? K = bnfinit(a^4-3*a^2+253009); B = bnrinit(K,9); B.cyc
%1 = [8400, 12, 6, 3]
? G = nfgaloisconj(K)
%2 = [-a, a, -1/503*a^3 + 3/503*a, 1/503*a^3 - 3/503*a]~
? bnrgaloismatrix(B, G[2]) \\ G[2] = Id ...
%3 =
[1 0 0 0]
[0 1 0 0]
[0 0 1 0]
[0 0 0 1]
? bnrgaloismatrix(B, G[3]) \\ automorphism of order 2
```

```

%4 =
[799 0 0 2800]
[ 0 7 0 4]
[ 4 0 5 2]
[ 0 0 0 2]
? M = %^2; for (i=1, #B.cyc, M[i,] %= B.cyc[i]); M
%5 = \\ acts on ray class group as automorphism of order 2
[1 0 0 0]
[0 1 0 0]
[0 0 1 0]
[0 0 0 1]

```

See `bnrisgalois` for further examples.

The library syntax is `GEN bnr_galoismatrix(GEN bnr, GEN aut)`. When *aut* is a polynomial or an algebraic number, `GEN bnr_autmatrix(GEN bnr, GEN aut)` is available.

3.13.36 `bnrinit`(*bnf*, *f*, {*flag* = 0}, {*cycmod*}). *bnf* is as output by `bnfinit` (including fundamental units), *f* is a modulus, initializes data linked to the ray class group structure corresponding to this module, a so-called `bnr` structure. One can input the attached *bid* with generators for *f* instead of the module itself, saving some time. (As in `idealstar`, the finite part of the conductor may be given by a factorization into prime ideals, as produced by `idealfactor`.)

If the positive integer *cycmod* is present, only compute the ray class group modulo *cycmod*, which may save a lot of time when some maximal ideals in *f* have a huge residue field. In applications, we are given a congruence subgroup *H* and study the class field attached to Cl_f/H . If that finite Abelian group has an exponent which divides *cycmod*, then we have changed nothing theoretically, while trivializing expensive discrete logs in residue fields (since computations can be made modulo *cycmod*-th powers). This is useful in `bnrclassfield`, for instance when computing *p*-elementary extensions.

The following member functions are available on the result: `.bnf` is the underlying *bnf*, `.mod` the modulus, `.bid` the *bid* structure attached to the modulus; finally, `.clgp`, `.no`, `.cyc`, `.gen` refer to the ray class group (as a finite abelian group), its cardinality, its elementary divisors, its generators (only computed if *flag* = 1).

The last group of functions are different from the members of the underlying *bnf*, which refer to the class group; use `bnr.bnf.xxx` to access these, e.g. `bnr.bnf.cyc` to get the cyclic decomposition of the class group.

They are also different from the members of the underlying *bid*, which refer to $(\mathbf{Z}_K/f)^*$; use `bnr.bid.xxx` to access these, e.g. `bnr.bid.no` to get $\phi(f)$.

If *flag* = 0 (default), the generators of the ray class group are not explicitly computed, which saves time. Hence `bnr.gen` would produce an error. Note that implicit generators are still fixed and stored in the *bnr* (and guaranteed to be the same for fixed *bnf* and *bid* inputs), in terms of `bnr.bnf.gen` and `bnr.bid.gen`. The computation which is not performed is the expansion of such products in the ray class group so as to fix explicit ideal representatives.

If *flag* = 1, as the default, except that generators are computed.

The library syntax is `GEN bnrinitmod(GEN bnf, GEN f, long flag, GEN cycmod = NULL)`. Instead of the above hardcoded numerical flags, one should rather use `GEN Buchraymod(GEN bnf, GEN module, long flag, GEN cycmod)` where an omitted `cycmod` is coded as `NULL` and `flag` is an or-ed combination of `nf_GEN` (include generators) and `nf_INIT` (if omitted, return just the cardinality of the ray class group and its structure), possibly 0. Or simply `GEN Buchray(GEN bnf, GEN module, long flag)` when `cycmod` is `NULL`.

3.13.37 bnrconductor($A, \{B\}, \{C\}$). Fast variant of `bnrconductor( $A, B, C$ )`; A, B, C represent an extension of the base field, given by class field theory (see Section 3.13.5). Outputs 1 if this modulus is the conductor, and 0 otherwise. This is slightly faster than `bnrconductor` when the character or subgroup is not primitive.

The library syntax is `long bnrconductor0(GEN A, GEN B = NULL, GEN C = NULL)`.

3.13.38 bnrsgalois(bnr, gal, H). Check whether the class field attached to the subgroup H is Galois over the subfield of `bnr.nf` fixed by the group gal , which can be given as output by `galoisinit`, or as a matrix or a vector of matrices as output by `bnrgaloismatrix`, the second option being preferable, since it saves the recomputation of the matrices. Note: The function assumes that the ray class field attached to bnr is Galois, which is not checked.

In the following example, we list the congruence subgroups of subextension of degree at most 3 of the ray class field of conductor 9 which are Galois over the rationals.

```
? K = bnfinit(a^4-3*a^2+253009); B = bnrinit(K,9); G = galoisinit(K);
? [H | H<-subgrouplist(B,3), bnrsgalois(B,G,H)];
time = 160 ms.
? M = bnrgaloismatrix(B,G);
? [H | H<-subgrouplist(B,3), bnrsgalois(B,M,H)]
time = 1 ms.
```

The second computation is much faster since `bnrgaloismatrix(B,G)` is computed only once.

The library syntax is `long bnrsgalois(GEN bnr, GEN gal, GEN H)`.

3.13.39 bnrprincipal($bnr, x, \{flag = 1\}$). Let bnr be the ray class group data output by `bnrinit(,1)` and let x be an ideal in any form, coprime to the modulus $f = \text{bnr.mod}$. Solves the discrete logarithm problem in the ray class group, with respect to the generators `bnr.gen`, in a way similar to `bnfisprincipal`. If x is not coprime to the modulus of bnr the result is undefined. Note that bnr need not contain the ray class group generators, i.e. it may be created with `bnrinit(,0)`; in that case, although `bnr.gen` is undefined, we can still fix natural generators for the ray class group (in terms of the generators in `bnr.bnf.gen` and `bnr.bid.gen`) and compute with respect to them.

The binary digits of `flag` (default `flag = 1`) mean:

- 1: If set returns a 2-component vector $[e, \alpha]$ where e is the vector of components of x on the ray class group generators, α is an element congruent to $1 \bmod^* f$ such that $x = \alpha \prod_i g_i^{e_i}$. If unset, returns only e .

- 4: If set, returns $[e, \alpha]$ where α is given in factored form (compact representation). This is orders of magnitude faster.

```
? K = bnfinit(x^2 - 30); bnr = bnrinit(K, [4, [1,1]]);
? bnr.clgp \ ray class group is isomorphic to Z/4 x Z/2 x Z/2
```

```

%2 = [16, [4, 2, 2]]
? P = idealprimedec(K, 3)[1]; \\ the ramified prime ideal above 3
? bnrprincipal(bnr,P) \\ bnr.gen undefined !
%5 = [[3, 0, 0]~, 9]
? bnrprincipal(bnr,P, 0) \\ omit principal part
%5 = [3, 0, 0]~
? bnr = bnrinit(bnr, bnr.bid, 1); \\ include explicit generators
? bnrprincipal(bnr,P) \\ ... alpha is different !
%7 = [[3, 0, 0]~, 1/128625]

```

It may be surprising that the generator α is different although the underlying *bnf* and *bid* are the same. This defines unique generators for the ray class group as ideal *classes*, whether we use `bnrinit(,0)` or `bnrinit(,1)`. But the actual ideal representatives (implicit if *flag* = 0, computed and stored in the *bnr* if *flag* = 1) are in general different and this is what happens here. Indeed, the implicit generators are naturally expressed in terms of `bnr.bnf.gen` and `bnr.bid.gen` and *then* expanded and simplified (in the same ideal class) so that we obtain ideal representatives for `bnr.gen` which are as simple as possible. And indeed the quotient of the two α found is 1 modulo the conductor (and positive at the infinite places it contains), and this is the only guaranteed property.

Beware that, when *bnr* is generated using `bnrinit(, cycmod)`, the results are given in Cl_f modulo *cycmod*-th powers:

```

? bnr2 = bnrinit(K, bnr.mod,, 2); \\ modulo squares
? bnr2.clgp
%9 = [8, [2, 2, 2]] \\ bnr.clgp tensored by Z/2Z
? bnrprincipal(bnr2,P, 0)
%10 = [1, 0, 0]~

```

The library syntax is `GEN bnrprincipal(GEN bnr, GEN x, long flag)`. Instead of hard-coded numerical flags, one should rather use `GEN isprincipalray(GEN bnr, GEN x)` for *flag* = 0, and if you want generators:

```
bnrprincipal(bnr, x, nf_GEN)
```

Also available is `GEN bnrprincipalmod(GEN bnr, GEN x, GEN mod, long flag)` that returns the discrete logarithm of *x* modulo the `t_INT mod`; the value `mod = NULL` is treated as 0 (full discrete logarithm), and *flag* = 1 is not allowed if *mod* is set.

3.13.40 bnrmap(*A*, *B*). This function has two different uses:

- if *A* and *B* are *bnr* structures for the same *bnf* attached to moduli m_A and m_B with $m_B \mid m_A$, return the canonical surjection from *A* to *B*, i.e. from the ray class group modulo m_A to the ray class group modulo m_B . The map is coded by a triple $[M, \text{cyc}_A, \text{cyc}_B]$: *M* gives the image of the fixed ray class group generators of *A* in terms of the ones in *B*, *cyc_A* and *cyc_B* are the cyclic structures `A.cyc` and `B.cyc` respectively. Note that this function does *not* need *A* or *B* to contain explicit generators for the ray class groups: they may be created using `bnrinit(,0)`.

If *B* is only known modulo *N*-th powers (from `bnrinit(,N)`), the result is correct provided *N* is a multiple of the exponent of *A*.

- if *A* is a projection map as above and *B* is either a congruence subgroup *H*, or a ray class character χ , or a discrete logarithm (from `bnrprincipal`) modulo m_A whose conductor divides

m_B , return the image of the subgroup (resp. the character, the discrete logarithm) as defined modulo m_B . The main use of this variant is to compute the primitive subgroup or character attached to a *bnr* modulo their conductor. This is more efficient than `bnrconductor` in two respects: the *bnr* attached to the conductor need only be computed once and, most importantly, the ray class group can be computed modulo N -th powers, where N is a multiple of the exponent of Cl_{m_A}/H (resp. of the order of χ). Whereas `bnrconductor` is specified to return a *bnr* attached to the full ray class group, which may lead to untractable discrete logarithms in the full ray class group instead of a tiny quotient.

The library syntax is `GEN bnrmap(GEN A, GEN B)`.

3.13.41 `bnrrootnumber`(*bnr*, *chi*, {*flag* = 0}). If $\chi = \textit{chi}$ is a character over *bnr*, not necessarily primitive, let $L(s, \chi) = \sum_{id} \chi(id) N(id)^{-s}$ be the attached Artin L-function. Returns the so-called Artin root number, i.e. the complex number $W(\chi)$ of modulus 1 such that

$$\Lambda(1-s, \chi) = W(\chi) \Lambda(s, \bar{\chi})$$

where $\Lambda(s, \chi) = A(\chi)^{s/2} \gamma_\chi(s) L(s, \chi)$ is the enlarged L-function attached to L .

You can set *flag* = 1 if the character is known to be primitive. Example:

```
bnf = bnfinit(x^2 - x - 57);
bnr = bnrinit(bnf, [7, [1, 1]]);
bnrrootnumber(bnr, [2, 1])
```

returns the root number of the character χ of $\text{Cl}_{7\infty_1\infty_2}(\mathbf{Q}(\sqrt{229}))$ defined by $\chi(g_1^a g_2^b) = \zeta_1^{2a} \zeta_2^b$. Here g_1, g_2 are the generators of the ray-class group given by `bnr.gen` and $\zeta_1 = e^{2i\pi/N_1}, \zeta_2 = e^{2i\pi/N_2}$ where N_1, N_2 are the orders of g_1 and g_2 respectively ($N_1 = 6$ and $N_2 = 3$ as `bnr.cyc` readily tells us).

The library syntax is `GEN bnrrootnumber(GEN bnr, GEN chi, long flag, long prec)`.

3.13.42 `bnrstark`(*bnr*, {*subgroup*}). *bnr* being as output by `bnrinit`, finds a relative equation for the class field corresponding to the modulus in *bnr* and the given congruence subgroup (as usual, omit *subgroup* if you want the whole ray class group).

The main variable of *bnr* must not be x , and the ground field and the class field must be totally real. When the base field is $\mathbf{Q}$, the vastly simpler `galoissubcyclo` is used instead. Here is an example:

```
bnf = bnfinit(y^2 - 3);
bnr = bnrinit(bnf, 5);
bnrstark(bnr)
```

returns the ray class field of $\mathbf{Q}(\sqrt{3})$ modulo 5. Usually, one wants to apply to the result one of

```
rnfpolredbest(bnf, pol)    \\ compute a reduced relative polynomial
rnfpolredbest(bnf, pol, 2) \\ compute a reduced absolute polynomial
```

The routine uses Stark units and needs to find a suitable auxiliary conductor, which may not exist when the class field is not cyclic over the base. In this case `bnrstark` is allowed to return a vector of polynomials defining *independent* relative extensions, whose compositum is the requested

class field. We decided that it was useful to keep the extra information thus made available, hence the user has to take the compositum herself, see `nfcompositum`.

Even if it exists, the auxiliary conductor may be so large that later computations become unfeasible. (And of course, Stark's conjecture may simply be wrong.) In case of difficulties, try `bnrclassfield`:

```
? bnr = bnrinit(bnfinit(y^8-12*y^6+36*y^4-36*y^2+9,1), 2);
? bnrstark(bnr)
***   at top-level: bnrstark(bnr)
***   ^-----
*** bnrstark: need 3919350809720744 coefficients in initzeta.
*** Computation impossible.
? bnrclassfield(bnr)
time = 20 ms.
%2 = [x^2 + (-2/3*y^6 + 7*y^4 - 14*y^2 + 3)]
```

The library syntax is GEN `bnrstark(GEN bnr, GEN subgroup = NULL, long prec)`.

3.13.43 `bnrstarkunit`(*bnr*, {*subgroup*}). *bnr* being as output by `bnrinit`, returns the characteristic polynomial of the (conjectural) Stark unit corresponding to the modulus in *bnr* and the given congruence subgroup (as usual, omit *subgroup* if you want the whole ray class group).

The ground field attached to *bnr* must be totally real and all but one infinite place must become complex in the class field, which must be a quadratic extension of its totally real subfield. Finally, the output is given as a polynomial in *x*, so the main variable of *bnr* must not be *x*. Here is an example:

```
? bnf = bnfinit(y^2 - 2);
? bnr = bnrinit(bnf, [15, [1,0]]);
? lift(bnrstarkunit(bnr))
%3 = x^8 + (-9000*y - 12728)*x^7 + (57877380*y + 81850978)*x^6 + ... + 1
```

The library syntax is GEN `bnrstarkunit(GEN bnr, GEN subgroup = NULL)`.

3.13.44 `dirzetak`(*nf*, *b*). Gives as a vector the first *b* coefficients of the Dedekind zeta function of the number field *nf* considered as a Dirichlet series.

The library syntax is GEN `dirzetak(GEN nf, GEN b)`.

3.13.45 `factornf`(*x*, *t*). This function is obsolete, use `nffactor`.

factorization of the univariate polynomial *x* over the number field defined by the (univariate) polynomial *t*. *x* may have coefficients in $\mathbf{Q}$ or in the number field. The algorithm reduces to factorization over $\mathbf{Q}$ (Trager's trick). The direct approach of `nffactor`, which uses van Hoeij's method in a relative setting, is in general faster.

The main variable of *t* must be of *lower* priority than that of *x* (see Section 2.5.3). However if nonrational number field elements occur (as polmods or polynomials) as coefficients of *x*, the variable of these polmods *must* be the same as the main variable of *t*. For example

```
? factornf(x^2 + Mod(y, y^2+1), y^2+1);
? factornf(x^2 + y, y^2+1); \\ these two are OK
? factornf(x^2 + Mod(z,z^2+1), y^2+1)
```

```

*** at top-level: factornf(x^2+Mod(z,z
***
*** factornf: inconsistent data in rnf function.
? factornf(x^2 + z, y^2+1)
*** at top-level: factornf(x^2+z,y^2+1
***
*** factornf: incorrect variable in rnf function.

```

The library syntax is GEN polfnf(GEN x, GEN t).

3.13.46 galoischarDET(gal, chi, {o = 1}). Let G be the group attached to the `galoisinit` structure gal , and let χ be the character of some representation ρ of the group G , where a polynomial variable is to be interpreted as an o -th root of 1. For instance, if `[T,o] = galoischartable(gal)` the characters χ are input as the columns of T .

Return the degree-1 character $\det \rho$ as the list of $\det \rho(g)$, where g runs through representatives of the conjugacy classes in `galoisconjclasses(gal)`, with the same ordering.

```

? P = x^5 - x^4 - 5*x^3 + 4*x^2 + 3*x - 1;
? polgalois(P)
%2 = [10, 1, 1, "D(5) = 5:2"]
? K = nfsplitting(P);
? gal = galoisinit(K); \\ dihedral of order 10
? [T,o] = galoischartable(gal);
? chi = T[,1]; \\ trivial character
? galoischarDET(gal, chi, o)
%7 = [1, 1, 1, 1]~
? [galoischarDET(gal, T[,i], o) | i <- [1..#T]] \\ all characters
%8 = [[1, 1, 1, 1]~, [1, 1, -1, 1]~, [1, 1, -1, 1]~, [1, 1, -1, 1]~]

```

The library syntax is GEN galoischarDET(GEN gal, GEN chi, long o).

3.13.47 galoischarpoly(gal, chi, {o = 1}). Let G be the group attached to the `galoisinit` structure gal , and let χ be the character of some representation ρ of the group G , where a polynomial variable is to be interpreted as an o -th root of 1, e.g., if `[T,o] = galoischartable(gal)` and χ is a column of T . Return the list of characteristic polynomials $\det(1 - \rho(g)T)$, where g runs through representatives of the conjugacy classes in `galoisconjclasses(gal)`, with the same ordering.

```

? T = x^5 - x^4 - 5*x^3 + 4*x^2 + 3*x - 1;
? polgalois(T)
%2 = [10, 1, 1, "D(5) = 5:2"]
? K = nfsplitting(T);
? gal = galoisinit(K); \\ dihedral of order 10
? [T,o] = galoischartable(gal);
? o
%5 = 5
? galoischarpoly(gal, T[,1], o) \\ T[,1] is the trivial character
%6 = [-x + 1, -x + 1, -x + 1, -x + 1]~
? galoischarpoly(gal, T[,3], o)
%7 = [x^2 - 2*x + 1,
      x^2 + (y^3 + y^2 + 1)*x + 1,

```

```

-x^2 + 1,
x^2 + (-y^3 - y^2)*x + 1]~

```

The library syntax is `GEN galoischarpoly(GEN gal, GEN chi, long o)`.

3.13.48 galoischartable(*gal*). Compute the character table of G , where G is the underlying group of the `galoisinit` structure *gal*. The input *gal* is also allowed to be a `t_VEC` of permutations that is closed under products. Let N be the number of conjugacy classes of G . Return a `t_VEC` $[M, e]$ where $e \geq 1$ is an integer and M is a square `t_MAT` of size N giving the character table of G .

- Each column corresponds to an irreducible character; the characters are ordered by increasing dimension and the first column is the trivial character (hence contains only 1's).

- Each row corresponds to a conjugacy class; the conjugacy classes are ordered as specified by `galoisconjclasses(gal)`, in particular the first row corresponds to the identity and gives the dimension $\chi(1)$ of the irreducible representation attached to the successive characters χ .

The value $M[i, j]$ of the character j at the conjugacy class i is represented by a polynomial in y whose variable should be interpreted as an e -th root of unity, i.e. as the lift of

```
Mod(y, polcyclo(e, 'y))
```

(Note that M is the transpose of the usual orientation for character tables.)

The integer e divides the exponent of the group G and is chosen as small as possible; for instance $e = 1$ when the characters are all defined over $\mathbf{Q}$, as is the case for S_n . Examples:

```

? K = nfsplitting(x^4+x+1);
? gal = galoisinit(K);
? [M,e] = galoischartable(gal);
? M~  \\ take the transpose to get the usual orientation
%4 =
[1  1  1  1  1]
[1 -1 -1  1  1]
[2  0  0 -1  2]
[3 -1  1  0 -1]
[3  1 -1  0 -1]
? e
%5 = 1
? {G = [Vecsmall([1, 2, 3, 4, 5]), Vecsmall([1, 5, 4, 3, 2]),
      Vecsmall([2, 1, 5, 4, 3]), Vecsmall([2, 3, 4, 5, 1]),
      Vecsmall([3, 2, 1, 5, 4]), Vecsmall([3, 4, 5, 1, 2]),
      Vecsmall([4, 3, 2, 1, 5]), Vecsmall([4, 5, 1, 2, 3]),
      Vecsmall([5, 1, 2, 3, 4]), Vecsmall([5, 4, 3, 2, 1])];}
  \\G = D10
? [M,e] = galoischartable(G);
? M~
%8 =
[1  1          1          1]
[1 -1          1          1]
[2  0 -y^3 - y^2 - 1      y^3 + y^2]

```

```

[2 0      y^3 + y^2 -y^3 - y^2 - 1]
? e
%9 = 5

```

The library syntax is `GEN galoischartable(GEN gal)`.

3.13.49 galoisconjclasses(*gal*). *gal* being output by `galoisinit`, return the list of conjugacy classes of the underlying group. The ordering of the classes is consistent with `galoischartable` and the trivial class comes first.

```

? G = galoisinit(x^6+108);
? galoisidentify(G)
%2 = [6, 1]  \\ S_3
? S = galoisconjclasses(G)
%3 = [[Vecsmall([1,2,3,4,5,6])],
      [Vecsmall([3,1,2,6,4,5]),Vecsmall([2,3,1,5,6,4])],
      [Vecsmall([6,5,4,3,2,1]),Vecsmall([5,4,6,2,1,3]),
      Vecsmall([4,6,5,1,3,2])]]
? [[permorder(c[1]),#c] | c <- S ]
%4 = [[1,1], [3,2], [2,3]]

```

This command also accepts subgroups returned by `galoissubgroups`:

```

? subs = galoissubgroups(G); H = subs[5];
? galoisidentify(H)
%2 = [2, 1]  \\ Z/2
? S = galoisconjclasses(subgroups_of_G[5]);
? [[permorder(c[1]),#c] | c <- S ]
%4 = [[1,1], [2,1]]

```

The library syntax is `GEN galoisconjclasses(GEN gal)`.

3.13.50 galoisexport(*gal*, {*flag*}). *gal* being be a Galois group as output by `galoisinit`, export the underlying permutation group as a string suitable for (no flags or *flag* = 0) GAP or (*flag* = 1) Magma. The following example compute the index of the underlying abstract group in the GAP library:

```

? G = galoisinit(x^6+108);
? s = galoisexport(G)
%2 = "Group((1, 2, 3)(4, 5, 6), (1, 4)(2, 6)(3, 5))"
? extern("echo \"IdGroup(\"s\");\" | gap -q")
%3 = [6, 1]
? galoisidentify(G)
%4 = [6, 1]

```

This command also accepts subgroups returned by `galoissubgroups`.

To *import* a GAP permutation into gp (for `galoissubfields` for instance), the following GAP function may be useful:

```

PermToGP := function(p, n)
  return Permuted([1..n],p);
end;

```

```

gap> p:= (1,26)(2,5)(3,17)(4,32)(6,9)(7,11)(8,24)(10,13)(12,15)(14,27)
      (16,22)(18,28)(19,20)(21,29)(23,31)(25,30)
gap> PermToGP(p,32);
[ 26, 5, 17, 32, 2, 9, 11, 24, 6, 13, 7, 15, 10, 27, 12, 22, 3, 28, 20, 19,
  29, 16, 31, 8, 30, 1, 14, 18, 21, 25, 23, 4 ]

```

The library syntax is `GEN galoisexport(GEN gal, long flag)`.

3.13.51 galoisfixedfield(*gal*, *perm*, {*flag*}, {*v* = *y*}). *gal* being be a Galois group as output by `galoisinit` and *perm* an element of *gal.group*, a vector of such elements or a subgroup of *gal* as returned by `galoissubgroups`, computes the fixed field of *gal* by the automorphism defined by the permutations *perm* of the roots *gal.roots*. *P* is guaranteed to be squarefree modulo *gal.p*.

If no flags or *flag* = 0, output format is the same as for `nfsubfield`, returning [*P*, *x*] such that *P* is a polynomial defining the fixed field, and *x* is a root of *P* expressed as a polmod in *gal.pol*.

If *flag* = 1 return only the polynomial *P*.

If *flag* = 2 return [*P*, *x*, *F*] where *P* and *x* are as above and *F* is the factorization of *gal.pol* over the field defined by *P*, where variable *v* (*y* by default) stands for a root of *P*. The priority of *v* must be less than the priority of the variable of *gal.pol* (see Section 2.5.3). In this case, *P* is also expressed in the variable *v* for compatibility with *F*. Example:

```

? G = galoisinit(x^4+1);
? galoisfixedfield(G,G.group[2],2)
%2 = [y^2 - 2, Mod(- x^3 + x, x^4 + 1), [x^2 - y*x + 1, x^2 + y*x + 1]]

```

computes the factorization $x^4 + 1 = (x^2 - \sqrt{2}x + 1)(x^2 + \sqrt{2}x + 1)$

The library syntax is `GEN galoisfixedfield(GEN gal, GEN perm, long flag, long v = -1)` where *v* is a variable number.

3.13.52 galoisgetgroup(*a*, {*b*}). Query the `galpol` package for a group of order *a* with index *b* in the GAP4 Small Group library, by Hans Ulrich Besche, Bettina Eick and Eamonn O'Brien.

The current version of `galpol` supports groups of order $a \leq 143$. If *b* is omitted, return the number of isomorphism classes of groups of order *a*.

The library syntax is `GEN galoisgetgroup(long a, long b)`. Also available is `GEN galoisnbpol(long a)` when *b* is omitted.

3.13.53 galoisgetname(*a*, *b*). Query the `galpol` package for a string describing the group of order *a* with index *b* in the GAP4 Small Group library, by Hans Ulrich Besche, Bettina Eick and Eamonn O'Brien. The strings were generated using the GAP4 function `StructureDescription`. The command below outputs the names of all abstract groups of order 12:

```

? o = 12; N = galoisgetgroup(o); \ # of abstract groups of order 12
? for(i=1, N, print(i, ". ", galoisgetname(o,i)))
1. C3 : C4
2. C12
3. A4
4. D12
5. C6 x C2

```

The current version of `galpol` supports groups of order $a \leq 143$. For $a \geq 16$, it is possible for different groups to have the same name:

```
? o = 20; N = galoisgetgroup(o);
? for(i=1, N, print(i, ". ", galoisgetname(o,i)))
1. C5 : C4
2. C20
3. C5 : C4
4. D20
5. C10 x C2
```

The library syntax is `GEN galoisgetname(long a, long b)`.

3.13.54 `galoisgetpol`($a, \{b\}, \{s\}$). Query the `galpol` package for a polynomial with Galois group isomorphic to `GAP4(a,b)`, totally real if $s = 1$ (default) and totally complex if $s = 2$. The current version of `galpol` supports groups of order $a \leq 143$. The output is a vector `[pol, den]` where

- `pol` is the polynomial of degree a
- `den` is the denominator of `nfgaloisconj(pol)`. Pass it as an optional argument to `galoisinit` or `nfgaloisconj` to speed them up:

```
? [pol,den] = galoisgetpol(64,4,1);
? G = galoisinit(pol);
time = 352ms
? galoisinit(pol, den); \\ passing 'den' speeds up the computation
time = 264ms
? % == %'
%4 = 1 \\ same answer
```

If b and s are omitted, return the number of isomorphism classes of groups of order a .

The library syntax is `GEN galoisgetpol(long a, long b, long s)`. Also available is `GEN galoisnbpol(long a)` when b and s are omitted.

3.13.55 `galoisidentify`(gal). gal being be a Galois group as output by `galoisinit`, output the isomorphism class of the underlying abstract group as a two-components vector `[o,i]`, where o is the group order, and i is the group index in the GAP4 Small Group library, by Hans Ulrich Besche, Bettina Eick and Eamonn O'Brien.

This command also accepts subgroups returned by `galoissubgroups`.

The current implementation is limited to degree less or equal to 127. Some larger “easy” orders are also supported.

The output is similar to the output of the function `IdGroup` in GAP4. Note that GAP4 `IdGroup` handles all groups of order less than 2000 except 1024, so you can use `galoisexport` and GAP4 to identify large Galois groups.

The library syntax is `GEN galoisidentify(GEN gal)`.

3.13.56 galoisinit(*pol*, {*den*}). Computes the Galois group and all necessary information for computing the fixed fields of the Galois extension $K/\mathbf{Q}$ where K is the number field defined by *pol* (monic irreducible polynomial in $\mathbf{Z}[X]$ or a number field as output by **nfinit**). The extension $K/\mathbf{Q}$ must be Galois with Galois group “weakly” super-solvable, see below; returns 0 otherwise. Hence this permits to quickly check whether a polynomial of order strictly less than 48 is Galois or not.

The algorithm used is an improved version of the paper “An efficient algorithm for the computation of Galois automorphisms”, Bill Allombert, Math. Comp, vol. 73, 245, 2001, pp. 359–375.

A group G is said to be “weakly” super-solvable if there exists a normal series

$$\{1\} = H_0 \triangleleft H_1 \triangleleft \cdots \triangleleft H_{n-1} \triangleleft H_n$$

such that each H_i is normal in G and for $i < n$, each quotient group H_{i+1}/H_i is cyclic, and either $H_n = G$ (then G is super-solvable) or G/H_n is isomorphic to either A_4 , S_4 or the group $(3 \times 3) : 4$ (**GAP4(36,9)**).

In practice, almost all small groups are WKSS, the exceptions having order 48(2), 56(1), 60(1), 72(3), 75(1), 80(1), 96(10), 112(1), 120(3) and ≥ 144 .

This function is a prerequisite for most of the **galoisxxx** routines. For instance:

```
P = x^6 + 108;
G = galoisinit(P);
L = galoissubgroups(G);
vector(#L, i, galoisisabelian(L[i],1))
vector(#L, i, galoisidentify(L[i]))
```

The output is an 8-component vector *gal*.

gal[1] contains the polynomial *pol* (*gal.pol*).

gal[2] is a three-components vector $[p, e, q]$ where p is a prime number (*gal.p*) such that *pol* totally split modulo p , e is an integer and $q = p^e$ (*gal.mod*) is the modulus of the roots in *gal.roots*.

gal[3] is a vector L containing the p -adic roots of *pol* as integers implicitly modulo *gal.mod*. (*gal.roots*).

gal[4] is the inverse of the Vandermonde matrix of the p -adic roots of *pol*, multiplied by *gal*[5].

gal[5] is a multiple of the least common denominator of the automorphisms expressed as polynomial in a root of *pol*.

gal[6] is the Galois group G expressed as a vector of permutations of L (*gal.group*).

gal[7] is a generating subset $S = [s_1, \dots, s_g]$ of G expressed as a vector of permutations of L (*gal.gen*).

gal[8] contains the relative orders $[o_1, \dots, o_g]$ of the generators of S (*gal.orders*).

Let H_n be as above, we have the following properties:

- if $G/H_n \simeq A_4$ then $[o_1, \dots, o_g]$ ends by $[2, 2, 3]$.
- if $G/H_n \simeq S_4$ then $[o_1, \dots, o_g]$ ends by $[2, 2, 3, 2]$.
- if $G/H_n \simeq (3 \times 3) : 4$ (**GAP4(36,9)**) then $[o_1, \dots, o_g]$ ends by $[3, 3, 4]$.

- for $1 \leq i \leq g$ the subgroup of G generated by $[s_1, \dots, s_i]$ is normal, with the exception of $i = g - 2$ in the A_4 and $(3 \times 3) : 4$ cases and of $i = g - 3$ in the S_4 case.

- the relative order o_i of s_i is its order in the quotient group $G/\langle s_1, \dots, s_{i-1} \rangle$, with the same exceptions.

- for any $x \in G$ there exists a unique family $[e_1, \dots, e_g]$ such that (no exceptions):

- for $1 \leq i \leq g$ we have $0 \leq e_i < o_i$

- $x = g_1^{e_1} g_2^{e_2} \dots g_n^{e_n}$

If present *den* must be a suitable value for *gal*[5].

The library syntax is `GEN galoisinit(GEN pol, GEN den = NULL)`.

3.13.57 galoisisabelian(*gal*, {*flag* = 0}). *gal* being as output by `galoisinit`, return 0 if *gal* is not an abelian group, and the HNF matrix of *gal* over `gal.gen` if *flag* = 0, 1 if *flag* = 1, and the SNF matrix of *gal* if *flag* = 2.

This command also accepts subgroups returned by `galoissubgroups`.

The library syntax is `GEN galoisisabelian(GEN gal, long flag)`.

3.13.58 galoisisnormal(*gal*, *subgrp*). *gal* being as output by `galoisinit`, and *subgrp* a subgroup of *gal* as output by `galoissubgroups`, return 1 if *subgrp* is a normal subgroup of *gal*, else return 0.

This command also accepts subgroups returned by `galoissubgroups`.

The library syntax is `long galoisisnormal(GEN gal, GEN subgrp)`.

3.13.59 galoispermtopol(*gal*, *perm*). *gal* being a Galois group as output by `galoisinit` and *perm* a element of *gal.group*, return the polynomial defining the Galois automorphism, as output by `nfgaloisconj`, attached to the permutation *perm* of the roots *gal.roots*. *perm* can also be a vector or matrix, in this case, `galoispermtopol` is applied to all components recursively.

Note that

```
G = galoisinit(pol);
galoispermtopol(G, G[6])~
```

is equivalent to `nfgaloisconj(pol)`, if degree of *pol* is greater or equal to 2.

The library syntax is `GEN galoispermtopol(GEN gal, GEN perm)`.

3.13.60 galoissplittinginit(*P*, {*d*}). Compute the Galois group over Q of the splitting field of P , that is the smallest field over which P is totally split. P is assumed to be integral, monic and irreducible; it can also be given by a `nf` structure. If d is given, it must be a multiple of the splitting field degree. The output is compatible with functions expecting a `galoisinit` structure.

The library syntax is `GEN galoissplittinginit(GEN P, GEN d = NULL)`.

3.13.61 galoissubcyclo($N, H, \{flag = 0\}, \{v\}$). Computes the subextension L of $\mathbf{Q}(\zeta_n)$ fixed by the subgroup $H \subset (\mathbf{Z}/n\mathbf{Z})^*$. By the Kronecker-Weber theorem, all abelian number fields can be generated in this way (uniquely if n is taken to be minimal). This function output is somewhat canonical, as it returns the minimal polynomial of a Gaussian period $\text{Tr}_{\mathbf{Q}(\zeta_f)/L}(\zeta_f)$, where f is the smallest integer so that $\zeta_n \in \mathbf{Q}(\zeta_f)$.

The pair (n, H) is deduced from the parameters (N, H) as follows

- N an integer: then $n = N$; H is a generator, i.e. an integer or an integer modulo n ; or a vector of generators.
- N the output of `znstar`(n) or `znstar`($n, 1$). H as in the first case above, or a matrix, taken to be a HNF left divisor of the SNF for $(\mathbf{Z}/n\mathbf{Z})^*$ (`N.cyc`), giving the generators of H in terms of $N.gen$.
- N the output of `bnrinit`(`bnfinit`(y), m) where m is a module. H as in the first case, or a matrix taken to be a HNF left divisor of the SNF for the ray class group modulo m (of type `N.cyc`), giving the generators of H in terms of $N.bid.gen$ ($= N.gen$ if N includes generators).

In this last case, beware that H is understood relatively to N ; in particular, if the infinite place does not divide the module, e.g if m is an integer, then it is not a subgroup of $(\mathbf{Z}/n\mathbf{Z})^*$, but of its quotient by $\{\pm 1\}$.

If $flag = 0$, computes a polynomial (in the variable v) defining the subfield of $\mathbf{Q}(\zeta_n)$ fixed by the subgroup H of $(\mathbf{Z}/n\mathbf{Z})^*$.

If $flag = 1$, computes only the conductor of the abelian extension, as a module.

If $flag = 2$, outputs $[pol, N]$, where pol is the polynomial as output when $flag = 0$ and N the conductor as output when $flag = 1$.

If $flag = 3$; outputs `galoisinit`(`pol`).

The following function can be used to compute all subfields of $\mathbf{Q}(\zeta_n)$ (of exact degree d , if d is set):

```
subcyclo(n, d = -1)=
{ my(bnr,L,IndexBound);
  IndexBound = if (d < 0, n, [d]);
  bnr = bnrinit(bnfinit(y), [n,[1]]);
  L = subgrouplist(bnr, IndexBound, 1);
  vector(#L,i, galoissubcyclo(bnr,L[i]));
}
```

Setting `L = subgrouplist(bnr, IndexBound)` would produce subfields of exact conductor $n\infty$.

The library syntax is `GEN galoissubcyclo(GEN N, GEN H = NULL, long flag, long v = -1)` where v is a variable number.

3.13.62 galoissubfields($G, \{flag = 0\}, \{v\}$). Outputs all the subfields of the Galois group G , as a vector. This works by applying `galoisfixedfield` to all subgroups. The meaning of $flag$ is the same as for `galoisfixedfield`.

The library syntax is `GEN galoissubfields(GEN G, long flag, long v = -1)` where v is a variable number.

3.13.63 galoissubgroups(G). Outputs all the subgroups of the Galois group gal . A subgroup is a vector $[\text{gen}, \text{orders}]$, with the same meaning as for gal.gen and gal.orders . Hence gen is a vector of permutations generating the subgroup, and orders is the relative orders of the generators. The cardinality of a subgroup is the product of the relative orders. Such subgroup can be used instead of a Galois group in the following command: `galoisisabelian`, `galoissubgroups`, `galoisexport` and `galoisidentify`.

To get the subfield fixed by a subgroup sub of gal , use

```
galoisfixedfield(gal,sub[1])
```

The library syntax is GEN `galoissubgroups(GEN G)`.

3.13.64 gcharalgebraic($gc, \{\text{type}\}$). gc being the structure returned by `gcharinit`, returns a `t_MAT` whose columns form a basis of the subgroup of algebraic Grossencharacters in gc (Weil type A0). The last component is interpreted as a power of the norm.

If type is a `t_VEC` of length $gc.r1 + gc.r2$, containing a pair of integers $[p_\tau, q_\tau]$ for each complex embedding τ , returns a `t_VEC` containing a character whose infinity type at τ is

$$z \mapsto z^{-p_\tau} \bar{z}^{-q_\tau}$$

if such a character exists, or empty otherwise. The full set of characters of that infinity type is obtained by multiplying by the group of finite order characters.

```
? bnf = bnfinit(x^4-2*x^3+23*x^2-22*x+6,1);
? gc = gcharinit(bnf,1);
? gc.cyc
% = [6, 0, 0, 0, 0.E-57]
? gcharalgebraic(gc)
% =
[1 0 0 0]
[0 1 0 0]
[0 0 1 0]
[0 0 0 0]
[0 0 -1/2 -1]
? gcharalgebraic(gc,[[1,1],[0,1]])
% = [] \ \ p_\tau + q_\tau must be constant for an algebraic character to exist
? chi = gcharalgebraic(gc,[[1,1],[0,2]])[1]
% = [0, 1, 2, 0, -1]~
? for(i=0,5,print(lfun Euler([gc,chi+[i,0,0,0,0]~],3)));
\ \ all characters with this infinity type: multiply by finite order characters
```

When the torsion subgroup is not cyclic, we can enumerate the characters of a given type with `forvec`.

```
? bnf = bnfinit(x^4+15*x^2+45,1);
? gc = gcharinit(bnf,1);
? gc.cyc
% = [2, 2, 0, 0, 0, 0.E-57]
? [chi] = gcharalgebraic(gc,[[2,0],[2,0]]);
? {forvec(v=vectorv(2,i,[0,gc.cyc[i]-1]),
  print(round(lfunan([gc,chi+concat(v,[0,0,0,0]~)],20)));}
```

```

    });
    [1, 0, 0, 4, -5, 0, 0, 0, -9, 0, 16, 0, 0, 0, 0, 16, 0, 0, 16, -20]
    [1, 0, 0, -4, 5, 0, 0, 0, 9, 0, 16, 0, 0, 0, 0, 16, 0, 0, -16, -20]
    [1, 0, 0, 4, 5, 0, 0, 0, 9, 0, -16, 0, 0, 0, 0, 16, 0, 0, 16, 20]
    [1, 0, 0, -4, -5, 0, 0, 0, -9, 0, -16, 0, 0, 0, 0, 16, 0, 0, -16, 20]

```

Some algebraic Hecke characters are related to CM Abelian varieties. We first show an example with an elliptic curve.

```

? E = ellinit([0, 0, 1, -270, -1708]); \\ elliptic curve with potential CM by  $\mathbf{Q}(\sqrt{-3})$ 
? bnf = bnfinit(x^2+3,1);
? p3 = idealprimedec(bnf,3)[1];
? gc = gcharinit(bnf,Mat([p3,2]));
? gc.cyc
% = [0, 0.E-57]
? [chi] = gcharalgebraic(gc,[[1,0]])
% = [[-1, -1/2]~]
? LE = lfuncreate(E);
? lfunan(LE,20)
% = [1, 0, 0, -2, 0, 0, -1, 0, 0, 0, 0, 0, 5, 0, 0, 4, 0, 0, -7, 0]
? Lchi = lfuncreate([gc,chi]);
? round(lfunan(Lchi,20))
% = [1, 0, 0, -2, 0, 0, -1, 0, 0, 0, 0, 0, 5, 0, 0, 4, 0, 0, -7, 0]

```

Here is an example with a CM Abelian surface.

```

? L = lfungenus2([-2*x^4 - 2*x^3 + 2*x^2 + 3*x - 2, x^3]);
? bnf = bnfinit(a^4 - a^3 + 2*a^2 + 4*a + 3, 1);
? pr = idealprimedec(bnf,13)[1];
? gc = gcharinit(bnf,pr);
? gc.cyc
% = [3, 0, 0, 0, 0.E-57]
? chitors = [1,0,0,0,0]~;
? typ = [[1,0],[1,0]];
? [chi0] = gcharalgebraic(gc,typ);
? igood = oo; nbgood = 0;
? {for(i=0,gc.cyc[1]-1,
    chi = chi0 + i*chitors;
    Lchi = lfuncreate([gc,chi]);
    if(lfunparams(L) == lfunparams(Lchi)
        && exponent(lfunan(L,10) - lfunan(Lchi,10)) < -50,
        igood=i; nbgood++
    );
});
? nbgood
% = 1
? chi = chi0 + igood*chitors;
? Lchi = lfuncreate([gc,chi]);
? lfunan(L,30)
% = [1, 0, -3, 0, 0, 0, 0, 0, 4, 0, 0, 0, 0, 0, 0, -4, 0, 0,
    0, 0, 0, 0, 0, 0, 0, 0, -6, 0, -3, 0]

```

```
? round(lfunan(Lchi,30))
% = [1, 0, -3, 0, 0, 0, 0, 0, 4, 0, 0, 0, 0, 0, 0, -4, 0, 0,
     0, 0, 0, 0, 0, 0, 0, 0, -6, 0, -3, 0]
```

The library syntax is GEN `gcharalgebraic(GEN gc, GEN type = NULL)`.

3.13.65 gcharconductor(*gc, chi*). Returns the conductor of *chi*, as a modulus over *gc.bnf*. This is the minimum modulus *m* such that $U(\mathfrak{m}) \subset \ker(chi)$ indicating the exact ramification of *chi*.

- for a real place *v*, $v \mid \mathfrak{m}$ iff $\chi_v(-1) = -1$.
- for a finite place *p*, the prime power $\mathfrak{p}^e$ divides exactly *m* if $e \geq 0$ is the smallest integer such that $\chi_{\mathfrak{p}}(U_e) = 1$ where $U_0 = \mathbf{Z}_{\mathfrak{p}}^{\times}$ and $U_i = 1 + \mathfrak{p}^i \mathbf{Z}_{\mathfrak{p}}$ for $i > 0$.

```
? bnf = bnfinit(x^2-5,1);
? gc = gcharinit(bnf,[(13*19)^2,[1,1]]);
? gc.cyc
% = [8892, 6, 2, 0, 0.E-57]
? chi = [0,0,1,1]~;
? gcharconductor(gc,chi)
% = [[61009, 7267; 0, 169], [1, 0]]
? gcharconductor(gc,13*chi)
% = [[4693, 559; 0, 13], [1, 0]]
? gcharconductor(gc,13*19*chi)
% = [[247, 65; 0, 13], [1, 0]]
? gcharconductor(gc,13*19*168*chi)
% = [[19, 5; 0, 1], [0, 0]]
```

The library syntax is GEN `gchar_conductor(GEN gc, GEN chi)`.

3.13.66 gcharduallog(*gc, chi*). Returns internal logarithm vector of character *chi* as a *t_VEC* in $\mathbf{R}^n$, so that for all *x*, `gchareval(gc,chi,x,0)` is equal to `gcharduallog(gc,chi) * gcharlog(gc,x)` in $\mathbf{R}/\mathbf{Z}$.

The components are organized as follows:

- the first *ns* components are in $\mathbf{R}$ and describe the character on the class group generators: θ encodes $\mathfrak{p} \mapsto \exp(2i\pi\theta)$,
- the next *nc* components are in $\mathbf{R}$ and describe the `idealstar` group character via its image on generators: θ encodes the image $\exp(2i\pi\theta)$,
- the next $r_1 + r_2$ components are in $\mathbf{R}$ and correspond to characters of $\mathbf{R}$ for each infinite place: φ encodes $x \mapsto |x|^{i\varphi}$ in the real case and $z \mapsto |z|^{2i\varphi}$ in the complex case,
- the last r_2 components are in $\mathbf{Z}$ and correspond to characters of $\mathbf{R}/\mathbf{Z}$ for each complex place: k encodes $z \mapsto (z/|z|)^k$.
- the last component *s* is in $\mathbf{C}$ and corresponds to a power $\|\cdot\|^s$ of the adélic norm.

See also `gcharlog`.

```
? bnf = bnfinit(x^3+4*x-1,1);
? gc = gcharinit(bnf,[1,[1]]);
? gc.cyc
% = [2, 0, 0, 0.E-57]
```

```
? chi = [0,1,0]~;
? f = gcharduallog(gc,chi)
% = [0.153497221319231, 1/2, 0.776369647248353, -0.388184823624176, 1, 0]
? pr = idealprimedec(bnf,2)[1];
? v = gcharlog(gc,pr);
? exp(2*I*Pi*f*v)
% = -0.569867696226731232993110144 - 0.821736459454756074068598760*I
? gchareval(gc,chi,pr)
% = -0.569867696226731232993110144 - 0.821736459454756074068598760*I
```

The library syntax is GEN gcharduallog(GEN gc, GEN chi).

3.13.67 gchareval(*gc*, *chi*, *x*, {*flag* = 1}). *gc* being the structure returned by gcharinit, *chi* a character in *gc*, and *x* an ideal of the base field, returns the value $\chi(x)$. If *flag* = 1 (default), returns a value in $\mathbf{C}^\times$; if *flag* = 0, returns a value in $\mathbf{C}/\mathbf{Z}$, normalized so that the real part is between $-1/2$ and $1/2$.

```
? bnf = bnfinit(x^2-5);
? gc = gcharinit(bnf,1);
? chi = [1]~;
? pr = idealprimedec(bnf,11)[1];
? a = gchareval(gc,chi,pr)
% = -0.3804107379142448929315340886 - 0.9248176417432464199580504588*I
? b = gchareval(gc,chi,pr,0)
% = -0.3121086861831031476247589216
? a == exp(2*Pi*I*b)
%7 = 1
```

The library syntax is GEN gchareval(GEN gc, GEN chi, GEN x, long flag).

3.13.68 gcharidentify(*gc*, *Lv*, *Lchiv*). *gc* being a Grossencharacter group as output by gcharinit, *Lv* being a **t_VEC** of places *v* encoded by a **t_INT** (infinite place) or a prime ideal structure representing a prime not dividing the modulus of *gc* (finite place), and *Lchiv* being a **t_VEC** of local characters χ_v encoded by [*k*, φ] with *k* a **t_INT** and φ a **t_REAL** or **t_COMPLEX** representing $x \mapsto \text{sign}(x)^k |x|^{i\varphi}$ (real place) or $z \mapsto (z/|z|)^k |z|^{2i\varphi}$ (complex place) or by a **t_REAL** or **t_COMPLEX** θ representing $\mathfrak{p} \mapsto \exp(2i\pi\theta)$ (finite place), returns a Grossencharacter ψ belonging to *g* such that $\psi_v \approx \chi_v$ for all *v*. At finite places, in place of a scalar one can provide a **t_VEC** whose last component is θ , as output by gcharlocal. To ensure proper identification, it is recommended to provide all infinite places together with a set of primes that generate the ray class group of modulus *gc.mod*.

```
? bnf = bnfinit(x^2-5,1);
? gc = gcharinit(bnf,1);
? chi = gcharidentify(gc,[2],[[0,13.]]);
? gcharlocal(gc,chi,2)
% = [0, 13.057005210545987626926134713745179631]
? pr = idealprimedec(bnf,11)[1];
? chi = gcharidentify(gc,[pr],[0.3]);
? gchareval(gc,chi,pr,0)
% = 0.30000006229129706787363344444425752636
```

If you know only few digits, it may be a good idea to reduce the current precision to obtain a meaningful result.

```

? bnf = bnfinit(x^2-5,1);
? gc = gcharinit(bnf,1);
? pr = idealprimedec(bnf,11)[1];
? chi = gcharidentify(gc,[pr],[0.184760])
% = [-420226]~ \\ unlikely to be meaningful
? gchareval(gc,chi,pr,0)
% = 0.18475998070331376194260927294721168954
? \p 10
  realprecision = 19 significant digits (10 digits displayed)
? chi = gcharidentify(gc,[pr],[0.184760])
% = [-7]~ \\ probably what we were looking for
? gchareval(gc,chi,pr,0)
% = 0.1847608033
? \p 38
  realprecision = 38 significant digits
? gchareval(gc,chi,pr,0)
% = 0.18476080328172203337331245154966763237

```

The output may be a quasi-character.

```

? bnf = bnfinit(x^2-2,1);
? gc = gcharinit(bnf,1); gc.cyc
% = [0, 0.E-57]
? gcharidentify(gc,[1,2],[[0,3.5+1/3*I],[0,-3.5+1/3*I]])
% = [-1, 1/3]~

```

The library syntax is GEN gchar_identify(GEN gc, GEN Lv, GEN Lchiv, long prec)

3.13.69 `gcharinit(bnf, f)`. *bnf* being a number field output by `bnfinit` (including fundamental units), *f* a modulus, initializes a structure (`gc`) describing the group of Hecke Grossencharacters of modulus *f*. (As in `idealstar`, the finite part of the conductor may be given by a factorization into prime ideals, as produced by `idealfactor`.)

The following member functions are available on the result: `.bnf` is the underlying *bnf*, `.mod` the modulus, `.cyc` its elementary divisors.

The internal representation uses a logarithm map on ideals $\mathcal{L} : I \rightarrow \mathbf{R}^n$, so that a Hecke Grossencharacter χ can be described by a *n* components vector *v* via $\chi : a \in I \mapsto \exp(2i\pi v \cdot \mathcal{L}(a))$.

See `gcharlog` for more details on the map $\mathcal{L}$.

```

? bnf = bnfinit(polcyclo(5),1); \\ initializes number field Q(ζ5)
? pr = idealprimedec(bnf,5)[1]; \\ prime p = (1 - ζ5) above 5
? gc = gcharinit(bnf,idealpow(bnf,pr,2)); \\ characters of modulus dividing p2
? gc.cyc \\ structure as an abelian group
% = [0,0,0,0.E-57]
? chi = [1,1,-1,0]~; \\ a character
? gcharconductor(gc,chi)[1]
% =
[5 4 1 4]
[0 1 0 0]

```

```
[0 0 1 0]
[0 0 0 1]
```

Currently, `gc` is a row vector with 11 components:

`gc[1]` is a matrix whose rows describe a system of generators of the characters as vectors of $\mathbf{R}^n$, under the above description.

`gc[2]` contains the underlying number field `bnf (gc.bnf)`.

`gc[3]` contains the underlying number field `nf (gc.nf)`, possibly stored at higher precision than `bnf`.

`gc[4]` contains data for computing in $(\mathbf{Z}_K/f)^\times$.

`gc[5]` is a vector S of prime ideals which generate the class group.

`gc[6]` contains data to compute discrete logarithms with respect to S in the class group.

`gc[7]` is a vector `[Sunits,m]`, where `Sunits` describes the S -units of `bnf` and m is a relation matrix for internal usage.

`gc[8]` is `[Vecsmall([evalprec,prec,nfprec]), Vecsmall([ntors,nfree,nalg])]` caching precisions and various dimensions.

`gc[9]` is a vector describing `gc` as a $\mathbf{Z}$ -module via its SNF invariants (`gc.cyc`), the last component representing the norm character.

`gc[10]` is a vector `[R,U,Ui]` allowing to convert characters from SNF basis to internal combination of generators.

Specifically, a character `chi` in SNF basis has coordinates `chi*Ui` in internal basis (the rows of `gc[1]`).

`gc[11] = m` is the matrix of $\mathcal{L}(v)$ for all S -units v .

`gc[12] = u` is an integral base change matrix such that `gc[1]` corresponds to $(mu)^{-1}$.

The library syntax is `GEN gcharinit(GEN bnf, GEN f, long prec)`.

3.13.70 gcharisalgebraic(*gc*, *chi*, {&*type*}). *gc* being the structure returned by `gcharinit` and *chi* a character on *gc*, returns 1 if and only if *chi* is an algebraic (Weil type A0) character, so that its infinity type at every complex embedding τ can be written

$$z \mapsto z^{-p_\tau} \bar{z}^{-q_\tau}$$

for some pair of integers (p_τ, q_τ) .

If *type* is given, it is set to the `t_VEC` of exponents $[p_\tau, q_\tau]$.

```
? bnf = bnfinit(x^4+1,1);
? gc = gcharinit(bnf,1);
? gc.cyc
% = [0, 0, 0, 0.E-57]
? chi1 = [0,0,1]~;
? gcharisalgebraic(gc,chi1)
% = 0
? gcharlocal(gc,chi1,1)
```

```
% = [-3, -0.89110698909568455588720672648627467040]
? chi2 = [1,0,0,-3]~;
? gcharisalgebraic(gc,chi2,&typ)
% = 1
? typ
% = [[6, 0], [2, 4]]
? gcharlocal(gc,chi2,1)
% = [-6, 3*I]
```

The library syntax is `int gcharisalgebraic(GEN gc, GEN chi, GEN *type = NULL)`.

3.13.71 `gcharlocal(gc, chi, v, {&BID})`. `gc` being a `gchar` structure initialised by `gcharinit`, returns the local component χ_v , where v is either an integer between 1 and $r_1 + r_2$ encoding an infinite place, or a prime ideal structure encoding a finite place.

- if v is a real place, $\chi_v(x) = \text{sign}(x)^k |x|^{i\varphi}$ is encoded as $[k, \varphi]$;
- if v is a complex place, $\chi_v(z) = (z/|z|)^k |z|^{2i\varphi}$ is encoded as $[k, \varphi]$;
- if $v = \mathfrak{p}$ is a finite place not dividing `gc.mod`, $\chi_v(\pi_v) = \exp(2i\pi\theta)$ is encoded as $[\theta]$;

• if $v = \mathfrak{p}$ is a finite place dividing `gc.mod`, we can define a *bid* structure attached to the multiplicative group $G = (\mathbf{Z}_K/\mathfrak{p}^k)^*$, where $\mathfrak{p}^k$ divides exactly `gc.mod` (see `idealstar`). Then χ_v is encoded as $[c_1, \dots, c_n, \theta]$ where $[c_1, \dots, c_n]$ defines a character on G (see `gchareval`) and $\chi_v(\pi_v) = \exp(2i\pi\theta)$. This *bid* structure only depends on `gc` and v (and not on the character χ); it can be recovered through the optional argument *BID*.

```
? bnf = bnfinit(x^3-x-1);
? gc = gcharinit(bnf,1);
? gc.cyc
% = [0, 0, 0.E-57]
? chi = [0,1,1/3]~;
? pr = idealprimedec(bnf,5)[1];
? gcharlocal(gc,chi,1)
% = [0, -4.8839310048284836274074581373242545693 - 1/3*I]
? gcharlocal(gc,chi,2)
% = [6, 2.4419655024142418137037290686621272847 - 1/3*I]
? gcharlocal(gc,chi,pr)
% = [0.115465135184293124024408915 + 0.0853833331211293579127218326*I]
? bnf = bnfinit(x^2+1,1);
? pr3 = idealprimedec(bnf,3)[1];
? pr5 = idealprimedec(bnf,5)[1];
? gc = gcharinit(bnf,[pr3,2;pr5,3]);
? gc.cyc
% = [600, 3, 0, 0.E-57]
? chi = [1,1,1]~;
? gcharlocal(gc,chi,pr3,&bid)
% = [1, 1, -21/50]
? bid.cyc
% = [24, 3]
? gcharlocal(gc,chi,pr5,&bid)
% = [98, -0.30120819117478336291229946188762973702]
```

```
? bid.cyc
% = [100]
```

The library syntax is `GEN gcharlocal(GEN gc, GEN chi, GEN v, long prec, GEN *BID = NULL)`.

3.13.72 gcharlog(gc, x). Returns the internal (logarithmic) representation of the ideal x suitable for computations in gc , as a `t_COL` in $\mathbf{R}^n$.

Its $n = \mathbf{ns} + \mathbf{nc} + (r_1 + r_2) + r_2 + 1$ components correspond to a logarithm map on the group of fractional ideals $\mathcal{L} : I \rightarrow \mathbf{R}^n$, see `gcharinit`.

More precisely, let $x = (\alpha) \prod \mathfrak{p}_i^{a_i}$ a principalization of x on a set S of primes generating the class group (see `bnfisprincipal`), then the logarithm of x is the `t_COL`

$$\mathcal{L}(x) = \left[(a_i), \log_f(\alpha), \frac{\log |x/\alpha|_\tau}{2\pi}, \frac{\arg(x/\alpha)_\tau}{2\pi}, \frac{\log N(x)}{2\pi} \cdot i \right]$$

where

- the exponent vector (a_i) has $\mathbf{ns}$ components, where $\mathbf{ns} = \#S$ is the number of prime ideals used to generate the class group,
- $\log_f(\alpha)$ is a discrete logarithm of α in the `idealstar` group $(\mathbf{Z}_K/f)^\times$, with $\mathbf{nc}$ components,
- $\log |x/\alpha|_\tau$ has $r_1 + r_2$ components, one for each real embedding and pair of complex embeddings $\tau: K \rightarrow \mathbf{C}$ (and $|z|_\tau = |z|^2$ for complex τ).
- $\arg(x/\alpha)_\tau$ has r_2 components, one for each pair of complex embeddings $\tau: K \rightarrow \mathbf{C}$.
- $N(x)$ is the norm of the ideal x .

```
? bnf = bnfinit(x^3-x^2+5*x+1,1);
? gc = gcharinit(bnf,3);
? gc.cyc
% = [3, 0, 0, 0.E-57]
? chi = [1,1,0,-1]~;
? f = gcharduallog(gc,chi);
? pr = idealprimedec(bnf,5)[1];
? v = gcharlog(gc,pr)
% = [2, -5, -1, 0.0188115475004995312411, -0.0188115475004995312411,
      -0.840176314833856764413, 0.256149999363388073738*I]~
? exp(2*I*Pi*f*v)
% = -4.5285995080704456583673312 + 2.1193835177957097598574507*I
? gchareval(gc,chi,pr)
% = -4.5285995080704456583673312 + 2.1193835177957097598574507*I
```

The library syntax is `GEN gcharlog(GEN gc, GEN x, long prec)`.

3.13.73 gcharnewprec(gc). gc being a Grossencharacter group output by `gcharinit`, recomputes its archimedean components ensuring accurate computations to current precision.

It is advisable to increase the precision before computing several values at large ideals.

The library syntax is `GEN gcharnewprec(GEN gc, long prec)`.

3.13.74 idealadd(nf, x, y). Sum of the two ideals x and y in the number field nf . The result is given in HNF.

```
? K = nfinit(x^2 + 1);
? a = idealadd(K, 2, x + 1)  \\ ideal generated by 2 and 1+I
%2 =
[2 1]

[0 1]
? pr = idealprimedec(K, 5)[1];  \\ a prime ideal above 5
? idealadd(K, a, pr)  \\ coprime, as expected
%4 =
[1 0]

[0 1]
```

This function cannot be used to add arbitrary $\mathbf{Z}$ -modules, since it assumes that its arguments are ideals:

```
? b = Mat([1,0]~);
? idealadd(K, b, b)  \\ only square t_MATs represent ideals
*** idealadd: nonsquare t_MAT in idealtyp.
? c = [2, 0; 2, 0]; idealadd(K, c, c)  \\ nonsense
%6 =
[2 0]

[0 2]
? d = [1, 0; 0, 2]; idealadd(K, d, d)  \\ nonsense
%7 =
[1 0]

[0 1]
```

In the last two examples, we get wrong results since the matrices c and d do not correspond to an ideal: the $\mathbf{Z}$ -span of their columns (as usual interpreted as coordinates with respect to the integer basis $K.zk$) is not an $\mathbf{Z}_K$ -module. To add arbitrary $\mathbf{Z}$ -modules generated by the columns of matrices A and B , use `mathnf(concat(A,B))`.

The library syntax is `GEN idealadd(GEN nf, GEN x, GEN y)`.

3.13.75 idealaddtoone($nf, x, \{y\}$). x and y being two co-prime integral ideals (given in any form), this gives a two-component row vector $[a, b]$ such that $a \in x$, $b \in y$ and $a + b = 1$.

The alternative syntax `idealaddtoone(nf, v)`, is supported, where v is a k -component vector of ideals (given in any form) which sum to $\mathbf{Z}_K$. This outputs a k -component vector e such that $e[i] \in x[i]$ for $1 \leq i \leq k$ and $\sum_{1 \leq i \leq k} e[i] = 1$.

The library syntax is `GEN idealaddtoone0(GEN nf, GEN x, GEN y = NULL)`.

3.13.76 idealappr($nf, x, \{flag\}$). If x is a fractional ideal (given in any form), gives an element α in nf such that for all prime ideals $\mathfrak{p}$ such that the valuation of x at $\mathfrak{p}$ is nonzero, we have $v_{\mathfrak{p}}(\alpha) = v_{\mathfrak{p}}(x)$, and $v_{\mathfrak{p}}(\alpha) \geq 0$ for all other $\mathfrak{p}$.

The argument x may also be given as a prime ideal factorization, as output by **idealfactor**, but allowing zero exponents. This yields an element α such that for all prime ideals $\mathfrak{p}$ occurring in x , $v_{\mathfrak{p}}(\alpha) = v_{\mathfrak{p}}(x)$; for all other prime ideals, $v_{\mathfrak{p}}(\alpha) \geq 0$.

flag is deprecated (ignored), kept for backward compatibility.

The library syntax is **GEN idealappr0**(GEN *nf*, GEN *x*, long *flag*). Use directly **GEN idealappr**(GEN *nf*, GEN *x*) since *flag* is ignored.

3.13.77 idealchinese($nf, x, \{y\}$). x being a prime ideal factorization (i.e. a 2-columns matrix whose first column contains prime ideals and the second column contains integral exponents), y a vector of elements in nf indexed by the ideals in x , computes an element b such that

$v_{\mathfrak{p}}(b - y_{\mathfrak{p}}) \geq v_{\mathfrak{p}}(x)$ for all prime ideals in x and $v_{\mathfrak{p}}(b) \geq 0$ for all other $\mathfrak{p}$.

```
? K = nfinit(t^2-2);
? x = idealfactor(K, 2^2*3)
%2 =
[[2, [0, 1]~, 2, 1, [0, 2; 1, 0]] 4]
[
      [3, [3, 0]~, 1, 2, 1] 1]
? y = [t,1];
? idealchinese(K, x, y)
%4 = [4, -3]~
```

The argument x may also be of the form $[x, s]$ where the first component is as above and s is a vector of signs, with r_1 components s_i in $\{-1, 0, 1\}$: if σ_i denotes the i -th real embedding of the number field, the element b returned satisfies further **sign**($\sigma_i(b)$) = s_i for all i such that $s_i = \pm 1$. In other words, the sign is fixed to s_i at the i -th embedding whenever s_i is nonzero.

```
? idealchinese(K, [x, [1,1]], y)
%5 = [16, -3]~
? idealchinese(K, [x, [-1,-1]], y)
%6 = [-20, -3]~
? idealchinese(K, [x, [1,-1]], y)
%7 = [4, -3]~
```

If y is omitted, return a data structure which can be used in place of x in later calls and allows to solve many chinese remainder problems for a given x more efficiently. In this case, the right hand side y is not allowed to have denominators, unless they are coprime to x .

```
? C = idealchinese(K, [x, [1,1]]);
? idealchinese(K, C, y) \\ as above
%9 = [16, -3]~
? for(i=1,10^4, idealchinese(K,C,y)) \\ ... but faster !
time = 80 ms.
? for(i=1,10^4, idealchinese(K,[x,[1,1]],y))
time = 224 ms.
```

Finally, this structure is itself allowed in place of x , the new s overriding the one already present in the structure. This allows to initialize for different sign conditions more efficiently when the underlying ideal factorization remains the same.

```
? D = idealchinese(K, [C, [1,-1]]);  \\ replaces [1,1]
? idealchinese(K, D, y)
%13 = [4, -3]~
? for(i=1,10^4,idealchinese(K,[C,[1,-1]]))
time = 40 ms.  \\ faster than starting from scratch
? for(i=1,10^4,idealchinese(K,[x,[1,-1]]))
time = 128 ms.
```

The library syntax is GEN idealchinese(GEN nf, GEN x, GEN y = NULL). Also available is GEN idealchineseinit(GEN nf, GEN x) when $y = \text{NULL}$.

3.13.78 idealcoprime(nf, x, y). Given two integral ideals x and y in the number field nf , returns a β in the field, such that $\beta \cdot x$ is an integral ideal coprime to y . In fact, β is also guaranteed to be integral outside primes dividing y .

The library syntax is GEN idealcoprime(GEN nf, GEN x, GEN y).

3.13.79 idealdiv($nf, x, y, \{flag = 0\}$). Quotient $x \cdot y^{-1}$ of the two ideals x and y in the number field nf . The result is given in HNF.

If $flag$ is nonzero, the quotient $x \cdot y^{-1}$ is assumed to be an integral ideal. This can be much faster when the norm of the quotient is small even though the norms of x and y are large. More precisely, the algorithm cheaply removes all maximal ideals above rational primes such that $v_p(Nx) = v_p(Ny)$.

The library syntax is GEN idealdiv0(GEN nf, GEN x, GEN y, long flag). Also available are GEN idealdiv(GEN nf, GEN x, GEN y) ($flag = 0$) and GEN idealdivexact(GEN nf, GEN x, GEN y) ($flag = 1$).

3.13.80 idealdown(nf, x). Let nf be a number field as output by **nfinit**, and x a fractional ideal. This function returns the nonnegative rational generator of $x \cap \mathbf{Q}$. If x is an extended ideal, the extended part is ignored.

```
? nf = nfinit(y^2+1);
? idealdown(nf, -1/2)
%2 = 1/2
? idealdown(nf, (y+1)/3)
%3 = 2/3
? idealdown(nf, [2, 11]~)
%4 = 125
? x = idealprimedec(nf, 2)[1]; idealdown(nf, x)
%5 = 2
? idealdown(nf, [130, 94; 0, 2])
%6 = 130
```

The library syntax is GEN idealdown(GEN nf, GEN x).

3.13.81 idealfactor($nf, x, \{lim\}$). Factors into prime ideal powers the ideal x in the number field nf . The output format is similar to the **factor** function, and the prime ideals are represented in the form output by the **idealprimedec** function. If lim is set, return partial factorization, including only prime ideals above rational primes $< lim$.

```
? nf = nfinit(x^3-2);
? idealfactor(nf, x) \\ a prime ideal above 2
%2 =
[[2, [0, 1, 0]~, 3, 1, ...] 1]
? A = idealhnf(nf, 6*x, 4+2*x+x^2)
%3 =
[6 0 4]
[0 6 2]
[0 0 1]
? idealfactor(nf, A)
%4 =
[[2, [0, 1, 0]~, 3, 1, ...] 2]
[[3, [1, 1, 0]~, 3, 1, ...] 2]
? idealfactor(nf, A, 3) \\ restrict to primes above p < 3
%5 =
[[2, [0, 1, 0]~, 3, 1, ...] 2]
```

The library syntax is GEN **gpidealfactor**(GEN nf , GEN x , GEN $lim = \text{NULL}$). This function should only be used by the **gp** interface. Use directly GEN **idealfactor**(GEN nf , GEN x) or GEN **idealfactor_limit**(GEN nf , GEN x , ulong lim).

3.13.82 idealfactorback($nf, f, \{e\}, \{flag = 0\}$). Gives back the ideal corresponding to a factorization. The integer 1 corresponds to the empty factorization. If e is present, e and f must be vectors of the same length (e being integral), and the corresponding factorization is the product of the $f[i]^{e[i]}$.

If not, and f is vector, it is understood as in the preceding case with e a vector of 1s: we return the product of the $f[i]$. Finally, f can be a regular factorization, as produced by **idealfactor**.

```
? nf = nfinit(y^2+1); idealfactor(nf, 4 + 2*y)
%1 =
[[2, [1, 1]~, 2, 1, [1, 1]~] 2]
[[5, [2, 1]~, 1, 1, [-2, 1]~] 1]
? idealfactorback(nf, %)
%2 =
[10 4]
[0 2]
? f = %1[,1]; e = %1[,2]; idealfactorback(nf, f, e)
%3 =
[10 4]
[0 2]
? % == idealhnf(nf, 4 + 2*y)
```

```
%4 = 1
```

If *flag* is nonzero, perform ideal reductions (`idealred`) along the way. This is most useful if the ideals involved are all *extended* ideals (for instance with trivial principal part), so that the principal parts extracted by `idealred` are not lost. Here is an example:

```
? f = vector(#f, i, [f[i], [;]]); \\ transform to extended ideals
? idealfactorback(nf, f, e, 1)
%6 = [[1, 0; 0, 1], [2, 1; [2, 1]~, 1]]
? nffactorback(nf, %6[2])
%7 = [4, 2]~
```

The extended ideal returned in %6 is the trivial ideal 1, extended with a principal generator given in factored form. We use `nffactorback` to recover it in standard form.

The library syntax is `GEN idealfactorback(GEN nf, GEN f, GEN e = NULL, long flag)`

3.13.83 idealfrobenius(*nf, gal, pr*). Let K be the number field defined by *nf* and assume $K/\mathbb{Q}$ be a Galois extension with Galois group given `gal=galoisinit(nf)`, and that *pr* is an unramified prime ideal $\mathfrak{p}$ in `prid` format. This function returns a permutation of `gal.group` which defines the Frobenius element $\text{Frob}_{\mathfrak{p}}$ attached to $\mathfrak{p}$. If p is the unique prime number in $\mathfrak{p}$, then $\text{Frob}(x) \equiv x^p \pmod{\mathfrak{p}}$ for all $x \in \mathbb{Z}_K$.

```
? nf = nfinit(polcyclo(31));
? gal = galoisinit(nf);
? pr = idealprimedec(nf, 101)[1];
? g = idealfrobenius(nf, gal, pr);
? galoispermtopol(gal, g)
%5 = x^8
```

This is correct since $101 \equiv 8 \pmod{31}$.

The library syntax is `GEN idealfrobenius(GEN nf, GEN gal, GEN pr)`.

3.13.84 idealhnf(*nf, u, {v}*). Gives the Hermite normal form of the ideal $u\mathbb{Z}_K + v\mathbb{Z}_K$, where u and v are elements of the number field K defined by *nf*.

```
? nf = nfinit(y^3 - 2);
? idealhnf(nf, 2, y+1)
%2 =
[1 0 0]
[0 1 0]
[0 0 1]
? idealhnf(nf, y/2, [0,0,1/3]~)
%3 =
[1/3 0 0]
[0 1/6 0]
[0 0 1/6]
```

If v is omitted, returns the HNF of the ideal defined by u : u may be an algebraic number (defining a principal ideal), a maximal ideal (as given by `idealprimedec` or `idealfactor`), or a

matrix whose columns give generators for the ideal. This last format is a little complicated, but useful to reduce general modules to the canonical form once in a while:

- if strictly less than $N = [K : \mathbf{Q}]$ generators are given, u is the $\mathbf{Z}_K$ -module they generate,
- if N or more are given, it is *assumed* that they form a $\mathbf{Z}$ -basis of the ideal, in particular that the matrix has maximal rank N . This acts as `mathnf` since the $\mathbf{Z}_K$ -module structure is (taken for granted hence) not taken into account in this case.

```
? idealhnf(nf, idealprimedec(nf,2)[1])
%4 =
[2 0 0]
[0 1 0]
[0 0 1]
? idealhnf(nf, [1,2;2,3;3,4])
%5 =
[1 0 0]
[0 1 0]
[0 0 1]
```

Finally, when K is quadratic with discriminant D_K , we allow $u = \text{Qfb}(\mathbf{a}, \mathbf{b}, \mathbf{c})$, provided $b^2 - 4ac = D_K$. As usual, this represents the ideal $a\mathbf{Z} + (1/2)(-b + \sqrt{D_K})\mathbf{Z}$.

```
? K = nfinit(x^2 - 60); K.disc
%1 = 60
? idealhnf(K, qfbprimeform(60,2))
%2 =
[2 1]
[0 1]
? idealhnf(K, Qfb(1,2,3))
*** at top-level: idealhnf(K,Qfb(1,2,3
*** ^-----
*** idealhnf: Qfb(1, 2, 3) has discriminant != 60 in idealhnf.
```

The library syntax is `GEN idealhnf0(GEN nf, GEN u, GEN v = NULL)`. Also available is `GEN idealhnf(GEN nf, GEN a)`, where nf is a true nf structure.

3.13.85 idealintersect(nf, A, B). Intersection of the two ideals A and B in the number field nf . The result is given in HNF.

```
? nf = nfinit(x^2+1);
? idealintersect(nf, 2, x+1)
%2 =
[2 0]
[0 2]
```

This function does not apply to general $\mathbf{Z}$ -modules, e.g. orders, since its arguments are replaced by the ideals they generate. The following script intersects $\mathbf{Z}$ -modules A and B given by matrices of compatible dimensions with integer coefficients:

```
ZM_intersect(A,B) =
```

```
{ my(Ker = matkerint(concat(A,B)));
  mathnf( A * Ker[1..#A,] )
}
```

The library syntax is `GEN idealintersect(GEN nf, GEN A, GEN B)`.

3.13.86 idealinv(nf, x). Inverse of the ideal x in the number field nf , given in HNF. If x is an extended ideal, its principal part is suitably updated: i.e. inverting $[I, t]$, yields $[I^{-1}, 1/t]$.

The library syntax is `GEN idealinv(GEN nf, GEN x)`.

3.13.87 idealismaximal(nf, x). Given nf a number field as output by `nfinit` and an ideal x , return 0 if x is not a maximal ideal. Otherwise return a `prid` structure nf attached to the ideal. This function uses `ispseudoprime` and may return a wrong result in case the underlying rational pseudoprime is not an actual prime number: apply `isprime(pr.p)` to guarantee correctness. If x is an extended ideal, the extended part is ignored.

```
? K = nfinit(y^2 + 1);
? idealismaximal(K, 3) \\ 3 is inert
%2 = [3, [3, 0]~, 1, 2, 1]
? idealismaximal(K, 5) \\ 5 is not
%3 = 0
? pr = idealprimedec(K,5)[1] \\ already a prid
%4 = [5, [-2, 1]~, 1, 1, [2, -1; 1, 2]]
? idealismaximal(K, pr) \\ trivial check
%5 = [5, [-2, 1]~, 1, 1, [2, -1; 1, 2]]
? x = idealhnf(K, pr)
%6 =
[5 3]
[0 1]
? idealismaximal(K, x) \\ converts from matrix form to prid
%7 = [5, [-2, 1]~, 1, 1, [2, -1; 1, 2]]
```

This function is noticeably faster than `idealfactor` since it never involves an actual factorization, in particular when $x \cap \mathbf{Z}$ is not a prime number.

The library syntax is `GEN idealismaximal(GEN nf, GEN x)`.

3.13.88 idealispower($nf, A, n, \{&B\}$). Let nf be a number field and $n > 0$ be a positive integer. Return 1 if the fractional ideal $A = B^n$ is an n -th power and 0 otherwise. If the argument B is present, set it to the n -th root of A , in HNF.

```
? K = nfinit(x^3 - 2);
? A = [46875, 30966, 9573; 0, 3, 0; 0, 0, 3];
? idealispower(K, A, 3, &B)
%3 = 1
? B
%4 =
[75 22 41]
[ 0  1  0]
[ 0  0  1]
```

```
? A = [9375, 2841, 198; 0, 3, 0; 0, 0, 3];
? idealispower(K, A, 3)
%5 = 0
```

The library syntax is `long idealispower(GEN nf, GEN A, long n, GEN *B = NULL)`.

3.13.89 ideallist(*nf*, *bound*, {*flag* = 4}). Computes the list of all ideals of norm less or equal to *bound* in the number field *nf*. The result is a row vector with exactly *bound* components. Each component is itself a row vector containing the information about ideals of a given norm, in no specific order. The information is inferred from local data and Chinese remainders and less expensive than computing than a direct global computation.

The binary digits of *flag* mean:

- 1: if the ideals are given by a *bid*, include generators; otherwise don't.
- 2: if this bit is set, *nf* must be a *bnf* with units. Each component is of the form $[bid, U]$, where *bid* is attached to an ideal *f* and *U* is a vector of discrete logarithms of the units in $(\mathbf{Z}_K/f)^*$. More precisely, *U* gives the `ideallogs` with respect to *bid* of $(\zeta, u_1, \dots, u_r)$ where ζ is the torsion unit generator `bnf.tu[2]` and (u_i) are the fundamental units in `bnf.fu`. This structure is technical, meant to be used in conjunction with `bnrclassnolist` or `bnrdisclist`.

- 4: give only the ideal (in HNF), else a *bid*.
- 8: omit ideals which cannot be conductors, i.e. divisible exactly by a prime ideal of norm 2.

```
? nf = nfinit(x^2+1);
? L = ideallist(nf, 100);
? L[1]
%3 = [[1, 0; 0, 1]] \ A single ideal of norm 1
? #L[65]
%4 = 4 \ There are 4 ideals of norm 65 in  $\mathbf{Z}[i]$ 
```

If one wants more information:

```
? L = ideallist(nf, 100, 0);
? l = L[25]; vector(#l, i, l[i].clgp)
%6 = [[20, [20]], [16, [4, 4]], [20, [20]]]
? l[1].mod
%7 = [[25, 18; 0, 1], []]
? l[2].mod
%8 = [[5, 0; 0, 5], []]
? l[3].mod
%9 = [[25, 7; 0, 1], []]
```

where we ask for the structures of the $(\mathbf{Z}[i]/f)^*$ for all three ideals of norm 25. In fact, for all moduli with finite part of norm 25 and trivial Archimedean part, as the last 3 commands show. See `ideallistarch` to treat general moduli.

Finally, one can input a negative **bound**. The function then returns the ideals of norm $|\text{bound}|$, given by their factorization matrix. The only valid value of *flag* is then the default. If needed, one can obtain their HNF using `idealfactorback`, and the corresponding *bid* structures using `idealstar` (which accepts ideals in factored form).

The library syntax is `GEN gideallist(GEN nf, GEN bound, long flag)`. Also available is `GEN ideallist0(GEN nf, long bound, long flag)` for a non-negative bound.

3.13.90 ideallistarch(*nf*, *list*, *arch*). *list* is a vector of vectors of bid's, as output by **ideallist** with flag 0 to 3. Return a vector of vectors with the same number of components as the original *list*. The leaves give information about moduli whose finite part is as in original list, in the same order, and Archimedean part is now *arch* (it was originally trivial). The information contained is of the same kind as was present in the input; see **ideallist**, in particular the meaning of *flag*.

```
? bnf = bnfinit(x^2-2);
? bnf.sign
%2 = [2, 0] \\ two places at infinity
? L = ideallist(bnf, 100, 0);
? l = L[98]; vector(#l, i, l[i].clgp)
%4 = [[42, [42]], [36, [6, 6]], [42, [42]]]
? La = ideallistarch(bnf, L, [1,1]); \\ add them to the modulus
? l = La[98]; vector(#l, i, l[i].clgp)
%6 = [[168, [42, 2, 2]], [144, [6, 6, 2, 2]], [168, [42, 2, 2]]]
```

Of course, the results above are obvious: adding *t* places at infinity will add *t* copies of $\mathbf{Z}/2\mathbf{Z}$ to $(\mathbf{Z}_K/f)^*$. The following application is more typical:

```
? L = ideallist(bnf, 100, 2); \\ units are required now
? La = ideallistarch(bnf, L, [1,1]);
? H = bnrclassnolist(bnf, La);
? H[98];
%4 = [2, 12, 2]
```

The library syntax is GEN **ideallistarch**(GEN *nf*, GEN *list*, GEN *arch*).

3.13.91 ideallog($\{nf\}, x, bid$). *nf* is a number field, *bid* is as output by **idealstar**(*nf*, *D*, ...) and *x* an element of *nf* which must have valuation equal to 0 at all prime ideals in the support of *D* and need not be integral. This function computes the discrete logarithm of *x* on the generators given in *bid.gen*. In other words, if g_i are these generators, of orders d_i respectively, the result is a column vector of integers (x_i) such that $0 \leq x_i < d_i$ and

$$x \equiv \prod_i g_i^{x_i} \pmod{*D}.$$

Note that when the support of *D* contains places at infinity, this congruence implies also sign conditions on the attached real embeddings. See **znlog** for the limitations of the underlying discrete log algorithms.

When *nf* is omitted, take it to be the rational number field. In that case, *x* must be a **t_INT** and *bid* must have been initialized by **znstar**(*N*,1).

The library syntax is GEN **ideallog**(GEN *nf* = NULL, GEN *x*, GEN *bid*). Also available are GEN **Zideallog**(GEN *bid*, GEN *x*) when *nf* is NULL, and GEN **ideallogmod**(GEN *nf*, GEN *x*, GEN *bid*, GEN *mod*) that returns the discrete logarithm of *x* modulo the **t_INT** *mod*; the value *mod* = NULL is treated as 0 (full discrete logarithm), but *nf* = NULL is not implemented with nonzero *mod*.

3.13.92 idealmin(*nf*, *ix*, {*vdir*}). This function is useless and kept for backward compatibility only, use **idealred**. Computes a pseudo-minimum of the ideal *x* in the direction *vdir* in the number field *nf*.

The library syntax is GEN **idealmin**(GEN *nf*, GEN *ix*, GEN *vdir* = NULL).

3.13.93 idealmul($nf, x, y, \{flag = 0\}$). Ideal multiplication of the ideals x and y in the number field nf ; the result is the ideal product in HNF. If either x or y are extended ideals, their principal part is suitably updated: i.e. multiplying $[I, t]$, $[J, u]$ yields $[IJ, tu]$; multiplying I and $[J, u]$ yields $[IJ, u]$.

```
? nf = nfinit(x^2 + 1);
? idealmul(nf, 2, x+1)
%2 =
[4 2]
[0 2]
? idealmul(nf, [2, x], x+1)      \\ extended ideal * ideal
%3 = [[4, 2; 0, 2], x]
? idealmul(nf, [2, x], [x+1, x])  \\ two extended ideals
%4 = [[4, 2; 0, 2], [-1, 0]~]
```

If $flag$ is nonzero, reduce the result using `idealred`.

The library syntax is `GEN idealmul0(GEN nf, GEN x, GEN y, long flag)`.

See also `GEN idealmul(GEN nf, GEN x, GEN y)` ($flag = 0$) and `GEN idealmulred(GEN nf, GEN x, GEN y)` ($flag \neq 0$).

3.13.94 idealnrm(nf, x). Computes the norm of the ideal x in the number field nf .

The library syntax is `GEN idealnrm(GEN nf, GEN x)`.

3.13.95 idealnumden(nf, x). Returns $[A, B]$, where A, B are coprime integer ideals such that $x = A/B$, in the number field nf .

```
? nf = nfinit(x^2+1);
? idealnumden(nf, (x+1)/2)
%2 = [[1, 0; 0, 1], [2, 1; 0, 1]]
```

The library syntax is `GEN idealnumden(GEN nf, GEN x)`.

3.13.96 idealpow($nf, x, k, \{flag = 0\}$). Computes the k -th power of the ideal x in the number field nf ; $k \in \mathbf{Z}$. If x is an extended ideal, its principal part is suitably updated: i.e. raising $[I, t]$ to the k -th power, yields $[I^k, t^k]$.

If $flag$ is nonzero, reduce the result using `idealred`, *throughout the (binary) powering process*; in particular, this is *not* the same as `idealpow`(nf, x, k) followed by reduction.

The library syntax is `GEN idealpow0(GEN nf, GEN x, GEN k, long flag)`.

See also `GEN idealpow(GEN nf, GEN x, GEN k)` and `GEN idealpows(GEN nf, GEN x, long k)` ($flag = 0$). Corresponding to $flag = 1$ is `GEN idealpowred(GEN nf, GEN x, GEN k)`.

3.13.97 idealprimedec($nf, p, \{f = 0\}$). Computes the prime ideal decomposition of the (positive) prime number p in the number field K represented by nf . If a nonprime p is given the result is undefined. If f is present and nonzero, restrict the result to primes of residue degree $\leq f$.

The result is a vector of *prid* structures, each representing one of the prime ideals above p in the number field nf . The representation $\mathbf{pr} = [p, a, e, f, mb]$ of a prime ideal means the following: a is an algebraic integer in the maximal order $\mathbf{Z}_K$ and the prime ideal is equal to $\mathfrak{p} = p\mathbf{Z}_K + a\mathbf{Z}_K$; e is the ramification index; f is the residual index; finally, mb is the multiplication table attached to an algebraic integer b such that $\mathfrak{p}^{-1} = \mathbf{Z}_K + b/p\mathbf{Z}_K$, which is used internally to compute valuations. In other words if p is inert, then mb is the integer 1, and otherwise it is a square $\mathbf{t_MAT}$ whose j -th column is $b \cdot \mathbf{nf.zk}[j]$.

The algebraic number a is guaranteed to have a valuation equal to 1 at the prime ideal (this is automatic if $e > 1$).

The components of $\mathbf{pr}$ should be accessed by member functions: $\mathbf{pr.p}$, $\mathbf{pr.e}$, $\mathbf{pr.f}$, and $\mathbf{pr.gen}$ (returns the vector $[p, a]$):

```
? K = nfinit(x^3-2);
? P = idealprimedec(K, 5);
? #P          \ 2 primes above 5 in Q(2^(1/3))
%3 = 2
? [p1,p2] = P;
? [p1.e, p1.f]    \ the first is unramified of degree 1
%5 = [1, 1]
? [p2.e, p2.f]    \ the second is unramified of degree 2
%6 = [1, 2]
? p1.gen
%7 = [5, [2, 1, 0]~]
? nfbasistoalg(K, %[2]) \ a uniformizer for p1
%8 = Mod(x + 2, x^3 - 2)
? #idealprimedec(K, 5, 1) \ restrict to f = 1
%9 = 1          \ now only p1
```

The library syntax is GEN idealprimedec_limit_f(GEN nf, GEN p, long f).

3.13.98 idealprincipalunits(nf, pr, k). Given a prime ideal in *idealprimedec* format, returns the multiplicative group $(1 + pr)/(1 + pr^k)$ as an abelian group. This function is much faster than *idealstar* when the norm of pr is large, since it avoids (useless) work in the multiplicative group of the residue field.

```
? K = nfinit(y^2+1);
? P = idealprimedec(K,2)[1];
? G = idealprincipalunits(K, P, 20);
? G.cyc
%4 = [512, 256, 4] \ Z/512 x Z/256 x Z/4
? G.gen
%5 = [[-1, -2]~, 1021, [0, -1]~] \ minimal generators of given order
```

The library syntax is GEN idealprincipalunits(GEN nf, GEN pr, long k).

3.13.99 idealramgroups(*nf*, *gal*, *pr*). Let K be the number field defined by *nf* and assume that $K/\mathbf{Q}$ is Galois with Galois group G given by `gal=galoisinit(nf)`. Let *pr* be the prime ideal $\mathfrak{P}$ in `prid` format. This function returns a vector g of subgroups of *gal* as follows:

- $g[1]$ is the decomposition group of $\mathfrak{P}$,
- $g[2]$ is $G_0(\mathfrak{P})$, the inertia group of $\mathfrak{P}$,
- and for $i \geq 2$,
- $g[i]$ is $G_{i-2}(\mathfrak{P})$, the $i-2$ -th ramification group of $\mathfrak{P}$.

The length of g is the number of nontrivial groups in the sequence, thus is 0 if $e = 1$ and $f = 1$, and 1 if $f > 1$ and $e = 1$. The following function computes the cardinality of a subgroup of G , as given by the components of g :

```
card(H) =my(o=H[2]); prod(i=1,#o,o[i]);

? nf=nfinit(x^6+3); gal=galoisinit(nf); pr=idealprimedec(nf,3)[1];
? g = idealramgroups(nf, gal, pr);
? apply(card,g)
%3 = [6, 6, 3, 3, 3] \\ cardinalities of the G_i

? nf=nfinit(x^6+108); gal=galoisinit(nf); pr=idealprimedec(nf,2)[1];
? iso=idealramgroups(nf,gal,pr)[2]
%5 = [[Vecsmall([2, 3, 1, 5, 6, 4])], Vecsmall([3])]
? nfdisc(galoisfixedfield(gal,iso,1))
%6 = -3
```

The field fixed by the inertia group of 2 is not ramified at 2.

The library syntax is `GEN idealramgroups(GEN nf, GEN gal, GEN pr)`.

3.13.100 idealred(*nf*, $I, \{v = 0\}$). LLL reduction of the ideal I in the number field K attached to *nf*, along the direction v . The v parameter is best left omitted, but if it is present, it must be an `nf.r1 + nf.r2`-component vector of *nonnegative* integers. (What counts is the relative magnitude of the entries: if all entries are equal, the effect is the same as if the vector had been omitted.)

This function finds an $a \in K^*$ such that $J = (a)I$ is “small” and integral (see the end for technical details). The result is the Hermite normal form of the “reduced” ideal J .

```
? K = nfinit(y^2+1);
? P = idealprimedec(K,5)[1];
? idealred(K, P)
%3 =
[1 0]
[0 1]
```

More often than not, a principal ideal yields the unit ideal as above. This is a quick and dirty way to check if ideals are principal, but it is not a necessary condition: a nontrivial result does not prove that the ideal is nonprincipal. For guaranteed results, see `bnfisprincipal`, which requires the computation of a full `bnf` structure.

If the input is an extended ideal $[I, s]$, the output is $[J, sa]$; in this way, one keeps track of the principal ideal part:

```
? idealred(K, [P, 1])
%5 = [[1, 0; 0, 1], [2, -1]~]
```

meaning that P is generated by $[2, -1]$. The number field element in the extended part is an algebraic number in any form *or* a factorization matrix (in terms of number field elements, not ideals!). In the latter case, elements stay in factored form, which is a convenient way to avoid coefficient explosion; see also `idealpow`.

Technical note. The routine computes an LLL-reduced basis for the lattice I^{-1} equipped with the quadratic form

$$||x||_v^2 = \sum_{i=1}^{r_1+r_2} 2^{v_i} \varepsilon_i |\sigma_i(x)|^2,$$

where as usual the σ_i are the (real and) complex embeddings and $\varepsilon_i = 1$, resp. 2, for a real, resp. complex place. The element a is simply the first vector in the LLL basis. The only reason you may want to try to change some directions and set some $v_i \neq 0$ is to randomize the elements found for a fixed ideal, which is heuristically useful in index calculus algorithms like `bnfinit` and `bnfisprincipal`.

Even more technical note. In fact, the above is a white lie. We do not use $|| \cdot ||_v$ exactly but a rescaled rounded variant which gets us faster and simpler LLLs. There's no harm since we are not using any theoretical property of a after all, except that it belongs to I^{-1} and that aI is “expected to be small”.

The library syntax is `GEN idealred0(GEN nf, GEN I, GEN v = NULL)`.

3.13.101 idealredmodpower($nf, x, n, \{B = factorlimit\}$). Let nf be a number field, x an ideal in nf and $n > 0$ be a positive integer. Return a number field element b such that $xb^n = v$ is small. If x is integral, then v is also integral.

More precisely, `idealnumden` reduces the problem to x integral. Then, factoring out the prime ideals dividing a rational prime $p \leq B$, we rewrite $x = IJ^n$ where the ideals I and J are both integral and I is B -smooth. Then we return a small element b in J^{-1} .

The bound B avoids a costly complete factorization of x ; as soon as the n -core of x is B -smooth (i.e., as soon as I is n -power free), then J is as large as possible and so is the expected reduction.

```
? T = x^6+108; nf = nfinit(T); a = Mod(x,T);
? setrand(1); u = (2*a^2+a+3)*random(2^1000*x^6)^6;
? sizebyte(u)
%3 = 4864
? b = idealredmodpower(nf,u,2);
? v2 = nfeltmul(nf,u, nfeltpow(nf,b,2))
%5 = [34, 47, 15, 35, 9, 3]~
? b = idealredmodpower(nf,u,6);
? v6 = nfeltmul(nf,u, nfeltpow(nf,b,6))
%7 = [3, 0, 2, 6, -7, 1]~
```

The last element `v6`, obtained by reducing modulo 6-th powers instead of squares, looks smaller than `v2` but its norm is actually a little larger:

```
? idealnorm(nf,v2)
%8 = 81309
```

```
? idealnorm(nf,v6)
%9 = 731781
```

The library syntax is `GEN idealredmodpower(GEN nf, GEN x, ulong n, ulong B)`.

3.13.102 idealstar($\{nf\}, N, \{flag = 1\}, \{cycmod\}$). Outputs a `bid` structure, necessary for computing in the finite abelian group $G = (\mathbf{Z}_K/N)^*$. Here, nf is a number field and N is a *modulus*: either an ideal in any form, or a row vector whose first component is an ideal and whose second component is a row vector of r_1 0 or 1. Ideals can also be given by a factorization into prime ideals, as produced by `idealfactor`.

If the positive integer `cycmod` is present, only compute the group modulo `cycmod`-th powers, which may save a lot of time when some maximal ideals in the modulus have a huge residue field. Whereas you might only be interested in quadratic or cubic residuosity; see also `bnrinit` for applications in class field theory.

This `bid` is used in `ideallog` to compute discrete logarithms. It also contains useful information which can be conveniently retrieved as `bid.mod` (the modulus), `bid.clgp` (G as a finite abelian group), `bid.no` (the cardinality of G), `bid.cyc` (elementary divisors) and `bid.gen` (generators).

If `flag = 1` (default), the result is a `bid` structure without generators: they are well defined but not explicitly computed, which saves time.

If `flag = 2`, as `flag = 1`, but including generators.

If `flag = 0`, only outputs $(\mathbf{Z}_K/N)^*$ as an abelian group, i.e as a 3-component vector $[h, d, g]$: h is the order, d is the vector of SNF cyclic components and g the corresponding generators.

If nf is omitted, we take it to be the rational number fields, N must be an integer and we return the structure of $(\mathbf{Z}/N\mathbf{Z})^*$. In other words `idealstar(, N, flag)` is short for

```
idealstar(nfinit(x), N, flag)
```

but faster. The alternative syntax `znstar(N, flag)` is also available for an analogous effect but, due to an unfortunate historical oversight, the default value of `flag` is different in the two functions (`znstar` does not initialize by default, you probably want `znstar(N,1)`).

The library syntax is `GEN idealstarmod(GEN nf = NULL, GEN N, long flag, GEN cycmod = NULL)`. Instead the above hardcoded numerical flags, one should rather use `GEN Idealstarmod(GEN nf, GEN ideal, long flag, GEN cycmod)` or `GEN Idealstar(GEN nf, GEN ideal, long flag)` (`cycmod` is `NULL`), where `flag` is an or-ed combination of `nf_GEN` (include generators) and `nf_INIT` (return a full `bid`, not a group), possibly 0. This offers one more combination: `gen`, but no `init`. The `nf` argument must be a true nf structure.

3.13.103 idealtwoelt($nf, x, \{a\}$). Computes a two-element representation of the ideal x in the number field nf , combining a random search and an approximation theorem; x is an ideal in any form (possibly an extended ideal, whose principal part is ignored)

- When called as `idealtwoelt(nf,x)`, the result is a row vector $[a, \alpha]$ with two components such that $x = a\mathbf{Z}_K + \alpha\mathbf{Z}_K$ and a is chosen to be the positive generator of $x \cap \mathbf{Z}$, unless x was given as a principal ideal in which case we may choose $a = 0$. The algorithm uses a fast lazy factorization of $x \cap \mathbf{Z}$ and runs in randomized polynomial time.

```
? K = nfinit(t^5-23);
? x = idealhnf(K, t^2*(t+1), t^3*(t+1))
```

```

%2 = \\ some random ideal of norm 552*23
[552 23 23 529 23]
[ 0 23 0 0 0]
[ 0 0 1 0 0]
[ 0 0 0 1 0]
[ 0 0 0 0 1]
? [a,alpha] = idealtwoelt(K, x)
%3 = [552, [23, 0, 1, 0, 0]~]
? nfbasistoalg(K, alpha)
%4 = Mod(t^2 + 23, t^5 - 23)

```

• When called as `idealtwoelt(nf,x,a)` with an explicit nonzero a supplied as third argument, the function assumes that $a \in x$ and returns $\alpha \in x$ such that $x = a\mathbf{Z}_K + \alpha\mathbf{Z}_K$. Note that we must factor a in this case, and the algorithm is generally slower than the default variant and gives larger generators:

```

? alpha2 = idealtwoelt(K, x, 552)
%5 = [-161, -161, -183, -207, 0]~
? idealhnf(K, 552, alpha2) == x
%6 = 1

```

Note that, in both cases, the return value is *not* recognized as an ideal by GP functions; one must use `idealhnf` as above to recover a valid ideal structure from the two-element representation.

The library syntax is `GEN idealtwoelt0(GEN nf, GEN x, GEN a = NULL)`. Also available are `GEN idealtwoelt(GEN nf, GEN x)` and `GEN idealtwoelt2(GEN nf, GEN x, GEN a)`.

3.13.104 idealval(nf, x, pr). Gives the valuation of the ideal x at the prime ideal pr in the number field nf , where pr is in `idealprimedec` format. The valuation of the 0 ideal is `+oo`.

The library syntax is `GEN gpidealval(GEN nf, GEN x, GEN pr)`. Also available is `long idealval(GEN nf, GEN x, GEN pr)`, which returns `LONG_MAX` if $x = 0$ and the valuation as a long integer.

3.13.105 matalgtobasis(nf, x). This function is deprecated, use `apply`.

nf being a number field in `nfinit` format, and x a (row or column) vector or matrix, apply `nfalgtobasis` to each entry of x .

The library syntax is `GEN matalgtobasis(GEN nf, GEN x)`.

3.13.106 matbasistoalg(nf, x). This function is deprecated, use `apply`.

nf being a number field in `nfinit` format, and x a (row or column) vector or matrix, apply `nfbasistoalg` to each entry of x .

The library syntax is `GEN matbasistoalg(GEN nf, GEN x)`.

3.13.107 modreverse(z). Let $z = \text{Mod}(A, T)$ be a polmod, and Q be its minimal polynomial, which must satisfy $\deg(Q) = \deg(T)$. Returns a “reverse polmod” $\text{Mod}(B, Q)$, which is a root of T .

This is quite useful when one changes the generating element in algebraic extensions:

```
? u = Mod(x, x^3 - x - 1); v = u^5;
? w = modreverse(v)
%2 = Mod(x^2 - 4*x + 1, x^3 - 5*x^2 + 4*x - 1)
```

which means that $x^3 - 5x^2 + 4x - 1$ is another defining polynomial for the cubic field

$$\mathbf{Q}(u) = \mathbf{Q}[x]/(x^3 - x - 1) = \mathbf{Q}[x]/(x^3 - 5x^2 + 4x - 1) = \mathbf{Q}(v),$$

and that $u \rightarrow v^2 - 4v + 1$ gives an explicit isomorphism. From this, it is easy to convert elements between the $A(u) \in \mathbf{Q}(u)$ and $B(v) \in \mathbf{Q}(v)$ representations:

```
? A = u^2 + 2*u + 3; subst(lift(A), 'x, w)
%3 = Mod(x^2 - 3*x + 3, x^3 - 5*x^2 + 4*x - 1)
? B = v^2 + v + 1; subst(lift(B), 'x, v)
%4 = Mod(26*x^2 + 31*x + 26, x^3 - x - 1)
```

If the minimal polynomial of z has lower degree than expected, the routine fails

```
? u = Mod(-x^3 + 9*x, x^4 - 10*x^2 + 1)
? modreverse(u)
*** modreverse: domain error in modreverse: deg(minpoly(z)) < 4
*** Break loop: type 'break' to go back to GP prompt
break> Vec( dbg_err() ) \\ ask for more info
["e_DOMAIN", "modreverse", "deg(minpoly(z))", "<", 4,
  Mod(-x^3 + 9*x, x^4 - 10*x^2 + 1)]
break> minpoly(u)
x^2 - 8
```

The library syntax is `GEN modreverse(GEN z)`.

3.13.108 newtonpoly(x, p). Gives the vector of the slopes of the Newton polygon of the polynomial x with respect to the prime number p . The n components of the vector are in decreasing order, where n is equal to the degree of x . Vertical slopes occur iff the constant coefficient of x is zero and are denoted by $+\infty$.

The library syntax is `GEN newtonpoly(GEN x, GEN p)`.

3.13.109 nfalgtobasis(nf, x). Given an algebraic number x in the number field nf , transforms it to a column vector on the integral basis $nf.zk$.

```
? nf = nfinit(y^2 + 4);
? nf.zk
%2 = [1, 1/2*y]
? nfalgtobasis(nf, [1,1]~)
%3 = [1, 1]~
? nfalgtobasis(nf, y)
%4 = [0, 2]~
? nfalgtobasis(nf, Mod(y, y^2+4))
%5 = [0, 2]~
```

This is the inverse function of `nfbasistoalg`.

The library syntax is `GEN algtobasis(GEN nf, GEN x)`.

3.13.110 nfbasis($T, \{&dK\}$). Let $T(X)$ be an irreducible polynomial with integral coefficients. This function returns an integral basis of the number field defined by T , that is a $\mathbf{Z}$ -basis of its maximal order. If present, `dK` is set to the discriminant of the returned order. The basis elements are given as elements in $K = \mathbf{Q}[X]/(T)$, in Hermite normal form with respect to the $\mathbf{Q}$ -basis $(1, X, \dots, X^{\deg T-1})$ of K , lifted to $\mathbf{Q}[X]$. In particular its first element is always 1 and its i -th element is a polynomial of degree $i-1$ whose leading coefficient is the inverse of an integer: the product of those integers is the index of $\mathbf{Z}[X]/(T)$ in the maximal order $\mathbf{Z}_K$:

```
? nfbasis(x^2 + 4) \\ Z[X]/(T) has index 2 in Z_K
%1 = [1, x/2]
? nfbasis(x^2 + 4, &D)
%2 = [1, x/2]
? D
%3 = -4
```

This function uses a modified version of the round 4 algorithm, due to David Ford, Sebastian Pauli and Xavier Roblot.

Local basis, orders maximal at certain primes.

Obtaining the maximal order is hard: it requires factoring the discriminant D of T . Obtaining an order which is maximal at a finite explicit set of primes is easy, but it may then be a strict suborder of the maximal order. To specify that we are interested in a given set of places only, we can replace the argument T by an argument $[T, listP]$, where $listP$ encodes the primes we are interested in: it must be a factorization matrix, a vector of integers or a single integer.

- Vector: we assume that it contains distinct *prime* numbers.
- Matrix: we assume that it is a two-column matrix of a (partial) factorization of D ; namely the first column contains distinct *primes* and the second one the valuation of D at each of these primes.
- Integer B : this is replaced by the vector of primes up to B . Note that the function will use at least $O(B)$ time: a small value, about 10^5 , should be enough for most applications. Values larger than 2^{32} are not supported.

In all these cases, the primes may or may not divide the discriminant D of T . The function then returns a $\mathbf{Z}$ -basis of an order whose index is not divisible by any of these prime numbers.

The result may actually be a global integral basis, in particular if all the prime divisors of the *field* discriminant are included, but this is not guaranteed! Note that `nfinit` has built-in support for such a check:

```
? K = nfinit([T, listP]);
? nfcertify(K)    \\ we computed an actual maximal order
%2 = [];
```

The first line initializes a number field structure incorporating `nfbasis([T, listP])` in place of a proven integral basis. The second line certifies that the resulting structure is correct. This allows to create an `nf` structure attached to the number field $K = \mathbf{Q}[X]/(T)$, when the discriminant of T cannot be factored completely, whereas the prime divisors of $\text{disc}K$ are known. If present, the argument `dK` is set to the discriminant of the returned order, and is equal to the field discriminant if and only if the order is maximal.

Of course, if *listP* contains a single prime number p , the function returns a local integral basis for $\mathbf{Z}_p[X]/(T)$:

```
? nfbasis(x^2+x-1001)
%1 = [1, 1/3*x - 1/3]
? nfbasis([x^2+x-1001, [2]])
%2 = [1, x]
```

The following function computes the index i_T of $\mathbf{Z}[X]/(T)$ in the order generated by the $\mathbf{Z}$ -basis B :

```
nfbasisindex(T, B) = vecprod([denominator(pollead(Q)) | Q <- B]);
```

In particular, B is a basis of the maximal order if and only if $\text{poldisc}(T)/i_T^2$ is equal to the field discriminant. More generally, this formula gives the square of index of the order given by B in $\mathbf{Z}_K$. For instance, assume that P is a vector of prime numbers containing (at least) all prime divisors of the field discriminant, then the following construct allows to provably compute the field discriminant and to check whether the returned basis is actually a basis of the maximal order

```
? B = nfbasis([T, P], &D);
? dK = sign(D) * vecprod([p^valuation(D,p) | p<-P]);
? dK * nbasisindex(T, B)^2 == poldisc(T)
```

The variable `dK` contains the field discriminant and the last command returns 1 if and only if B is a $\mathbf{Z}$ -basis of the maximal order. Of course, the `nfinit` / `nfcertify` approach is simpler, but it is also more costly.

The Buchmann-Lenstra algorithm.

We now complicate the picture: it is in fact allowed to include *composite* numbers instead of primes in `listP` (Vector or Matrix case), provided they are pairwise coprime. The result may still be a correct integral basis if the field discriminant factors completely over the actual primes in the list; again, this is not guaranteed. Adding a composite C such that C^2 divides D may help because when we consider C as a prime and run the algorithm, two good things can happen: either we succeed in proving that no prime dividing C can divide the index (without actually needing to find those primes), or the computation exhibits a nontrivial zero divisor, thereby factoring C and we go on with the refined factorization. (Note that including a C such that C^2 does not divide D is useless.) If neither happen, then the computed basis need not generate the maximal order. Here is an example:

```
? B = 10^5;
? listP = factor(poldisc(T), B); \\ primes <= B dividing D + cofactor
? basis = nfbasis([T, listP], &D)
```

If the computed discriminant D factors completely over the primes less than B (together with the primes contained in the `addprimes` table), then everything is certified: D is the field discriminant and `basis` generates the maximal order. This can be tested as follows:

```
F = factor(D, B); P = F[,1]; E = F[,2];
for (i = 1, #P,
  if (P[i] > B && !isprime(P[i]), warning("nf may be incorrect")));
```

This is a sufficient but not a necessary condition, hence the warning, instead of an error.

The function `nfcertify` speeds up and automates the above process:

```
? B = 10^5;
? nf = nfinit([T, B]);
? nfcertify(nf)
%3 = [] \\ nf is unconditionally correct
? [basis, disc] = [nf.zk, nf.disc];
```

The library syntax is `GEN nfbasis(GEN T, GEN *dK = NULL)`.

3.13.111 `nfbasistoalg`(nf, x). Given an algebraic number x in the number field nf , transforms it into `t_POLMOD` form.

```
? nf = nfinit(y^2 + 4);
? nf.zk
%2 = [1, 1/2*y]
? nfbasistoalg(nf, [1,1]~)
%3 = Mod(1/2*y + 1, y^2 + 4)
? nfbasistoalg(nf, y)
%4 = Mod(y, y^2 + 4)
? nfbasistoalg(nf, Mod(y, y^2+4))
%5 = Mod(y, y^2 + 4)
```

This is the inverse function of `nfalgtobasis`.

The library syntax is `GEN basistoalg(GEN nf, GEN x)`.

3.13.112 nfcertify(*nf*). *nf* being as output by **nfinit**, checks whether the integer basis is known unconditionally. This is in particular useful when the argument to **nfinit** was of the form $[T, \text{listP}]$, specifying a finite list of primes when p -maximality had to be proven, or a list of coprime integers to which Buchmann-Lenstra algorithm was to be applied.

The function returns a vector of coprime composite integers. If this vector is empty, then **nf.zk** and **nf.disc** are correct. Otherwise, the result is dubious. In order to obtain a certified result, one must completely factor each of the given integers, then **addprime** each of their prime factors, then check whether **nfdisc**(**nf.pol**) is equal to **nf.disc**.

The library syntax is **GEN nfcertify**(**GEN nf**).

3.13.113 nfcompositum(*nf*, P , Q , {*flag* = 0}). Let *nf* be a number field structure attached to the field K and let P and Q be squarefree polynomials in $K[X]$ in the same variable. Outputs the simple factors of the étale K -algebra $A = K[X, Y]/(P(X), Q(Y))$. The factors are given by a list of polynomials R in $K[X]$, attached to the number field $K[X]/(R)$, and sorted by increasing degree (with respect to lexicographic ordering for factors of equal degrees). Returns an error if one of the polynomials is not squarefree.

Note that it is more efficient to reduce to the case where P and Q are irreducible first. The routine will not perform this for you, since it may be expensive, and the inputs are irreducible in most applications anyway. In this case, there will be a single factor R if and only if the number fields defined by P and Q are linearly disjoint (their intersection is K).

The binary digits of *flag* mean

1: outputs a vector of 4-component vectors $[R, a, b, k]$, where R ranges through the list of all possible compositums as above, and a (resp. b) expresses the root of P (resp. Q) as an element of $K[X]/(R)$. Finally, k is a small integer such that $b + ka = X$ modulo R .

2: assume that P and Q define number fields that are linearly disjoint: both polynomials are irreducible and the corresponding number fields have no common subfield besides K . This allows to save a costly factorization over K . In this case return the single simple factor instead of a vector with one element.

A compositum is often defined by a complicated polynomial, which it is advisable to reduce before further work. Here is an example involving the field $K(\zeta_5, 5^{1/10})$, $K = \mathbf{Q}(\sqrt{5})$:

```
? K = nfinit(y^2-5);
? L = nfcompositum(K, x^5 - y, polcyclo(5), 1); \\ list of [R, a, b, k]
? [R, a] = L[1]; \\ pick the single factor, extract R, a (ignore b, k)
? lift(R) \\ defines the compositum
%4 = x^10 + (-5/2*y + 5/2)*x^9 + (-5*y + 20)*x^8 + (-20*y + 30)*x^7 + \
(-45/2*y + 145/2)*x^6 + (-71/2*y + 121/2)*x^5 + (-20*y + 60)*x^4 + \
(-25*y + 5)*x^3 + 45*x^2 + (-5*y + 15)*x + (-2*y + 6)
? a^5 - y \\ a fifth root of y
%5 = 0
? [T, X] = rnfpolredbest(K, R, 1);
? lift(T) \\ simpler defining polynomial for K[x]/(R)
%7 = x^10 + (-11/2*y + 25/2)
? liftall(X) \\ root of R in K[x]/(T(x))
%8 = (3/4*y + 7/4)*x^7 + (-1/2*y - 1)*x^5 + 1/2*x^2 + (1/4*y - 1/4)
? a = subst(a.pol, 'x, X); \\ a in the new coordinates
```

```
? liftall(a)
%10 = (-3/4*y - 7/4)*x^7 - 1/2*x^2
? a^5 - y
%11 = 0
```

The main variables of P and Q must be the same and have higher priority than that of nf (see `varhigher` and `varlower`).

The library syntax is `GEN nfcompositum(GEN nf, GEN P, GEN Q, long flag)`.

3.13.114 `nfdetint( $nf, x$ )`. Given a pseudo-matrix x , computes a nonzero ideal contained in (i.e. multiple of) the determinant of x . This is particularly useful in conjunction with `nfhnfmod`.

The library syntax is `GEN nfdetint(GEN nf, GEN x)`.

3.13.115 `nfdisc( $T$ )`. field discriminant of the number field defined by the integral, preferably monic, irreducible polynomial $T(X)$. Returns the discriminant of the number field $\mathbf{Q}[X]/(T)$, using the Round 4 algorithm.

Local discriminants, valuations at certain primes.

As in `nfbasis`, the argument T can be replaced by $[T, \text{listP}]$, where `listP` is as in `nfbasis`: a vector of pairwise coprime integers (usually distinct primes), a factorization matrix, or a single integer. In that case, the function returns the discriminant of an order whose basis is given by `nfbasis(T, listP)`, which need not be the maximal order, and whose valuation at a prime entry in `listP` is the same as the valuation of the field discriminant.

In particular, if `listP` is $[p]$ for a prime p , we can return the p -adic discriminant of the maximal order of $\mathbf{Z}_p[X]/(T)$, as a power of p , as follows:

```
? padicdisc(T,p) = p^valuation(nfdisc([T,[p]]), p);
? nfdisc(x^2 + 6)
%2 = -24
? padicdisc(x^2 + 6, 2)
%3 = 8
? padicdisc(x^2 + 6, 3)
%4 = 3
```

The following function computes the discriminant of the maximal order under the assumption that P is a vector of prime numbers containing (at least) all prime divisors of the field discriminant:

```
globaldisc(T, P) =
{ my (D = nfdisc([T, P]));
  sign(D) * vecprod([p^valuation(D,p) | p <-P]);
}
? globaldisc(x^2 + 6, [2, 3, 5])
%1 = -24
```

The library syntax is `nfdisc(GEN T)`. Also available is `GEN nfbasis(GEN T, GEN *d)`, which returns the order basis, and where `*d` receives the order discriminant.

3.13.116 nfdiscfactors(T). Given a polynomial T with integer coefficients, return $[D, faD]$ where D is `nfdisc( $T$ )` and faD is the factorization of $|D|$. All the variants `[ $T$ , listP]` are allowed (see `??nfdisc`), in which case faD is the factorization of the discriminant underlying order (which need not be maximal at the primes not specified by `listP`) and the factorization may contain large composites.

```
? T = x^3 - 6021021*x^2 + 12072210077769*x - 8092423140177664432;
? [D,faD] = nfdiscfactors(T); print(faD); D
[3, 3; 500009, 2]
%2 = -6750243002187]

? T = x^3 + 9*x^2 + 27*x - 125014250689643346789780229390526092263790263725;
? [D,faD] = nfdiscfactors(T); print(faD); D
[3, 3; 1000003, 2]
%4 = -27000162000243

? [D,faD] = nfdiscfactors([T, 10^3]); print(faD)
[3, 3; 125007125141751093502187, 2]
```

In the final example, we only get a partial factorization, which is only guaranteed correct at primes $\leq 10^3$.

The function also accept number field structures, for instance as output by `nfinit`, and returns the field discriminant and its factorization:

```
? T = x^3 + 9*x^2 + 27*x - 125014250689643346789780229390526092263790263725;
? nf = nfinit(T); [D,faD] = nfdiscfactors(T); print(faD); D
%2 = -27000162000243
? nf.disc
%3 = -27000162000243
```

The library syntax is `GEN nfdiscfactors(GEN T)`.

3.13.117 nfeltadd(nf, x, y). Given two elements x and y in nf , computes their sum $x + y$ in the number field nf .

```
? nf = nfinit(1+x^2);
? nfeltadd(nf, 1, x) \\ 1 + I
%2 = [1, 1]~
```

The library syntax is `GEN nfadd(GEN nf, GEN x, GEN y)`.

3.13.118 nfeltdiv(nf, x, y). Given two elements x and y in nf , computes their quotient x/y in the number field nf .

The library syntax is `GEN nfdiv(GEN nf, GEN x, GEN y)`.

3.13.119 nfeltdivuuc(nf, x, y). Given two elements x and y in nf , computes an algebraic integer q in the number field nf such that the components of $x - qy$ are reasonably small. In fact, this is functionally identical to `round(nfdiv( $nf, x, y$ ))`.

The library syntax is `GEN nfdivuuc(GEN nf, GEN x, GEN y)`.

3.13.120 nfeltdivmodpr(nf, x, y, pr). This function is obsolete, use **nfmodpr**.

Given two elements x and y in nf and pr a prime ideal in **modpr** format (see **nfmodprinit**), computes their quotient x/y modulo the prime ideal pr .

The library syntax is **GEN nfdivmodpr**(**GEN nf**, **GEN x**, **GEN y**, **GEN pr**). This function is normally useless in library mode. Project your inputs to the residue field using **nf_to_Fq**, then work there.

3.13.121 nfeltdivrem(nf, x, y). Given two elements x and y in nf , gives a two-element row vector $[q, r]$ such that $x = qy + r$, q is an algebraic integer in nf , and the components of r are reasonably small.

The library syntax is **GEN nfdivrem**(**GEN nf**, **GEN x**, **GEN y**).

3.13.122 nfeltembed($nf, x, \{pl\}$). Given an element x in the number field nf , return the (real or) complex embeddings of x specified by optional argument pl , at the current **realprecision**:

- pl omitted: return the vector of embeddings at all $r_1 + r_2$ places;
- pl an integer between 1 and $r_1 + r_2$: return the i -th embedding of x , attached to the i -th root of **nf.pol**, i.e. **nf.roots**[i];
- pl a vector or **t_VECSMALL**: return the vector of embeddings; the i -th entry gives the embedding at the place attached to the $pl[i]$ -th real root of **nf.pol**.

```
? nf = nfinit('y^3 - 2);
? nf.sign
%2 = [1, 1]
? nfeltembed(nf, 'y)
%3 = [1.25992[...], -0.62996[...], 1.09112[...]*I]
? nfeltembed(nf, 'y, 1)
%4 = 1.25992[...]
? nfeltembed(nf, 'y, 3) \\ there are only 2 arch. places
*** at top-level: nfeltembed(nf,'y,3)
*** ^-----
*** nfeltembed: domain error in nfeltembed: index > 2
```

The library syntax is **GEN nfeltembed**(**GEN nf**, **GEN x**, **GEN pl** = **NULL**, **long prec**).

3.13.123 nfeltispower($nf, x, n, \{&y\}$). Returns 1 if x is an n -th power in the number field **nf** (and sets y to an n -th root if the argument is present), else returns 0.

```
? nf = nfinit(1+x^2);
? nfeltispower(nf, -4, 4, &y)
%2 = 1
? y
%3 = [-1, -1]~
```

The library syntax is **long nfispower**(**GEN nf**, **GEN x**, **long n**, **GEN *y** = **NULL**).

3.13.124 nfeltissquare($nf, x, \{&y\}$). Returns 1 if x is a square in nf (and sets y to a square root if the argument is present), else returns 0.

```
? nf = nfinit(1+x^2);
? nfeltissquare(nf, -1, &y)
%2 = 1
? y
%3 = [0, -1]~
```

The library syntax is `long nfissquare(GEN nf, GEN x, GEN *y = NULL)`.

3.13.125 nfeltmod(nf, x, y). Given two elements x and y in nf , computes an element r of nf of the form $r = x - qy$ with q algebraic integer, and such that r is small. This is functionally identical to

$$x - \text{nfmul}(nf, \text{round}(\text{nfdiv}(nf, x, y)), y).$$

The library syntax is `GEN nfmod(GEN nf, GEN x, GEN y)`.

3.13.126 nfeltmul(nf, x, y). Given two elements x and y in nf , computes their product $x * y$ in the number field nf .

The library syntax is `GEN nfmul(GEN nf, GEN x, GEN y)`.

3.13.127 nfeltmulmodpr(nf, x, y, pr). This function is obsolete, use `nfmodpr`.

Given two elements x and y in nf and pr a prime ideal in `modpr` format (see `nfmodprinit`), computes their product $x * y$ modulo the prime ideal pr .

The library syntax is `GEN nfmulmodpr(GEN nf, GEN x, GEN y, GEN pr)`. This function is normally useless in library mode. Project your inputs to the residue field using `nf_to_Fq`, then work there.

3.13.128 nfeltnorm(nf, x). Returns the absolute norm of x .

The library syntax is `GEN nfnorm(GEN nf, GEN x)`.

3.13.129 nfeltpow(nf, x, k). Given an element x in nf , and a positive or negative integer k , computes x^k in the number field nf .

The library syntax is `GEN nfpow(GEN nf, GEN x, GEN k)`. `GEN nfinv(GEN nf, GEN x)` correspond to $k = -1$, and `GEN nfsqr(GEN nf, GEN x)` to $k = 2$.

3.13.130 nfeltpowmodpr(nf, x, k, pr). This function is obsolete, use `nfmodpr`.

Given an element x in nf , an integer k and a prime ideal pr in `modpr` format (see `nfmodprinit`), computes x^k modulo the prime ideal pr .

The library syntax is `GEN nfpowmodpr(GEN nf, GEN x, GEN k, GEN pr)`. This function is normally useless in library mode. Project your inputs to the residue field using `nf_to_Fq`, then work there.

3.13.131 nfeltreduce(nf, a, id). Given an ideal id in Hermite normal form and an element a of the number field nf , finds an element r in nf such that $a - r$ belongs to the ideal and r is small.

The library syntax is `GEN nfreduce(GEN nf, GEN a, GEN id)`.

3.13.132 nfeltreducemodpr(nf, x, pr). This function is obsolete, use `nfmodpr`.

Given an element x of the number field nf and a prime ideal pr in `modpr` format compute a canonical representative for the class of x modulo pr .

The library syntax is `GEN nfreducemodpr(GEN nf, GEN x, GEN pr)`. This function is normally useless in library mode. Project your inputs to the residue field using `nf_to_Fq`, then work there.

3.13.133 nfeltsign($nf, x, \{pl\}$). Given an element x in the number field nf , returns the signs of the real embeddings of x specified by optional argument pl :

- pl omitted: return the vector of signs at all r_1 real places;
- pl an integer between 1 and r_1 : return the sign of the i -th embedding of x , attached to the i -th real root of `nf.pol`, i.e. `nf.roots[i]`;
- pl a vector or `t_VECSMALL`: return the vector of signs; the i -th entry gives the sign at the real place attached to the $pl[i]$ -th real root of `nf.pol`.

```
? nf = nfinit(polsubcyclo(11,5,'y)); \\ Q(cos(2 pi/11))
? nf.sign
%2 = [5, 0]
? x = Mod('y, nf.pol);
? nfeltsign(nf, x)
%4 = [-1, -1, -1, 1, 1]
? nfeltsign(nf, x, 1)
%5 = -1
? nfeltsign(nf, x, [1..4])
%6 = [-1, -1, -1, 1]
? nfeltsign(nf, x, 6) \\ there are only 5 real embeddings
*** at top-level: nfeltsign(nf,x,6)
*** ^-----
*** nfeltsign: domain error in nfeltsign: index > 5
```

The library syntax is `GEN nfeltsign(GEN nf, GEN x, GEN pl = NULL)`.

3.13.134 nfelttrace(nf, x). Returns the absolute trace of x .

The library syntax is `GEN nftrace(GEN nf, GEN x)`.

3.13.135 nfeltval($nf, x, pr, \{\&y\}$). Given an element x in nf and a prime ideal pr in the format output by `idealprimedec`, computes the valuation v at pr of the element x . The valuation of 0 is `+oo`.

```
? nf = nfinit(x^2 + 1);
? P = idealprimedec(nf, 2)[1];
? nfeltval(nf, x+1, P)
%3 = 1
```

This particular valuation can also be obtained using `idealval(nf, x, pr)`, since x is then converted to a principal ideal.

If the y argument is present, sets $y = x\tau^v$, where τ is a fixed “anti-uniformizer” for pr : its valuation at pr is -1 ; its valuation is 0 at other prime ideals dividing $pr.p$ and nonnegative at all other primes. In other words y is the part of x coprime to pr . If x is an algebraic integer, so is y .

```
? nfeltval(nf, x+1, P, &y); y
%4 = [0, 1]~
```

For instance if $x = \prod_i x_i^{e_i}$ is known to be coprime to pr , where the x_i are algebraic integers and $e_i \in \mathbf{Z}$ then, if $v_i = \text{nfeltval}(nf, x_i, pr, \&y_i)$, we still have $x = \prod_i y_i^{e_i}$, where the y_i are still algebraic integers but now all of them are coprime to pr . They can then be mapped to the residue field of pr more efficiently than if the product had been expanded beforehand: we can reduce mod pr after each ring operation.

The library syntax is `GEN gpnfvalrem(GEN nf, GEN x, GEN pr, GEN *y = NULL)`. Also available are `long nfvalrem(GEN nf, GEN x, GEN pr, GEN *y = NULL)`, which returns `LONG_MAX` if $x = 0$ and the valuation as a long integer, and `long nfval(GEN nf, GEN x, GEN pr)`, which only returns the valuation ($y = \text{NULL}$).

3.13.136 nffactor(nf, T). Factorization of the univariate polynomial (or rational function) T over the number field nf given by `nfinit`; T has coefficients in nf (i.e. either scalar, polmod, polynomial or column vector). The factors are sorted by increasing degree.

The main variable of nf must be of *lower* priority than that of T , see Section 2.5.3. However if the polynomial defining the number field occurs explicitly in the coefficients of T as modulus of a `t_POLMOD` or as a `t_POL` coefficient, its main variable must be *the same* as the main variable of T . For example,

```
? nf = nfinit(y^2 + 1);
? nffactor(nf, x^2 + y); \\ OK
? nffactor(nf, x^2 + Mod(y, y^2+1)); \\ OK
? nffactor(nf, x^2 + Mod(z, z^2+1)); \\ WRONG
```

It is possible to input a defining polynomial for nf instead, but this is in general less efficient since parts of an `nf` structure will then be computed internally. This is useful in two situations: when you do not need the `nf` elsewhere, or when you cannot initialize an `nf` due to integer factorization difficulties when attempting to compute the field discriminant and maximal order. In all cases, the function runs in polynomial time using Belabas's variant of van Hoeij's algorithm, which copes with hundreds of modular factors.

Caveat. `nfinit([T, listP])` allows to compute in polynomial time a conditional nf structure, which sets `nf.zk` to an order which is not guaranteed to be maximal at all primes. Always either use `nfcertify` first (which may not run in polynomial time) or make sure to input `nf.pol` instead of the conditional nf : `nffactor` is able to recover in polynomial time in this case, instead of potentially missing a factor.

The library syntax is `GEN nffactor(GEN nf, GEN T)`.

3.13.137 nffactorback($nf, f, \{e\}$). Gives back the nf element corresponding to a factorization. The integer 1 corresponds to the empty factorization.

If e is present, e and f must be vectors of the same length (e being integral), and the corresponding factorization is the product of the $f[i]^{e[i]}$.

If not, and f is vector, it is understood as in the preceding case with e a vector of 1s: we return the product of the $f[i]$. Finally, f can be a regular factorization matrix.

```
? nf = nfinit(y^2+1);
? nffactorback(nf, [3, y+1, [1,2]~], [1, 2, 3])
```

```
%2 = [12, -66]~
? 3 * (I+1)^2 * (1+2*I)^3
%3 = 12 - 66*I
```

The library syntax is GEN `nffactorback`(GEN `nf`, GEN `f`, GEN `e` = NULL).

3.13.138 nffactormod(*nf*, *Q*, *pr*). This routine is obsolete, use `nfmodpr` and `factormod`.

Factors the univariate polynomial Q modulo the prime ideal pr in the number field nf . The coefficients of Q belong to the number field (scalar, polmod, polynomial, even column vector) and the main variable of nf must be of lower priority than that of Q (see Section 2.5.3). The prime ideal pr is either in `idealprimedec` or (preferred) `modprinit` format. The coefficients of the polynomial factors are lifted to elements of nf :

```
? K = nfinit(y^2+1);
? P = idealprimedec(K, 3)[1];
? nffactormod(K, x^2 + y*x + 18*y+1, P)
%3 =
[x + (2*y + 1) 1]
[x + (2*y + 2) 1]
? P = nfmodprinit(K, P); \\ convert to nfmodprinit format
? nffactormod(K, x^2 + y*x + 18*y+1)
%5 =
[x + (2*y + 1) 1]
[x + (2*y + 2) 1]
```

Same result, of course, here about 10% faster due to the precomputation.

The library syntax is GEN `nffactormod`(GEN `nf`, GEN `Q`, GEN `pr`).

3.13.139 nfgaloisapply(*nf*, *aut*, *x*). Let nf be a number field as output by `nfinit`, and let aut be a Galois automorphism of nf expressed by its image on the field generator (such automorphisms can be found using `nfgaloisconj`). The function computes the action of the automorphism aut on the object x in the number field; x can be a number field element, or an ideal (possibly extended). Because of possible confusion with elements and ideals, other vector or matrix arguments are forbidden.

```
? nf = nfinit(x^2+1);
? L = nfgaloisconj(nf)
%2 = [-x, x]~
? aut = L[1]; /* the nontrivial automorphism */
? nfgaloisapply(nf, aut, x)
%4 = Mod(-x, x^2 + 1)
? P = idealprimedec(nf,5); /* prime ideals above 5 */
? nfgaloisapply(nf, aut, P[2]) == P[1]
%6 = 0 \\ !!!!
? idealval(nf, nfgaloisapply(nf, aut, P[2]), P[1])
%7 = 1
```

The surprising failure of the equality test (%7) is due to the fact that although the corresponding prime ideals are equal, their representations are not. (A prime ideal is specified by a uniformizer,

and there is no guarantee that applying automorphisms yields the same elements as a direct `idealprimedec` call.)

The automorphism can also be given as a column vector, representing the image of `Mod(x, nf.pol)` as an algebraic number. This last representation is more efficient and should be preferred if a given automorphism must be used in many such calls.

```
? nf = nfinit(x^3 - 37*x^2 + 74*x - 37);
? aut = nfgaloisconj(nf)[2]; \\ an automorphism in basistoalg form
%2 = -31/11*x^2 + 1109/11*x - 925/11
? AUT = nfalgtobasis(nf, aut); \\ same in algtobasis form
%3 = [16, -6, 5]~
? v = [1, 2, 3]~; nfgaloisapply(nf, aut, v) == nfgaloisapply(nf, AUT, v)
%4 = 1 \\ same result...
? for (i=1,10^5, nfgaloisapply(nf, aut, v))
time = 463 ms.
? for (i=1,10^5, nfgaloisapply(nf, AUT, v))
time = 343 ms. \\ but the latter is faster
```

The library syntax is `GEN galoisapply(GEN nf, GEN aut, GEN x)`.

3.13.140 nfgaloisconj(*nf*, {*flag* = 0}, {*d*}). *nf* being a number field as output by `nfinit`, computes the conjugates of a root *r* of the nonconstant polynomial $x = nf[1]$ expressed as polynomials in *r*. This also makes sense when the number field is not Galois since some conjugates may lie in the field. *nf* can simply be a polynomial.

If no flags or *flag* = 0, use a combination of flag 4 and 1 and the result is always complete. There is no point whatsoever in using the other flags.

If *flag* = 1, use `nfroots`: a little slow, but guaranteed to work in polynomial time.

If *flag* = 4, use `galoisinit`: very fast, but only applies to (most) Galois fields. If the field is Galois with weakly super-solvable Galois group (see `galoisinit`), return the complete list of automorphisms, else only the identity element. If present, *d* is assumed to be a multiple of the least common denominator of the conjugates expressed as polynomial in a root of *pol*.

This routine can only compute $\mathbf{Q}$ -automorphisms, but it may be used to get *K*-automorphism for any base field *K* as follows:

```
rnfgaloisconj(nfK, R) = \\ K-automorphisms of L = K[X] / (R)
{
  my(polabs, N, al, S, ala, k, vR);
  R *= Mod(1, nfK.pol); \\ convert coeffs to polmod elts of K
  vR = variable(R);
  al = Mod(variable(nfK.pol), nfK.pol);
  [polabs, ala, k] = rnfequation(nfK, R, 1);
  Rt = if(k==0, R, subst(R, vR, vR-al*k));
  N = nfgaloisconj(polabs) % Rt; \\ Q-automorphisms of L
  S = select(s->subst(Rt, vR, Mod(s, Rt)) == 0, N);
  if (k==0, S, apply(s->subst(s, vR, vR+k*al)-k*al, S));
}
K = nfinit(y^2 + 7);
rnfgaloisconj(K, x^4 - y*x^3 - 3*x^2 + y*x + 1) \\ K-automorphisms of L
```

The library syntax is `GEN galoisconj0(GEN nf, long flag, GEN d = NULL, long prec)`. Use directly `GEN galoisconj(GEN nf, GEN d)`, corresponding to $flag = 0$, the others only have historical interest.

3.13.141 nfgrunwaldwang($nf, Lpr, Ld, pl, \{v = 'x\}$). Given nf a number field in nf or bnf format, a `t_VEC` Lpr of primes of nf and a `t_VEC` Ld of positive integers of the same length, a `t_VECSMALL` pl of length r_1 the number of real places of nf , computes a polynomial with coefficients in nf defining a cyclic extension of nf of minimal degree satisfying certain local conditions:

- at the prime $Lpr[i]$, the extension has local degree a multiple of $Ld[i]$;
- at the i -th real place of nf , it is complex if $pl[i] = -1$ (no condition if $pl[i] = 0$).

The extension has degree the LCM of the local degrees. Currently, the degree is restricted to be a prime power for the search, and to be prime for the construction because of the `rnfkummer` restrictions.

When nf is $\mathbf{Q}$, prime integers are accepted instead of `prid` structures. However, their primality is not checked and the behavior is undefined if you provide a composite number.

Warning. If the number field nf does not contain the n -th roots of unity where n is the degree of the extension to be computed, the function triggers the computation of the bnf of $nf(\zeta_n)$, which may be costly.

```
? nf = nfinit(y^2-5);
? pr = idealprimedec(nf,13)[1];
? pol = nfgrunwaldwang(nf, [pr], [2], [0,-1], 'x)
%3 = x^2 + Mod(3/2*y + 13/2, y^2 - 5)
```

The library syntax is `GEN nfgrunwaldwang(GEN nf, GEN Lpr, GEN Ld, GEN pl, long v = -1)` where v is a variable number.

3.13.142 nfhilbert($nf, a, b, \{pr\}$). If pr is omitted, compute the global quadratic Hilbert symbol (a, b) in nf , that is 1 if $x^2 - ay^2 - bz^2$ has a non trivial solution (x, y, z) in nf , and -1 otherwise. Otherwise compute the local symbol modulo the prime ideal pr , as output by `idealprimedec`.

The library syntax is `long nfhilbert0(GEN nf, GEN a, GEN b, GEN pr = NULL)`.

Also available is `long nfhilbert(GEN nf, GEN a, GEN b)` (global quadratic Hilbert symbol), where nf is a true nf structure.

3.13.143 nfhnf($nf, x, \{flag = 0\}$). Given a pseudo-matrix (A, I) , finds a pseudo-basis (B, J) in Hermite normal form of the module it generates. If $flag$ is nonzero, also return the transformation matrix U such that $AU = [0|B]$.

The library syntax is `GEN nfhnf0(GEN nf, GEN x, long flag)`. Also available:

`GEN nfhnf(GEN nf, GEN x)` ($flag = 0$).

`GEN rnfsimplifybasis(GEN bnf, GEN x)` simplifies the pseudo-basis $x = (A, I)$, returning a pseudo-basis (B, J) . The ideals in the list J are integral, primitive and either trivial (equal to the full ring of integer) or nonprincipal.

3.13.144 nfhnfmod(*nf*, *x*, *detx*). Given a pseudo-matrix (A, I) and an ideal *detx* which is contained in (read integral multiple of) the determinant of (A, I) , finds a pseudo-basis in Hermite normal form of the module generated by (A, I) . This avoids coefficient explosion. *detx* can be computed using the function `nfdetint`.

The library syntax is `GEN nfhnfmod(GEN nf, GEN x, GEN detx)`.

3.13.145 nfinit(*pol*, {*flag* = 0}). *pol* being a nonconstant irreducible polynomial in $\mathbf{Q}[X]$, preferably monic and integral, initializes a *number field* (or *nf*) structure attached to the field K defined by *pol*. As such, it's a technical object passed as the first argument to most `nfxxx` functions, but it contains some information which may be directly useful. Access to this information via *member functions* is preferred since the specific data organization given below may change in the future. Currently, `nf` is a row vector with 9 components:

`nf[1]` contains the polynomial *pol* (`nf.pol`).

`nf[2]` contains [*r1*, *r2*] (`nf.sign`, `nf.r1`, `nf.r2`), the number of real and complex places of K .

`nf[3]` contains the discriminant $d(K)$ (`nf.disc`) of K .

`nf[4]` contains the index of `nf[1]` (`nf.index`), i.e. $[\mathbf{Z}_K : \mathbf{Z}[\theta]]$, where θ is any root of `nf[1]`.

`nf[5]` is a vector containing 7 matrices M , G , *roundG*, T , MD , TI , MDI and a vector vP defined as follows:

- M is the $(r1+r2) \times n$ matrix whose columns represent the numerical values of the conjugates of the elements of the integral basis.

- G is an $n \times n$ matrix such that $T2 = {}^tGG$, where $T2$ is the quadratic form $T_2(x) = \sum |\sigma(x)|^2$, σ running over the embeddings of K into $\mathbf{C}$.

- *roundG* is a rescaled copy of G , rounded to nearest integers.

- T is the $n \times n$ matrix whose coefficients are $\text{Tr}(\omega_i \omega_j)$ where the ω_i are the elements of the integral basis. Note also that $\det(T)$ is equal to the discriminant of the field K . Also, when understood as an ideal, the matrix T^{-1} generates the codifferent ideal.

- The columns of MD (`nf.diff`) express a $\mathbf{Z}$ -basis of the different of K on the integral basis.

- TI is equal to the primitive part of T^{-1} , which has integral coefficients.

- MDI is a two-element representation (for faster ideal product) of $d(K)$ times the codifferent ideal (`nf.disc*nf.codiff`, which is an integral ideal). This is used in `idealinv`.

- vP is the list of prime divisors of the field discriminant, i.e, the ramified primes (`nf.p`); `nfdiscfactors(nf)` is the preferred way to access that information.

`nf[6]` is the vector containing the $r1+r2$ roots (`nf.roots`) of `nf[1]` corresponding to the $r1+r2$ embeddings of the number field into $\mathbf{C}$ (the first $r1$ components are real, the next $r2$ have positive imaginary part).

`nf[7]` is a $\mathbf{Z}$ -basis for $d\mathbf{Z}_K$, where $d = [\mathbf{Z}_K : \mathbf{Z}(\theta)]$, expressed on the powers of θ . The multiplication by d ensures that all polynomials have integral coefficients and `nf[7]/d` (`nf.zk`) is an integral basis for $\mathbf{Z}_K$. Its first element is guaranteed to be 1. This basis is LLL-reduced with respect to T_2 (strictly speaking, it is a permutation of such a basis, due to the condition that the first element be 1).

`nf[8]` is the $n \times n$ integral matrix expressing the power basis in terms of the integral basis, and finally

`nf[9]` is the $n \times n^2$ matrix giving the multiplication table of the integral basis.

If a non monic or non integral polynomial is input, `nfinit` will transform it, and return a structure attached to the new (monic integral) polynomial together with the attached change of variables, see `flag = 3`. It is allowed, though not very useful given the existence of `nfnewprec`, to input a `nf` or a `bnf` instead of a polynomial. It is also allowed to input a `rnf`, in which case an `nf` structure attached to the absolute defining polynomial `polabs` is returned (`flagis` then ignored).

```
? nf = nfinit(x^3 - 12); \\ initialize number field Q[X] / (X^3 - 12)
? nf.pol    \\ defining polynomial
%2 = x^3 - 12
? nf.disc   \\ field discriminant
%3 = -972
? nf.index  \\ index of power basis order in maximal order
%4 = 2
? nf.zk     \\ integer basis, lifted to Q[X]
%5 = [1, x, 1/2*x^2]
? nf.sign   \\ signature
%6 = [1, 1]
? factor(abs(nf.disc)) \\ determines ramified primes
%7 =
[2 2]
[3 5]
? idealfactor(nf, 2)
%8 =
[[2, [0, 0, -1]~, 3, 1, [0, 1, 0]~] 3] \\ p_2^3
```

Huge discriminants, helping `nfdisc`.

In case `pol` has a huge discriminant which is difficult to factor, it is hard to compute from scratch the maximal order. The following special input formats are also accepted:

- `[pol, B]` where `pol` is a monic integral polynomial and `B` is the lift of an integer basis, as would be computed by `nfbasis`: a vector of polynomials with first element 1 (implicitly modulo `pol`). This is useful if the maximal order is known in advance.

- `[pol, B, P]` where `pol` and `B` are as above (a monic integral polynomial and the lift of an integer basis), and `P` is the list of ramified primes in the extension.

- `[pol, listP]` where `pol` is a rational polynomial and `listP` specifies a list of primes as in `nfbasis`. Instead of the maximal order, `nfinit` then computes an order which is maximal at these particular primes as well as the primes contained in the private prime table, see `addprimes`. The result has a good chance of being correct when the discriminant `nf.disc` factors completely over this set of primes but this is not guaranteed. The function `nfcertify` automates this:

```
? pol = polcompositum(x^5 - 101, polcyclo(7))[1];
? nf = nfinit([pol, 10^3]);
? nfcertify(nf)
%3 = []
```

A priori, `nf.zk` defines an order which is only known to be maximal at all primes $\leq 10^3$ (no prime $\leq 10^3$ divides `nf.index`). The certification step proves the correctness of the computation. Had it failed, that particular `nf` structure could not have been trusted and may have caused routines using it to fail randomly. One particular function that remains trustworthy in all cases is `idealprimedec` when applied to a prime included in the above list of primes or, more generally, a prime not dividing any entry in `nfcertify` output.

In order to explain the meaning of *flag*, let $P = \text{polredbest}(pol)$, a polynomial defining the same number field obtained using the LLL algorithm on the lattice $(\mathbf{Z}_K, T_2)$, which may be equal to pol but is usually different and simpler. Binary digits of *flag* mean:

- 1: return $[nf, \text{Mod}(a, P)]$, where nf is `nfinit(P)` and $\text{Mod}(a, P) = \text{Mod}(x, pol)$ gives the change of variables. If only this bit is set, the behaviour is useless since we have $P = pol$.

- 2: return `nfinit(P)`.

Both flags are set automatically when pol is not monic or not integral: first a linear change of variables is performed, to get a monic integral polynomial, then `polredbest`.

- 4: do not LLL-reduce `nf.zk`, which saves time in large degrees, you may expect to gain a factor 2 or so in degree $n \geq 100$ or more, at the expense of *possibly* slowing down later uses of the nf structure. Use this flag if you only need basic arithmetic (the `nfelt*`, `nfmodpr*` and `ideal*` functions); or if you expect the natural basis of the maximal order to contain small elements, this will be the case for cyclotomic fields for instance. On the other hand, functions involving LLL reduction of rank n lattices should be avoided since each call will be about as costly as the initial LLL reduction that the flag prevents and may become more costly because of this missing initial reduction. In particular it is silly to use this flag in addition to the first two, although GP will not protest.

```
? T = polcyclo(307);
? K = nfinit(T);
time = 19,390 ms.
? a = idealhnf(K,1-x);
time = 477ms
? idealfactor(K, a)
time = 294ms

? Kno = nfinit(T, 4);
time = 11,256 ms.
? ano = idealhnf(Kno,1-x); \\ no slowdown, even slightly faster
time = 460ms
? idealfactor(Kno, ano)
time = 264ms

? nfinit(T, 2); \\ polredbest is very slow in high degree
time = 4min, 34,870 ms.
? norml2(%.pol) == norml2(T) \\ and gains nothing here
%9 = 1
```

The library syntax is `GEN nfinit0(GEN pol, long flag, long prec)`. Also available are `GEN nfinit(GEN x, long prec)` (*flag* = 0), `GEN nfinitred(GEN x, long prec)` (*flag* = 2), `GEN nfinitred2(GEN x, long prec)` (*flag* = 3). Instead of the above hardcoded numerical flags in `nfinit0`, one should rather use an or-ed combination of

- `nf_RED`: find a simpler defining polynomial,
- `nf_ORIG`: also return the change of variable,
- `nf_NOLLL`: do not LLL-reduce the maximal order $\mathbf{Z}$ -basis.

3.13.146 `nfisideal(nf, x)`. Returns 1 if x is an ideal in the number field nf , 0 otherwise.

The library syntax is `long isideal(GEN nf, GEN x)`.

3.13.147 `nfisincl(f, g, {flag = 0})`. Let f and g define number fields, where f and g are irreducible polynomials in $\mathbf{Q}[X]$ and nf structures as output by `nfinit`. If either f or g is not irreducible, the result is undefined. Tests whether the number field f is conjugate to a subfield of the field g . If not, the output is the integer 0; if it is, the output depends on the value of *flag*:

- *flag* = 0 (default): return a vector of polynomials $[a_1, \dots, a_n]$ with rational coefficients, representing all distinct embeddings: we have $g \mid f \circ a_i$ for all i .

- *flag* = 1: return a single polynomial a representing a single embedding; this can be n times faster than the default when the embeddings have huge coefficients.

- *flag* = 2: return a vector of rational functions $[r_1, \dots, r_n]$ whose denominators are coprime to g and such that $r_i \% g$ is the polynomial a_i from *flag* = 0. This variant is always faster than *flag* = 0 but produces results which are harder to use. If the denominators are hard to invert in $\mathbf{Q}[X]/(g)$, this may be even faster than *flag* = 1.

```
? T = x^6 + 3*x^4 - 6*x^3 + 3*x^2 + 18*x + 10;
? U = x^3 + 3*x^2 + 3*x - 2;
? nfisincl(U, T)
%3 = [24/179*x^5-27/179*x^4+80/179*x^3-234/179*x^2+380/179*x+94/179]
? a = nfisincl(U, T, 1)
%4 = 24/179*x^5-27/179*x^4+80/179*x^3-234/179*x^2+380/179*x+94/179
? subst(U, x, Mod(a,T))
%5 = Mod(0, x^6 + 3*x^4 - 6*x^3 + 3*x^2 + 18*x + 10)
? nfisincl(U, T, 2) \\ a as a t_RFRAC
%6 = [(2*x^3 - 3*x^2 + 2*x + 4)/(3*x^2 - 1)]
? (a - %[1]) % T
%7 = 0
? #nfisincl(x^2+1, T) \\ two embeddings
%8 = 2

\\ same result with nf structures
? L = nfinit(T); K = nfinit(U); v = [a];
? nfisincl(U, L) == v
%10 = 1
? nfisincl(K, T) == v
%11 = 1
? nfisincl(K, L) == v
%12 = 1

\\ comparative bench: an nf is a little faster, esp. for the subfield
? B = 2000;
? for (i=1, B, nfisincl(U,T))
time = 1,364 ms.
```

```

? for (i=1, B, nfisincl(K,T))
time = 988 ms.
? for (i=1, B, nfisincl(U,L))
time = 1,341 ms.
? for (i=1, B, nfisincl(K,L))
time = 880 ms.

```

Using an *nf* structure for the tentative subfield is faster if the structure is already available. On the other hand, the gain in *nfisincl* is usually not sufficient to make it worthwhile to initialize only for that purpose.

```

? for (i=1, B, nfinit(U))
time = 590 ms.

```

A final more complicated example

```

? f = x^8 - 72*x^6 + 1944*x^4 - 30228*x^2 - 62100*x - 34749;
? g = nfsplitting(f); poldegree(g)
%2 = 96
? #nfisincl(f, g)
time = 559 ms.
%3 = 8
? nfisincl(f,g,1);
time = 172 ms.
? v = nfisincl(f,g,2);
time = 199 ms.
? apply(x->poldegree(denominator(x)), v)
%6 = [81, 81, 81, 81, 81, 81, 80, 81]
? v % g;
time = 407 ms.

```

This final example shows that mapping rational functions to $\mathbf{Q}[X]/(g)$ can be more costly than that the rest of the algorithm. Note that *nfsplitting* also admits a *flag* yielding an embedding.

The library syntax is `GEN nfisincl0(GEN f, GEN g, long flag)`. Also available is `GEN nfisisom(GEN a, GEN b) (flag = 0)`.

3.13.148 *nfisisom*(*f*, *g*). As *nfisincl*, but tests for isomorphism. More efficient if *f* or *g* is a number field structure.

```

? f = x^6 + 30*x^5 + 495*x^4 + 1870*x^3 + 16317*x^2 - 22560*x + 59648;
? g = x^6 + 42*x^5 + 999*x^4 + 8966*x^3 + 36117*x^2 + 21768*x + 159332;
? h = x^6 + 30*x^5 + 351*x^4 + 2240*x^3 + 10311*x^2 + 35466*x + 58321;
? #nfisisom(f,g) \\ two isomorphisms
%3 = 2
? nfisisom(f,h) \\ not isomorphic
%4 = 0
\\ comparative bench
? K = nfinit(f); L = nfinit(g); B = 10^3;
? for (i=1, B, nfisisom(f,g))
time = 6,124 ms.
? for (i=1, B, nfisisom(K,g))

```

```

time = 3,356 ms.
? for (i=1, B, nfisisom(f,L))
time = 3,204 ms.
? for (i=1, B, nfisisom(K,L))
time = 3,173 ms.

```

The function is usually very fast when the fields are nonisomorphic, whenever the fields can be distinguished via a simple invariant such as degree, signature or discriminant. It may be slower when the fields share all invariants, but still faster than computing actual isomorphisms:

```

\\ usually very fast when the answer is 'no':
? for (i=1, B, nfisisom(f,h))
time = 32 ms.

\\ but not always
? u = x^6 + 12*x^5 + 6*x^4 - 377*x^3 - 714*x^2 + 5304*x + 15379
? v = x^6 + 12*x^5 + 60*x^4 + 166*x^3 + 708*x^2 + 6600*x + 23353
? nfisisom(u,v)
%13 = 0
? polsturm(u) == polsturm(v)
%14 = 1
? nfdisc(u) == nfdisc(v)
%15 = 1
? for(i=1,B, nfisisom(u,v))
time = 1,821 ms.
? K = nfinit(u); L = nfinit(v);
? for(i=1,B, nfisisom(K,v))
time = 232 ms.

```

The library syntax is GEN nfisisom(GEN f, GEN g).

3.13.149 nfislocalpower(*nf*, *pr*, *a*, *n*). Let *nf* be a *nf* structure attached to a number field *K*, let *a* ∈ *K* and let *pr* be a *prid* structure attached to a maximal ideal *v*. Return 1 if *a* is an *n*-th power in the completed local field K_v , and 0 otherwise.

```

? K = nfinit(y^2+1);
? P = idealprimedec(K,2)[1]; \\ the ramified prime above 2
? nfislocalpower(K,P,-1, 2) \\ -1 is a square
%3 = 1
? nfislocalpower(K,P,-1, 4) \\ ... but not a 4-th power
%4 = 0
? nfislocalpower(K,P,2, 2) \\ 2 is not a square
%5 = 0
? Q = idealprimedec(K,5)[1]; \\ a prime above 5
? nfislocalpower(K,Q, [0, 32]~, 30) \\ 32*I is locally a 30-th power
%7 = 1

```

The library syntax is long nfislocalpower(GEN nf, GEN pr, GEN a, GEN n).

3.13.150 nfkernelmodpr(nf, x, pr). This function is obsolete, use **nfmodpr**.

Kernel of the matrix a in $\mathbf{Z}_K/pr$, where pr is in **modpr** format (see **nfmodprinit**).

The library syntax is **GEN nfkernelmodpr**(**GEN nf**, **GEN x**, **GEN pr**). This function is normally useless in library mode. Project your inputs to the residue field using **nfM_to_FqM**, then work there.

3.13.151 nflist($G, \{N\}, \{s = -1\}, \{F\}$). Finds number fields (up to isomorphism) with Galois group of Galois closure isomorphic to G with s complex places. The number fields are given by polynomials. This function supports the following groups:

- degree 2: $C_2 = 2T1$;
- degree 3: $C_3 = 3T1$ and $S_3 = 3T2$;
- degree 4: $C_4 = 4T1$, $V_4 = 4T2$, $D_4 = 4T3$, $A_4 = 4T4$ and $S_4 = 4T5$;
- degree 5: $C_5 = 5T1$, $D_5 = 5T2$, $F_5 = M_{20} = 5T3$ and $A_5 = 5T4$;
- degree 6: $C_6 = 6T1$, $S_3(6) = D_6(6) = 6T2$, $D_6(12) = 6T3$, $A_4(6) = 6T4$, $S_3 \times C_3 = 6T5$, $A_4(6) \times C_2 = 6T6$, $S_4(6)^+ = 6T7$, $S_4(6)^- = 6T8$, $S_3^2 = 6T9$, $C_3^2 : C_4 = 6T10$, $S_4(6) \times C_2 = 6T11$, $A_5(6) = PSL_2(5) = 6T12$ and $C_3^2 : D_4 = 6T13$;
- degree 7: $C_7 = 7T1$, $D_7 = 7T2$, $M_{21} = 7T3$ and $M_{42} = 7T4$;
- degree 9: $C_9 = 9T1$, $C_3 \times C_3 = 9T2$ and $D_9 = 9T3$;
- degree ℓ with ℓ prime: $C_\ell = \ell T1$ and $D_\ell = \ell T2$.

The groups A_5 and $A_5(6)$ require the optional package **nflistdata**.

In addition, if N is a polynomial, all transitive subgroups of S_n with $n \leq 15$, as well as alternating groups A_n and the full symmetric group S_n for all n (see below for details and explanations).

The groups are coded as $[n, k]$ using the **nTk** format where n is the degree and k is the T -number, the index in the classification of transitive subgroups of S_n .

Alternatively, the groups C_n , D_n , A_n , S_n , V_4 , $F_5 = M_{20}$, M_{21} and M_{42} can be input as character strings exactly as written, lifting subscripts; for instance "S4" or "M21". If the group is not recognized or is unsupported the function raises an exception.

The number fields are computed on the fly (and not from a preexisting table) using a variety of algorithms, with the exception of A_5 and $A_5(6)$ which are obtained by table lookup. The algorithms are recursive and use the following ingredients: build distinguished subfields (or resolvent fields in Galois closures) of smaller degrees, use class field theory to build abelian extensions over a known base, select subfields using Galois theory. Because of our use of class field theory, and ultimately **bnfinit**, all results depend on the GRH in degree $n > 3$.

To avoid wasting time, the output polynomials defining the number fields are usually not the simplest possible, use **polredbest** or **polredabs** to reduce them.

The non-negative integer s specifies the number of complex places, between 0 and $n/2$. Additional supported values are:

- $s = -1$ (default) all signatures;
- $s = -2$ all signatures, given by increasing number of complex places; in degree n , this means a vector with $1 + \text{floor}(n/2)$ components: the i -th entry corresponds to $s = i - 1$.

If the irreducible monic polynomial $F \in \mathbf{Z}[X]$ is specified, gives only number fields having $\mathbf{Q}[X]/(F)$ as a subfield, or in the case of S_3 , D_ℓ , A_4 , S_4 , F_5 , M_{21} and M_{42} , as a resolvent field (see also the function `nfresolvent` for these cases).

The parameter N can be the following:

- a positive integer: finds all fields with absolute discriminant N (recall that the discriminant over $\mathbf{Q}$ is $(-1)^s N$).
- a pair of non-negative real numbers $[a, b]$ specifying a real interval: finds all fields with absolute value of discriminant between a and b . For most Galois groups, this is faster than iterating on individual N .
- omitted (default): a few fields of small discriminant (not always those with smallest absolute discriminant) are output with given G and s ; usually about 10, less if too difficult to find. The parameter F is ignored.
- a polynomial with main variable, say t , of priority lower than x . The program outputs a *regular* polynomial in $\mathbf{Q}(t)[x]$ (in fact in $\mathbf{Z}[x, t]$) with the given Galois group. By Hilbert irreducibility, almost all specializations of t will give suitable polynomials. The parameters s and F are ignored. This is implemented for all transitive subgroups of S_n with $n \leq 15$ as well as for the alternating and symmetric groups A_n and S_n for all n . Polynomials for A_n were inspired by J.-F. Mestre, a few polynomials in degree ≤ 8 come from G. W. Smith, “Some polynomials over $\mathbf{Q}(t)$ and their Galois groups”, *Math. Comp.*, **69** (230), 1999, pp. 775–796 most others in degree ≤ 11 were provided by J. Klüners and G. Malle (see G. Malle and B. H. Matzat, *Inverse Galois Theory*, Springer, 1999) and T. Dokchitser completed the list up to degree 15. But for A_n and S_n , subgroups of S_n for $n > 7$ require the optional `nflistdata` package.

Complexity. : For a positive integer N , the complexity is subexponential in $\log N$ (and involves factoring N). For an interval $[a, b]$, the complexity is roughly as follows, ignoring terms which are subexponential in $\log b$. It is usually linear in the output size.

- C_n : $O(b^{1/\phi(n)})$ for $n = 2, 4, 6, 9$ or any odd prime;
- D_n : $O(b^{2/\phi(n)})$ for $n = 4$ or any odd prime;
- V_4 , A_4 : $O(b^{1/2})$, S_4 : $O(b)$; N.B. The subexponential terms are expensive for A_4 and S_4 .
- M_{20} : $O(b)$.
- $S_4(6)^-$, $S_4(6)^+$, $A_4(6) \times C_2$, $S_3 \times S_3$, $S_4(6) \times C_2$: $O(b)$, $D_6(12)$, $A_4(6)$, $S_3(6)$, $S_3 \times C_3$, C_3^2 : C_4 : $O(b^{1/2})$.
- M_{21} , M_{42} : $O(b)$.
- $C_3 \times C_3$: $O(b^{1/3})$, D_9 : $O(b^{5/12})$.

```
? #nflist("S3", [1, 10^5]) \\ S3 cubic fields
%1 = 21794
? #nflist("S3", [1, 10^5], 0) \\ real S3 cubic fields (0 complex place)
%2 = 4753
? #nflist("S3", [1, 10^5], 1) \\ complex cubic fields (1 complex place)
%3 = 17041
? v = nflist("S3", [1, 10^5], -2); apply(length,v)
%4 = [4753, 17041]
? nflist("S4") \\ a few S4 fields
```

```

%5 = [x^4 + 12*x^2 - 8*x + 16, x^4 - 2*x^2 - 8*x + 25, ...]
? nflist("S4",,0) \\ a few real S4 fields
%6 = [x^4 - 52*x^2 - 56*x + 48, x^4 - 26*x^2 - 8*x + 1, ...]
? nflist("S4",,-2) \\ a few real S4 fields, by signature
%7 = [[x^4 - 52*x^2 - 56*x + 48, ...],
      [x^4 - 8*x - 16, ... ],
      [x^4 + 138*x^2 - 8*x + 4541, ...]]
? nflist("S3",,,x^2+23) \\ a few cubic fields with resolvent Q(sqrt(-23))
%8 = [x^3 + x + 1, x^3 + 2*x + 1, ...]
? nflist("C3", 3969) \\ C3 fields of given discriminant
%9 = [x^3 - 21*x + 28, x^3 - 21*x - 35]
? nflist([3,1], 3969) \\ C3 fields, using nTt label
%10 = [x^3 - 21*x + 28, x^3 - 21*x - 35]
? P = nflist([8,12],t) \\ geometric 8T12 polynomial
%11 = x^8 + (-t^2 - 803)*x^6 + (264*t^2 + 165528)*x^4
      + (-2064*t^2 - 1724976)*x^2 + 4096*t^2
? polgalois(subst(P, t, 11))
%12 = [24, 1, 12, "2A_4(8)=[2]A(4)=SL(2,3)"]
? nflist("S11")
*** at top-level: nflist("S11")
*** ^-----
*** nflist: unsupported group (S11). Use one of
"C1"=[1,1];
"C2"=[2,1];
"C3"=[3,1], "S3"=[3,2];
"C4"=[4,1], "V4"=[4,2], "D4"=[4,3], "A4"=[4,4], "S4"=[4,5];
"C5"=[5,1], "D5"=[5,2], "F5"="M20"=[5,3], "A5"=[5,4];
"C6"=[6,1], "D6"=[6,2], [6,3], ..., [6,13];
"C7"=[7,1], "D7"=[7,2], "M21"=[7,3], "M42"=[7,4];
"C9"=[9,1], [9,2], "D9"=[9,3].
Also supported are "Cp"=[p,1] and "Dp"=[p,2] for any odd prime p.
? nflist("S25", 't)
%13 = x^25 + x*t + 1

```

The library syntax is `GEN nflist(GEN G, GEN N = NULL, long s, GEN F = NULL)`.

3.13.152 nfmodpr(*nf*, *x*, *pr*). Map *x* to a `t_FFELT` in the residue field modulo *pr*. The argument *pr* is either a maximal ideal in `idealprimedec` format or, preferably, a *modpr* structure from `nfmodprinit`. The function `nfmodprlift` allows to lift back to $\mathbf{Z}_K$.

Note that the function applies to number field elements and not to vector / matrices / polynomials of such. Use `apply` to convert recursive structures.

```

? K = nfinit(y^3-250);
? P = idealprimedec(K, 5)[2];
? modP = nfmodprinit(K, P, 't);
? K.zk
%4 = [1, 1/5*y, 1/25*y^2]
? apply(t->nfmodpr(K,t,modP), K.zk)
%5 = [1, t, 2*t + 1]

```

```
? %[1].mod
%6 = t^2 + 3*t + 4
? K.index
%7 = 125
```

For clarity, we represent elements in the residue field $\mathbf{F}_5[t]/(T)$ as polynomials in the variable t . Whenever the underlying rational prime does not divide `K.index`, it is actually the case that t is the reduction of y in $\mathbf{Q}[y]/(K.pol)$ modulo an irreducible factor of `K.pol` over $\mathbf{F}_p$. In the above example, 5 divides the index and t is actually the reduction of $y/5$.

The library syntax is `GEN nfmodpr(GEN nf, GEN x, GEN pr)`.

3.13.153 nfmodprinit(*nf*, *pr*, {*v* = *variable*(*nf.pol*)}). Transforms the prime ideal *pr* into `modpr` format necessary for all operations modulo *pr* in the number field *nf*. The functions `nfmodpr` and `nfmodprlift` allow to project to and lift from the residue field. The variable *v* is used to display finite field elements (see `ffgen`).

```
? K = nfinit(y^3-250);
? P = idealprimedec(K, 5)[2];
? modP = nfmodprinit(K, P, 't);
? K.zk
%4 = [1, 1/5*y, 1/25*y^2]
? apply(t->nfmodpr(K,t,modP), K.zk)
%5 = [1, t, 2*t + 1]
? %[1].mod
%6 = t^2 + 3*t + 4
? K.index
%7 = 125
```

For clarity, we represent elements in the residue field $\mathbf{F}_5[t]/(T)$ as polynomials in the variable t . Whenever the underlying rational prime does not divide `K.index`, it is actually the case that t is the reduction of y in $\mathbf{Q}[y]/(K.pol)$ modulo an irreducible factor of `K.pol` over $\mathbf{F}_p$. In the above example, 5 divides the index and t is actually the reduction of $y/5$.

The library syntax is `GEN nfmodprinit0(GEN nf, GEN pr, long v) = -1` where *v* is a variable number.

3.13.154 nfmodprlift(*nf*, *x*, *pr*). Lift the `t_FFELT` *x* (from `nfmodpr`) in the residue field modulo *pr* to the ring of integers. Vectors and matrices are also supported. For polynomials, use `apply` and the `present` function.

The argument *pr* is either a maximal ideal in `idealprimedec` format or, preferably, a `modpr` structure from `nfmodprinit`. There are no compatibility checks to try and decide whether *x* is attached the same residue field as defined by *pr*: the result is undefined if not.

The function `nfmodpr` allows to reduce to the residue field.

```
? K = nfinit(y^3-250);
? P = idealprimedec(K, 5)[2];
? modP = nfmodprinit(K,P);
? K.zk
%4 = [1, 1/5*y, 1/25*y^2]
? apply(t->nfmodpr(K,t,modP), K.zk)
```

```
%5 = [1, y, 2*y + 1]
? nfmodprlift(K, %, modP)
%6 = [1, 1/5*y, 2/5*y + 1]
? nfeltval(K, %[3] - K.zk[3], P)
%7 = 1
```

The library syntax is GEN nfmodprlift(GEN nf, GEN x, GEN pr).

3.13.155 nfnewprec(*nf*). Transforms the number field *nf* into the corresponding data using current (usually larger) precision. This function works as expected if *nf* is in fact a *bnf*, a *bnr* or a *rnf* (update structure to current precision). If the original *bnf* structure was *not* computed by **bnfinit**(,1), then this may be quite slow and even fail: many generators of principal ideals have to be computed and the algorithm may fail because the accuracy is not sufficient to bootstrap the required generators and fundamental units.

The library syntax is GEN nfnewprec(GEN nf, long prec). See also GEN bnfnewprec(GEN bnf, long prec) and GEN bnrnewprec(GEN bnr, long prec).

3.13.156 nfpolsturm(*nf*, *T*, {*pl*}). Given a polynomial *T* with coefficients in the number field *nf*, returns the number of real roots of the *s*(*T*) where *s* runs through the real embeddings of the field specified by optional argument *pl*:

- *pl* omitted: all r_1 real places;
- *pl* an integer between 1 and r_1 : the embedding attached to the *i*-th real root of **nf.pol**, i.e. **nf.roots**[*i*];

- *pl* a vector or **t_VECSMALL**: the embeddings attached to the *pl*[*i*]-th real roots of **nf.pol**.

```
? nf = nfinit('y^2 - 2);
? nf.sign
%2 = [2, 0]
? nf.roots
%3 = [-1.414..., 1.414...]
? T = x^2 + 'y;
? nfpolsturm(nf, T, 1) \\ subst(T,y,sqrt(2)) has two real roots
%5 = 2
? nfpolsturm(nf, T, 2) \\ subst(T,y,-sqrt(2)) has no real root
%6 = 0
? nfpolsturm(nf, T) \\ all embeddings together
%7 = [2, 0]
? nfpolsturm(nf, T, [2,1]) \\ second then first embedding
%8 = [0, 2]
? nfpolsturm(nf, x^3) \\ number of distinct roots !
%9 = [1, 1]
? nfpolsturm(nf, x, 6) \\ there are only 2 real embeddings !
*** at top-level: nfpolsturm(nf,x,6)
*** ^-----
*** nfpolsturm: domain error in nfpolsturm: index > 2
```

The library syntax is GEN nfpolsturm(GEN nf, GEN T, GEN pl = NULL).

3.13.157 nfresolvent(*pol*, {*flag* = 0}). Let *pol* be an irreducible integral polynomial defining a number field K with Galois closure $\tilde{K}$. This function is limited to the Galois groups supported by **nflist**; in the following ℓ denotes an odd prime. If $\text{Gal}(\tilde{K}/\mathbf{Q})$ is D_ℓ , A_4 , S_4 , F_5 (M_{20}), A_5 , M_{21} or M_{42} , returns a polynomial R defining the corresponding resolvent field (quadratic for D_ℓ , cyclic cubic for A_4 and M_{21} , noncyclic cubic for S_4 , cyclic quartic for F_5 , $A_5(6)$ sextic for A_5 , and cyclic sextic for M_{42}). In the $A_5(6)$ case, returns the A_5 field of which it is the resolvent. Otherwise, gives a “canonical” subfield, or 0 if the Galois group is not supported.

The binary digits of *flag* correspond to 1: returns a pair $[R, f]$ where f is a “conductor” whose definition is specific to each group and given below; 2: returns all “canonical” subfields.

Let D be the discriminant of the resolvent field **nfdisc**(R):

- In cases C_ℓ , D_ℓ , A_4 , or S_4 , $\text{disc}(K) = (Df^2)^m$ with $m = (\ell - 1)/2$ in the first two cases, and 1 in the last two.
- In cases where K is abelian over the resolvent subfield, the conductor of the relative extension.
- In case F_5 , $\text{disc}(K) = Df^4$ if $f > 0$ or $5^2 Df^4$ if $f < 0$.
- In cases M_{21} or M_{42} , $\text{disc}(K) = D^m f^6$ if $f > 0$ or $7^3 D^m f^6$ if $f < 0$, where $m = 2$ for M_{21} and $m = 1$ for M_{42} .
- In cases A_5 and $A_5(6)$, *flag* is currently ignored.

```
? pol = x^6-3*x^5+7*x^4-9*x^3+7*x^2-3*x+1; \\ Galois closure D_6
? nfresolvent(pol)
%2 = x^3 + x - 1
? nfresolvent(pol,1)
%3 = [x^3 + x - 1, [[31, 21, 3; 0, 1, 0; 0, 0, 1], [1]]]
```

The library syntax is **GEN nfresolvent(GEN pol, long flag)**.

3.13.158 nfroots({*nf*}, x). Roots of the polynomial x in the number field *nf* given by **nfinit** without multiplicity (in $\mathbf{Q}$ if *nf* is omitted). x has coefficients in the number field (scalar, polmod, polynomial, column vector). The main variable of *nf* must be of lower priority than that of x (see Section 2.5.3). However if the coefficients of the number field occur explicitly (as polmods) as coefficients of x , the variable of these polmods *must* be the same as the main variable of t (see **nfactor**).

It is possible to input a defining polynomial for *nf* instead, but this is in general less efficient since parts of an **nf** structure will then be computed internally. This is useful in two situations: when you do not need the **nf** elsewhere, or when you cannot initialize an **nf** due to integer factorization difficulties when attempting to compute the field discriminant and maximal order.

Caveat. `nfinit([T, listP])` allows to compute in polynomial time a conditional nf structure, which sets `nf.zk` to an order which is not guaranteed to be maximal at all primes. Always either use `nfcertify` first (which may not run in polynomial time) or make sure to input `nf.pol` instead of the conditional nf : `nfroots` is able to recover in polynomial time in this case, instead of potentially missing a factor.

The library syntax is `GEN nfroots(GEN nf = NULL, GEN x)`. See also `GEN nfrootsQ(GEN x)`, corresponding to `nf = NULL`.

3.13.159 nfrootsof1(nf). Returns a two-component vector $[w, z]$ where w is the number of roots of unity in the number field nf , and z is a primitive w -th root of unity. It is possible to input a defining polynomial for nf instead.

```
? K = nfinit(polcyclo(11));
? nfrootsof1(K)
%2 = [22, [0, 0, 0, 0, 0, -1, 0, 0, 0, 0]~]
? z = nfbasistoalg(K, %[2])    \\ in algebraic form
%3 = Mod(-x^5, x^10 + x^9 + x^8 + x^7 + x^6 + x^5 + x^4 + x^3 + x^2 + x + 1)
? [lift(z^11), lift(z^2)]      \\ proves that the order of z is 22
%4 = [-1, -x^9 - x^8 - x^7 - x^6 - x^5 - x^4 - x^3 - x^2 - x - 1]
```

This function guesses the number w as the gcd of the $\#k(v)^*$ for unramified v above odd primes, then computes the roots in nf of the w -th cyclotomic polynomial. The algorithm is polynomial time with respect to the field degree and the bitsize of the multiplication table in nf (both of them polynomially bounded in terms of the size of the discriminant). Fields of degree up to 100 or so should require less than one minute.

The library syntax is `GEN nfrootsof1(GEN nf)`.

3.13.160 nfsnf($nf, x, \{flag = 0\}$). Given a torsion $\mathbf{Z}_K$ -module x attached to the square integral invertible pseudo-matrix (A, I, J) , returns an ideal list $D = [d_1, \dots, d_n]$ which is the Smith normal form of x . In other words, x is isomorphic to $\mathbf{Z}_K/d_1 \oplus \dots \oplus \mathbf{Z}_K/d_n$ and d_i divides d_{i-1} for $i \geq 2$. If $flag$ is nonzero return $[D, U, V]$, where UAV is the identity.

See Section 3.13.4 for the definition of integral pseudo-matrix; briefly, it is input as a 3-component row vector $[A, I, J]$ where $I = [b_1, \dots, b_n]$ and $J = [a_1, \dots, a_n]$ are two ideal lists, and A is a square $n \times n$ matrix with columns $(A_1, \dots, A_n)$, seen as elements in K^n (with canonical basis $(e_1, \dots, e_n)$). This data defines the $\mathbf{Z}_K$ module x given by

$$(b_1 e_1 \oplus \dots \oplus b_n e_n) / (a_1 A_1 \oplus \dots \oplus a_n A_n) ,$$

The integrality condition is $a_{i,j} \in b_i a_j^{-1}$ for all i, j . If it is not satisfied, then the d_i will not be integral. Note that every finitely generated torsion module is isomorphic to a module of this form and even with $b_i = \mathbf{Z}_K$ for all i .

The library syntax is `GEN nfsnf0(GEN nf, GEN x, long flag)`. Also available:

`GEN nfsnf(GEN nf, GEN x) (flag = 0)`.

3.13.161 nfsolvemodpr(*nf*, *a*, *b*, *P*). This function is obsolete, use **nfmodpr**.

Let P be a prime ideal in **modpr** format (see **nfmodprinit**), let a be a matrix, invertible over the residue field, and let b be a column vector or matrix. This function returns a solution of $a \cdot x = b$; the coefficients of x are lifted to *nf* elements.

```
? K = nfinit(y^2+1);
? P = idealprimedec(K, 3)[1];
? P = nfmodprinit(K, P);
? a = [y+1, y; y, 0]; b = [1, y]~
? nfsolvemodpr(K, a,b, P)
%5 = [1, 2]~
```

The library syntax is **GEN nfsolvemodpr**(**GEN nf**, **GEN a**, **GEN b**, **GEN P**). This function is normally useless in library mode. Project your inputs to the residue field using **nfM.to_FqM**, then work there.

3.13.162 nfsplitting($P, \{d\}, \{fl\}$). Defining polynomial S over $\mathbf{Q}$ for the splitting field of $P \in \mathbf{Q}[x]$, that is the smallest field over which P is totally split. If irreducible, the polynomial P can also be given by a **nf** structure, which is more efficient. If d is given, it must be a multiple of the splitting field degree. Note that if P is reducible the splitting field degree can be smaller than the degree of P .

If *flag* is non-zero, we assume P to be monic, integral and irreducible and the return value depends on *flag*:

- *flag* = 1: return $[S, C]$ where S is as before and C is an embedding of $\mathbf{Q}[x]/(P)$ in its splitting field given by a polynomial (implicitly modulo S , as in **nfisincl**).

- *flag* = 2: return $[S, C]$ where C is vector of rational functions whose image in $\mathbf{Q}[x]/(S)$ yields the embedding; this avoids inverting the denominator, which is costly. when the degree of the splitting field is huge.

- *flag* = 3: return $[S, v, p]$ a data structure allowing to quickly compute the Galois group of the splitting field, which is used by **galoissplittinginit**; more precisely, p is a prime splitting completely in the splitting field and v is a vector with $\deg S$ elements describing the automorphisms of S acting on the roots of S modulo p .

```
? K = nfinit(x^3 - 2);
? nfsplitting(K)
%2 = x^6 + 108
? nfsplitting(x^8 - 2)
%3 = x^16 + 272*x^8 + 64
? S = nfsplitting(x^6 - 8) \\ reducible
%4 = x^4 + 2*x^2 + 4
? lift(nfroots(subst(S,x,a),x^6-8))
%5 = [-a, a, -1/2*a^3 - a, -1/2*a^3, 1/2*a^3, 1/2*a^3 + a]
? P = x^8-2;
? [S,C] = nfsplitting(P,,1)
%7 = [x^16 + 272*x^8 + 64, -7/768*x^13 - 239/96*x^5 + 1/2*x]
? subst(P, x, Mod(C,S))
%8 = Mod(0, x^16 + 272*x^8 + 64)
```

Specifying the degree d of the splitting field can make the computation faster; if d is not a multiple of the true degree, it will be ignored with a warning.

```
? nfsplitting(x^17-123);
time = 3,607 ms.
? poldegree(%)
%2 = 272
? nfsplitting(x^17-123,272);
time = 150 ms.
? nfsplitting(x^17-123,273);
*** nfsplitting: Warning: ignoring incorrect degree bound 273
time = 3,611 ms.
```

The complexity of the algorithm is polynomial in the degree d of the splitting field and the bitsize of T ; if d is large the result will likely be unusable, e.g. `nfinit` will not be an option:

```
? nfsplitting(x^6-x-1)
[... degree 720 polynomial deleted ...]
time = 11,020 ms.
```

Variant: Also available is `GEN nfsplitting(GEN T, GEN D)` for $flag = 0$.

The library syntax is `GEN nfsplitting0(GEN P, GEN d = NULL, long fl)`.

3.13.163 nfsubfields($pol, \{d = 0\}, \{flag = 0\}$). Finds all subfields of degree d of the number field defined by the (monic, integral) polynomial pol (all subfields if d is null or omitted). The result is a vector of subfields, each being given by $[g, h]$ (default) or simply g ($flag = 1$), where g is an absolute equation and h expresses one of the roots of g in terms of the root x of the polynomial defining nf . This routine uses

- Allombert's `galoissubfields` when nf is Galois (with weakly supersolvable Galois group).
- Klüners's or van Hoeij–Klüners–Novocin algorithm in the general case. The latter runs in polynomial time and is generally superior unless there exists a small unramified prime p such that pol has few irreducible factors modulo p .

An input of the form `[nf, fa]` is also allowed, where `fa` is the factorisation of $nf.pol$ over nf , expressed as a famat of polynomials with coefficients in the variable of `nf`, in which case the van Hoeij–Klüners–Novocin algorithm is used.

```
? pol = x^4 - x^3 - x^2 + x + 1;
? nfsubfields(pol)
%2 = [[x, 0], [x^2 - x + 1, x^3 - x^2 + 1], [x^4 - x^3 - x^2 + x + 1, x]]
? nfsubfields(pol,,1)
%2 = [x, x^2 - x + 1, x^4 - x^3 - x^2 + x + 1]
? y=varhigher("y"); fa = nffactor(pol,subst(pol,x,y));
? #nfsubfields([pol,fa])
%5 = 3
```

The library syntax is `GEN nfsubfields0(GEN pol, long d, long flag)`. Also available is `GEN nfsubfields(GEN nf, long d)`, corresponding to $flag = 0$.

3.13.164 nfsubfieldscm(*nf*, {*flag* = 0}). Computes the maximal CM subfield of *nf*. Returns 0 if *nf* does not have a CM subfield, otherwise returns [*g*, *h*] (default) or *g* (*flag* = 1) where *g* is an absolute equation and *h* expresses a root of *g* in terms of the generator of *nf*. Moreover, the CM involution is given by $X \bmod g(X) \mapsto -X \bmod g(X)$, i.e. $X \bmod g(X)$ is a totally imaginary element.

An input of the form [*nf*, *fa*] is also allowed, where *fa* is the factorisation of *nf.pol* over *nf*, and *nf* is also allowed to be a monic defining polynomial for the number field.

```
? nf = nfinit(x^8 + 20*x^6 + 10*x^4 - 4*x^2 + 9);
? nfsubfieldscm(nf)
%2 = [x^4 + 4480*x^2 + 3612672, 3*x^5 + 58*x^3 + 5*x]
? pol = y^16-8*y^14+29*y^12-60*y^10+74*y^8-48*y^6+8*y^4+4*y^2+1;
? fa = nffactor(pol, subst(pol,y,x));
? nfsubfieldscm([pol,fa])
%5 = [y^8 + ... , ...]
```

The library syntax is GEN nfsubfieldscm(GEN nf, long flag).

3.13.165 nfsubfieldsmax(*nf*, {*flag* = 0}). Computes the list of maximal subfields of *nf*. The result is a vector as in nfsubfields.

An input of the form [*nf*, *fa*] is also allowed, where *fa* is the factorisation of *nf.pol* over *nf*, and *nf* is also allowed to be a monic defining polynomial for the number field.

The library syntax is GEN nfsubfieldsmax(GEN nf, long flag).

3.13.166 nfweilheight(*nf*, *v*). Let *nf* be attached to a number field *K*, let *v* be a vector of elements of *K*, not all of them 0, seen as element of the projective space of dimension #*v* - 1. Return the absolute logarithmic Weil height of that element, which does not depend on the number field used to compute it.

When the entries of *v* are rational, the height is $\log(\text{normlp}(v / \text{content}(v), \infty))$.

```
? v = [1, 2, -3, 101]; Q = nfinit(x); Qi = nfinit(x^2 + 1);
? exponent(nfweilheight(Q, v) - log(101))
%2 = -125
? exponent(nfweilheight(Qi, v) - log(101))
%3 = -125
```

The library syntax is GEN nfweilheight(GEN nf, GEN v, long prec).

3.13.167 polcompositum(*P*, *Q*, {*flag* = 0}). *P* and *Q* being squarefree polynomials in $\mathbf{Z}[X]$ in the same variable, outputs the simple factors of the étale $\mathbf{Q}$ -algebra $A = \mathbf{Q}(X, Y)/(P(X), Q(Y))$. The factors are given by a list of polynomials *R* in $\mathbf{Z}[X]$, attached to the number field $\mathbf{Q}(X)/(R)$, and sorted by increasing degree (with respect to lexicographic ordering for factors of equal degrees). Returns an error if one of the polynomials is not squarefree.

Note that it is more efficient to reduce to the case where *P* and *Q* are irreducible first. The routine will not perform this for you, since it may be expensive, and the inputs are irreducible in most applications anyway. In this case, there will be a single factor *R* if and only if the number fields defined by *P* and *Q* are linearly disjoint (their intersection is $\mathbf{Q}$).

Assuming P is irreducible (of smaller degree than Q for efficiency), it is in general much faster to proceed as follows

```
nf = nfinit(P); L = nffactor(nf, Q)[,1];
vector(#L, i, rnfequation(nf, L[i]))
```

to obtain the same result. If you are only interested in the degrees of the simple factors, the `rnfequation` instruction can be replaced by a trivial `poldegree(P) * poldegree(L[i])`.

The binary digits of *flag* mean

1: outputs a vector of 4-component vectors $[R, a, b, k]$, where R ranges through the list of all possible compositums as above, and a (resp. b) expresses the root of P (resp. Q) as an element of $\mathbf{Q}(X)/(R)$. Finally, k is a small integer such that $b + ka = X$ modulo R .

2: assume that P and Q define number fields which are linearly disjoint: both polynomials are irreducible and the corresponding number fields have no common subfield besides $\mathbf{Q}$. This allows to save a costly factorization over $\mathbf{Q}$. In this case return the single simple factor instead of a vector with one element.

A compositum is often defined by a complicated polynomial, which it is advisable to reduce before further work. Here is an example involving the field $\mathbf{Q}(\zeta_5, 5^{1/5})$:

```
? L = polcompositum(x^5 - 5, polcyclo(5), 1); \\ list of [R, a, b, k]
? [R, a] = L[1]; \\ pick the single factor, extract R, a (ignore b, k)
? R \\ defines the compositum
%3 = x^20 + 5*x^19 + 15*x^18 + 35*x^17 + 70*x^16 + 141*x^15 + 260*x^14\
+ 355*x^13 + 95*x^12 - 1460*x^11 - 3279*x^10 - 3660*x^9 - 2005*x^8 \
+ 705*x^7 + 9210*x^6 + 13506*x^5 + 7145*x^4 - 2740*x^3 + 1040*x^2 \
- 320*x + 256
? a^5 - 5 \\ a fifth root of 5
%4 = 0
? [T, X] = polredbest(R, 1);
? T \\ simpler defining polynomial for Q[x]/(R)
%6 = x^20 + 25*x^10 + 5
? X \\ root of R in Q[y]/(T(y))
%7 = Mod(-1/11*x^15 - 1/11*x^14 + 1/22*x^10 - 47/22*x^5 - 29/11*x^4 + 7/22,\
x^20 + 25*x^10 + 5)
? a = subst(a.pol, 'x, X) \\ a in the new coordinates
%8 = Mod(1/11*x^14 + 29/11*x^4, x^20 + 25*x^10 + 5)
? a^5 - 5
%9 = 0
```

In the above example, $x^5 - 5$ and the 5-th cyclotomic polynomial are irreducible over $\mathbf{Q}$; they have coprime degrees so define linearly disjoint extensions and we could have started by

```
? [R,a] = polcompositum(x^5 - 5, polcyclo(5), 3); \\ [R, a, b, k]
```

The library syntax is GEN `polcompositum0(GEN P, GEN Q, long flag)`. Also available are GEN `compositum(GEN P, GEN Q)` (*flag* = 0) and GEN `compositum2(GEN P, GEN Q)` (*flag* = 1).

3.13.168 polgalois(T). Galois group of the nonconstant polynomial $T \in \mathbf{Q}[X]$. In the present version 2.17.4, T must be irreducible and the degree d of T must be less than or equal to 7. If the **galdata** package has been installed, degrees 8, 9, 10 and 11 are also implemented. By definition, if $K = \mathbf{Q}[x]/(T)$, this computes the action of the Galois group of the Galois closure of K on the d distinct roots of T , up to conjugacy (corresponding to different root orderings).

The output is a 4-component vector $[n, s, k, \text{name}]$ with the following meaning: n is the cardinality of the group, s is its signature ($s = 1$ if the group is a subgroup of the alternating group A_d , $s = -1$ otherwise) and name is a character string containing name of the transitive group according to the GAP 4 transitive groups library by Alexander Hulpke.

k is more arbitrary and the choice made up to version 2.2.3 of PARI is rather unfortunate: for $d > 7$, k is the numbering of the group among all transitive subgroups of S_d , as given in “The transitive groups of degree up to eleven”, G. Butler and J. McKay, *Communications in Algebra*, vol. 11, 1983, pp. 863–911 (group k is denoted T_k there). And for $d \leq 7$, it was ad hoc, so as to ensure that a given triple would denote a unique group. Specifically, for polynomials of degree $d \leq 7$, the groups are coded as follows, using standard notations

In degree 1: $S_1 = [1, 1, 1]$.

In degree 2: $S_2 = [2, -1, 1]$.

In degree 3: $A_3 = C_3 = [3, 1, 1]$, $S_3 = [6, -1, 1]$.

In degree 4: $C_4 = [4, -1, 1]$, $V_4 = [4, 1, 1]$, $D_4 = [8, -1, 1]$, $A_4 = [12, 1, 1]$, $S_4 = [24, -1, 1]$.

In degree 5: $C_5 = [5, 1, 1]$, $D_5 = [10, 1, 1]$, $M_{20} = [20, -1, 1]$, $A_5 = [60, 1, 1]$, $S_5 = [120, -1, 1]$.

In degree 6: $C_6 = [6, -1, 1]$, $S_3 = [6, -1, 2]$, $D_6 = [12, -1, 1]$, $A_4 = [12, 1, 1]$, $G_{18} = [18, -1, 1]$, $S_4^- = [24, -1, 1]$, $A_4 \times C_2 = [24, -1, 2]$, $S_4^+ = [24, 1, 1]$, $G_{36}^- = [36, -1, 1]$, $G_{36}^+ = [36, 1, 1]$, $S_4 \times C_2 = [48, -1, 1]$, $A_5 = PSL_2(5) = [60, 1, 1]$, $G_{72} = [72, -1, 1]$, $S_5 = PGL_2(5) = [120, -1, 1]$, $A_6 = [360, 1, 1]$, $S_6 = [720, -1, 1]$.

In degree 7: $C_7 = [7, 1, 1]$, $D_7 = [14, -1, 1]$, $M_{21} = [21, 1, 1]$, $M_{42} = [42, -1, 1]$, $PSL_2(7) = PSL_3(2) = [168, 1, 1]$, $A_7 = [2520, 1, 1]$, $S_7 = [5040, -1, 1]$.

This is deprecated and obsolete, but for reasons of backward compatibility, we cannot change this behavior yet. So you can use the default **new_galois_format** to switch to a consistent naming scheme, namely k is always the standard numbering of the group among all transitive subgroups of S_n . If this default is in effect, the above groups will be coded as:

In degree 1: $S_1 = [1, 1, 1]$.

In degree 2: $S_2 = [2, -1, 1]$.

In degree 3: $A_3 = C_3 = [3, 1, 1]$, $S_3 = [6, -1, 2]$.

In degree 4: $C_4 = [4, -1, 1]$, $V_4 = [4, 1, 2]$, $D_4 = [8, -1, 3]$, $A_4 = [12, 1, 4]$, $S_4 = [24, -1, 5]$.

In degree 5: $C_5 = [5, 1, 1]$, $D_5 = [10, 1, 2]$, $M_{20} = [20, -1, 3]$, $A_5 = [60, 1, 4]$, $S_5 = [120, -1, 5]$.

In degree 6: $C_6 = [6, -1, 1]$, $S_3 = [6, -1, 2]$, $D_6 = [12, -1, 3]$, $A_4 = [12, 1, 4]$, $G_{18} = [18, -1, 5]$, $A_4 \times C_2 = [24, -1, 6]$, $S_4^+ = [24, 1, 7]$, $S_4^- = [24, -1, 8]$, $G_{36}^- = [36, -1, 9]$, $G_{36}^+ = [36, 1, 10]$, $S_4 \times$

$C_2 = [48, -1, 11]$, $A_5 = PSL_2(5) = [60, 1, 12]$, $G_{72} = [72, -1, 13]$, $S_5 = PGL_2(5) = [120, -1, 14]$, $A_6 = [360, 1, 15]$, $S_6 = [720, -1, 16]$.

In degree 7: $C_7 = [7, 1, 1]$, $D_7 = [14, -1, 2]$, $M_{21} = [21, 1, 3]$, $M_{42} = [42, -1, 4]$, $PSL_2(7) = PSL_3(2) = [168, 1, 5]$, $A_7 = [2520, 1, 6]$, $S_7 = [5040, -1, 7]$.

Warning. The method used is that of resolvent polynomials and is sensitive to the current precision. The precision is updated internally but, in very rare cases, a wrong result may be returned if the initial precision was not sufficient.

The library syntax is `GEN polgalois(GEN T, long prec)`. To enable the new format in library mode, set the global variable `new_galois_format` to 1.

3.13.169 polred($T, \{flag = 0\}$). This function is *deprecated*, use `polredbest` instead. Finds polynomials with reasonably small coefficients defining subfields of the number field defined by T . One of the polynomials always defines $\mathbf{Q}$ (hence has degree 1), and another always defines the same number field as T if T is irreducible.

All T accepted by `nfinit` are also allowed here; in particular, the format `[T, listP]` is recommended, e.g. with `listP = 105` or a vector containing all ramified primes. Otherwise, the maximal order of $\mathbf{Q}[x]/(T)$ must be computed.

The following binary digits of *flag* are significant:

1: Possibly use a suborder of the maximal order. The primes dividing the index of the order chosen are larger than `primelimit` or divide integers stored in the `addprimes` table. This flag is *deprecated*, the `[T, listP]` format is more flexible.

2: gives also elements. The result is a two-column matrix, the first column giving primitive elements defining these subfields, the second giving the corresponding minimal polynomials.

```
? M = polred(x^4 + 8, 2)
%1 =
[          1          x - 1]
[ 1/2*x^2 + 1 x^2 - 2*x + 3]
[-1/2*x^2 + 1 x^2 - 2*x + 3]
[      1/2*x^2      x^2 + 2]
[      1/4*x^3      x^4 + 2]
? minpoly(Mod(M[4,1], x^4+8))
%2 = x^2 + 2
```

The library syntax is `polred(GEN T) (flag = 0)`. Also available is `GEN polred2(GEN T) (flag = 2)`. The function `polred0` is deprecated, provided for backward compatibility.

3.13.170 polredabs($T, \{flag = 0\}$). Returns a canonical defining polynomial P for the number field $\mathbf{Q}[X]/(T)$ defined by T , such that the sum of the squares of the modulus of the roots (i.e. the T_2 -norm) is minimal. Different T defining isomorphic number fields will yield the same P . All T accepted by **bnfinit** are also allowed here, e.g. nonmonic polynomials, or pairs $[T, \text{listP}]$ specifying that a nonmaximal order may be used. For convenience, any number field structure ($nf, bnf, \dots$) can also be used instead of T .

```
? polredabs(x^2 + 16)
%1 = x^2 + 1
? K = bnfinit(x^2 + 16); polredabs(K)
%2 = x^2 + 1
```

Warning 1. Using a **t_POL** T requires computing and fully factoring the discriminant d_K of the maximal order which may be very hard. You can use the format $[T, \text{listP}]$, where **listP** encodes a list of known coprime divisors of $\text{disc}(T)$ (see **??nfbasis**), to help the routine, thereby replacing this part of the algorithm by a polynomial time computation. But this may only compute a suborder of the maximal order, when the divisors are not squarefree or do not include all primes dividing d_K . The routine attempts to certify the result independently of this order computation as per **nfcertify**: we try to prove that the computed order is maximal. If the certification fails, the routine then fully factors the integers returned by **nfcertify**. You can also use **polredbest** to avoid this factorization step; in this case, the result is small but no longer canonical.

Warning 2. Apart from the factorization of the discriminant of T , this routine runs in polynomial time for a *fixed* degree. But the complexity is exponential in the degree: this routine may be exceedingly slow when the number field has many subfields, hence a lot of elements of small T_2 -norm. If you do not need a canonical polynomial, the function **polredbest** is in general much faster (it runs in polynomial time), and tends to return polynomials with smaller discriminants.

The binary digits of *flag* mean

1: outputs a two-component row vector $[P, a]$, where P is the default output and $\text{Mod}(a, P)$ is a root of the original T .

4: gives *all* polynomials of minimal T_2 norm; of the two polynomials $P(x)$ and $\pm P(-x)$, only one is given.

16: (OBSOLETE) Possibly use a suborder of the maximal order, *without* attempting to certify the result as in Warning 1. This makes **polredabs** behave like **polredbest**. Just use the latter.

```
? T = x^16 - 136*x^14 + 6476*x^12 - 141912*x^10 + 1513334*x^8 \
    - 7453176*x^6 + 13950764*x^4 - 5596840*x^2 + 46225
? T1 = polredabs(T); T2 = polredbest(T);
? [ norml2(polroots(T1)), norml2(polroots(T2)) ]
%3 = [88.0000000, 120.0000000]
? [ sizedigit(poldisc(T1)), sizedigit(poldisc(T2)) ]
%4 = [75, 67]
```

The precise definition of the output of **polredabs** is as follows.

- Consider the finite list of characteristic polynomials of primitive elements of K that are in $\mathbf{Z}_K$ and minimal for the T_2 norm; now remove from the list the polynomials whose discriminant do not have minimal absolute value. Note that this condition is restricted to the original list of polynomials with minimal T_2 norm and does not imply that the defining polynomial for the field with smallest discriminant belongs to the list !

- To a polynomial $P(x) = x^n + \dots + a_n \in \mathbf{R}[x]$ we attach the sequence $S(P)$ given by $|a_1|, a_1, \dots, |a_n|, a_n$. Order the polynomials P by the lexicographic order on the coefficient vectors $S(P)$. Then the output of **polredabs** is the smallest polynomial in the above list for that order. In other words, the monic polynomial which is lexicographically smallest with respect to the absolute values of coefficients, favouring negative coefficients to break ties, i.e. choosing $x^3 - 2$ rather than $x^3 + 2$.

The library syntax is **GEN polredabs0(GEN T, long flag)**. Instead of the above hardcoded numerical flags, one should use an or-ed combination of

- **nf_PARTIALFACT** (OBSOLETE): possibly use a suborder of the maximal order, *without* attempting to certify the result.

- **nf_ORIG**: return $[P, a]$, where $\text{Mod}(a, P)$ is a root of T .

- **nf_RAW**: return $[P, b]$, where $\text{Mod}(b, T)$ is a root of P . The algebraic integer b is the raw result produced by the small vectors enumeration in the maximal order; P was computed as the characteristic polynomial of $\text{Mod}(b, T)$. $\text{Mod}(a, P)$ as in **nf_ORIG** is obtained with **modreverse**.

- **nf_ADDZK**: if r is the result produced with some of the above flags (of the form P or $[P, c]$), return $[r, \text{zk}]$, where **zk** is a **Z**-basis for the maximal order of $\mathbf{Q}[X]/(P)$.

- **nf_ALL**: return a vector of results of the above form, for all polynomials of minimal T_2 -norm.

3.13.171 polredbest($T, \{flag = 0\}$). Finds a polynomial with reasonably small coefficients defining the same number field as T . All T accepted by **nfinit** are also allowed here (e.g. nonmonic polynomials, **nf**, **bnf**, **[T,Z_K.basis]**). Contrary to **polredabs**, this routine runs in polynomial time, but it offers no guarantee as to the minimality of its result.

This routine computes an LLL-reduced basis for an order in $\mathbf{Q}[X]/(T)$, then examines small linear combinations of the basis vectors, computing their characteristic polynomials. It returns the *separable* polynomial P of smallest discriminant, the one with lexicographically smallest **abs**(**Vec**(**P**)) in case of ties. This is a good candidate for subsequent number field computations since it guarantees that the denominators of algebraic integers, when expressed in the power basis, are reasonably small. With no claim of minimality, though.

It can happen that iterating this functions yields better and better polynomials, until it stabilizes:

```
? \p5
? P = X^12+8*X^8-50*X^6+16*X^4-3069*X^2+625;
? poldisc(P)*1.
%2 = 1.2622 E55
? P = polredbest(P);
? poldisc(P)*1.
%4 = 2.9012 E51
? P = polredbest(P);
? poldisc(P)*1.
%6 = 8.8704 E44
```

In this example, the initial polynomial P is the one returned by **polredabs**, and the last one is stable.

If $flag = 1$: outputs a two-component row vector $[P, a]$, where P is the default output and **a**, a **t_POLMOD** modulo **P**, is a root of the original T .

```
? [P,a] = polredbest(x^4 + 8, 1)
%1 = [x^4 + 2, Mod(x^3, x^4 + 2)]
? charpoly(a)
%2 = x^4 + 8
```

In particular, the map $\mathbf{Q}[x]/(T) \rightarrow \mathbf{Q}[x]/(P)$, $x \mapsto \mathbf{a}$ defines an isomorphism of number fields, which can be computed as

```
subst(lift(Q), 'x, a)
```

if Q is a `t_POLMOD` modulo T ; `b = modreverse(a)` returns a `t_POLMOD` giving the inverse of the above map (which should be useless since $\mathbf{Q}[x]/(P)$ is a priori a better representation for the number field and its elements).

The library syntax is `GEN polredbest(GEN T, long flag)`.

3.13.172 polredord(x). This function is obsolete, use `polredbest`.

The library syntax is `GEN polredord(GEN x)`.

3.13.173 poltschirnhaus(x). Applies a random Tschirnhausen transformation to the polynomial x , which is assumed to be nonconstant and separable, so as to obtain a new equation for the étale algebra defined by x . This is for instance useful when computing resolvents, hence is used by the `polgalois` function.

The library syntax is `GEN tschirnhaus(GEN x)`.

3.13.174 rnfalgtobasis(rnf, x). Expresses x on the relative integral basis. Here, rnf is a relative number field extension L/K as output by `rnfinit`, and x an element of L in absolute form, i.e. expressed as a polynomial or polmod with polmod coefficients, *not* on the relative integral basis.

The library syntax is `GEN rnfalgtobasis(GEN rnf, GEN x)`.

3.13.175 rnfbasis(bnf, M). Let K the field represented by bnf , as output by `bnfinit`. M is a projective $\mathbf{Z}_K$ -module of rank n ($M \otimes K$ is an n -dimensional K -vector space), given by a pseudo-basis of size n . The routine returns either a true $\mathbf{Z}_K$ -basis of M (of size n) if it exists, or an $n + 1$ -element generating set of M if not.

It is allowed to use a monic irreducible polynomial P in $K[X]$ instead of M , in which case, M is defined as the ring of integers of $K[X]/(P)$, viewed as a $\mathbf{Z}_K$ -module.

Huge discriminants, helping rnfdisc. The format $[T, B]$ is also accepted instead of T and computes an order which is maximal at all maximal ideals specified by B , see `??rnfinit`: the valuation of D is then correct at all such maximal ideals but may be incorrect at other primes.

The library syntax is `GEN rnfbasis(GEN bnf, GEN M)`.

3.13.176 rnfbasistoalg(rnf, x). Computes the representation of x as a polmod with polmods coefficients. Here, rnf is a relative number field extension L/K as output by `rnfinit`, and x an element of L expressed on the relative integral basis.

The library syntax is `GEN rnfbasistoalg(GEN rnf, GEN x)`.

3.13.177 rnfcharpoly($nf, T, a, \{var = 'x\}$). Characteristic polynomial of a over nf , where a belongs to the algebra defined by T over nf , i.e. $nf[X]/(T)$. Returns a polynomial in variable v (x by default).

```
? nf = nfinit(y^2+1);
? rnfcharpoly(nf, x^2+y*x+1, x+y)
%2 = x^2 + Mod(-y, y^2 + 1)*x + 1
```

The library syntax is GEN rnfcharpoly(GEN nf, GEN T, GEN a, long var = -1) where var is a variable number.

3.13.178 rnfconductor($bnf, T, \{flag = 0\}$). Given a bnf structure attached to a number field K , as produced by **bnfinit**, and T an irreducible polynomial in $K[x]$ defining an Abelian extension $L = K[x]/(T)$, computes the class field theory conductor of this Abelian extension. If T does not define an Abelian extension over K , the result is undefined; it may be the integer 0 (in which case the extension is definitely not Abelian) or a wrong result.

The result is a 3-component vector $[f, bnr, H]$, where f is the conductor of the extension given as a 2-component row vector $[f_0, f_\infty]$, bnr is the attached **bnr** structure and H is a matrix in HNF defining the subgroup of the ray class group on the ray class group generators **bnr.gen**; in particular, it is a left divisor of the diagonal matrix attached to **bnr.cyc** and $|\det H| = N = \deg T$.

- If $flag$ is 1, return $[f, bnrmod, H]$, where **bnrmod** is now attached to Cl_f/Cl_f^N , and H is as before since it contains the N -th powers. This is useful when f contains a maximal ideal with huge residue field, since the corresponding tough discrete logarithms are trivialized: in the quotient group, all elements have small order dividing N . This allows to work in Cl_f/H but no longer in Cl_f .

- If $flag$ is 2, only return $[f, fa]$ where **fa** is the factorization of the conductor finite part ($= f[1]$).

Huge discriminants, helping rnfdisc. The format $[T, B]$ is also accepted instead of T and computes the conductor of the extension provided it factors completely over the maximal ideals specified by B , see **??rnfinit**: the valuation of f_0 is then correct at all such maximal ideals but may be incorrect at other primes.

The library syntax is GEN rnfconductor0(GEN bnf, GEN T, long flag). Also available is GEN rnfconductor(GEN bnf, GEN T) when $flag = 0$.

3.13.179 rnfdedekind($nf, pol, \{pr\}, \{flag = 0\}$). Given a number field K coded by nf and a monic polynomial $P \in \mathbf{Z}_K[X]$, irreducible over K and thus defining a relative extension L of K , applies Dedekind's criterion to the order $\mathbf{Z}_K[X]/(P)$, at the prime ideal pr . It is possible to set pr to a vector of prime ideals (test maximality at all primes in the vector), or to omit altogether, in which case maximality at *all* primes is tested; in this situation $flag$ is automatically set to 1.

The default historic behavior ($flag$ is 0 or omitted and pr is a single prime ideal) is not so useful since **rnfpsudobasis** gives more information and is generally not that much slower. It returns a 3-component vector $[max, basis, v]$:

- $basis$ is a pseudo-basis of an enlarged order O produced by Dedekind's criterion, containing the original order $\mathbf{Z}_K[X]/(P)$ with index a power of pr . Possibly equal to the original order.

- max is a flag equal to 1 if the enlarged order O could be proven to be pr -maximal and to 0 otherwise; it may still be maximal in the latter case if pr is ramified in L ,

- v is the valuation at pr of the order discriminant.

If *flag* is nonzero, on the other hand, we just return 1 if the order $\mathbf{Z}_K[X]/(P)$ is *pr*-maximal (resp. maximal at all relevant primes, as described above), and 0 if not. This is much faster than the default, since the enlarged order is not computed.

```
? nf = nfinit(y^2-3); P = x^3 - 2*y;
? pr3 = idealprimedec(nf,3)[1];
? rnfdekind(nf, P, pr3)
%3 = [1, [[1, 0, 0; 0, 1, 0; 0, 0, 1], [1, 1, 1]], 8]
? rnfdekind(nf, P, pr3, 1)
%4 = 1
```

In this example, *pr3* is the ramified ideal above 3, and the order generated by the cube roots of y is already *pr3*-maximal. The order-discriminant has valuation 8. On the other hand, the order is not maximal at the prime above 2:

```
? pr2 = idealprimedec(nf,2)[1];
? rnfdekind(nf, P, pr2, 1)
%6 = 0
? rnfdekind(nf, P, pr2)
%7 = [0, [[2, 0, 0; 0, 1, 0; 0, 0, 1], [[1, 0; 0, 1], [1, 0; 0, 1],
      [1, 1/2; 0, 1/2]]], 2]
```

The enlarged order is not proven to be *pr2*-maximal yet. In fact, it is; it is in fact the maximal order:

```
? B = rnfpsudobasis(nf, P)
%8 = [[1, 0, 0; 0, 1, 0; 0, 0, 1], [1, 1, [1, 1/2; 0, 1/2]],
      [162, 0; 0, 162], -1]
? idealval(nf,B[3], pr2)
%9 = 2
```

It is possible to use this routine with nonmonic $P = \sum_{i \leq n} p_i X^i \in \mathbf{Z}_K[X]$ if *flag* = 1; in this case, we test maximality of Dedekind's order generated by

$$1, p_n \alpha, p_n \alpha^2 + p_{n-1} \alpha, \dots, p_n \alpha^{n-1} + p_{n-1} \alpha^{n-2} + \dots + p_1 \alpha.$$

The routine will fail if P vanishes on the projective line over the residue field $\mathbf{Z}_K/\mathfrak{pr}$ (FIXME).

The library syntax is GEN rnfdekind(GEN nf, GEN pol, GEN pr = NULL, long flag)

3.13.180 rnfDET(*nf*, *M*). Given a pseudo-matrix M over the maximal order of *nf*, computes its determinant.

The library syntax is GEN rnfDET(GEN nf, GEN M).

3.13.181 rnfDISC(*nf*, *T*). Given an *nf* structure attached to a number field K , as output by *nfinit*, and a monic irreducible polynomial $T \in K[x]$ defining a relative extension $L = K[x]/(T)$, compute the relative discriminant of L . This is a vector $[D, d]$, where D is the relative ideal discriminant and d is the relative discriminant considered as an element of K^*/K^{*2} . The main variable of *nf* must be of lower priority than that of T , see Section 2.5.3.

Huge discriminants, helping rnfdisc. The format $[T, B]$ is also accepted instead of T and computes an order which is maximal at all maximal ideals specified by B , see `??rnfinit`: the valuation of D is then correct at all such maximal ideals but may be incorrect at other primes.

The library syntax is `GEN rnfdiscf(GEN nf, GEN T)`.

3.13.182 rnfeltabstorel(*rnf*, *x*). Let *rnf* be a relative number field extension L/K as output by `rnfinit` and let *x* be an element of L expressed either

- as a polynomial modulo the absolute equation *rnf*.polabs,
- or in terms of the absolute $\mathbf{Z}$ -basis for $\mathbf{Z}_L$ if *rnf* contains one (as in `rnfinit(nf, pol, 1)`, or after a call to `nfinit(rnf)`).

Computes *x* as an element of the relative extension L/K as a polmod with polmod coefficients. If *x* is actually rational, return it as a rational number:

```
? K = nfinit(y^2+1); L = rnfinit(K, x^2-y);
? L.polabs
%2 = x^4 + 1
? rnfeltabstorel(L, Mod(x, L.polabs))
%3 = Mod(x, x^2 + Mod(-y, y^2 + 1))
? rnfeltabstorel(L, 1/3)
%4 = 1/3
? rnfeltabstorel(L, Mod(x, x^2-y))
%5 = Mod(x, x^2 + Mod(-y, y^2 + 1))
? rnfeltabstorel(L, [0,0,0,1]~) \\ Z_L not initialized yet
*** at top-level: rnfeltabstorel(L,[0,
*** ^-----
*** rnfeltabstorel: incorrect type in rnfeltabstorel, apply nfinit(rnf).
? nfinit(L); \\ initialize now
? rnfeltabstorel(L, [0,0,0,1]~)
%6 = Mod(Mod(y, y^2 + 1)*x, x^2 + Mod(-y, y^2 + 1))
? rnfeltabstorel(L, [1,0,0,0]~)
%7 = 1
```

The library syntax is `GEN rnfeltabstorel(GEN rnf, GEN x)`.

3.13.183 rnfeltdown(*rnf*, *x*, {*flag* = 0}). *rnf* being a relative number field extension L/K as output by `rnfinit` and *x* being an element of L expressed as a polynomial or polmod with polmod coefficients (or as a `t_COL` on `nfinit(rnf).zk`), computes *x* as an element of K as a `t_POLMOD` if *flag* = 0 and as a `t_COL` otherwise. If *x* is not in K , a domain error occurs. Note that if *x* is in fact rational, it is returned as a rational number, ignoring *flag*.

```
? K = nfinit(y^2+1); L = rnfinit(K, x^2-y);
? L.pol
%2 = x^4 + 1
? rnfeltdown(L, Mod(x^2, L.pol))
%3 = Mod(y, y^2 + 1)
? rnfeltdown(L, Mod(x^2, L.pol), 1)
%4 = [0, 1]~
? rnfeltdown(L, Mod(y, x^2-y))
```

```

%5 = Mod(y, y^2 + 1)
? rnfeltdown(L, Mod(y,K.pol))
%6 = Mod(y, y^2 + 1)
? rnfeltdown(L, Mod(x, L.pol))
*** at top-level: rnfeltdown(L,Mod(x,x
*** ^-----
*** rnfeltdown: domain error in rnfeltdown: element not in the base field
? rnfeltdown(L, Mod(y, x^2-y), 1) \\ as a t_COL
%7 = [0, 1]~
? rnfeltdown(L, [0,0,1,0]~) \\ not allowed without absolute nf struct
*** rnfeltdown: incorrect type in rnfeltdown (t_COL).
? nfinit(L); \\ add absolute nf structure to L
? rnfeltdown(L, [0,0,1,0]~) \\ now OK
%8 = Mod(y, y^2 + 1)

```

If we had started with $L = \text{rnfin}(K, x^2-y, 1)$, then the final command would have worked directly.

The library syntax is `GEN rnfeltdown0(GEN rnf, GEN x, long flag)`. Also available is `GEN rnfeltdown(GEN rnf, GEN x)` (*flag* = 0).

3.13.184 rnfeltnorm(*rnf*, *x*). *rnf* being a relative number field extension L/K as output by `rnfin` and *x* being an element of L , returns the relative norm $N_{L/K}(x)$ as an element of K .

```

? K = nfinit(y^2+1); L = rnfin(K, x^2-y);
? rnfeltnorm(L, Mod(x, L.pol))
%2 = Mod(x, x^2 + Mod(-y, y^2 + 1))
? rnfeltnorm(L, 2)
%3 = 4

```

The library syntax is `GEN rnfeltnorm(GEN rnf, GEN x)`.

3.13.185 rnfeltreltoabs(*rnf*, *x*). *rnf* being a relative number field extension L/K as output by `rnfin` and *x* being an element of L expressed as a polynomial or polmod with polmod coefficients, computes *x* as an element of the absolute extension $L/\mathbb{Q}$ as a polynomial modulo the absolute equation *rnf*.polabs.

```

? K = nfinit(y^2+1); L = rnfin(K, x^2-y);
? L.polabs
%2 = x^4 + 1
? rnfeltreltoabs(L, Mod(x, L.pol))
%3 = Mod(x, x^4 + 1)
? rnfeltreltoabs(L, Mod(y, x^2-y))
%4 = Mod(x^2, x^4 + 1)
? rnfeltreltoabs(L, Mod(y,K.pol))
%5 = Mod(x^2, x^4 + 1)

```

If the input is actually rational, then `rnfeltreltoabs` returns it as a rational number instead of a `t_POLMOD`:

```

? rnfeltreltoabs(L, Mod(2, K.pol))
%6 = 2

```

The library syntax is `GEN rnfeltreltoabs(GEN rnf, GEN x)`.

3.13.186 rnfelttrace(*rnf*, *x*). *rnf* being a relative number field extension L/K as output by `rnfini`t and *x* being an element of L , returns the relative trace $Tr_{L/K}(x)$ as an element of K .

```
? K = nfinit(y^2+1); L = rnfini(K, x^2-y);
? rnfelttrace(L, Mod(x, L.pol))
%2 = 0
? rnfelttrace(L, 2)
%3 = 4
```

The library syntax is `GEN rnfelttrace(GEN rnf, GEN x)`.

3.13.187 rnfeltup(*rnf*, *x*, {*flag* = 0}). *rnf* being a relative number field extension L/K as output by `rnfini`t and *x* being an element of K , computes *x* as an element of the absolute extension $L/\mathbf{Q}$. As a `t_POLMOD` modulo *rnf*.pol if *flag* = 0 and as a `t_COL` on the absolute field integer basis if *flag* = 1. Note that if *x* is in fact rational, it is returned as a rational number, ignoring *flag*.

```
? K = nfinit(y^2+1); L = rnfini(K, x^2-y);
? L.pol
%2 = x^4 + 1
? rnfeltup(L, Mod(y, K.pol))
%3 = Mod(x^2, x^4 + 1)
? rnfeltup(L, y)
%4 = Mod(x^2, x^4 + 1)
? rnfeltup(L, [1,2]~) \\ in terms of K.zk
%5 = Mod(2*x^2 + 1, x^4 + 1)
? rnfeltup(L, y, 1) \\ in terms of nfinit(L).zk
%6 = [0, 1, 0, 0]~
? rnfeltup(L, [1,2]~, 1)
%7 = [1, 2, 0, 0]~
? rnfeltup(L, [1,0]~) \\ rational
%8 = 1
```

The library syntax is `GEN rnfeltup0(GEN rnf, GEN x, long flag)`. Also available is `GEN rnfeltup(GEN rnf, GEN x)` (*flag* = 0).

3.13.188 rnfequation(*nf*, *pol*, {*flag* = 0}). Given a number field *nf* as output by `nfinit` (or simply a monic irreducible integral polynomial defining the field) and a polynomial *pol* with coefficients in *nf* defining a relative extension L of *nf*, computes an absolute equation of L over $\mathbf{Q}$.

The main variable of *nf* *must* be of lower priority than that of *pol* (see Section 2.5.3). Note that for efficiency, this does not check whether the relative equation is irreducible over *nf*, but only if it is squarefree. If it is reducible but squarefree, the result will be the absolute equation of the étale algebra defined by *pol*. If *pol* is not squarefree, raise an `e_DOMAIN` exception.

```
? rnfequation(y^2+1, x^2 - y)
%1 = x^4 + 1
? T = y^3-2; rnfequation(nfinit(T), (x^3-2)/(x-Mod(y,T)))
%2 = x^6 + 108 \\ Galois closure of Q(2^(1/3))
```

If *flag* is nonzero, outputs a 3-component row vector $[z, a, k]$, where

- *z* is the absolute equation of L over $\mathbf{Q}$, as in the default behavior,
- *a* expresses as a `t_POLMOD` modulo *z* a root α of the polynomial defining the base field *nf*,

• k is a small integer such that $\theta = \beta + k\alpha$ is a root of z , where β is a root of pol . It is guaranteed that $k = 0$ whenever $\mathbf{Q}(\beta) = L$.

```
? T = y^3-2; pol = x^2 + x*y + y^2;
? [z,a,k] = rnfequation(T, pol, 1);
? z
%3 = x^6 + 108
? subst(T, y, a)
%4 = 0
? alpha= Mod(y, T);
? beta = Mod(x*Mod(1,T), pol);
? subst(z, x, beta + k*alpha)
%7 = 0
```

The library syntax is `GEN rnfequation0(GEN nf, GEN pol, long flag)`. Also available are `GEN rnfequation(GEN nf, GEN pol)` ($flag = 0$) and `GEN rnfequation2(GEN nf, GEN pol)` ($flag = 1$).

3.13.189 rnfhnfbasis(bnf, M). Given a bnf attached to a number field K and a projective $\mathbf{Z}_K$ -module M given by a pseudo-matrix, returns either a true HNF basis of M if one exists, or zero otherwise. If M is a polynomial with coefficients in K , replace it by the pseudo-matrix returned by `rnfpsudobasis`.

The library syntax is `GEN rnfhnfbasis(GEN bnf, GEN M)`.

3.13.190 rnfidealabstorel(rnf, x). Let rnf be a relative number field extension L/K as output by `rnfinit` and let x be an ideal of the absolute extension $L/\mathbf{Q}$. Returns the relative pseudo-matrix in HNF giving the ideal x considered as an ideal of the relative extension L/K , i.e. as a $\mathbf{Z}_K$ -module.

Let `Labs` be an (absolute) `nf` structure attached to L , obtained via `Labs = nfinit(rnf)`. Then `rnf` “knows” about `Labs` and x may be given in any format attached to `Labs`, e.g. a prime ideal or an ideal in HNF wrt. `Labs.zk`:

```
? K = nfinit(y^2+1); rnf = rnfinit(K, x^2-y); Labs = nfinit(rnf);
? m = idealhnf(Labs, 17, x^3+2); \\ some ideal in HNF wrt. Labs.zk
? B = rnfidealabstorel(rnf, m)
%3 = [[1, 8; 0, 1], [[17, 4; 0, 1], 1]] \\ pseudo-basis for m as Z_K-module
? A = rnfidealreltoabs(rnf, B)
%4 = [17, x^2 + 4, x + 8, x^3 + 8*x^2] \\ Z-basis for m in Q[x]/(rnf.polabs)
? mathnf(matalgtobasis(Labs, A)) == m
%5 = 1
```

If on the other hand, we do not have a `Labs` at hand, because it would be too expensive to compute, but we nevertheless have a $\mathbf{Z}$ -basis for x , then we can use the function with this basis as argument. The entries of x may be given either modulo `rnf.polabs` (absolute form, possibly lifted) or modulo `rnf.pol` (relative form as `t_POLMODs`):

```
? K = nfinit(y^2+1); rnf = rnfinit(K, x^2-y);
? rnfidealabstorel(rnf, [17, x^2 + 4, x + 8, x^3 + 8*x^2])
%2 = [[1, 8; 0, 1], [[17, 4; 0, 1], 1]]
? rnfidealabstorel(rnf, Mod([17, y + 4, x + 8, y*x + 8*y], x^2-y))
%3 = [[1, 8; 0, 1], [[17, 4; 0, 1], 1]]
```

The library syntax is `GEN rnfidealabstorel(GEN rnf, GEN x)`.

3.13.191 rnfidealdown(rnf, x). Let rnf be a relative number field extension L/K as output by `rnfinit`, and x an ideal of L , given either in relative form or by a $\mathbf{Z}$ -basis of elements of L (see Section 3.13.190). This function returns the ideal of K below x , i.e. the intersection of x with K .

The library syntax is `GEN rnfidealdown(GEN rnf, GEN x)`.

3.13.192 rnfidealfactor(rnf, x). Factor into prime ideal powers the ideal x in the attached absolute number field $L = \text{nfininit}(rnf)$. The output format is similar to the `factor` function, and the prime ideals are represented in the form output by the `idealprimedec` function for L .

```
? rnf = rnfinit(nfinit(y^2+1), x^2-y+1);
? rnfidealfactor(rnf, y+1) \\ P_2^2
%2 =
[[2, [0,0,1,0]~, 4, 1, [0,0,0,2;0,0,-2,0;-1,-1,0,0;1,-1,0,0]] 2]
? rnfidealfactor(rnf, x) \\ P_2
%3 =
[[2, [0,0,1,0]~, 4, 1, [0,0,0,2;0,0,-2,0;-1,-1,0,0;1,-1,0,0]] 1]
? L = nfinit(rnf);
? id = idealhnf(L, idealhnf(L, 25, (x+1)^2));
? idealfactor(L, id) == rnfidealfactor(rnf, id)
%6 = 1
```

Note that ideals of the base field K must be explicitly lifted to L via `rnfidealup` before they can be factored.

The library syntax is `GEN rnfidealfactor(GEN rnf, GEN x)`.

3.13.193 rnfidealhnf(rnf, x). rnf being a relative number field extension L/K as output by `rnfinit` and x being a relative ideal (which can be, as in the absolute case, of many different types, including of course elements), computes the HNF pseudo-matrix attached to x , viewed as a $\mathbf{Z}_K$ -module.

The library syntax is `GEN rnfidealhnf(GEN rnf, GEN x)`.

3.13.194 rnfidealmul(rnf, x, y). rnf being a relative number field extension L/K as output by `rnfinit` and x and y being ideals of the relative extension L/K given by pseudo-matrices, outputs the ideal product, again as a relative ideal.

The library syntax is `GEN rnfidealmul(GEN rnf, GEN x, GEN y)`.

3.13.195 rnfidealnrmabs(rnf, x). Let rnf be a relative number field extension L/K as output by `rnfinit` and let x be a relative ideal (which can be, as in the absolute case, of many different types, including of course elements). This function computes the norm of the x considered as an ideal of the absolute extension $L/\mathbf{Q}$. This is identical to

```
idealnrm(rnf, rnfidealnrmrel(rnf, x))
```

but faster.

The library syntax is `GEN rnfidealnrmabs(GEN rnf, GEN x)`.

3.13.196 rnfidealnormrel(*rnf*, *x*). Let *rnf* be a relative number field extension L/K as output by **rnfinit** and let *x* be a relative ideal (which can be, as in the absolute case, of many different types, including of course elements). This function computes the relative norm of *x* as an ideal of K in HNF.

The library syntax is **GEN rnfidealnormrel**(**GEN rnf**, **GEN x**).

3.13.197 rnfidealprimedec(*rnf*, *pr*). Let *rnf* be a relative number field extension L/K as output by **rnfinit**, and *pr* a maximal ideal of K (*prid*), this function completes the *rnf* with a *nf* structure attached to L (see Section 3.13.201) and returns the vector *S* of prime ideals of $\mathbf{Z}_L$ above *pr*.

```
? K = nfinit(y^2+1); rnf = rnfinit(K, x^3+y+1);
? pr = idealprimedec(K, 2)[1];
? S = rnfidealprimedec(rnf, pr);
? #S
%4 = 1
```

The relative ramification indices and residue degrees can be obtained as **PR.e / pr.e** and **PR.f / PR.f**, if **PR** is an element of *S*.

The argument *pr* is also allowed to be a prime number *p*, in which case the function returns a pair of vectors [**SK**,**SL**], where **SK** contains the primes of K above *p* and **SL**[*i*] is the vector of primes of L above **SK**[*i*].

```
? [SK,SL] = rnfidealprimedec(rnf, 5);
? [#SK, vector(#SL,i,#SL[i])]
%6 = [2, [2, 2]]
```

The library syntax is **GEN rnfidealprimedec**(**GEN rnf**, **GEN pr**).

3.13.198 rnfidealreltoabs(*rnf*, *x*, {*flag* = 0}). Let *rnf* be a relative number field extension L/K as output by **rnfinit** and let *x* be a relative ideal, given as a $\mathbf{Z}_K$ -module by a pseudo matrix $[A, I]$. This function returns the ideal *x* as an absolute ideal of $L/\mathbf{Q}$. If *flag* = 0, the result is given by a vector of **t_POLMODs** modulo **rnf.pol** forming a $\mathbf{Z}$ -basis; if *flag* = 1, it is given in HNF in terms of the fixed $\mathbf{Z}$ -basis for $\mathbf{Z}_L$, see Section 3.13.201.

```
? K = nfinit(y^2+1); rnf = rnfinit(K, x^2-y);
? P = idealprimedec(K,2)[1];
? P = rnfidealup(rnf, P)
%3 = [2, x^2 + 1, 2*x, x^3 + x]
? Prel = rnfidealhnf(rnf, P)
%4 = [[1, 0; 0, 1], [[2, 1; 0, 1], [2, 1; 0, 1]]]
? rnfidealreltoabs(rnf,Prel)
%5 = [2, x^2 + 1, 2*x, x^3 + x]
? rnfidealreltoabs(rnf,Prel,1)
%6 =
[2 1 0 0]
[0 1 0 0]
[0 0 2 1]
[0 0 0 1]
```

The reason why we do not return by default ($flag = 0$) the customary HNF in terms of a fixed $\mathbf{Z}_L$ -basis for $\mathbf{Z}_L$ is precisely because a rnf does not contain such a basis by default. Completing the structure so that it contains a nf structure for L is polynomial time but costly when the absolute degree is large, thus it is not done by default. Note that setting $flag = 1$ will complete the rnf .

The library syntax is `GEN rnfidealreltoabs0(GEN rnf, GEN x, long flag)`. Also available is `GEN rnfidealreltoabs(GEN rnf, GEN x) (flag = 0)`.

3.13.199 rnfidealtwoelt(rnf, x). rnf being a relative number field extension L/K as output by `rnfini`t and x being an ideal of the relative extension L/K given by a pseudo-matrix, gives a vector of two generators of x over $\mathbf{Z}_L$ expressed as polmods with polmod coefficients.

The library syntax is `GEN rnfidealtwoelement(GEN rnf, GEN x)`.

3.13.200 rnfidealup($rnf, x, \{flag = 0\}$). Let rnf be a relative number field extension L/K as output by `rnfini`t and let x be an ideal of K . This function returns the ideal $x\mathbf{Z}_L$ as an absolute ideal of $L/\mathbf{Q}$, in the form of a $\mathbf{Z}$ -basis. If $flag = 0$, the result is given by a vector of polynomials (modulo `rnf.pol`); if $flag = 1$, it is given in HNF in terms of the fixed $\mathbf{Z}$ -basis for $\mathbf{Z}_L$, see Section 3.13.201.

```
? K = nfinit(y^2+1); rnf = rnfini(K, x^2-y);
? P = idealprimedec(K,2)[1];
? rnfidealup(rnf, P)
%3 = [2, x^2 + 1, 2*x, x^3 + x]
? rnfidealup(rnf, P,1)
%4 =
[2 1 0 0]
[0 1 0 0]
[0 0 2 1]
[0 0 0 1]
```

The reason why we do not return by default ($flag = 0$) the customary HNF in terms of a fixed $\mathbf{Z}$ -basis for $\mathbf{Z}_L$ is precisely because a rnf does not contain such a basis by default. Completing the structure so that it contains a nf structure for L is polynomial time but costly when the absolute degree is large, thus it is not done by default. Note that setting $flag = 1$ will complete the rnf .

The library syntax is `GEN rnfidealup0(GEN rnf, GEN x, long flag)`. Also available is `GEN rnfidealup(GEN rnf, GEN x) (flag = 0)`.

3.13.201 rnfinit($nf, T, \{flag = 0\}$). Given an nf structure attached to a number field K , as output by `nfinit`, and a monic irreducible polynomial T in $\mathbf{Z}_K[x]$ defining a relative extension $L = K[x]/(T)$, this computes data to work in L/K . The main variable of T must be of higher priority (see Section 2.5.3) than that of nf , and the coefficients of T must be in K .

The result is a row vector, whose components are technical. We let $m = [K : \mathbf{Q}]$ the degree of the base field, $n = [L : K]$ the relative degree, r_1 and r_2 the number of real and complex places of K . Access to this information via *member functions* is preferred since the specific data organization specified below will change in the future.

If $flag = 1$, add an nf structure attached to L to rnf . This is likely to be very expensive if the absolute degree mn is large, but fixes an integer basis for $\mathbf{Z}_L$ as a $\mathbf{Z}$ -module and allows to

input and output elements of L in absolute form: as `t_COL` for elements, as `t_MAT` in HNF for ideals, as `prid` for prime ideals. Without such a call, elements of L are represented as `t_POLMOD`, etc. Note that a subsequent `nfinit(rnf)` will also explicitly add such a component, and so will the following functions `rnfidealmul`, `rnfidealtwoelt`, `rnfidealprimedec`, `rnfidealup` (with flag 1) and `rnfidealreltoabs` (with flag 1). The absolute nf structure attached to L can be recovered using `nfinit(rnf)`.

`rnf[1](rnf.pol)` contains the relative polynomial T .

`rnf[2]` contains the integer basis $[A, d]$ of K , as (integral) elements of $L/\mathbf{Q}$. More precisely, A is a vector of polynomial with integer coefficients, d is a denominator, and the integer basis is given by A/d .

`rnf[3](rnf.disc)` is a two-component row vector $[\mathfrak{d}(L/K), s]$ where $\mathfrak{d}(L/K)$ is the relative ideal discriminant of L/K and s is the discriminant of L/K viewed as an element of $K^*/(K^*)^2$, in other words it is the output of `rnfdisc`.

`rnf[4](rnf.index)` is the ideal index $\mathfrak{f}$, i.e. such that $d(T)\mathbf{Z}_K = \mathfrak{f}^2\mathfrak{d}(L/K)$.

`rnf[5](rnf.p)` is the list of rational primes dividing the norm of the relative discriminant ideal.

`rnf[7](rnf.zk)` is the pseudo-basis (A, I) for the maximal order $\mathbf{Z}_L$ as a $\mathbf{Z}_K$ -module: A is the relative integral pseudo basis expressed as polynomials (in the variable of T) with polmod coefficients in nf , and the second component I is the ideal list of the pseudobasis in HNF.

`rnf[8]` is the inverse matrix of the integral basis matrix, with coefficients polmods in nf .

`rnf[9]` is currently unused.

`rnf[10](rnf.nf)` is nf .

`rnf[11]` is an extension of `rnfequation(K, T, 1)`. Namely, a vector $[P, a, k, K.pol, T]$ describing the *absolute* extension $L/\mathbf{Q}$: P is an absolute equation, more conveniently obtained as `rnf.polabs`; a expresses the generator $\alpha = y \bmod K.pol$ of the number field K as an element of L , i.e. a polynomial modulo the absolute equation P ;

k is a small integer such that, if β is an abstract root of T and α the generator of K given above, then $P(\beta + k\alpha) = 0$. It is guaranteed that $k = 0$ if $\mathbf{Q}(\beta) = L$.

Caveat. Be careful if $k \neq 0$ when dealing simultaneously with absolute and relative quantities since $L = \mathbf{Q}(\beta + k\alpha) = K(\alpha)$, and the generator chosen for the absolute extension is not the same as for the relative one. If this happens, one can of course go on working, but we advise to change the relative polynomial so that its root becomes $\beta + k\alpha$. Typical GP instructions would be

```
[P,a,k] = rnfequation(K, T, 1);
if (k, T = subst(T, x, x - k*Mod(y, K.pol)));
L = nfinit(K, T);
```

`rnf[12]` is by default unused and set equal to 0. This field is used to store further information about the field as it becomes available (which is rarely needed, hence would be too expensive to compute during the initial `nfinit` call).

Huge discriminants, helping rnfdisc. When T has a discriminant which is difficult to factor, it is hard to compute $\mathbf{Z}_L$. As in `nfinit`, the special input format $[T, B]$ is also accepted, where T is a polynomial as above and B specifies a list of maximal ideals. The following formats are recognized for B :

- an integer: the list of all maximal ideals above a rational prime $p < B$.
- a vector of rational primes or prime ideals: the list of all maximal ideals dividing an element in the list.

Instead of $\mathbf{Z}_L$, this produces an order which is maximal at all such maximal ideals primes. The result may actually be a complete and correct *rnf* structure if the relative ideal discriminant factors completely over this list of maximal ideals but this is not guaranteed. In general, the order may not be maximal at primes $\mathfrak{p}$ not in the list such that $\mathfrak{p}^2$ divides the relative ideal discriminant.

The library syntax is `GEN rnfinit0(GEN nf, GEN T, long flag)`. Also available is `GEN rnfinit(GEN nf, GEN T)` ($flag = 0$).

3.13.202 rnfisabelian(*nf*, *T*). T being a relative polynomial with coefficients in nf , return 1 if it defines an abelian extension, and 0 otherwise.

```
? K = nfinit(y^2 + 23);
? rnfisabelian(K, x^3 - 3*x - y)
%2 = 1
```

The library syntax is `long rnfisabelian(GEN nf, GEN T)`.

3.13.203 rnfisfree(*bnf*, *M*). Given a *bnf* attached to a number field K and a projective $\mathbf{Z}_K$ -module M given by a pseudo-matrix, return true (1) if M is free else return false (0). If M is a polynomial with coefficients in K , replace it by the pseudo-matrix returned by `rnfpsudobasis`.

The library syntax is `long rnfisfree(GEN bnf, GEN M)`.

3.13.204 rnfislocalcyclo(*rnf*). Let *rnf* be a relative number field extension L/K as output by `rnfinit` whose degree $[L : K]$ is a power of a prime ℓ . Return 1 if the ℓ -extension is locally cyclotomic (locally contained in the cyclotomic $\mathbf{Z}_\ell$ -extension of K_v at all places $v|\ell$), and 0 if not.

```
? K = nfinit(y^2 + y + 1);
? L = rnfinit(K, x^3 - y); /* = K(zeta_9), globally cyclotomic */
? rnfislocalcyclo(L)
%3 = 1
\\ we expect 3-adic continuity by Krasner's lemma
? vector(5, i, rnfislocalcyclo(rnfinit(K, x^3 - y + 3^i)))
%5 = [0, 1, 1, 1, 1]
```

The library syntax is `long rnfislocalcyclo(GEN rnf)`.

3.13.205 rnfisnorm($T, a, \{flag = 0\}$). Similar to **bnfisnorm** but in the relative case. T is as output by **rnfisnorminit** applied to the extension L/K . This tries to decide whether the element a in K is the norm of some x in the extension L/K .

The output is a vector $[x, q]$, where $a = \text{Norm}(x) * q$. The algorithm looks for a solution x which is an S -integer, with S a list of places of K containing at least the ramified primes, the generators of the class group of L , as well as those primes dividing a . If L/K is Galois, then this is enough but you may want to add more primes to S to produce different elements, possibly smaller; otherwise, $flag$ is used to add more primes to S : all the places above the primes $p \leq flag$ (resp. $p \nmid flag$) if $flag > 0$ (resp. $flag < 0$).

The answer is guaranteed (i.e. a is a norm iff $q = 1$) if the field is Galois, or, under GRH, if S contains all primes less than $4 \log^2 |\text{disc}(M)|$, where M is the normal closure of L/K .

If **rnfisnorminit** has determined (or was told) that L/K is Galois, and $flag \neq 0$, a Warning is issued (so that you can set $flag = 1$ to check whether L/K is known to be Galois, according to T). Example:

```
bnf = bnfinit(y^3 + y^2 - 2*y - 1);
p = x^2 + Mod(y^2 + 2*y + 1, bnf.pol);
T = rnfisnorminit(bnf, p);
rnfisnorm(T, 17)
```

checks whether 17 is a norm in the Galois extension $\mathbf{Q}(\beta)/\mathbf{Q}(\alpha)$, where $\alpha^3 + \alpha^2 - 2\alpha - 1 = 0$ and $\beta^2 + \alpha^2 + 2\alpha + 1 = 0$ (it is).

The library syntax is GEN **rnfisnorm**(GEN T , GEN a , long $flag$).

3.13.206 rnfisnorminit($pol, polrel, \{flag = 2\}$). Let K be defined by a root of pol , and L/K the extension defined by the polynomial $polrel$. As usual, pol can in fact be an nf , or bnf , etc; if pol has degree 1 (the base field is $\mathbf{Q}$), $polrel$ is also allowed to be an nf , etc. Computes technical data needed by **rnfisnorm** to solve norm equations $Nx = a$, for x in L , and a in K .

If $flag = 0$, do not care whether L/K is Galois or not.

If $flag = 1$, L/K is assumed to be Galois (unchecked), which speeds up **rnfisnorm**.

If $flag = 2$, let the routine determine whether L/K is Galois.

The library syntax is GEN **rnfisnorminit**(GEN pol , GEN $polrel$, long $flag$).

3.13.207 rnfkummer($bnr, \{subgp\}$). This function is deprecated, use **bnrclassfield**.

The library syntax is GEN **rnfkummer**(GEN bnr , GEN $subgp = \text{NULL}$, long $prec$).

3.13.208 rnflllgram($nf, pol, order$). Given a polynomial pol with coefficients in nf defining a relative extension L and a suborder $order$ of L (of maximal rank), as output by **rnfpsudobasis**(nf, pol) or similar, gives $[[neworder], U]$, where $neworder$ is a reduced order and U is the unimodular transformation matrix.

The library syntax is GEN **rnflllgram**(GEN nf , GEN pol , GEN $order$, long $prec$).

3.13.209 rnfnormgroup(*bnr*, *pol*). *bnr* being a big ray class field as output by **bnrinit** and *pol* a relative polynomial defining an Abelian extension, computes the norm group (alias Artin or Takagi group) corresponding to the Abelian extension of $bnf = \mathbf{bnr}.bnf$ defined by *pol*, where the module corresponding to *bnr* is assumed to be a multiple of the conductor (i.e. *pol* defines a subextension of *bnr*). The result is the HNF defining the norm group on the given generators of **bnr.gen**. Note that neither the fact that *pol* defines an Abelian extension nor the fact that the module is a multiple of the conductor is checked. The result is undefined if the assumption is not correct, but the function will return the empty matrix `[]` if it detects a problem; it may also not detect the problem and return a wrong result.

The library syntax is `GEN rnfnormgroup(GEN bnr, GEN pol)`.

3.13.210 rnfpolred(*nf*, *pol*). This function is obsolete: use **rnfpolredbest** instead. Relative version of **polred**. Given a monic polynomial *pol* with coefficients in *nf*, finds a list of relative polynomials defining some subfields, hopefully simpler and containing the original field. In the present version 2.17.4, this is slower and less efficient than **rnfpolredbest**.

Remark. This function is based on an incomplete reduction theory of lattices over number fields, implemented by **rnfl11gram**, which deserves to be improved.

The library syntax is `GEN rnfpolred(GEN nf, GEN pol, long prec)`.

3.13.211 rnfpolredabs(*nf*, *pol*, {*flag* = 0}). Relative version of **polredabs**. Given an irreducible monic polynomial *pol* with coefficients in the maximal order of *nf*, finds a canonical relative polynomial defining the same field, hopefully with small coefficients. Note that the equation is only canonical for a fixed *nf*, using a different defining polynomial in the *nf* structure will produce a different relative equation.

The binary digits of *flag* correspond to 1: add information to convert elements to the new representation, 2: absolute polynomial, instead of relative, 16: possibly use a suborder of the maximal order. More precisely:

0: default, return *P*

1: returns $[P, a]$ where *P* is the default output and *a*, a `t_POLMOD` modulo *P*, is a root of *pol*.

2: returns *Pabs*, an absolute, instead of a relative, polynomial. This polynomial is canonical and does not depend on the *nf* structure. Same as but faster than

`polredabs(rnfequation(nf, pol))`

3: returns $[Pabs, a, b]$, where *Pabs* is an absolute polynomial as above, *a*, *b* are `t_POLMOD` modulo *Pabs*, roots of **nf.pol** and *pol* respectively.

16: (OBSOLETE) possibly use a suborder of the maximal order. This makes **rnfpolredabs** behave as **rnfpolredbest**. Just use the latter.

Warning. The complexity of `rnfpolredabs` is exponential in the absolute degree. The function `rnfpolredbest` runs in polynomial time, and tends to return polynomials with smaller discriminants. It also supports polynomials with arbitrary coefficients in nf , neither integral nor necessarily monic.

The library syntax is `GEN rnfpolredabs(GEN nf, GEN pol, long flag)`.

3.13.212 rnfpolredbest($nf, pol, \{flag = 0\}$). Relative version of `polredbest`. Given a polynomial pol with coefficients in nf , finds a simpler relative polynomial P defining the same field. As opposed to `rnfpolredabs` this function does not return a *smallest* (canonical) polynomial with respect to some measure, but it does run in polynomial time.

The binary digits of $flag$ correspond to 1: add information to convert elements to the new representation, 2: absolute polynomial, instead of relative. More precisely:

0: default, return P

1: returns $[P, a]$ where P is the default output and a , a `t_POLMOD` modulo P , is a root of pol .

2: returns $Pabs$, an absolute, instead of a relative, polynomial. Same as but faster than

`rnfequation(nf, rnfpolredbest(nf, pol))`

3: returns $[Pabs, a, b]$, where $Pabs$ is an absolute polynomial as above, a, b are `t_POLMOD` modulo $Pabs$, roots of $nf.pol$ and pol respectively.

```
? K = nfinit(y^3-2); pol = x^2 +x*y + y^2;
? [P, a] = rnfpolredbest(K, pol, 1);
? P
%3 = x^2 - x + Mod(y - 1, y^3 - 2)
? a
%4 = Mod(Mod(2*y^2+3*y+4, y^3-2)*x + Mod(-y^2-2*y-2, y^3-2),
          x^2 - x + Mod(y-1, y^3-2))
? subst(K.pol, y, a)
%5 = 0
? [Pabs, a, b] = rnfpolredbest(K, pol, 3);
? Pabs
%7 = x^6 - 3*x^5 + 5*x^3 - 3*x + 1
? a
%8 = Mod(-x^2+x+1, x^6-3*x^5+5*x^3-3*x+1)
? b
%9 = Mod(2*x^5-5*x^4-3*x^3+10*x^2+5*x-5, x^6-3*x^5+5*x^3-3*x+1)
? subst(K.pol, y, a)
%10 = 0
? substvec(pol, [x, y], [a, b])
%11 = 0
```

The library syntax is `GEN rnfpolredbest(GEN nf, GEN pol, long flag)`.

3.13.213 rnfpsudobasis(*nf*, *T*). Given an *nf* structure attached to a number field K , as output by `nfinit`, and a monic irreducible polynomial T in $\mathbf{Z}_K[x]$ defining a relative extension $L = K[x]/(T)$, computes the relative discriminant of L and a pseudo-basis (A, J) for the maximal order $\mathbf{Z}_L$ viewed as a $\mathbf{Z}_K$ -module. This is output as a vector $[A, J, D, d]$, where D is the relative ideal discriminant and d is the relative discriminant considered as an element of K^*/K^{*2} .

```
? K = nfinit(y^2+1);
? [A,J,D,d] = rnfpsudobasis(K, x^2+y);
? A
%3 =
[1 0]
[0 1]
? J
%4 = [1, 1]
? D
%5 = [0, -4]~
? d
%6 = [0, -1]~
```

Huge discriminants, helping rnfdisc. The format $[T, B]$ is also accepted instead of T and produce an order which is maximal at all prime ideals specified by B , see `??rnfininit`.

```
? p = 585403248812100232206609398101;
? q = 711171340236468512951957953369;
? T = x^2 + 3*(p*q)^2;
? [A,J,D,d] = V = rnfpsudobasis(K, T); D
time = 22,178 ms.
%10 = 3
? [A,J,D,d] = W = rnfpsudobasis(K, [T,100]); D
time = 5 ms.
%11 = 3
? V == W
%12 = 1
? [A,J,D,d] = W = rnfpsudobasis(K, [T, [3]]); D
%13 = 3
? V == W
%14 = 1
```

In this example, the results are identical since $D \cap \mathbf{Z}$ factors over primes less than 100 (and in fact, over 3). Had it not been the case, the order would have been guaranteed maximal at primes $\mathfrak{p}|p$ for $p \leq 100$ only (resp. $\mathfrak{p}|3$). And might have been nonmaximal at any other prime ideal $\mathfrak{p}$ such that $\mathfrak{p}^2$ divided D .

The library syntax is `GEN rnfpsudobasis(GEN nf, GEN T)`.

3.13.214 rnfsteinitz(nf, M). Given a nf attached to a number field K and a projective module M given by a pseudo-matrix, returns a pseudo-basis (A, I) (not in HNF in general) such that all the ideals of I except perhaps the last one are equal to the ring of integers of nf . If M is a polynomial with coefficients in K , replace it by the pseudo-matrix returned by **rnfpsudobasis** and return the four-component row vector $[A, I, D, d]$ where (A, I) are as before and (D, d) are discriminants as returned by **rnfpsudobasis**. The ideal class of the last ideal of I is well defined; it is the Steinitz class of M (its image in $SK_0(\mathbf{Z}_K)$).

The library syntax is **GEN rnfsteinitz**(**GEN nf**, **GEN M**).

3.13.215 subcyclohminus($fH, \{p = 0\}$). Let F be the abelian number field contained in $\mathbf{Q}(\zeta_f)$ corresponding to the subgroup H of $(\mathbf{Z}/f\mathbf{Z})^*$. Computes the relative class number $h^-(F) = h(F)/h(F^+)$ of F . The argument fH encodes F and the data $[f, H]$ as follows:

- $fH = [f, H]$, where H is given by a vector of integral generators,
- $fH = [bnr, H]$, where bnr is attached to $\text{Cl}_f(\mathbf{Q})$ and H is a congruence subgroup,
- $fH = [G, H]$, where G is **idealstar**($f, 1$), and H is a subgroup of $(\mathbf{Z}/f\mathbf{Z})^\times$,
- $fH = f$, where we assume that $H = \{1\}$, i.e., $F = \mathbf{Q}(\zeta_f)$,
- an irreducible integral polynomial defining a primitive element for F .

The algorithm is based on an analytic class number formula:

$$h^-(F) = Q(F)w(F) \prod_{K \subset F} N_{\mathbf{Q}(\zeta_d)/\mathbf{Q}}(-B_{1,\chi}/2),$$

where $Q(F)$ is the unit index of F , $w(F)$ is the number of roots of unity contained in F and K runs through all imaginary cyclic subfields of F . For each K , d is the degree $[K : \mathbf{Q}]$, χ is an arbitrary injective character of $G(K/\mathbf{Q})$ to $\mathbf{C}^\times$ and the Bernoulli number is given by

$$B_{1,\chi} = (1/f_\chi) \sum_{a=1}^{f_\chi} a\chi(a) = -(1/(2 - \bar{\chi}(2))) \sum_{1 \leq a \leq f_\chi/2} \chi(a),$$

where f_χ is the conductor of χ , namely the conductor of K . The unit index $Q \in \{1, 2\}$ is difficult to determine in general. If it could be computed, the function returns $[a, b] = [h^-, Q]$; else it returns $[2h^-/Q, 0]$. More precisely, the second component is 0 unless we are in one of the following cases:

- If $f = p^a$ with a prime number p , then $Q = 1$.
- If $F = \mathbf{Q}(\zeta_f)$, then $Q = 1$ if and only if $f = p^a$.
- If $f = 4p^a$ or $p^a q^b$ with odd prime numbers p, q , then $Q = 1$ if and only if $[\mathbf{Q}(\zeta_f) : F]$ is even.

Finally, the optional parameter p is an *odd* prime number. If p is given, then **subcyclohminus** outputs the valuation at p of $h^-(F)$, in other words the maximal integer e such that $p^e \mid h^-(F)$ by evaluating p -adic valuations of Bernoulli numbers. Since p is odd and $Q \in \{1, 2\}$, the latter can be disregarded and the result is the same as **valuation**(**subcyclohminus**(**f,H**) [1], **p**), but adding this argument p can be much faster when p does not divide $[F : \mathbf{Q}]$ or if a high power of p divides $[F : \mathbf{Q}]$.

? [a,b] = **subcyclohminus**(22220); b

```

%1 = 2 \\ = Q
? sizedigit(a)
%2 = 4306 \\ huge...
? valuation(a, 101)
%3 = 41
? subcycloHminus(22220, 101) \\ directly compute the valuation
%4 = 41

```

shows that 101^{41} divides $h^-(\mathbf{Q}(\zeta_{22220}))$ exactly. Let k_n be the n -th layer of the cyclotomic $\mathbf{Z}_3$ -extension of $k = \mathbf{Q}(\sqrt{-1501391})$; the following computes e_n for $1 \leq n \leq 3$, where 3^{e_n} is the 3-part of the relative class number $h^-(k_n)$:

```

? d = 1501391;
? subcycloHminus([9*d, [28,10,8]], 3)
%1 = 5
? subcycloHminus([27*d, [28,188,53]], 3)
%2 = 12
? subcycloHminus([81*d, [161,80,242]], 3)
%3 = 26

```

Note that $h^+(k_n)$ is prime to 3 for all $n \geq 0$.

The following example computes the 3-part of $h^-(F)$, where F is the subfield of the 7860079-th cyclotomic field with degree $2 \cdot 3^8$.

```

? p=7860079; a=znprimroot(p)^(2*3^8);
? valuation(subcycloHminus([p,a])[1], 3)
time = 1min, 47,896 ms.
%2 = 65
? subcycloHminus([p,a], 3)
time = 1,290 ms.
%3 = 65

```

The library syntax is `GEN subcycloHminus(GEN fH, GEN p = NULL)`.

3.13.216 subcycloIwasawa($fH, p, \{n = 0\}$). Let F be the abelian number field contained in $\mathbf{Q}(\zeta_f)$ corresponding to the subgroup H of $(\mathbf{Z}/f\mathbf{Z})^*$, let $p > 2$ be an odd prime not dividing $[F : \mathbf{Q}]$, let F_∞ be the cyclotomic $\mathbf{Z}_p$ -extension of F and let F_n be its n -th layer. Computes the minus part of Iwasawa polynomials and λ -invariants attached to F_∞ , using the Stickelberger elements ξ_n^χ belonging to F_n .

The function is only implemented when p , n and f are relatively small: all of p^4 , p^{n+1} and f must fit into an `unsigned long` integer. The argument `fH` encodes the data $[f, H]$ as follows:

- $fH = [f, H]$, where H is given by a vector of integral generators,
- $fH = [bnr, H]$, where bnr is attached to $\text{Cl}_f(\mathbf{Q})$ and H is a congruence subgroup,
- $fH = [G, H]$, where G is `idealstar`($f, 1$), and H is a subgroup of $(\mathbf{Z}/f\mathbf{Z})^\times$,
- $fH = f$, where we assume that $H = \{1\}$, i.e., $F = \mathbf{Q}(\zeta_f)$,
- an irreducible integral polynomial defining a primitive element for F .

If F is quadratic, we also allow $p = 2$ and more data is output (see below).

For a number field K , we write K_n for the n -th layer of the cyclotomic $\mathbf{Z}_p$ -extension of K . The algorithm considers all cyclic subfields K of F and all injective odd characters $\chi : \text{Gal}(K/\mathbf{Q}) \rightarrow \overline{\mathbf{Q}}_p^\times$. Let $\Sigma_n = \text{Gal}(K_n/K)$, which is cyclic generated by the Frobenius automorphism σ ; we write $K_\chi = \mathbf{Q}_p(\chi)$, $\mathcal{O}_\chi = \mathbf{Z}_p[\chi]$ with maximal ideal $\mathfrak{p}$. The Stickelberger element ξ_n^χ belongs to $\mathcal{O}_\chi[\Sigma_n]$; the polynomial $f_n^\chi(x) \in \mathcal{O}_\chi[x]$ is constructed from ξ_n^χ by the correspondence $\sigma \mapsto 1+x$. If n is sufficiently large, then $\mathfrak{p}$ does not divide $f_n^\chi(x)$ and the distinguished polynomial $g_n^\chi(x) \in \mathcal{O}_\chi[x]$ is uniquely determined by the relation $f_n^\chi(x) = u(x)g_n^\chi(x)$, $u(x) \in \mathcal{O}_\chi[x]^\times$. Owing to Iwasawa Main Conjecture proved by Mazur-Wiles, we can define the Iwasawa polynomial $g_\chi(x) = \lim_{n \rightarrow \infty} g_n^\chi(x) \in \mathcal{O}_\chi[x]$. If r is the smallest integer satisfying $\deg g_n^\chi \leq p^r$, then we have

$$g_\chi(x) \equiv g_n^\chi(x) \pmod{\mathfrak{p}^{n+1-r}}.$$

Applying the norm from K_χ down to $\mathbf{Q}_p$, we obtain polynomials $G_\chi(x), G_n^\chi(x) \in \mathbf{Z}_p[x]$ satisfying the congruence

$$G_\chi(x) \equiv G_n^\chi(x) \pmod{p^{n+1-r}}.$$

Note that $\lambda_p^-(F) = \sum_{K, \chi} \deg G_\chi(x)$ is the Iwasawa λ^- -invariant of F , while the μ -invariant $\mu_p(F)$ is known to be zero by the theorem of Ferrero-Washington.

If $n = 0$, the function returns $[\lambda_p^-(F)]$ (the vector may contain further useful components, see below); for positive n , it returns all non-constant $G_n^\chi(x) \pmod{p^{n+1-r}}$ as (K, χ) vary.

```
? subcycloiwasa(22220, 41)  \ \ f = 22220, H = {1}
%1 = [217]
? P = polcompositum(x^2 - 42853, polcyclo(5))[1];
? subcycloiwasa(P, 5)
%3 = [3]
? subcycloiwasa(P, 5, 4) \ \ the sum of the degrees is indeed 3
%4 = [T + 585, T^2 + 405*T]
```

The first example corresponds to $F = \mathbf{Q}(\zeta_{22220})$ and shows, that $\lambda_{41}^-(F) = 217$. The second one builds $F = \mathbf{Q}(\sqrt{42853}, \zeta_5)$ then lists the non-constant $G_4^\chi(x) \pmod{p^4}$ for $p = 5$. Note that in this case all degrees are ≤ 5 hence $r \leq 1$ and $n + 1 - r \geq n$; so the above also gives G_χ modulo p^4 .

We henceforth restrict to the quadratic case, where more information is available, and $p = 2$ is now allowed: we write $F = \mathbf{Q}(\sqrt{d})$ of discriminant d ($\neq 1$) and character χ .

Algorithm and output for $n = 0$, $F = \mathbf{Q}(\sqrt{d})$. Currently, only the case $d < 0$ (F quadratic imaginary, i.e. $\chi(-1) = -1$) is implemented.

- If $p > 2$, the function returns $[\lambda, \nu, [e_0, \dots, e_k]]$, where $\lambda = \lambda_p^-(F)$, p^{e_n} denotes the p -part of the class number of F_n and $e_n = \lambda n + \nu$ for all $n > k$. We use Gold's theorem (Acta Arith. vol.26 (1974), pp. 21–32, vol.26 (1975), pp. 233–240). Then as soon as $e_n - e_{n-1} < \varphi(p^n)$ for some $n \geq 1$, we have $\lambda_p(F) = e_n - e_{n-1}$; if $\chi(p) = 1$ we can weaken the hypothesis to $e_n - e_{n-1} \leq \varphi(p^n)$ for some $n \geq 1$ and obtain the same conclusion. To compute $e_n - e_{n-1}$ we use Bernoulli numbers (`subcyclohminus`) if $\chi(p) = 0$ and a much faster algorithm of Gold (Pacific J. Math. vol.40 (1972), pp.83–88) otherwise.

- For $p = 2$, we use Kida's formula (Tohoku Math. J. vol. 31 (1979), pp. 91–96) and only return $[\lambda^-]$.

When $d > 1$, `subcycloiwasa` should calculate $\lambda_p(F) = \lambda_p^+(F)$, which is conjectured to be zero. But this is not yet implemented.

```
? subcycloiwasaawa(x^2+11111, 2)
%1 = [5] /*  $\lambda_2(\mathbf{Q}(\sqrt{-11111})) = 5$  */
? subcycloiwasaawa(x^2+11111, 3)
%2 = [1, 0, []]
? subcycloiwasaawa(x^2+11111, 11)
%3 = [0, 0, []]
```

This shows that for $p = 3$, we have $\lambda = 1$, $\nu = 0$, and $e_n = n$ for all $n \geq 0$. And at $p = 11$, we have $e_n = 0$ for all $n \geq 0$.

```
? subcycloiwasaawa(x^2+1501391, 3)
time = 23 ms.
%4 = [14, -16, [2, 5]]
```

computes e_n by Gold's algorithm for $F = \mathbf{Q}(\sqrt{-1501391})$. This shows that at $p = 3$, we have $\lambda = 14$, $\nu = -16$, then $e_0 = 2$, $e_1 = 5$, and $e_n = 14n - 16$ for $n \geq 2$.

```
? subcycloiwasaawa(x^2+956238, 3)
time = 141 ms.
%5 = [14, -19, [1, 3]]
```

computes e_n using Bernoulli numbers for $F = \mathbf{Q}(\sqrt{-956238})$. This shows that $e_0 = 1$, $e_1 = 3$ and $e_n = 14n - 19$ for $n \geq 2$.

Algorithm and output for $n > 0$; $F = \mathbf{Q}(\sqrt{d})$.

- When $d < 0$ and $n \geq 1$, `subcycloiwasaawa` computes the Stickelberger element $\xi_n = \xi_n^\chi \in \mathbf{Z}_p[\Sigma_n]$ and the Iwasawa polynomial $g(x) = g_\chi(x) \in \mathbf{Z}_p[x]$ from the n -th layer F_n of the cyclotomic $\mathbf{Z}_p$ -extension of F . Let q be p (p odd) or 4 ($p = 2$) and let q_0 be the lcm of q and the discriminant d of F , and let $q_n = q_0 p^n$. Then $\Sigma_n = \text{Gal}(\mathbf{Q}_n/\mathbf{Q}) = \text{Gal}(F_n/F) = \langle s \rangle$, where s is the Frobenius automorphism $(\mathbf{Q}_n/\mathbf{Q}, 1 + q_0)$ and

$$\xi_n = q_n^{-1} \sum_{a=1, (a, q_n)=1}^{q_n} a \chi(a)^{-1} (\mathbf{Q}_n/\mathbf{Q}, a)^{-1}$$

is an element of $\mathbf{Q}[\Sigma_n]$. For $(p, d) = (2, -1), (2, -2), (2, -3), (2, -6), (3, -3)$, we know that $\lambda_p(F) = 0$ and there is nothing to do. For the other cases, it is proved that $(1/2)\xi_n \in \mathbf{Z}_p[\Sigma_n]$. The polynomial $f_n(x) \in \mathbf{Z}_p[x]$ is constructed from $(1/2)\xi_n$ by the correspondence $s \mapsto 1 + x$. If n is sufficiently large, then p does not divide $f_n(x)$ and the distinguished polynomial $g_n(x) \in \mathbf{Z}_p[x]$ is uniquely determined by the relation $f_n(x) = u(x)g_n(x)$, $u(x) \in \mathbf{Z}_p[[x]]^\times$. The Iwasawa polynomial $g(x)$ is defined by $g(x) = \lim_{n \rightarrow \infty} g_n(x)$; if r is the smallest integer satisfying $\deg g = \lambda_p(F) \leq p^r$, then we have $g(x) \equiv g_n(x) \pmod{p^{n+1-r}}$ when $p > 2$ and modulo 2^{n-r} otherwise.

Conjecturally, we have further

1. case $q_0 = p$: $\xi_n \in \mathbf{Z}[\Sigma_n]$.
2. case $d = -1$ and $\chi(p) = -1$: $\xi_n \in \mathbf{Z}[\Sigma_n]$.
3. case $d = -3$ and $\chi(p) = -1$: $(3/2)\xi_n \in \mathbf{Z}[\Sigma_n]$.
4. other cases: $(1/2)\xi_n \in \mathbf{Z}[\Sigma_n]$.

Finally, `subcycloiwasaawa` outputs $[g]$ where g is $g_n(x) \bmod p^{n+1-r}$ (p odd) or $\bmod 2^{n-r}$ ($p = 2$).

```
? subcycloiwasawa(x^2+239, 3, 10)
%6 = [x^6 + 18780*x^5 + 14526*x^4 + 18168*x^3 + 3951*x^2 + 1128*x]
```

This is $g(x) \bmod 3^9$. Indeed, $n = 10$, $\lambda = 6$ (the degree), hence $r = 2$ and $n + 1 - r = 2$.

• When $d > 1$ and $n \geq 1$, $\xi_n^* \in \mathbf{Q}[\Sigma_n]$ is constructed from $\chi^* = \chi^{-1}\omega$, where χ is the character of $F = \mathbf{Q}(\sqrt{d})$ and ω is the Teichmüller character mod q . Next we construct $f_n^*(x) \in \mathbf{Z}_p[x]$ from $(1/2)\xi_n^*$ by the correspondence $s^{-1} \mapsto (1+x)(1+q_0)^{-1}$ and define the distinguished polynomial $g_n^*(x) \in \mathbf{Z}_p[x]$ using $f_n^*(x)$. Then $g^*(x) = \lim_{n \rightarrow \infty} g_n^*(x)$ is the Iwasawa polynomial, which has a connection with Greenberg conjecture for F . Let r be the smallest integer satisfying $\deg g^* \leq p^r$, then we have $g^*(x) \equiv g_n^*(x) \pmod{p^{n+1-r}}$ when $p > 2$ and $g^*(x) \equiv g_n^*(x) \pmod{2^{n-r}}$ when $p = 2$. Finally, `subcycloiwasawa` outputs $[g^*]$ where g^* is $g_n^*(x) \bmod p^{n+1-r}$ (p odd) or $\bmod 2^{n-r}$ ($p = 2$).

```
? subcycloiwasawa(x^2-13841, 2, 19)
time = 1min, 17,238 ms.
%7 = [x^3 + 30644*x^2 + 126772*x + 44128]
```

This is $g^*(x) \bmod 2^{17}$ ($r = 2$), the distinguished polynomial treated in a paper of T. Fukuda, K. Komatsu, M. Ozaki and T. Tsuji (Funct. Approx. Comment. Math. vol.54.1, pp. 7–17, 2016).

The library syntax is `GEN subcycloiwasawa(GEN fH, GEN p, long n)`.

3.13.217 subcyclopclgp($fH, p, \{flag = 0\}$). Let F be the abelian number field contained in $\mathbf{Q}(\zeta_f)$ corresponding to the subgroup H of $(\mathbf{Z}/f\mathbf{Z})^*$, let $p > 2$ be an odd prime not dividing $[F : \mathbf{Q}]$. Computes the p -Sylow subgroup A_F of the ideal class group using an unconditional algorithm of M. Aoki and T. Fukuda (LNCS. vol.4076, pp.56–71, 2006).

The argument `fH` encodes the data $[f, H]$ as follows:

- `fH` = $[f, H]$, where H is given by a vector of integral generators,
- `fH` = $[bnr, H]$, where bnr is attached to $\text{Cl}_f(\mathbf{Q})$ and H is a congruence subgroup,
- `fH` = $[G, H]$, where G is `idealstar`($f, 1$), and H is a subgroup of $(\mathbf{Z}/f\mathbf{Z})^\times$,
- `fH` = f , where we assume that $H = \{1\}$, i.e., $F = \mathbf{Q}(\zeta_f)$,
- an irreducible integral polynomial defining a primitive element for F .

The result is a 6-component vector v , and components 2 or 3 can be left empty or only partially computed to save time (see `flag` below):

$v[1]$ is p .

$v[2]$ contains $[E, [e_1, \dots, e_k]]$ with $E = \sum_i e_i$, meaning that the order of A_F^+ is p^E and its cyclic structure is $\mathbf{Z}/p^{e_1}\mathbf{Z} \times \dots \times \mathbf{Z}/p^{e_k}\mathbf{Z}$

$v[3]$ similarly describes the order and the structure of A_F^- .

$v[4]$ contains the structure of $\text{Gal}(F/\mathbf{Q})$ as a product of cyclic groups (elementary divisors).

$v[5]$ is the number of cyclic subfields K of F except for $\mathbf{Q}$.

$v[6]$ is the number of $\mathbf{Q}_p$ -conjugacy classes of injective characters $\chi : \text{Gal}(K/\mathbf{Q}) \rightarrow \overline{\mathbf{Q}_p}^\times$.

A vector of primes p is also accepted and the result is then a vector of vectors as above, in the same order as the primes.

The group A_F is the direct sum of A_F^+ and A_F^- ; each of A_F^+ and A_F^- is decomposed into χ -parts A_χ . By default, the function computes only $|A_F^-|$ and an upper bound for $|A_F^+|$ (expected to be equal to $|A_F^+|$) separately with different algorithms. This is expected to be fast. The behavior is controlled by the binary digits of *flag*:

1: if $|A_F^+|$ or $|A_F^-|$ is computed, also determines its group structure and guarantees informations about A_F^+ . This last part is usually costly.

2: do not compute quantities related to A_F^+ (the corresponding (e_i) in $v[2]$ is replaced with a dummy empty vector).

4: do not compute quantities related to A_F^- (the corresponding (e_i) in $v[3]$ is replaced with a dummy empty vector).

8: ignores proper subfields of F . This is motivated by the following kind of problems: let $\mathbf{Q}(p^k)$ be the k -th layer of the cyclotomic $\mathbf{Z}_p$ -extension of $\mathbf{Q}$ and define $\mathbf{Q}(n) = \mathbf{Q}(p_1^{e_1}) \cdots \mathbf{Q}(p_r^{e_r})$ when n factors as $n = p_1^{e_1} \cdots p_r^{e_r}$, which is a real cyclic field of degree n satisfying $\mathbf{Q}(n) \subset \mathbf{Q}(m)$ when $n \mid m$. What are the prime factors of the class number $h(n)$ of $\mathbf{Q}(n)$? The new prime factors of $h(n)$, not occurring in a lower level, will all be present when using this *flag*.

The other values are technical and only useful when bit 1 (certification and structure) is set; do not set them unless you run into difficulties with default parameters.

16: when this bit is set, the function tries to save memory, sacrificing speed; this typically uses half the memory for a slowdown of a factor 2.

32: likely to speed up the algorithm when the rank of A_χ is large and to create a minor slowdown otherwise. Though the effect is restricted, the 3-class group of $\mathbf{Q}(\sqrt{15338}, \zeta_5)$ is computed 4 times faster when this bit is set (see below).

Examples. With default $flag = 0$, the function (quickly) determines the exact value of $|A_F^-|$ and a rigorous upper bound of $|A_F^+|$ which is expected to be equal to $|A_F^+|$; of course, when the upper bound is 0, we know for sure that A_F^+ is trivial. With $flag = 1$ we obtain the group structure of A_F completely and guarantee the informations about A_F^+ (slow).

[illegible]

This computes the 101-part A_F of the ideal class group of $F = \mathbf{Q}(\zeta_{22220})$. The output says that $A_F^+ = 0$, which is rigorous (since trivial), and $|A_F^-| = 101^{41}$, more precisely A_F^- is isomorphic to $(\mathbf{Z}/101\mathbf{Z})^{41}$ which is also rigorous (since the description of A_F^- is always rigorous). The Galois group $\text{Gal}(F/\mathbf{Q})$ is $\mathbf{Z}/100\mathbf{Z} \oplus \mathbf{Z}/20\mathbf{Z} \oplus \mathbf{Z}/2\mathbf{Z} \oplus \mathbf{Z}/2\mathbf{Z}$. The field F has 479 cyclic subfields different from $\mathbf{Q}$ and there are 7999 $\mathbf{Q}_{101}$ -conjugacy classes of injective characters $\chi : \text{Gal}(K/\mathbf{Q}) \rightarrow \overline{\mathbf{Q}}_{101}^\times$.

```
? subcycloplgp(22220, 11)
time = 83 ms.
%2 = [11, [2, [1, 1]], [16, []], [100, 20, 2, 2], 479, 1799]
```

This computes the 11-part A_F for the same F . The result says that $|A_F^+| = 11^2$, A_F^+ is isomorphic to $(\mathbf{Z}/11\mathbf{Z})^2$ which is not rigorous and is only an upper bound, and $|A_F^-| = 11^{16}$ which is rigorous. The group structure of A_F^- is unknown.

```
? subcyclopcclgp(22220, 11, 1)
time = 185 ms.
%3 = [11, [2, [1, 1]], [16, [2, 1, 1, 1, 1, 1, 1, 1, 1, 1, 1, 1, 1, 1, 1]],
[100, 20, 2, 2], 479, 1799]
```

now guarantees that A_F^+ is isomorphic to $(\mathbf{Z}/11\mathbf{Z})^2$ and determines that A_F^- is isomorphic to $\mathbf{Z}/11^2\mathbf{Z} \oplus (\mathbf{Z}/11\mathbf{Z})^{14}$, at the expense of slightly increasing the running time.

We now try a much harder example: $F = \mathbf{Q}(\sqrt{36322}, \zeta_5)$, which we could define using $f = 726440$ and $H = [41, 61, 111, 131]$ (prove it!). We will use a defining polynomial instead:

```
? T = polcompositum(x^2-36322, polcyclo(5), 2);
? subcyclopcclgp(T, 5) \\ fast when non rigorous for A^+
time = 82 ms.
%4 = [5, [1, [1]], [4, []], [4, 2], 5, 7]
\\ try to certify; requires about 2GB of memory
? subcyclopcclgp(T, 5, 1)
*** subcyclopcclgp: the PARI stack overflows !
current stack size: 1000003072 (1907.352 Mbytes)
? default(parisizemax, "2G");
? subcyclopcclgp(T, 5, 1) \\ with more memory, we get an answer
time = 36,201 ms.
%6 = [5, [1, [1]], [4, [3, 1]], [4, 2], 5, 7]
\\ trying to reduce memory use does not work (still need 2GB); slower
? subcyclopcclgp(T, 5, 1+16)
time = 39,450 ms.
```

This shows that A_F^+ is isomorphic to $\mathbf{Z}/5\mathbf{Z}$ and A_F^- is isomorphic to $\mathbf{Z}/5^3\mathbf{Z} \oplus \mathbf{Z}/5\mathbf{Z}$ for $p = 5$. For this example, trying to reduce memory use with $flag = 1 + 16$ fails: the computation becomes slower and still needs 2GB; $flag = 1 + 16 + 32$ is a disaster: it requires about 8GB and 9 minutes of computation.

Here's a situation where the technical flags make a difference: let $F = \mathbf{Q}(\sqrt{15338}, \zeta_5)$.

```
? T = polcompositum(x^2-15338, polcyclo(5), 2);
? subcyclopcclgp(T, 3)
time = 123 ms.
%2 = [3, [1, [1]], [4, []], [4, 2], 5, 5]
? subcyclopcclgp(T, 3, 1) \\ requires a stack of 8GB
time = 4min, 47,822 ms.
%3 = [3, [1, [1]], [4, [1, 1, 1, 1]], [4, 2], 5, 5]
? subcyclopcclgp(T, 3, 1+16);
time = 7min, 20,876 ms. \\ works with 5GB, but slower
? subcyclopcclgp(T, 3, 1+32);
time = 1min, 11,424 ms. \\ also works with 5GB, 4 times faster than original
? subcyclopcclgp(T, 3, 1+16+32);
time = 1min, 47,285 ms. \\ now works with 2.5GB
```

Let $F = \mathbf{Q}(106)$ defined as above; namely, F is the composite field of $\mathbf{Q}(\sqrt{2})$ and the subfield of $\mathbf{Q}(\zeta_{53^2})$ with degree 53. This time we shall build the compositum using class field theory:

```
? Q = bnfinit(y);
? bnr1 = bnrinit(Q, 8); H1 = Mat(2);
```

```

? bnr2 = bnrinit(Q, [53^2, [1]]); H2 = Mat(53);
? [bnr,H] = bnrcompositum([bnr1, H1], [bnr2, H2]);
? subcyclopclgp([bnr,H], 107)
time = 10 ms.
%5 = [107, [1, [1]], [0, []], [106], 3, 105]
? subcyclopclgp([bnr,H], 107, 1) \\ requires 2.5GB
time = 15min, 13,537 ms.
%6 = [107, [1, [1]], [0, []], [106], 3, 105]

```

Both results are identical (and they were expected to be), but only the second is rigorous. Flag bit 32 has a minor impact in this case (reduces timings by 20 s.)

The library syntax is `GEN subcyclopclgp(GEN fH, GEN p, long flag)`.

3.13.218 subgrouplist(*cyc*, {*bound*}, {*flag* = 0}). *cyc* being a vector of positive integers giving the cyclic components for a finite Abelian group G (or any object which has a `.cyc` method), outputs the list of subgroups of G . Subgroups are given as HNF left divisors of the SNF matrix corresponding to G .

If *flag* = 0 (default) and *cyc* is a *bnr* structure output by `bnrinit`, gives only the subgroups whose conductor is the modulus `bnr.mod`. Otherwise, all subgroups are given.

If *bound* is present, and is a positive integer, restrict the output to subgroups of index less than *bound*. If *bound* is a vector containing a single positive integer B , then only subgroups of index exactly equal to B are computed. For instance

```

? subgrouplist([6,2])
%1 = [[6, 0; 0, 2], [2, 0; 0, 2], [6, 3; 0, 1], [2, 1; 0, 1], [3, 0; 0, 2],
[1, 0; 0, 2], [6, 0; 0, 1], [2, 0; 0, 1], [3, 0; 0, 1], [1, 0; 0, 1]]
? subgrouplist([6,2],3) \\ index less than 3
%2 = [[2, 1; 0, 1], [1, 0; 0, 2], [2, 0; 0, 1], [3, 0; 0, 1], [1, 0; 0, 1]]
? subgrouplist([6,2],[3]) \\ index 3
%3 = [[3, 0; 0, 1]]
? bnr = bnrinit(bnfinit(x), [120,[1]], 1);
? L = subgrouplist(bnr, [8]);

```

In the last example, L corresponds to the 24 subfields of $\mathbf{Q}(\zeta_{120})$, of degree 8 and conductor 120∞ (by setting *flag*, we see there are a total of 43 subgroups of degree 8).

```

? vector(#L, i, galoissubcyclo(bnr, L[i]))

```

will produce their equations. (For a general base field, you would have to rely on `bnrstark`, or `bnrclassfield`.)

Warning. This function requires factoring the exponent of G . If you are only interested in subgroups of index n (or dividing n), you may considerably speed up the function by computing the subgroups of G/G^n , whose cyclic components are `apply(x->gcd(n,x), C)` (where C gives the cyclic components of G). If you want the *bnr* variant, now is a good time to use `bnrinit(,, n)` as well, to directly compute the ray class group modulo n -th powers.

The library syntax is `GEN subgrouplist0(GEN cyc, GEN bound = NULL, long flag)`.

3.14 Associative and central simple algebras.

This section collects functions related to associative algebras and central simple algebras (CSA) over number fields.

3.14.1 Algebra definitions.

Let A be a finite-dimensional unital associative algebra over a field K . The algebra A is *central* if its center is K and it is *simple* if it has no nontrivial two-sided ideals.

We provide functions to handle associative algebras of finite dimension over $\mathbf{Q}$ or $\mathbf{F}_p$. We represent them by the left multiplication table on a basis over the prime subfield; the function `algtblinit` creates the object representing an associative algebra. We also provide functions to handle central simple algebras over a number field K . We represent them either by the left multiplication table on a basis over the center K or by a cyclic algebra (see below); the function `alginit` creates the object representing a central simple algebra.

The set of elements of an algebra A that annihilate every simple left A -module is a two-sided ideal, called the *Jacobson radical* of A . If the Jacobson radical is trivial, the algebra is *semisimple*: it is isomorphic to a direct product of simple algebras. The dimension of a CSA over its center K is always a square d^2 and the integer d is called the *degree* of the algebra over K . A CSA over a field K is always isomorphic to $M_k(D)$ for some integer k and some central division algebra D of degree e : the integer e is the *index* of the algebra.

Let L/K be a cyclic extension of degree d , let σ be a generator of $\text{Gal}(L/K)$ and let $b \in K^*$. Then the *cyclic algebra* $(L/K, \sigma, b)$ is the algebra $\bigoplus_{i=0}^{d-1} x^i L$ with $x^d = b$ and $\ell x = x\sigma(\ell)$ for all $\ell \in L$. The algebra $(L/K, \sigma, b)$ is a central simple K -algebra of degree d , and it is an L -vector space. Left multiplication is L -linear and induces a K -algebra isomorphism $(L/K, \sigma, b) \otimes_K L \rightarrow M_d(L)$.

Let K be a nonarchimedean local field with uniformizer π , and let L/K be the unique unramified extension of degree d . Then every central simple algebra A of degree d over K is isomorphic to $(L/K, \text{Frob}, \pi^h)$ for some integer h . The element $h/d \in \mathbf{Q}/\mathbf{Z}$ is called the *Hasse invariant* of A .

Let $\mathbf{H}$ be the Hamilton quaternion algebra, that is the 4-dimensional algebra over $\mathbf{R}$ with basis $1, i, j, ij$ and multiplication given by $i^2 = j^2 = -1$ and $ji = -ij$, which is also the cyclic algebra $(\mathbf{C}/\mathbf{R}, z \mapsto \bar{z}, -1)$. Every central simple algebra A of degree d over $\mathbf{R}$ is isomorphic to $M_d(\mathbf{R})$ or $M_{d/2}(\mathbf{H})$. We define the *Hasse invariant* of A to be $0 \in \mathbf{Q}/\mathbf{Z}$ in the first case and $1/2 \in \mathbf{Q}/\mathbf{Z}$ in the second case.

3.14.2 Orders in algebras.

Let A be an algebra of finite dimension over $\mathbf{Q}$. An *order* in A is a finitely generated $\mathbf{Z}$ -submodule $\mathcal{O}$ such that $\mathbf{Q}\mathcal{O} = A$, that is also a subring with unit. By default the data computed by `alginit` contains a $\mathbf{Z}$ -basis of a maximal order $\mathcal{O}_0$. We define natural orders in central simple algebras defined by a cyclic algebra or by a multiplication table over the center. Let $A = (L/K, \sigma, b) = \bigoplus_{i=0}^{d-1} x^i L$ be a cyclic algebra over a number field K of degree n with ring of integers $\mathbf{Z}_K$. Let $\mathbf{Z}_L$ be the ring of integers of L , and assume that b is integral. Then the submodule $\mathcal{O} = \bigoplus_{i=0}^{d-1} x^i \mathbf{Z}_L$ is an order in A , called the *natural order*. Let $\omega_0, \dots, \omega_{nd-1}$ be a $\mathbf{Z}$ -basis of $\mathbf{Z}_L$. The *natural basis* of $\mathcal{O}$ is $b_0, \dots, b_{nd^2-1}$ where $b_i = x^{i/(nd)} \omega_{(i \bmod nd)}$. Now let A be a central simple algebra of degree d over a number field K of degree n with ring of integers $\mathbf{Z}_K$. Let $e_0, \dots, e_{d^2-1}$ be a basis of A over K and assume that the left multiplication table of A on (e_i) is integral. Then the submodule $\mathcal{O} = \bigoplus_{i=0}^{d^2-1} \mathbf{Z}_K e_i$ is an order in A , called the *natural order*. Let $\omega_0, \dots, \omega_{n-1}$ be a $\mathbf{Z}$ -basis of $\mathbf{Z}_K$. The *natural basis* of $\mathcal{O}$ is $b_0, \dots, b_{nd^2-1}$ where $b_i = \omega_{(i \bmod n)} e_{i/n}$.

3.14.3 Lattices in algebras.

We also provide functions to handle full lattices in algebras over $\mathbf{Q}$. A full lattice $J \subset A$ is represented by a 2-component `t_VEC` $[I, t]$ representing $J = tI$, where

- I is an integral nonsingular upper-triangular matrix representing a sublattice of $\mathcal{O}_0$ expressed on the integral basis, and
- $t \in \mathbf{Q}_{>0}$ is a `t_INT` or `t_FRAC`.

For the sake of efficiency you should use matrices I that are primitive and in Hermite Normal Form; this makes the representation unique. No GP function uses this property, but all GP functions return lattices in this form. The prefix for lattice functions is `alglat`.

3.14.4 GP conventions for algebras.

As with number fields, we represent elements of central simple algebras in two ways, called the *algebraic representation* and the *basis representation*, and you can convert between the two with the functions `algalgtobasis` and `algbasistoalg`. In every central simple algebra object, we store a $\mathbf{Z}$ -basis of an order $\mathcal{O}_0$, and the basis representation is simply a `t_COL` with coefficients in $\mathbf{Q}$ expressing the element in that basis. If no maximal order was computed by `alginit`, then $\mathcal{O}_0$ is the natural order. If a maximal order was computed, then $\mathcal{O}_0$ is a maximal order containing the natural order. For a cyclic algebra $A = (L/K, \sigma, b)$, the algebraic representation is a `t_COL` with coefficients in L representing the element in the decomposition $A = \bigoplus_{i=0}^{d-1} x^i L$. For a central simple algebra defined by a multiplication table over its center K on a basis (e_i) , the algebraic representation is a `t_COL` with coefficients in K representing the element on the basis (e_i) .

Warning. The coefficients in the decomposition $A = \bigoplus_{i=0}^{d-1} x^i L$ are not the same as those in the decomposition $A = \bigoplus_{i=0}^{d-1} Lx^i$. The i -th coefficients are related by conjugating by x^i , which on L amounts to acting by σ^i .

Warning. For a central simple algebra over $\mathbf{Q}$ defined by a multiplication table, we cannot distinguish between the basis and the algebraic representations from the size of the vectors. The behavior is then to always interpret the column vector as a basis representation if the coefficients are `t_INT` or `t_FRAC`, and as an algebraic representation if the coefficients are `t_POL` or `t_POLMOD`.

An element of the Hamilton quaternion algebra $\mathbf{H}$ can be represented as a `t_REAL`, a `t_COMPLEX` representing an element of $\mathbf{C} = \mathbf{R} + \mathbf{R}i \subset \mathbf{H}$, or a 4 components `t_COL` of `t_REAL` encoding the coordinates on the basis $1, i, j, ij$.

3.14.5 `algadd`($\{al\}, x, y$). Given two elements x and y in al (Hamilton quaternions if omitted), computes their sum $x + y$ in the algebra al .

```
? A = alginit(nfinit(y), [-1,1]);
? algadd(A, [1,x]~, [1,2,3,4]~)
% = [2, 1, 1, 6]~
? algadd(sqrt(2+I), [-1,0,1,2]~)
% = [0.4553466902, 0.3435607497, 1, 2]~
```

Also accepts matrices with coefficients in al .

If x and y are given in the same format, then one should simply use `+` instead of `algadd`.

The library syntax is `GEN algadd(GEN al = NULL, GEN x, GEN y)`.

3.14.6 algalgtobasis(al, x). Given an element x in the central simple algebra al output by `alginit`, transforms it to a column vector on the integral basis of al . This is the inverse function of `algbasistoalg`.

```
? A = alginit(nfinit(y^2-5), [2,y]);
? algalgtobasis(A, [y,1]~)
%2 = [0, 2, 0, -1, 2, 0, 0, 0]~
? algbasistoalg(A, algalgtobasis(A, [y,1]~))
%3 = [Mod(Mod(y, y^2 - 5), x^2 - 2), 1]~
```

The library syntax is `GEN algalgtobasis(GEN al, GEN x)`.

3.14.7 algaut(al). Given a cyclic algebra $al = (L/K, \sigma, b)$ output by `alginit`, returns the automorphism σ .

```
? nf = nfinit(y);
? p = idealprimedec(nf, 7)[1];
? p2 = idealprimedec(nf, 11)[1];
? A = alginit(nf, [3, [[p, p2], [1/3, 2/3]], [0]]);
? algaut(A)
%5 = -1/3*x^2 + 1/3*x + 26/3
```

The library syntax is `GEN algaut(GEN al)`.

3.14.8 algb(al). Given a cyclic algebra $al = (L/K, \sigma, b)$ output by `alginit`, returns the element $b \in K$.

```
nf = nfinit(y);
? p = idealprimedec(nf, 7)[1];
? p2 = idealprimedec(nf, 11)[1];
? A = alginit(nf, [3, [[p, p2], [1/3, 2/3]], [0]]);
? algb(A)
%5 = Mod(-77, y)
```

The library syntax is `GEN algb(GEN al)`.

3.14.9 algbasis(al). Given a central simple algebra al output by `alginit`, returns a $\mathbf{Z}$ -basis of the order $\mathcal{O}_0$ stored in al with respect to the natural order in al . It is a maximal order if one has been computed.

```
A = alginit(nfinit(y), [-1,-1]);
? algbasis(A)
%2 =
[1 0 0 1/2]
[0 1 0 1/2]
[0 0 1 1/2]
[0 0 0 1/2]
```

The library syntax is `GEN algbasis(GEN al)`.

3.14.10 algbasistoalg(al, x). Given an element x in the central simple algebra al output by `alginit`, transforms it to its algebraic representation in al . This is the inverse function of `algalgtobasis`.

```
? A = alginit(nfinit(y^2-5), [2,y]);
? z = algbasistoalg(A, [0,1,0,0,2,-3,0,0]~);
? liftall(z)
%3 = [(-1/2*y - 2)*x + (-1/4*y + 5/4), -3/4*y + 7/4]~
? algalgtobasis(A,z)
%4 = [0, 1, 0, 0, 2, -3, 0, 0]~
```

The library syntax is `GEN algbasistoalg(GEN al, GEN x)`.

3.14.11 algcenter(al). If al is a table algebra output by `algtableinit`, returns a basis of the center of the algebra al over its prime field ($\mathbf{Q}$ or $\mathbf{F}_p$). If al is a central simple algebra output by `alginit`, returns the center of al , which is stored in al .

A simple example: the 2×2 upper triangular matrices over $\mathbf{Q}$, generated by I_2 , $a = [0, 1; 0, 0]$ and $b = [0, 0; 0, 1]$, such that $a^2 = 0$, $ab = a$, $ba = 0$, $b^2 = b$: the diagonal matrices form the center.

```
? mt = [matid(3), [0,0,0;1,0,1;0,0,0], [0,0,0;0,0,0;1,0,1]];
? A = algtableinit(mt);
? algcenter(A) \\ = (I_2)
%3 =
[1]
[0]
[0]
```

An example in the central simple case:

```
? nf = nfinit(y^3-y+1);
? A = alginit(nf, [-1,-1]);
? algcenter(A).pol
%3 = y^3 - y + 1
```

The library syntax is `GEN algcenter(GEN al)`.

3.14.12 algcentralproj($al, z, \{maps = 0\}$). Given a table algebra al output by `algtableinit` and a `t_VEC` $z = [z_1, \dots, z_n]$ of orthogonal central idempotents, returns a `t_VEC` $[al_1, \dots, al_n]$ of algebras such that $al_i = z_i al$. If $maps = 1$, each al_i is a `t_VEC` $[quo, proj, lift]$ where quo is the quotient algebra, $proj$ is a `t_MAT` representing the projection onto this quotient and $lift$ is a `t_MAT` representing a lift.

A simple example: $\mathbf{F}_2 \times \mathbf{F}_4$, generated by $1 = (1, 1)$, $e = (1, 0)$ and x such that $x^2 + x + 1 = 0$. We have $e^2 = e$, $x^2 = x + 1$ and $ex = 0$.

```
? mt = [matid(3), [0,0,0; 1,1,0; 0,0,0], [0,0,1; 0,0,0; 1,0,1]];
? A = algtableinit(mt,2);
? e = [0,1,0]~;
? e2 = algsub(A, [1,0,0]~, e);
? [a,a2] = algcentralproj(A, [e,e2]);
? algdim(a)
%6 = 1
```

```
? algdim(a2)
%7 = 2
```

The library syntax is GEN alg_centralproj(GEN al, GEN z, long maps).

3.14.13 algchar(*al*). Given an algebra *al* output by `alginit` or `algtbleinit`, returns the characteristic of *al*.

```
? mt = [matid(3), [0,0,0; 1,1,0; 0,0,0], [0,0,1; 0,0,0; 1,0,1]];
? A = algtbleinit(mt,13);
? algchar(A)
%3 = 13
```

The library syntax is GEN algchar(GEN al).

3.14.14 algcharpoly(*{al}*, *b*, *{v = 'x}*, *{abs = 0}*). Given an element *b* in *al* (Hamilton quaternions if omitted), returns its characteristic polynomial as a polynomial in the variable *v*. If *al* is a table algebra output by `algtbleinit` or if *abs* = 1, returns the absolute characteristic polynomial of *b*, which is an element of $\mathbf{F}_p[v]$, $\mathbf{Q}[v]$ or $\mathbf{R}[v]$; if *al* is omitted or a central simple algebra output by `alginit` and *abs* = 0, returns the reduced characteristic polynomial of *b*, which is an element of $K[v]$ where *K* is the center of *al*.

```
? al = alginit(nfinit(y), [-1,-1]); \\ (-1,-1)_Q
? algcharpoly(al, [0,1]~)
%2 = x^2 + 1
? algcharpoly(al, [0,1]~,1)
%3 = x^4 + 2*x^2 + 1
? nf = nfinit(y^2-5);
? al = alginit(nf, [-1,y]);
? a = [y,1+x]~*Mod(1,y^2-5)*Mod(1,x^2+1);
? P = lift(algcharpoly(al,a))
%7 = x^2 - 2*y*x + (-2*y + 5)
? algcharpoly(al,a,,1)
%8 = x^8 - 20*x^6 - 80*x^5 + 110*x^4 + 800*x^3 + 1500*x^2 - 400*x + 25
? lift(P*subst(P,y,-y)*Mod(1,y^2-5))^2
%9 = x^8 - 20*x^6 - 80*x^5 + 110*x^4 + 800*x^3 + 1500*x^2 - 400*x + 25
? algcharpoly([sqrt(2),-1,0,Pi]~) \\ Hamilton quaternions
%10 = x^2 - 2.8284271247*x + 12.8696044010
```

Also accepts a square matrix with coefficients in *al*.

The library syntax is GEN algcharpoly(GEN al = NULL, GEN b, long v = -1, long abs) where *v* is a variable number.

3.14.15 algdegree(*al*). Given a central simple algebra *al* output by `alginit`, returns the degree of *al*.

```
? nf = nfinit(y^3-y+1);
? A = alginit(nf, [-1,-1]);
? algdegree(A)
%3 = 2
```

The library syntax is long algdegree(GEN al).

3.14.16 algdim(*al*, {*abs* = 0}). If *al* is a table algebra output by `algtblinit` or if *abs* = 1, returns the dimension of *al* over its prime subfield ($\mathbf{Q}$ or $\mathbf{F}_p$) or over $\mathbf{R}$ for real algebras. If *al* is a central simple algebra output by `alginit` and *abs* = 0, returns the dimension of *al* over its center.

```
? nf = nfinit(y^3-y+1);
? A = alginit(nf, [-1,-1]);
? algdim(A)
%3 = 4
? algdim(A,1)
%4 = 12
? C = alginit(I,0); \\ complex numbers as a real algebra
? algdim(C,1)
%6 = 2
```

The library syntax is `long algdim(GEN al, long abs)`.

3.14.17 algdisc(*al*). Given a central simple algebra *al* output by `alginit`, computes the discriminant of the order $\mathcal{O}_0$ stored in *al*, that is the determinant of the trace form $\text{Tr} : \mathcal{O}_0 \times \mathcal{O}_0 \rightarrow \mathbf{Z}$.

```
? nf = nfinit(y^2-5);
? A = alginit(nf, [-3,1-y]);
? [PR,h] = alghassef(A)
%3 = [[[2, [2, 0]~, 1, 2, 1], [3, [3, 0]~, 1, 2, 1]], Vecsmall([0, 1])]
? n = algdegree(A);
? D = algdim(A,1);
? h = vector(#h, i, n - gcd(n,h[i]));
? n^D * nf.disc^(n^2) * ideallnorm(nf, idealfactorback(nf,PR,h))^n
%4 = 12960000
? algdisc(A)
%5 = 12960000
```

The library syntax is `GEN algdisc(GEN al)`.

3.14.18 algdivl({*al*}, *x*, *y*). Given two elements *x* and *y* in *al* (Hamilton quaternions if omitted), computes their left quotient $x \backslash y$ in the algebra *al*: an element *z* such that $xz = y$ (such an element is not unique when *x* is a zerodivisor). If *x* is invertible, this is the same as $x^{-1}y$. Assumes that *y* is left divisible by *x* (i.e. that *z* exists). Also accepts square matrices with coefficients in *al*.

```
? A = alginit(nfinit(y), [-1,1]);
? x = [1,1]~; alginv(A,x)
% = 0
? z = algmul(A,x,alrandom(A,2))
% = [0, 0, 0, 8]~
? algdivl(A,x,z)
% = [4, 4, 0, 0]~
```

The library syntax is `GEN algdivl(GEN al = NULL, GEN x, GEN y)`.

3.14.19 algdivr({*al*}, *x*, *y*). Given two elements *x* and *y* in *al* (Hamilton quaternions if omitted), returns xy^{-1} . Also accepts square matrices with coefficients in *al*.

The library syntax is `GEN algdivr(GEN al = NULL, GEN x, GEN y)`.

3.14.20 alggroup(*gal*, {*p* = 0}). Initializes the group algebra $K[G]$ over $K = \mathbf{Q}$ (*p* omitted) or $\mathbf{F}_p$ where G is the underlying group of the **galoisinit** structure *gal*. The input *gal* is also allowed to be a **t_VEC** of permutations that is closed under products.

Example:

```
? K = nfsplitting(x^3-x+1);
? gal = galoisinit(K);
? al = alggroup(gal);
? algisemisimple(al)
%4 = 1
? G = [Vecsmall([1,2,3]), Vecsmall([1,3,2])];
? al2 = alggroup(G, 2);
? algisemisimple(al2)
%8 = 0
```

The library syntax is **GEN alggroup**(**GEN gal**, **GEN p** = NULL).

3.14.21 alggroupcenter(*gal*, {*p* = 0}, {&*cc*}). Initializes the center $Z(K[G])$ of the group algebra $K[G]$ over $K = \mathbf{Q}$ (*p* = 0 or omitted) or $\mathbf{F}_p$ where G is the underlying group of the **galoisinit** structure *gal*. The input *gal* is also allowed to be a **t_VEC** of permutations that is closed under products. Sets *cc* to a **t_VEC** [*elts*, *conjclass*, *rep*, *flag*] where *elts* is a sorted **t_VEC** containing the list of elements of G , *conjclass* is a **t_VECSMALL** of the same length as *elts* containing the index of the conjugacy class of the corresponding element (an integer between 1 and the number of conjugacy classes), and *rep* is a **t_VECSMALL** of length the number of conjugacy classes giving for each conjugacy class the index in *elts* of a representative of this conjugacy class. Finally *flag* is 1 if and only if the permutation representation of G is transitive, in which case the *i*-th element of *elts* is characterized by $g[1] = i$; this is always the case when *gal* is a **galoisinit** structure. The basis of $Z(K[G])$ as output consists of the indicator functions of the conjugacy classes in the ordering given by *cc*. Example:

```
? K = nfsplitting(x^4+x+1);
? gal = galoisinit(K); \\ S4
? al = alggroupcenter(gal, &cc);
? algiscommutative(al)
%4 = 1
? #cc[3] \\ number of conjugacy classes of S4
%5 = 5
? gal = [Vecsmall([1,2,3]), Vecsmall([1,3,2])]; \\ C2
? al = alggroupcenter(gal, &cc);
? cc[3]
%8 = Vecsmall([1, 2])
? cc[4]
%9 = 0
```

The library syntax is **GEN alggroupcenter**(**GEN gal**, **GEN p** = NULL, **GEN *cc** = NULL)

3.14.22 alghasse(*al*, {*pl*}). Given a central simple algebra *al* output by `alginit` and a prime ideal or an integer between 1 and $r_1 + r_2$, returns a `t_FRAC` *h* : the local Hasse invariant of *al* at the place specified by *pl*. If *al* is an algebra over $\mathbf{R}$, returns the Hasse invariant of *al*

```
? nf = nfinit(y^2-5);
? A = alginit(nf, [-1,y]);
? alghasse(A, 1)
%3 = 1/2
? alghasse(A, 2)
%4 = 0
? alghasse(A, idealprimedec(nf,2)[1])
%5 = 1/2
? alghasse(A, idealprimedec(nf,5)[1])
%6 = 0
? H = alginit(1.,1/2); \\ Hamilton quaternion algebra
? alghasse(H)
%8 = 1/2
```

The library syntax is `GEN alghasse(GEN al, GEN pl = NULL)`.

3.14.23 alghassef(*al*). Given a central simple algebra *al* output by `alginit`, returns a `t_VEC` [*PR*, *h_f*] describing the local Hasse invariants at the finite places of the center: *PR* is a `t_VEC` of primes and *h_f* is a `t_VECSMALL` of integers modulo the degree *d* of *al*. The Hasse invariant of *al* at the primes outside *PR* is 0, but *PR* can include primes at which the Hasse invariant is 0.

```
? nf = nfinit(y^2-5);
? A = alginit(nf, [-1,2*y-1]);
? [PR,hf] = alghassef(A);
? PR
%4 = [[19, [10, 2]~, 1, 1, [-8, 2; 2, -10]], [2, [2, 0]~, 1, 2, 1]]
? hf
%5 = Vecsmall([1, 0])
```

The library syntax is `GEN alghassef(GEN al)`.

3.14.24 alghassei(*al*). Given a central simple algebra *al* output by `alginit`, returns a `t_VECSMALL` *h_i* of r_1 integers modulo the degree *d* of *al*, where r_1 is the number of real places of the center: the local Hasse invariants of *al* at infinite places.

```
? nf = nfinit(y^2-5);
? A = alginit(nf, [-1,y]);
? alghassei(A)
%3 = Vecsmall([1, 0])
```

The library syntax is `GEN alghassei(GEN al)`.

3.14.25 algindex(*al*, {*pl*}). Returns the index of the central simple algebra A over K (as output by `alginit`), that is the degree e of the unique central division algebra D over K such that A is isomorphic to some matrix algebra $M_k(D)$. If *pl* is set, it should be a prime ideal of K or an integer between 1 and $r_1 + r_2$, and in that case return the local index at the place *pl* instead.

```
? nf = nfinit(y^2-5);
? A = alginit(nf, [-1,y]);
? algindex(A, 1)
%3 = 2
? algindex(A, 2)
%4 = 1
? algindex(A, idealprimedec(nf,2)[1])
%5 = 2
? algindex(A, idealprimedec(nf,5)[1])
%6 = 1
? algindex(A)
%7 = 2
```

The library syntax is `long algindex(GEN al, GEN pl = NULL)`.

3.14.26 alginit($B, C, \{v\}, \{flag = 3\}$). Initializes the central simple algebra defined by data B, C and variable v , as follows.

- (multiplication table) B is the base number field K in `nfinit` form, C is a “multiplication table” over K . As a K -vector space, the algebra is generated by a basis $(e_1 = 1, \dots, e_n)$; the table is given as a `t_VEC` of n matrices in $M_n(K)$, giving the left multiplication by the basis elements e_i , in the given basis. Assumes that $e_1 = 1$, that the multiplication table is integral, and that $(\bigoplus_{i=1}^n Ke_i, C)$ describes a central simple algebra over K .

```
{ mi = [0,-1,0, 0;
        1, 0,0, 0;
        0, 0,0,-1;
        0, 0,1, 0];
  mj = [0, 0,-1,0;
        0, 0, 0,1;
        1, 0, 0,0;
        0,-1, 0,0];
  mk = [0, 0, 0, -1;
        0, 0,-1, 0;
        0, 1, 0, 0;
        1, 0, 0, 0];
  A = alginit(nfinit(y), [matid(4), mi,mj,mk], , 0); }
```

represents (in a complicated way) the quaternion algebra $(-1, -1)_{\mathbf{Q}}$. See below for a simpler solution.

- (cyclic algebra) B is an `rnf` structure attached to a cyclic number field extension L/K of degree d , C is a `t_VEC` [`sigma`, `b`] with 2 components: `sigma` is a `t_POLMOD` representing an automorphism generating $\text{Gal}(L/K)$, b is an element in K^* . This represents the cyclic algebra $(L/K, \sigma, b)$. Currently the element b has to be integral.

```
? Q = nfinit(y); T = polcyclo(5, 'x'); F = rnfinit(Q, T);
```

```
? A = alginit(F, [Mod(x^2,T), 3]);
```

defines the cyclic algebra $(L/\mathbf{Q}, \sigma, 3)$, where $L = \mathbf{Q}(\zeta_5)$ and $\sigma : \zeta \mapsto \zeta^2$ generates $\text{Gal}(L/\mathbf{Q})$.

- (quaternion algebra, special case of the above) B is an **nf** structure attached to a number field K , $C = [a, b]$ is a vector containing two elements of K^* with a not a square in K , returns the quaternion algebra $(a, b)_K$. The variable v ('x by default) must have higher priority than the variable of $K.\text{pol}$ and is used to represent elements in the splitting field $L = K[x]/(x^2 - a)$.

```
? Q = nfinit(y); A = alginit(Q, [-1,-1]); \\ (-1,-1)_{\mathbf{Q}}
```

- (algebra/ K defined by local Hasse invariants) B is an **nf** structure attached to a number field K , $C = [d, [\text{PR}, h_f], h_i]$ is a triple containing an integer $d > 1$, a pair $[\text{PR}, h_f]$ describing the Hasse invariants at finite places, and h_i the Hasse invariants at archimedean (real) places. A local Hasse invariant belongs to $(1/d)\mathbf{Z}/\mathbf{Z} \subset \mathbf{Q}/\mathbf{Z}$, and is given either as a **t_FRAC** (lift to $(1/d)\mathbf{Z}$), a **t_INT** or **t_INTMOD** modulo d (lift to $\mathbf{Z}/d\mathbf{Z}$); a whole vector of local invariants can also be given as a **t_VEC SMALL**, whose entries are handled as **t_INT**s. PR is a list of prime ideals (**prid** structures), and h_f is a vector of the same length giving the local invariants at those maximal ideals. The invariants at infinite real places are indexed by the real roots $K.\text{roots}$: if the Archimedean place v is attached to the j -th root, the value of h_v is given by $h_i[j]$, must be 0 or $1/2$ (or $d/2$ modulo d), and can be nonzero only if d is even.

By class field theory, provided the local invariants h_v sum to 0, up to Brauer equivalence, there is a unique central simple algebra over K with given local invariants and trivial invariant elsewhere. In particular, up to isomorphism, there is a unique such algebra A of degree d .

We realize A as a cyclic algebra through class field theory. The variable v ('x by default) must have higher priority than the variable of $K.\text{pol}$ and is used to represent elements in the (cyclic) splitting field extension L/K for A .

```
? nf = nfinit(y^2+1);
? PR = idealprimedec(nf,5); #PR
%2 = 2
? hi = [];
? hf = [PR, [1/3,-1/3]];
? A = alginit(nf, [3,hf,hi]);
? algsplittingfield(A).pol
%6 = x^3 - 21*x + 7
```

- (matrix algebra, toy example) B is an **nf** structure attached to a number field K , $C = d$ is a positive integer. Returns a cyclic algebra isomorphic to the matrix algebra $M_d(K)$.

- (algebras over $\mathbf{R}$) If B is a **t_REAL** and $C = 1/2$, returns a structure representing the Hamilton quaternion algebra $\mathbf{H} = (-1, -1)_{\mathbf{R}}$. If B is a **t_REAL** and $C = 0$, returns an algebra structure representing $\mathbf{R}$. If B is a **t_COMPLEX** and $C = 0$, returns an algebra structure representing $\mathbf{C}$.

In all cases over a number field, this function factors various discriminants and computes a maximal order for the algebra by default, which may require a lot of time. This can be controlled by *flag*, whose binary digits mean:

- 1: compute a maximal order.
- 2: fully factor the discriminants instead of using a lazy factorisation. If this digit of *flag* is set to 0, the local Hasse invariants are not computed.

If only a partial factorisation is known, the computed order is only guaranteed to be maximal at the known prime factors.

The pari object representing such an algebra A is a `t_VEC` with the following data:

- A splitting field L of A of the same degree over K as A , in `rnfinit` format, accessed with `algsplittingfield`.
- The Hasse invariants at the real places of K , accessed with `alghassei`.
- The Hasse invariants of A at the finite primes of K that ramify in the natural order of A , accessed with `alghassef`.
- A basis of an order $\mathcal{O}_0$ expressed on the basis of the natural order, accessed with `algbasis`.
- A basis of the natural order expressed on the basis of $\mathcal{O}_0$, accessed with `alginvbasis`.
- The left multiplication table of $\mathcal{O}_0$ on the previous basis, accessed with `algmultable`.
- The characteristic of A (always 0), accessed with `algchar`.
- The absolute traces of the elements of the basis of $\mathcal{O}_0$.
- If A was constructed as a cyclic algebra $(L/K, \sigma, b)$ of degree d , a `t_VEC` $[\sigma, \sigma^2, \dots, \sigma^{d-1}]$. The function `algaut` returns σ .

- If A was constructed as a cyclic algebra $(L/K, \sigma, b)$, the element b , accessed with `algb`.
- If A was constructed with its multiplication table mt over K , the `t_VEC` of `t_MAT` mt , accessed with `algrelmultable`.
- If A was constructed with its multiplication table mt over K , a `t_VEC` with three components: a `t_COL` representing an element of A generating the splitting field L as a maximal subfield of A , a `t_MAT` representing an L -basis $\mathcal{B}$ of A expressed on the $\mathbf{Z}$ -basis of $\mathcal{O}_0$, and a `t_MAT` representing the $\mathbf{Z}$ -basis of $\mathcal{O}_0$ expressed on $\mathcal{B}$. This data is accessed with `algsplittingdata`.

The library syntax is `GEN alginit(GEN B, GEN C, long v = -1, long flag)` where v is a variable number.

3.14.27 `alginv`($\{al\}, x$). Given an element x in al (Hamilton quaternions if omitted), computes its inverse x^{-1} in the algebra al . Assumes that x is invertible.

```
? A = alginit(nfinit(y), [-1,-1]);
? alginv(A,[1,1,0,0]~)
%2 = [1/2, 1/2, 0, 0]~
? alginv([1,0,Pi,sqrt(2)]~) \\ Hamilton quaternions
%3 = [0.0777024661, 0, -0.2441094967, -0.1098878814]~
```

Also accepts square matrices with coefficients in al .

The library syntax is `GEN alginv(GEN al = NULL, GEN x)`.

3.14.28 alginvbasis(*al*). Given an central simple algebra *al* output by `alginit`, returns a $\mathbf{Z}$ -basis of the natural order in *al* with respect to the order $\mathcal{O}_0$ stored in *al*.

```
A = alginit(nfinit(y), [-1,-1]);
? alginvbasis(A)
%2 =
[1 0 0 -1]
[0 1 0 -1]
[0 0 1 -1]
[0 0 0 2]
```

The library syntax is `GEN alginvbasis(GEN al)`.

3.14.29 algisassociative(*mt*, *p* = 0). Returns 1 if the multiplication table *mt* is suitable for `algtbleinit(mt, p)`, 0 otherwise. More precisely, *mt* should be a `t_VEC` of *n* matrices in $M_n(K)$, giving the left multiplications by the basis elements $e_1, \dots, e_n$ (structure constants). We check whether the first basis element e_1 is 1 and $e_i(e_j e_k) = (e_i e_j) e_k$ for all i, j, k .

```
? mt = [matid(3), [0,0,0;1,0,1;0,0,0], [0,0,0;0,0,0;1,0,1]];
? algisassociative(mt)
%2 = 1
```

May be used to check a posteriori an algebra: we also allow *mt* as output by `algtbleinit` (*p* is ignored in this case).

The library syntax is `int algisassociative(GEN mt, GEN p)`.

3.14.30 algiscommutative(*al*). *al* being a table algebra output by `algtbleinit` or a central simple algebra output by `alginit`, tests whether the algebra *al* is commutative.

```
? mt = [matid(3), [0,0,0;1,0,1;0,0,0], [0,0,0;0,0,0;1,0,1]];
? A = algtbleinit(mt);
? algiscommutative(A)
%3 = 0
? mt = [matid(3), [0,0,0; 1,1,0; 0,0,0], [0,0,1; 0,0,0; 1,0,1]];
? A = algtbleinit(mt,2);
? algiscommutative(A)
%6 = 1
```

The library syntax is `int algiscommutative(GEN al)`.

3.14.31 `algsdivision`(*al*, {*pl*}). Given a central simple algebra *al* output by `alginit`, tests whether *al* is a division algebra. If *pl* is set, it should be a prime ideal of *K* or an integer between 1 and $r_1 + r_2$, and in that case tests whether *al* is locally a division algebra at the place *pl* instead.

```
? nf = nfinit(y^2-5);
? A = alginit(nf, [-1,y]);
? algsdivision(A, 1)
%3 = 1
? algsdivision(A, 2)
%4 = 0
? algsdivision(A, idealprimedec(nf,2)[1])
%5 = 1
? algsdivision(A, idealprimedec(nf,5)[1])
%6 = 0
? algsdivision(A)
%7 = 1
```

The library syntax is `int algsdivision(GEN al, GEN pl = NULL)`.

3.14.32 `algsdivl`({*al*}, *x*, *y*, {&*z*}). Given two elements *x* and *y* in *al* (Hamilton quaternions if omitted), tests whether *y* is left divisible by *x*, that is whether there exists *z* in *al* such that $xz = y$, and sets *z* to this element if it exists.

```
? A = alginit(nfinit(y), [-1,1]);
? algsdivl(A, [x+2,-x-2]~, [x,1]~)
%2 = 0
? algsdivl(A, [x+2,-x-2]~, [-x,x]~, &z)
%3 = 1
? z
%4 = [Mod(-2/5*x - 1/5, x^2 + 1), 0]~
```

Also accepts square matrices with coefficients in *al*.

The library syntax is `int algsdivl(GEN al = NULL, GEN x, GEN y, GEN *z = NULL)`.

3.14.33 `algsinv`({*al*}, *x*, {&*ix*}). Given an element *x* in *al* (Hamilton quaternions if omitted), tests whether *x* is invertible, and sets *ix* to the inverse of *x*.

```
? A = alginit(nfinit(y), [-1,1]);
? algsinv(A, [-1,1]~)
%2 = 0
? algsinv(A, [1,2]~, &ix)
%3 = 1
? ix
%4 = [Mod(Mod(-1/3, y), x^2 + 1), Mod(Mod(2/3, y), x^2 + 1)]~
```

Also accepts square matrices with coefficients in *al*.

The library syntax is `int algsinv(GEN al = NULL, GEN x, GEN *ix = NULL)`.

3.14.34 `algisramified(al, {pl})`. Given a central simple algebra *al* output by `alginit`, tests whether *al* is ramified, i.e. not isomorphic to a matrix algebra over its center. If *pl* is set, it should be a prime ideal of *K* or an integer between 1 and $r_1 + r_2$, and in that case tests whether *al* is locally ramified at the place *pl* instead.

```
? nf = nfinit(y^2-5);
? A = alginit(nf, [-1,y]);
? algisramified(A, 1)
%3 = 1
? algisramified(A, 2)
%4 = 0
? algisramified(A, idealprimedec(nf,2)[1])
%5 = 1
? algisramified(A, idealprimedec(nf,5)[1])
%6 = 0
? algisramified(A)
%7 = 1
```

The library syntax is `int algisramified(GEN al, GEN pl = NULL)`.

3.14.35 `algissemisimple(al)`. *al* being a table algebra output by `alhtableinit` or a central simple algebra output by `alginit`, tests whether the algebra *al* is semisimple.

```
? mt = [matid(3), [0,0,0;1,0,1;0,0,0], [0,0,0;0,0,0;1,0,1]];
? A = alhtableinit(mt);
? algissemisimple(A)
%3 = 0
? m_i=[0,-1,0,0;1,0,0,0;0,0,0,-1;0,0,1,0]; \\ quaternion algebra (-1,-1)
? m_j=[0,0,-1,0;0,0,0,1;1,0,0,0;0,-1,0,0];
? m_k=[0,0,0,-1;0,0,-1,0;0,1,0,0;1,0,0,0];
? mt = [matid(4), m_i, m_j, m_k];
? A = alhtableinit(mt);
? algissemisimple(A)
%9 = 1
```

The library syntax is `int algissemisimple(GEN al)`.

3.14.36 `algissimple(al, {ss = 0})`. *al* being a table algebra output by `alhtableinit` or a central simple algebra output by `alginit`, tests whether the algebra *al* is simple. If *ss* = 1, assumes that the algebra *al* is semisimple without testing it.

```
? mt = [matid(3), [0,0,0;1,0,1;0,0,0], [0,0,0;0,0,0;1,0,1]];
? A = alhtableinit(mt); \\ matrices [*,*; 0,*]
? algissimple(A)
%3 = 0
? algissimple(A,1) \\ incorrectly assume that A is semisimple
%4 = 1
? m_i=[0,-1,0,0;1,0,0,0;0,0,0,-1;0,0,1,0];
? m_j=[0,0,-1,0;0,0,0,1;1,0,0,0;0,-1,0,0];
? m_k=[0,0,0,-1;0,0,b,0;0,1,0,0;1,0,0,0];
? mt = [matid(4), m_i, m_j, m_k];
```

```
? A = algtableinit(mt); \\ quaternion algebra (-1,-1)
? algissimple(A)
%10 = 1
? mt = [matid(3), [0,0,0; 1,1,0; 0,0,0], [0,0,1; 0,0,0; 1,0,1]];
? A = algtableinit(mt,2); \\ direct product F_4 x F_2
? algissimple(A)
%13 = 0
```

The library syntax is `int algissimple(GEN al, long ss)`.

3.14.37 algissplit(*al*, {*pl*}). Given a central simple algebra *al* output by `alginit`, tests whether *al* is split, i.e. isomorphic to a matrix algebra over its center. If *pl* is set, it should be a prime ideal of *K* or an integer between 1 and $r_1 + r_2$, and in that case tests whether *al* is locally split at the place *pl* instead.

```
? nf = nfinit(y^2-5);
? A = alginit(nf, [-1,y]);
? algissplit(A, 1)
%3 = 0
? algissplit(A, 2)
%4 = 1
? algissplit(A, idealprimedec(nf,2)[1])
%5 = 0
? algissplit(A, idealprimedec(nf,5)[1])
%6 = 1
? algissplit(A)
%7 = 0
```

The library syntax is `int algissplit(GEN al, GEN pl = NULL)`.

3.14.38 alglatadd(*al*, *lat1*, *lat2*, {&*ptinter*}). Given an algebra *al* and two lattices *lat1* and *lat2* in *al*, computes the sum $lat1 + lat2$. If *ptinter* is present, set it to the intersection $lat1 \cap lat2$.

```
? al = alginit(nfinit(y^2+7), [-1,-1]);
? lat1 = alglathnf(al, [1,1,0,0,0,0,0,0]~);
? lat2 = alglathnf(al, [1,0,1,0,0,0,0,0]~);
? latsum = alglatadd(al,lat1,lat2,&latinter);
? matdet(latsum[1])
%5 = 4
? matdet(latinter[1])
%6 = 64
```

The library syntax is `GEN alglatadd(GEN al, GEN lat1, GEN lat2, GEN *ptinter = NULL)`

3.14.39 `alglatcontains`(*al*, *lat*, *x*, {&*ptc*}). Given an algebra *al*, a lattice *lat* and *x* in *al*, tests whether $x \in lat$. If *ptc* is present, sets it to the `t_COL` of coordinates of *x* in the basis of *lat*.

```
? al = alginit(nfinit(y^2+7), [-1,-1]);
? a1 = [1,-1,0,1,2,0,1,2]~;
? lat1 = alglathnf(al,a1);
? alglatcontains(al,lat1,a1,&c)
%4 = 1
? c
%5 = [-1, -2, -1, 1, 2, 0, 1, 1]~
```

The library syntax is `int alglatcontains(GEN al, GEN lat, GEN x, GEN *ptc = NULL)`

3.14.40 `alglatelement`(*al*, *lat*, *c*). Given an algebra *al*, a lattice *lat* and a `t_COL` *c*, returns the element of *al* whose coordinates on the $\mathbf{Z}$ -basis of *lat* are given by *c*.

```
? al = alginit(nfinit(y^2+7), [-1,-1]);
? a1 = [1,-1,0,1,2,0,1,2]~;
? lat1 = alglathnf(al,a1);
? c = [1..8]~;
? elt = alglatelement(al,lat1,c);
? alglatcontains(al,lat1,elt,&c2)
%6 = 1
? c==c2
%7 = 1
```

The library syntax is `GEN alglatelement(GEN al, GEN lat, GEN c)`.

3.14.41 `alglathnf`(*al*, *m*, {*d* = 0}). Given an algebra *al* and a matrix *m* with columns representing elements of *al*, returns the lattice *L* generated by the columns of *m*. If provided, *d* must be a rational number such that *L* contains *d* times the natural basis of *al*. The argument *m* is also allowed to be a `t_VEC` of `t_MAT`, in which case *m* is replaced by the concatenation of the matrices, or a `t_COL`, in which case *m* is replaced by its left multiplication table as an element of *al*.

```
? al = alginit(nfinit(y^2+7), [-1,-1]);
? a = [1,1,-1/2,1,1/3,-1,1,1]~;
? mt = algtomatrix(al,a,1);
? lat = alglathnf(al,mt);
? lat[2]
%5 = 1/6
```

The library syntax is `GEN alglathnf(GEN al, GEN m, GEN d)`.

3.14.42 `alglatindex`(*al*, *lat1*, *lat2*). Given an algebra *al* and two lattices *lat1* and *lat2* in *al*, computes the generalized index of *lat1* relative to *lat2*, i.e. $|lat2/lat1 \cap lat2|/|lat1/lat1 \cap lat2|$.

```
? al = alginit(nfinit(y^2+7), [-1,-1]);
? lat1 = alglathnf(al, [1,1,0,0,0,0,0,0]~);
? lat2 = alglathnf(al, [1,0,1,0,0,0,0,0]~);
? alglatindex(al, lat1, lat2)
%4 = 1
? lat1==lat2
%5 = 0
```

The library syntax is GEN `alglatindex`(GEN *al*, GEN *lat1*, GEN *lat2*).

3.14.43 `alglatinter`(*al*, *lat1*, *lat2*, {&*ptsum*}). Given an algebra *al* and two lattices *lat1* and *lat2* in *al*, computes the intersection $lat1 \cap lat2$. If *ptsum* is present, sets it to the sum $lat1 + lat2$.

```
? al = alginit(nfinit(y^2+7), [-1,-1]);
? lat1 = alglathnf(al, [1,1,0,0,0,0,0,0]~);
? lat2 = alglathnf(al, [1,0,1,0,0,0,0,0]~);
? latinter = alglatinter(al, lat1, lat2, &latsum);
? matdet(latsum[1])
%5 = 4
? matdet(latinter[1])
%6 = 64
```

The library syntax is GEN `alglatinter`(GEN *al*, GEN *lat1*, GEN *lat2*, GEN **ptsum* = NULL)

3.14.44 `alglatlefttransporter`(*al*, *lat1*, *lat2*). Given an algebra *al* and two lattices *lat1* and *lat2* in *al*, computes the left transporter from *lat1* to *lat2*, i.e. the set of $x \in al$ such that $x \cdot lat1 \subset lat2$.

```
? al = alginit(nfinit(y^2+7), [-1,-1]);
? lat1 = alglathnf(al, [1,-1,0,1,2,0,5,2]~);
? lat2 = alglathnf(al, [0,1,-2,-1,0,0,3,1]~);
? tr = alglatlefttransporter(al, lat1, lat2);
? a = alglatelement(al, tr, [0,0,0,0,0,0,1,0]~);
? alglatsubset(al, alglatmul(al, a, lat1), lat2)
%6 = 1
? alglatsubset(al, alglatmul(al, lat1, a), lat2)
%7 = 0
```

The library syntax is GEN `alglatlefttransporter`(GEN *al*, GEN *lat1*, GEN *lat2*).

3.14.45 `alglatmul(al, lat1, lat2)`. Given an algebra al and two lattices $lat1$ and $lat2$ in al , computes the lattice generated by the products of elements of $lat1$ and $lat2$. One of $lat1$ and $lat2$ is also allowed to be an element of al ; in this case, computes the product of the element and the lattice.

```
? al = alginit(nfinit(y^2+7), [-1,-1]);
? a1 = [1,-1,0,1,2,0,1,2]~;
? a2 = [0,1,2,-1,0,0,3,1]~;
? lat1 = alglathnf(al,a1);
? lat2 = alglathnf(al,a2);
? lat3 = alglatmul(al,lat1,lat2);
? matdet(lat3[1])
%7 = 29584
? lat3 == alglathnf(al, algmul(al,a1,a2))
%8 = 0
? lat3 == alglatmul(al, lat1, a2)
%9 = 0
? lat3 == alglatmul(al, a1, lat2)
%10 = 0
```

The library syntax is GEN `alglatmul(GEN al, GEN lat1, GEN lat2)`.

3.14.46 `alglatrightrighttransporter(al, lat1, lat2)`. Given an algebra al and two lattices $lat1$ and $lat2$ in al , computes the right transporter from $lat1$ to $lat2$, i.e. the set of $x \in al$ such that $lat1 \cdot x \subset lat2$.

```
? al = alginit(nfinit(y^2+7), [-1,-1]);
? lat1 = alglathnf(al,matdiagonal([1,3,7,1,2,8,5,2]));
? lat2 = alglathnf(al,matdiagonal([5,3,8,1,9,8,7,1]));
? tr = alglatrightrighttransporter(al,lat1,lat2);
? a = alglatelement(al,tr,[0,0,0,0,0,0,0,1]~);
? alglatsubset(al,alglatmul(al,lat1,a),lat2)
%6 = 1
? alglatsubset(al,alglatmul(al,a,lat1),lat2)
%7 = 0
```

The library syntax is GEN `alglatrightrighttransporter(GEN al, GEN lat1, GEN lat2)`.

3.14.47 `alglatsubset(al, lat1, lat2, {&ptindex})`. Given an algebra al and two lattices $lat1$ and $lat2$ in al , tests whether $lat1 \subset lat2$. If it is true and $ptindex$ is present, sets it to the index of $lat1$ in $lat2$.

```
? al = alginit(nfinit(y^2+7), [-1,-1]);
? lat1 = alglathnf(al,[1,1,0,0,0,0,0,0]~);
? lat2 = alglathnf(al,[1,0,1,0,0,0,0,0]~);
? alglatsubset(al,lat1,lat2)
%4 = 0
? latsum = alglatadd(al,lat1,lat2);
? alglatsubset(al,lat1,latsum,&index)
%6 = 1
? index
%7 = 4
```

The library syntax is `int alglatsubset(GEN al, GEN lat1, GEN lat2, GEN *ptindex = NULL)`.

3.14.48 `algmakeintegral`(*mt*, {*maps* = 0}). *mt* being a multiplication table over $\mathbf{Q}$ in the same format as the input of `algtbleinit`, computes an integral multiplication table *mt2* for an isomorphic algebra. When *maps* = 1, returns a `t_VEC` [*mt2*, *S*, *T*] where *S* and *T* are matrices respectively representing the map from the algebra defined by *mt* to the one defined by *mt2* and its inverse.

```
? mt = [matid(2), [0, -1/4; 1, 0]];
? algtbleinit(mt);
***   at top-level: algtbleinit(mt)
***   ~~~~~
*** algtbleinit: domain error in algtbleinit: denominator(mt) != 1
? mt2 = algmakeintegral(mt);
? al = algtbleinit(mt2);
? algisassociative(al)
%4 = 1
? [mt2, S, T] = algmakeintegral(mt, 1);
? S
%6 =
[1  0]

[0 1/4]
? T
%7 =
[1 0]

[0 4]
? vector(#mt, i, S * (mt * T[,i]) * T) == mt2
%8 = 1
```

The library syntax is `GEN algmakeintegral(GEN mt, long maps)`.

3.14.49 `algmul`({*al*}, *x*, *y*). Given two elements *x* and *y* in *al* (Hamilton quaternions if omitted), computes their product *xy* in the algebra *al*.

```
? A = alginit(nfinit(y), [-1, -1]);
? algmul(A, [1, 1, 0, 0]~, [0, 0, 2, 1]~)
% = [2, 3, 5, -4]~
? algmul(, [1, 2, 3, 4]~, sqrt(I)) \\ Hamilton quaternions
% = [-0.7071067811, 2.1213203435, 4.9497474683, 0.7071067811]~
```

Also accepts matrices with coefficients in *al*.

The library syntax is `GEN algmul(GEN al = NULL, GEN x, GEN y)`.

3.14.50 `algmultable(al)`. Returns a multiplication table of al over its prime subfield ($\mathbf{Q}$ or $\mathbf{F}_p$) or over $\mathbf{R}$ for real algebras, as a `t_VEC` of `t_MAT`: the left multiplication tables of basis elements. If al was output by `algtabinit`, returns the multiplication table used to define al . If al was output by `alginit`, returns the multiplication table of the order $\mathcal{O}_0$ stored in al .

```
? A = alginit(nfinit(y), [-1,-1]);
? M = algmultable(A);
? #M
%3 = 4
? M[1]  \\ multiplication by e_1 = 1
%4 =
[1 0 0 0]
[0 1 0 0]
[0 0 1 0]
[0 0 0 1]
? M[2]
%5 =
[0 -1  1  0]
[1  0  1  1]
[0  0  1  1]
[0  0 -2 -1]
? H = alginit(1.,1/2); \\ Hamilton quaternions
? algmultable(H)[3]  \\ multiplication by j
%7 =
[0  0 -1 0]
[0  0  0 1]
[1  0  0 0]
[0 -1  0 0]
```

The library syntax is `GEN algmultable(GEN al)`.

3.14.51 `algneg({al},x)`. Given an element x in al , computes its opposite $-x$ in the algebra al (Hamilton quaternions if omitted).

```
? A = alginit(nfinit(y), [-1,-1]);
? algneg(A,[1,1,0,0]~)
%2 = [-1, -1, 0, 0]~
```

Also accepts matrices with coefficients in al .

The library syntax is `GEN algneg(GEN al = NULL, GEN x)`.

3.14.52 algnorm($\{al\}, x, \{abs = 0\}$). Given an element x in al (Hamilton quaternions if omitted), computes its norm. If al is a table algebra output by `algtbleinit` or if $abs = 1$, returns the absolute norm of x , which is an element of $\mathbf{F}_p$, $\mathbf{Q}$ or $\mathbf{R}$; if al is omitted or a central simple algebra output by `alginit` and $abs = 0$ (default), returns the reduced norm of x , which is an element of the center of al .

```
? mt = [matid(3), [0,0,0; 1,1,0; 0,0,0], [0,0,1; 0,0,0; 1,0,1]];
? A = algtbleinit(mt,19);
? algnorm(A,[0,-2,3]~)
%3 = 18
? nf = nfinit(y^2-5);
? B = alginit(nf,[-1,y]);
? b = [x,1]~;
? n = algnorm(B,b)
%7 = Mod(-y + 1, y^2 - 5)
? algnorm(B,b,1)
%8 = 16
? nfeltnorm(nf,n)^algdegree(B)
%9 = 16
? algnorm([0,sqrt(3),0,sqrt(2)]~) \\ Hamilton quaternions
%10 = 5.0000000000
```

Also accepts a square matrix with coefficients in al .

The library syntax is `GEN algnorm(GEN al = NULL, GEN x, long abs)`.

3.14.53 algpoleval($\{al\}, T, b$). Given an element b in al (Hamilton quaternions if omitted) and a polynomial T in $K[X]$, computes $T(b)$ in al . Here $K = \mathbf{Q}$ or $\mathbf{F}_p$ for a table algebra (output by `algtbleinit`) and K is the center of al for a central simple algebra (output by `alginit`). Also accepts as input a `t_VEC` $[b, mb]$ where mb is the left multiplication table of b .

```
? nf = nfinit(y^2-5);
? al = alginit(nf,[y,-1]);
? b = [1..8]~;
? pol = algcharpoly(al,b,,1); \\absolute characteristic polynomial
? algpoleval(al,pol,b)==0
%5 = 1
? mb = algtomatrix(al,b,1);
? algpoleval(al,pol,[b,mb]) == 0
%7 = 1
? pol = algcharpoly(al,b); \\reduced characteristic polynomial
? algpoleval(al,pol,b) == 0
%9 = 1
? algpoleval(polcyclo(8),[1,0,0,1]~/sqrt(2)) \\ Hamilton quaternions
%10 = [0.E-38, 0, 0, 0.E-38]~
```

The library syntax is `GEN algpoleval(GEN al = NULL, GEN T, GEN b)`.

3.14.54 algpow($\{al\}, x, n$). Given an element x in al (Hamilton quaternions if omitted) and an integer n , computes the power x^n in the algebra al .

```
? A = alginit(nfinit(y), [-1,-1]);
? algpow(A,[1,1,0,0]~,7)
%2 = [8, -8, 0, 0]~
? algpow([1,2,3,sqrt(3)]~,-3) \\ Hamilton quaternions
% = [-0.0095664563, 0.0052920822, 0.0079381233, 0.0045830776]~
```

Also accepts a square matrix with coefficients in al .

The library syntax is GEN algpow(GEN al = NULL, GEN x, GEN n).

3.14.55 algprimesubalg(al). al being the output of `algtblinit` representing a semisimple algebra of positive characteristic, returns a basis of the prime subalgebra of al . The prime subalgebra of al is the subalgebra fixed by the Frobenius automorphism of the center of al . It is abstractly isomorphic to a product of copies of $\mathbf{F}_p$.

```
? mt = [matid(3), [0,0,0; 1,1,0; 0,0,0], [0,0,1; 0,0,0; 1,0,1]];
? A = algtblinit(mt,2);
? algprimesubalg(A)
%3 =
[1 0]
[0 1]
[0 0]
```

The library syntax is GEN algprimesubalg(GEN al).

3.14.56 algquotient($al, I, \{maps = 0\}$). al being a table algebra output by `algtblinit` and I being a basis of a two-sided ideal of al represented by a matrix, returns the quotient al/I . When $maps = 1$, returns a `t_VEC` [$al/I, proj, lift$] where $proj$ and $lift$ are matrices respectively representing the projection map and a section of it.

```
? mt = [matid(3), [0,0,0; 1,1,0; 0,0,0], [0,0,1; 0,0,0; 1,0,1]];
? A = algtblinit(mt,2);
? AQ = algquotient(A,[0;1;0]);
? algdim(AQ)
%4 = 2
```

The library syntax is GEN alg_quotient(GEN al, GEN I, long maps).

3.14.57 algradical(al). al being a table algebra output by `algtblinit`, returns a basis of the Jacobson radical of the algebra al over its prime field ($\mathbf{Q}$ or $\mathbf{F}_p$).

Here is an example with $A = \mathbf{Q}[x]/(x^2)$, with the basis $(1, x)$:

```
? mt = [matid(2), [0,0;1,0]];
? A = algtblinit(mt);
? algradical(A) \\ = (x)
%3 =
[0]
[1]
```

Another one with 2×2 upper triangular matrices over $\mathbf{Q}$, with basis I_2 , $a = [0, 1; 0, 0]$ and $b = [0, 0; 0, 1]$, such that $a^2 = 0$, $ab = a$, $ba = 0$, $b^2 = b$:

```
? mt = [matid(3), [0,0,0;1,0,1;0,0,0], [0,0,0;0,0,0;1,0,1]];
? A = algtblinit(mt);
? algradical(A) \\ = (a)
%6 =
[0]
[1]
[0]
```

The library syntax is `GEN algradical(GEN al)`.

3.14.58 algramifiedplaces(*al*). Given a central simple algebra *al* output by `alginit`, returns a `t_VEC` containing the list of places of the center of *al* that are ramified in *al*. Each place is described as an integer between 1 and r_1 or as a prime ideal.

```
? nf = nfinit(y^2-5);
? A = alginit(nf, [-1,y]);
? algramifiedplaces(A)
%3 = [1, [2, [2, 0]~, 1, 2, 1]]
```

The library syntax is `GEN algramifiedplaces(GEN al)`.

3.14.59 algrandom(*{al}*, *b*). Given an algebra *al* and a nonnegative integer *b*, returns a random element in *al* with coefficients in $[-b, b]$.

```
? al = alginit(nfinit(y), [-1,-1]);
? algrandom(al,3)
% = [2, 0, 3, -1]~
```

If *al* is an algebra over $\mathbf{R}$ (Hamilton quaternions if omitted) and *b* is a positive `t_REAL`, returns a random element of *al* with coefficients in $[-b, b]$.

```
? algrandom(,1.)
% = [-0.1806334680, -0.2810504190, 0.5011479961, 0.9498643737]~
```

The library syntax is `GEN algrandom(GEN al = NULL, GEN b)`.

3.14.60 algremltable(*al*). Given a central simple algebra *al* output by `alginit` defined by a multiplication table over its center (a number field), returns this multiplication table.

```
? nf = nfinit(y^3-5); a = y; b = y^2;
? {m_i = [0,a,0,0;
          1,0,0,0;
          0,0,0,a;
          0,0,1,0];}
? {m_j = [0, 0,b, 0;
          0, 0,0,-b;
          1, 0,0, 0;
          0,-1,0, 0];}
? {m_k = [0, 0,0,-a*b;
          0, 0,b, 0;]
```

```

0,-a,0, 0;
1, 0,0, 0];}
? mt = [matid(4), m_i, m_j, m_k];
? A = alginit(nf,mt,'x');
? M = algrelmultable(A);
? M[2] == m_i
%8 = 1
? M[3] == m_j
%9 = 1
? M[4] == m_k
%10 = 1

```

The library syntax is `GEN algrelmultable(GEN al)`.

3.14.61 `algsimpledec(al, {maps = 0})`. al being the output of `algtbleinit`, returns a `t_VEC` $[J, [al_1, \dots, al_n]]$ where J is a basis of the Jacobson radical of al and al/J is isomorphic to the direct product of the simple algebras al_i . When $maps = 1$, each al_i is replaced with a `t_VEC` $[al_i, proj_i, lift_i]$ where $proj_i$ and $lift_i$ are matrices respectively representing the projection map $al \rightarrow al_i$ and a section of it. Modulo J , the images of the $lift_i$ form a direct sum in al/J , so that the images of $1 \in al_i$ under $lift_i$ are central primitive idempotents of al/J . The factors are sorted by increasing dimension, then increasing dimension of the center. This ensures that the ordering of the isomorphism classes of the factors is deterministic over finite fields, but not necessarily over $\mathbf{Q}$.

The library syntax is `GEN algsimpledec(GEN al, long maps)`.

3.14.62 `algsplit(al, {v = 'x'})`. If al is a table algebra over $\mathbf{F}_p$ output by `algtbleinit` that represents a simple algebra, computes an isomorphism between al and a matrix algebra $M_d(\mathbf{F}_{p^n})$ where $N = nd^2$ is the dimension of al . Returns a `t_VEC` $[map, mapi]$, where:

- map is a `t_VEC` of N matrices of size $d \times d$ with `t_FFELT` coefficients using the variable v , representing the image of the basis of al under the isomorphism.

- $mapi$ is an $N \times N$ matrix with `t_INT` coefficients, representing the image in al by the inverse isomorphism of the basis (b_i) of $M_d(\mathbf{F}_p[\alpha])$ (where α has degree n over $\mathbf{F}_p$) defined as follows: let $E_{i,j}$ be the matrix having all coefficients 0 except the (i,j) -th coefficient equal to 1, and define

$$b_{i_3+n(i_2+di_1)+1} = E_{i_1+1, i_2+1} \alpha^{i_3},$$

where $0 \leq i_1, i_2 < d$ and $0 \leq i_3 < n$.

Example:

```

? al0 = alginit(nfinit(y^2+7), [-1,-1]);
? al = algtbleinit(algmtable(al0), 3); \\ isomorphic to M_2(F_9)
? [map,mapi] = algsplit(al, 'a');
? x = [1,2,1,0,0,0,0,0]~; fx = map*x
%4 =
[2*a 0]
[ 0 2]
? y = [0,0,0,0,1,0,0,1]~; fy = map*y
%5 =
[1 2*a]

```

```

[2 a + 2]
? map*algmul(al,x,y) == fx*fy
%6 = 1
? map*mapi[,6]
%7 =
[0 0]
[a 0]

```

Warning. If al is not simple, `algsplit(al)` can trigger an error, but can also run into an infinite loop. Example:

```

? al = alginit(nfinit(y),[-1,-1]); \\ ramified at 2
? al2 = algtableinit(algmultable(al),2); \\ maximal order modulo 2
? algsplit(al2); \\ not semisimple, infinite loop

```

The library syntax is `GEN algsplit(GEN al, long v = -1)` where v is a variable number.

3.14.63 algsplittingdata(al). Given a central simple algebra al output by `alginit` defined by a multiplication table over its center K (a number field), returns data stored to compute a splitting of al over an extension. This data is a `t_VEC` `[t,Lbas,Lbasinv]` with 3 components:

- an element t of al such that $L = K(t)$ is a maximal subfield of al ;
- a matrix `Lbas` expressing a L -basis of al (given an L -vector space structure by multiplication on the right) on the integral basis of al ;
- a matrix `Lbasinv` expressing the integral basis of al on the previous L -basis.

```

? nf = nfinit(y^3-5); a = y; b = y^2;
? {m_i = [0,a,0,0;
          1,0,0,0;
          0,0,0,a;
          0,0,1,0];}
? {m_j = [0, 0,b, 0;
          0, 0,0,-b;
          1, 0,0, 0;
          0,-1,0, 0];}
? {m_k = [0, 0,0,-a*b;
          0, 0,b, 0;
          0,-a,0, 0;
          1, 0,0, 0];}
? mt = [matid(4), m_i, m_j, m_k];
? A = alginit(nf,mt,'x);
? [t,Lb,Lbi] = algsplittingdata(A);
? t
%8 = [0, 0, 0, 1, 0, 0, 0, 0, 0, 0, 0, 0]~;
? matsize(Lb)
%9 = [12, 2]
? matsize(Lbi)
%10 = [2, 12]

```

The library syntax is `GEN algsplittingdata(GEN al)`.

3.14.64 algsplittingfield(*al*). Given a central simple algebra *al* output by `alginit`, returns an `rnf` structure: the splitting field of *al* that is stored in *al*, as a relative extension of the center.

```
nf = nfinit(y^3-5);
a = y; b = y^2;
{m_i = [0,a,0,0;
        1,0,0,0;
        0,0,0,a;
        0,0,1,0];}
{m_j = [0, 0,b, 0;
        0, 0,0,-b;
        1, 0,0, 0;
        0,-1,0, 0];}
{m_k = [0, 0,0,-a*b;
        0, 0,b,  0;
        0,-a,0,  0;
        1, 0,0,  0];}
mt = [matid(4), m_i, m_j, m_k];
A = alginit(nf,mt,'x');
algsplittingfield(A).pol
%8 = x^2 - y
```

The library syntax is `GEN algsplittingfield(GEN al)`.

3.14.65 algsqr(*{al}*, *x*). Given an element *x* in *al* (Hamilton quaternions if omitted), computes its square x^2 in the algebra *al*.

```
? A = alginit(nfinit(y), [-1,-1]);
? algsqr(A,[1,0,2,0]~)
%2 = [-3, 0, 4, 0]~
? algsqr([0,0,0,Pi]~) \\ Hamilton quaternions
%3 = [-9.8696044010, 0, 0, 0]~
```

Also accepts a square matrix with coefficients in *al*.

The library syntax is `GEN algsqr(GEN al = NULL, GEN x)`.

3.14.66 algsub(*{al}*, *x*, *y*). Given two elements *x* and *y* in *al* (Hamilton quaternions if omitted), computes their difference $x - y$ in the algebra *al*.

```
? A = alginit(nfinit(y), [-1,-1]);
? algsub(A,[1,1,0,0]~,[1,0,1,0]~)
%2 = [0, 1, -1, 0]~
```

Also accepts matrices with coefficients in *al*.

If *x* and *y* are given in the same format, then one should simply use `-` instead of `algsub`.

The library syntax is `GEN algsub(GEN al = NULL, GEN x, GEN y)`.

3.14.67 `algsubalg(al, B)`. *al* being a table algebra output by `algtblinit` and *B* being a basis of a subalgebra of *al* represented by a matrix, computes an algebra *al2* isomorphic to *B*.

Returns [*al2*, *B2*] where *B2* is a possibly different basis of the subalgebra *al2*, with respect to which the multiplication table of *al2* is defined.

```
? mt = [matid(3), [0,0,0; 1,1,0; 0,0,0], [0,0,1; 0,0,0; 1,0,1]];
? A = algtblinit(mt,2);
? B = algsubalg(A,[1,0; 0,0; 0,1]);
? alldim(A)
%4 = 3
? alldim(B[1])
%5 = 2
? m = matcompanion(x^4+1);
? mt = [m^i | i <- [0..3]];
? al = algtblinit(mt);
? B = [1,0;0,0;0,1/2;0,0];
? al2 = algsubalg(al,B);
? alldim(al2[1])
? al2[2]
%13 =
[1 0]
[0 0]
[0 1]
[0 0]
```

The library syntax is `GEN algsubalg(GEN al, GEN B)`.

3.14.68 `algtblinit(mt, {p = 0})`. Initializes the associative algebra over $K = \mathbf{Q}$ (*p* omitted) or $\mathbf{F}_p$ defined by the multiplication table *mt*. As a *K*-vector space, the algebra is generated by a basis $(e_1 = 1, e_2, \dots, e_n)$; the table is given as a `t_VEC` of *n* matrices in $M_n(K)$, giving the left multiplication by the basis elements e_i , in the given basis. Assumes that $e_1 = 1$, that $Ke_1 \oplus \dots \oplus Ke_n$ describes an associative algebra over *K*, and in the case $K = \mathbf{Q}$ that the multiplication table is integral. If the algebra is already known to be central and simple, then the case $K = \mathbf{F}_p$ is useless, and one should use `algin` directly.

The point of this function is to input a finite dimensional *K*-algebra, so as to later compute its radical, then to split the quotient algebra as a product of simple algebras over *K*.

The pari object representing such an algebra *A* is a `t_VEC` with the following data:

- The characteristic of *A*, accessed with `algchar`.
- The multiplication table of *A*, accessed with `algmtable`.
- The traces of the elements of the basis.

A simple example: the 2×2 upper triangular matrices over $\mathbf{Q}$, generated by I_2 , $a = [0, 1; 0, 0]$ and $b = [0, 0; 0, 1]$, such that $a^2 = 0$, $ab = a$, $ba = 0$, $b^2 = b$:

```
? mt = [matid(2), [0,0,0;1,0,1;0,0,0], [0,0,0;0,0,0;1,0,1]];
? A = algtblinit(mt);
? algradical(A) \\ = (a)
```

```

%6 =
[0]
[1]
[0]
? algcenter(A) \\ = (I_2)
%7 =
[1]
[0]
[0]

```

The library syntax is `GEN algtableinit(GEN mt, GEN p = NULL)`.

3.14.69 algtensor(*al1*, *al2*, {*flag* = 3}). Given two algebras *al1* and *al2*, computes their tensor product. *flag* has the same meaning as in `alginit`.

Currently only implemented for cyclic algebras of coprime degree over the same center K , and the tensor product is over K .

The library syntax is `GEN algtensor(GEN al1, GEN al2, long flag)`.

3.14.70 algtomatrix({*al*}, *x*, {*abs* = 0}). Given an element x in *al* (Hamilton quaternions if omitted), returns the image of x under a homomorphism to a matrix algebra. If *al* is a table algebra output by `algtableinit` or if *abs* = 1, returns the left multiplication table on the integral basis; if *al* is a central simple algebra and *abs* = 0, returns $\phi(x)$ where $\phi : A \otimes_K L \rightarrow M_d(L)$ (where d is the degree of the algebra and L is an extension of K with $[L : K] = d$) is an isomorphism stored in *al*. Also accepts a square matrix with coefficients in *al*.

```

? A = alginit(nfinit(y), [-1,-1]);
? algtomatrix(A,[0,0,0,2]~)
%2 =
[Mod(x + 1, x^2 + 1) Mod(Mod(1, y)*x + Mod(-1, y), x^2 + 1)]
[Mod(x + 1, x^2 + 1) Mod(-x + 1, x^2 + 1)]
? algtomatrix(A,[0,1,0,0]~,1)
%2 =
[0 -1 1 0]
[1 0 1 1]
[0 0 1 1]
[0 0 -2 -1]
? algtomatrix(A,[0,x]~,1)
%3 =
[-1 0 0 -1]
[-1 0 1 0]
[-1 -1 0 -1]
[ 2 0 0 1]
? algtomatrix(,[1,2,3,4]~) \\ Hamilton quaternions
%4 =
[1 + 2*I -3 - 4*I]

```

```

[3 - 4*I  1 - 2*I]
? algtomatrix(,I,1)
%5 =
[0 -1 0  0]
[1  0 0  0]
[0  0 0 -1]
[0  0 1  0]

```

Also accepts matrices with coefficients in *al*.

The library syntax is GEN algtomatrix(GEN al = NULL, GEN x, long abs).

3.14.71 algtrace({*al*}, *x*, {*abs* = 0}). Given an element *x* in *al* (Hamilton quaternions if omitted), computes its trace. If *al* is a table algebra output by `algtabinit` or if *abs* = 1, returns the absolute trace of *x*, which is an element of $\mathbf{F}_p$, $\mathbf{Q}$ or $\mathbf{R}$; if *al* is omitted or the output of `alginit` and *abs* = 0 (default), returns the reduced trace of *x*, which is an element of the center of *al*.

```

? A = alginit(nfinit(y), [-1,-1]);
? algtrace(A,[5,0,0,1]~)
%2 = 11
? algtrace(A,[5,0,0,1]~,1)
%3 = 22
? nf = nfinit(y^2-5);
? A = alginit(nf,[-1,y]);
? a = [1+x+y,2*y]~*Mod(1,y^2-5)*Mod(1,x^2+1);
? t = algtrace(A,a)
%7 = Mod(2*y + 2, y^2 - 5)
? algtrace(A,a,1)
%8 = 8
? algdegree(A)*nfeltttrace(nf,t)
%9 = 8
? algtrace(,[1.,2,3,4]~) \\ Hamilton quaternions
%10 = 2.0000000000
? algtrace(,[1.,2,3,4]~,0)
%11 = 4.0000000000

```

Also accepts a square matrix with coefficients in *al*.

The library syntax is GEN algtrace(GEN al = NULL, GEN x, long abs).

3.14.72 algtype(*al*). Given an algebra *al* output by `alginit` or by `algtabinit`, returns an integer indicating the type of algebra:

- 0: not a valid algebra.
- 1: table algebra output by `algtabinit`.
- 2: central simple algebra output by `alginit` and represented by a multiplication table over its center.
- 3: central simple algebra output by `alginit` and represented by a cyclic algebra.
- 4: division algebra over $\mathbf{R}$ ($\mathbf{R}$, $\mathbf{C}$ or Hamilton quaternion algebra $\mathbf{H}$).

```

? algtype([])
%1 = 0
? mt = [matid(3), [0,0,0; 1,1,0; 0,0,0], [0,0,1; 0,0,0; 1,0,1]];
? A = algtableinit(mt,2);
? algtype(A)
%4 = 1
? nf = nfinit(y^3-5);
? a = y; b = y^2;
? {m_i = [0,a,0,0;
          1,0,0,0;
          0,0,0,a;
          0,0,1,0];}
? {m_j = [0, 0,b, 0;
          0, 0,0,-b;
          1, 0,0, 0;
          0,-1,0, 0];}
? {m_k = [0, 0,0,-a*b;
          0, 0,b, 0;
          0,-a,0, 0;
          1, 0,0, 0];}
? mt = [matid(4), m_i, m_j, m_k];
? A = alginit(nf,mt,'x');
? algtype(A)
%12 = 2
? A = alginit(nfinit(y), [-1,-1]);
? algtype(A)
%14 = 3
? H = alginit(1.,1/2);
? algtype(H)
%16 = 4

```

The library syntax is `long algtype(GEN al)`.

3.15 Elliptic curves.

3.15.1 Elliptic curve structures. An elliptic curve is given by a Weierstrass model

$$y^2 + a_1xy + a_3y = x^3 + a_2x^2 + a_4x + a_6,$$

whose discriminant is nonzero. One can also define an elliptic curve with a

$$y^2 = x^3 + a_4x + a_6.$$

Affine points on E are represented as two-component vectors $[x,y]$; the point at infinity, i.e. the identity element of the group law, is represented by the one-component vector $[0]$.

Given a vector of coefficients $[a_1, a_2, a_3, a_4, a_6]$ or $[a_4, a_6]$, the function `ellinit` initializes and returns an *ell* structure. An additional optional argument allows to specify the base field in case it cannot be inferred from the curve coefficients. This structure contains data needed by elliptic curve related functions, and is generally passed as a first argument. Expensive data are skipped on initialization: they will be dynamically computed when (and if) needed, and then inserted in the structure. The precise layout of the *ell* structure is left undefined and should never be used directly. The following member functions are available, depending on the underlying domain.

All domains.

- **a1, a2, a3, a4, a6**: coefficients of the elliptic curve.
- **b2, b4, b6, b8**: b -invariants of the curve; in characteristic $\neq 2$, for $Y = 2y + a_1x + a_3$, the curve equation becomes

$$Y^2 = 4x^3 + b_2x^2 + 2b_4x + b_6 =: g(x).$$

- **c4, c6**: c -invariants of the curve; in characteristic $\neq 2, 3$, for $X = x + b_2/12$ and $Y = 2y + a_1x + a_3$, the curve equation becomes

$$Y^2 = 4X^3 - (c_4/12)X - (c_6/216).$$

- **disc**: discriminant of the curve. This is only required to be nonzero, not necessarily a unit.
- **j**: j -invariant of the curve.

These are used as follows:

```
? E = ellinit([0,0,0, a4,a6]);  
? E.b4  
%2 = 2*a4  
? E.disc  
%3 = -64*a4^3 - 432*a6^2
```

Curves over $\mathbf{C}$.

This in particular includes curves defined over $\mathbf{Q}$. All member functions in this section return data, as it is currently stored in the structure, if present; and otherwise compute it to the default accuracy, that was fixed *at the time of ellinit* (via a `t_REAL D` domain argument, or `realprecision` by default). The function `ellperiods` allows to recompute (and cache) the following data to *current realprecision*.

- **area**: volume of the complex lattice defining E .
- **roots** is a vector whose three components contain the complex roots of the right hand side $g(x)$ of the attached b -model $Y^2 = g(x)$. If the roots are all real, they are ordered by decreasing value. If only one is real, it is the first component.
- **omega**: $[\omega_1, \omega_2]$, periods forming a basis of the complex lattice defining E . The first component ω_1 is the (positive) real period, in other words the integral of the Néron differential $dx/(2y + a_1x + a_3)$ over the connected component of the identity component of $E(\mathbf{R})$. The second component ω_2 is a complex period, such that $\tau = \omega_1/\omega_2$ belongs to Poincaré's half-plane (positive imaginary part); not necessarily to the standard fundamental domain. It is normalized so that $\Im(\omega_2) < 0$ and either $\Re(\omega_2) = 0$, when `E.disc` > 0 ($E(\mathbf{R})$ has two connected components), or $\Re(\omega_2) = \omega_1/2$.
- **eta** is a row vector containing the quasi-periods η_1 and η_2 such that $\eta_i = 2\zeta(\omega_i/2)$, where ζ is the Weierstrass zeta function attached to the period lattice; see `ellzeta`. In particular, the Legendre relation holds: $\eta_2\omega_1 - \eta_1\omega_2 = 2\pi i$.

Warning. As for the orientation of the basis of the period lattice, beware that many sources use the inverse convention where ω_2/ω_1 has positive imaginary part and our ω_2 is the conjugate of theirs. Our convention $\tau = \omega_1/\omega_2$ ensures that the action of PSL_2 is the natural one:

$$[a, b; c, d] \cdot \tau = (a\tau + b)/(c\tau + d) = (a\omega_1 + b\omega_2)/(c\omega_1 + d\omega_2),$$

instead of a twisted one. (Our τ is $-1/\tau$ in the above inverse convention.)

Curves over $\mathbf{Q}_p$.

We advise to input a model defined over $\mathbf{Q}$ for such curves. In any case, if you input an approximate model with `t_PADIC` coefficients, it will be replaced by a lift to $\mathbf{Q}$ (an exact model “close” to the one that was input) and all quantities will then be computed in terms of this lifted model.

For the time being only curves with multiplicative reduction (split or nonsplit), i.e. $v_p(j) < 0$, are supported by nontrivial functions. In this case the curve is analytically isomorphic to $\bar{\mathbf{Q}}_p^*/q^{\mathbf{Z}} := E_q(\bar{\mathbf{Q}}_p)$, for some p -adic integer q (the Tate period). In particular, we have $j(q) = j(E)$.

- `p` is the residual characteristic
- `roots` is a vector with a single component, equal to the p -adic root e_1 of the right hand side $g(x)$ of the attached b -model $Y^2 = g(x)$. The point $(e_1, 0)$ corresponds to $-1 \in \bar{\mathbf{Q}}_p^*/q^{\mathbf{Z}}$ under the Tate parametrization.
- `tate` returns $[u^2, u, q, [a, b], Ei, L]$ in the notation of Henniart-Mestre (CRAS t. 308, p. 391–395, 1989): q is as above, $u \in \mathbf{Q}_p(\sqrt{-c_6})$ is such that $\phi^*dx/(2y + a_1x + a_3) = udt/t$, where $\phi : E_q \rightarrow E$ is an isomorphism (well defined up to sign) and dt/t is the canonical invariant differential on the Tate curve; $u^2 \in \mathbf{Q}_p$ does not depend on ϕ . (Technicality: if $u \notin \mathbf{Q}_p$, it is stored as a quadratic `t_POLMOD`.) The parameters $[a, b]$ satisfy $4u^2b \cdot \mathrm{agm}(\sqrt{a/b}, 1)^2 = 1$ as in Theorem 2 (*loc. cit.*). `Ei` describes the sequence of 2-isogenous curves (with kernel generated by $[0, 0]$) $E_i : y^2 = x(x + A_i)(x + A_i - B_i)$ converging quadratically towards the singular curve E_∞ . Finally, L is Mazur-Tate-Teitelbaum’s $\mathcal{L}$ -invariant, equal to $\log_p q/v_p(q)$.

Curves over $\mathbf{F}_q$.

- `p` is the characteristic of $\mathbf{F}_q$.
- `no` is $\#E(\mathbf{F}_q)$.
- `cyc` gives the cycle structure of $E(\mathbf{F}_q)$.
- `gen` returns the generators of $E(\mathbf{F}_q)$.
- `group` returns `[no, cyc, gen]`, i.e. $E(\mathbf{F}_q)$ as an abelian group structure.

Curves over $\mathbf{Q}$.

All functions should return a correct result, whether the model is minimal or not, but it is a good idea to stick to minimal models whenever $\gcd(c_4, c_6)$ is easy to factor (minor speed-up). The construction

```
E = ellminimalmodel(E0, &v)
```

replaces the original model E_0 by a minimal model E , and the variable change v allows to go between the two models:

```
ellchangept(P0, v)
ellchangeptinv(P, v)
```

respectively map the point P_0 on E_0 to its image on E , and the point P on E to its pre-image on E_0 .

A few routines — namely `ellgenerators`, `ellidentify`, `ellsearch`, `forell` — require the optional package `elldata` (John Cremona's database) to be installed. In that case, the function `ellinit` will allow alternative inputs, e.g. `ellinit("11a1")`. Functions using this package need to load chunks of a large database in memory and require at least 2MB stack to avoid stack overflows.

- `gen` returns the generators of $E(\mathbf{Q})$, if known (from John Cremona's database)

Curves over number fields.

- `nf` return the *nf* structure attached to the number field over which E is defined.
- `bnf` return the *bnf* structure attached to the number field over which E is defined or raise an error (if only an *nf* is available).
- `omega`, `eta`, `area`: vectors of complex periods, quasi-periods and lattice areas attached to the complex embeddings of E , in the same order as `E.nf.roots`.

3.15.2 Reduction. Let E be a curve defined over $\mathbf{Q}_p$ given by a p -integral model; if the curve has good reduction at p , we may define its reduction $\tilde{E}$ over the finite field $\mathbf{F}_p$:

```
? E = ellinit([-3,1], 0(5^10));  \ \  E/\mathbf{Q}_5
? Et = ellinit(E, 5)
? ellcard(Et)  \ \  \tilde{E}/\mathbf{F}_5 has 7 points
%3 = 7
? ellinit(E, 7)
*** at top-level: ellinit(E,7)
*** ^-----
*** ellinit: inconsistent moduli in ellinit: 5 != 7
```

Likewise, if a curve is defined over a number field K and $\mathfrak{p}$ is a maximal ideal with finite residue field $\mathbf{F}_q$, we define the reduction $\tilde{E}/\mathbf{F}_q$ provided E has good reduction at $\mathfrak{p}$. $E/\mathbf{Q}$ is an important special case:

```
? E = ellinit([-3,1]);
? factor(E.disc)
%2 =
[2 4]
[3 4]
```

```
? Et = ellinit(E, 5);
? ellcard(Et) \\  $\tilde{E} / \mathbf{F}_5$  has 7 points
%4 = 7
? ellinit(E, 3) \\ bad reduction at 3
%5 = []
```

General number fields are similar:

```
? K = nfinit(x^2+1); E = ellinit([x,x+1], K);
? idealfactor(K, E.disc) \\ three primes of bad reduction
%2 =
[ [2, [1, 1]~, 2, 1, [1, -1; 1, 1]] 10]
[ [5, [-2, 1]~, 1, 1, [2, -1; 1, 2]] 2]
[[5, [2, 1]~, 1, 1, [-2, -1; 1, -2]] 2]
? P = idealprimedec(K, 3); \\ a prime of good reduction
? ideálnorm(K, P)
%4 = 9
? Et = ellinit(E, P);
? ellcard(Et) \\  $\tilde{E} / \mathbf{F}_9$  has 4 points
%6 = 4
```

If the model is not locally minimal at $\mathfrak{p}$, the above will fail: `elllocalred` and `ellchangecurve` allow to reduce to that case.

Some functions such as `ellap`, `ellcard`, `ellgroup` and `ellissupersingular` even implicitly replace the given equation by a local minimal model and consider the group of nonsingular points $\tilde{E}^{ns}$ so they make sense even when the curve has bad reduction.

3.15.3 `ell2cover(E)`. If E is an elliptic curve over $\mathbf{Q}$, returns a basis of the set of everywhere locally soluble 2-covers of the curve E . For each cover a pair $[R, P]$ is returned where $y^2 - R(x)$ is a quartic curve and P is a point on $E(k)$, where $k = \mathbf{Q}(x)[y]/(y^2 - R(x))$. E can also be given as the output of `ellrankinit(E)`, or as a pair $[e, f]$, where e is an elliptic curve given by `ellrankinit` and f is a quadratic twist of e . We then look for points on f .

```
? E = ellinit([-25,4]);
? C = ell2cover(E); #C
%2 = 2
? [R,P] = C[1]; R
%3 = 64*x^4+480*x^2-128*x+100
? P[1]
%4 = -320/y^2*x^4 + 256/y^2*x^3 + 800/y^2*x^2 - 320/y^2*x - 436/y^2
? ellisoncurve(E, Mod(P, y^2-R))
%5 = 1
? H = hyperellratpoints(R,10)
%6 = [[0,10], [0,-10], [1/5,242/25], [1/5,-242/25], [2/5,282/25],
      [2/5,-282/25]]
? A = substvec(P, [x,y], H[1])
%7 = [-109/25, 686/125]
```

The library syntax is `GEN ell2cover(GEN E, long prec)`.

3.15.4 ellL1($E, \{r = 0\}$). Returns the value at $s = 1$ of the derivative of order r of the L -function of the elliptic curve $E/\mathbf{Q}$.

```
? E = ellinit("11a1"); \\ order of vanishing is 0
? ellL1(E)
%2 = 0.2538418608559106843377589233
? E = ellinit("389a1"); \\ order of vanishing is 2
? ellL1(E)
%4 = -5.384067311837218089235032414 E-29
? ellL1(E, 1)
%5 = 0
? ellL1(E, 2)
%6 = 1.518633000576853540460385214
```

The main use of this function, after computing at *low* accuracy the order of vanishing using `ellanalyticrank`, is to compute the leading term at *high* accuracy to check (or use) the Birch and Swinnerton-Dyer conjecture:

```
? \p18
  realprecision = 18 significant digits
? E = ellinit("5077a1"); ellanalyticrank(E)
time = 8 ms.
%1 = [3, 10.3910994007158041]
? \p200
  realprecision = 202 significant digits (200 digits displayed)
? ellL1(E, 3)
time = 104 ms.
%3 = 10.3910994007158041387518505103609170697263563756570092797[...]
```

Analogous and more general functionalities for E defined over general number fields are available through `lfun`.

The library syntax is `GEN ellL1(GEN E, long r, long bitprec)`.

3.15.5 elladd($E, z1, z2$). Sum of the points $z1$ and $z2$ on the elliptic curve corresponding to E .

The library syntax is `GEN elladd(GEN E, GEN z1, GEN z2)`.

3.15.6 ellak(E, n). Computes the coefficient a_n of the L -function of the elliptic curve $E/\mathbf{Q}$, i.e. coefficients of a newform of weight 2 by the modularity theorem (Taniyama-Shimura-Weil conjecture). E must be an `ell` structure over $\mathbf{Q}$ as output by `ellinit`. E must be given by an integral model, not necessarily minimal, although a minimal model will make the function faster.

```
? E = ellinit([1,-1,0,4,3]);
? ellak(E, 10)
%2 = -3
? e = ellchangecurve(E, [1/5,0,0,0]); \\ made not minimal at 5
? ellak(e, 10) \\ wasteful but works
%3 = -3
? E = ellminimalmodel(e); \\ now minimal
? ellak(E, 5)
%5 = -3
```

If the model is not minimal at a number of bad primes, then the function will be slower on those n divisible by the bad primes. The speed should be comparable for other n :

```
? for(i=1,10^6, ellak(E,5))
time = 699 ms.
? for(i=1,10^6, ellak(e,5)) \\ 5 is bad, markedly slower
time = 1,079 ms.

? for(i=1,10^5,ellak(E,5*i))
time = 1,477 ms.
? for(i=1,10^5,ellak(e,5*i)) \\ still slower but not so much on average
time = 1,569 ms.
```

The library syntax is `GEN akell(GEN E, GEN n)`.

3.15.7 `ellan`(E, n). Computes the vector of the first n Fourier coefficients a_k corresponding to the elliptic curve E defined over a number field. If E is defined over $\mathbf{Q}$, the curve may be given by an arbitrary model, not necessarily minimal, although a minimal model will make the function faster. Over a more general number field, the model must be locally minimal at all primes above 2 and 3.

The library syntax is `GEN ellan(GEN E, long n)`. Also available is `GEN ellanQ_zv(GEN e, long n)`, which returns a `t_VECSMALL` instead of a `t_VEC`, saving on memory.

3.15.8 `ellanalyticrank`($E, \{eps\}$). Returns the order of vanishing at $s = 1$ of the L -function of the elliptic curve $E/\mathbf{Q}$ and the value of the first nonzero derivative. To determine this order, it is assumed that any value less than `eps` is zero. If `eps` is omitted, $2^{-b/2}$ is used, where b is the current bit precision.

```
? E = ellinit("11a1"); \\ rank 0
? ellanalyticrank(E)
%2 = [0, 0.2538418608559106843377589233]
? E = ellinit("37a1"); \\ rank 1
? ellanalyticrank(E)
%4 = [1, 0.3059997738340523018204836835]
? E = ellinit("389a1"); \\ rank 2
? ellanalyticrank(E)
%6 = [2, 1.518633000576853540460385214]
? E = ellinit("5077a1"); \\ rank 3
? ellanalyticrank(E)
%8 = [3, 10.39109940071580413875185035]
```

Analogous and more general functionalities for E defined over general number fields are available through `lfun` and `lfunorderzero`.

The library syntax is `GEN ellanalyticrank(GEN E, GEN eps = NULL, long bitprec)`.

3.15.9 ellap($E, \{p\}$). Let E be an `ell` structure as output by `ellinit`, attached to an elliptic curve E/K . If the field $K = \mathbf{F}_q$ is finite, return the trace of Frobenius t , defined by the equation $\#E(\mathbf{F}_q) = q + 1 - t$.

For other fields of definition and p defining a finite residue field $\mathbf{F}_q$, return the trace of Frobenius for the reduction of E : the argument p is best left omitted if $K = \mathbf{Q}_\ell$ (else we must have $p = \ell$) and must be a prime number ($K = \mathbf{Q}$) or prime ideal (K a general number field) with residue field $\mathbf{F}_q$ otherwise. The equation need not be minimal or even integral at p ; of course, a minimal model will be more efficient.

For a number field K , the trace of Frobenius is the a_p coefficient in the Euler product defining the curve L -series, whence the function name:

$$L(E/K, s) = \prod_{\text{bad } p} (1 - a_p(Np)^{-s})^{-1} \prod_{\text{good } p} (1 - a_p(Np)^{-s} + (Np)^{1-2s})^{-1}.$$

When the characteristic of the finite field is large, the availability of the `seadata` package will speed up the computation.

```
? E = ellinit([0,1]); \\ y^2 = x^3 + 0.x + 1, defined over Q
? ellap(E, 7) \\ 7 necessary here
%2 = -4      \\ #E(F_7) = 7+1-(-4) = 12
? ellcard(E, 7)
%3 = 12      \\ OK

? E = ellinit([0,1], 11); \\ defined over F_11
? ellap(E)      \\ no need to repeat 11
%4 = 0
? ellap(E, 11)  \\ ... but it also works
%5 = 0
? ellgroup(E, 13) \\ ouch, inconsistent input!
*** at top-level: ellap(E,13)
***      ^-----
*** ellap: inconsistent moduli in Rg_to_Fp:
    11
    13
? a = ffgen(ffinit(11,3), 'a); \\ defines F_q := F_{11^3}
? E = ellinit([a+1,a]); \\ y^2 = x^3 + (a+1)x + a, defined over F_q
? ellap(E)
%8 = -3
```

If the curve is defined over a more general number field than $\mathbf{Q}$, the maximal ideal p must be explicitly given in `idealprimedec` format. There is no assumption of local minimality at p .

```
? K = nfinit(a^2+1); E = ellinit([1+a,0,1,0,0], K);
? fa = idealfactor(K, E.disc)
%2 =
[ [5, [-2, 1]~, 1, 1, [2, -1; 1, 2]] 1]
[[13, [5, 1]~, 1, 1, [-5, -1; 1, -5]] 2]
? ellap(E, fa[1,1])
%3 = -1 \\ nonsplit multiplicative reduction
```

```

? ellap(E, fa[2,1])
%4 = 1 \\ split multiplicative reduction
? P17 = idealprimedec(K,17)[1];
? ellap(E, P17)
%6 = 6 \\ good reduction
? E2 = ellchangeurve(E, [17,0,0,0]);
? ellap(E2, P17)
%8 = 6 \\ same, starting from a nonminimal model

? P3 = idealprimedec(K,3)[1];
? ellap(E, P3) \\ OK: E is minimal at P3
%10 = -2
? E3 = ellchangeurve(E, [3,0,0,0]);
? ellap(E3, P3) \\ not integral at P3
*** at top-level: ellap(E3,P3)
*** ^-----
*** ellap: impossible inverse in Rg_to_ff: Mod(0, 3).

```

Algorithms used. If $E/\mathbf{F}_q$ has CM by a principal imaginary quadratic order we use a fast explicit formula (involving essentially Kronecker symbols and Cornacchia's algorithm), in $O(\log q)^2$ bit operations. Otherwise, we use Shanks-Mestre's baby-step/giant-step method, which runs in time $\tilde{O}(q^{1/4})$ using $\tilde{O}(q^{1/4})$ storage, hence becomes unreasonable when q has about 30 digits. Above this range, the SEA algorithm becomes available, heuristically in $\tilde{O}(\log q)^4$, and primes of the order of 200 digits become feasible. In small characteristic we use Mestre's ($p=2$), Kohel's ($p=3,5,7,13$), Satoh-Harley (all in $\tilde{O}(p^2 n^2)$) or Kedlaya's (in $\tilde{O}(pn^3)$) algorithms.

The library syntax is GEN ellap(GEN E, GEN p = NULL).

3.15.10 ellbil($E, z1, z2$). Deprecated alias for ellheight(E, P, Q).

The library syntax is GEN bilhell(GEN E, GEN z1, GEN z2, long prec).

3.15.11 ellbsd(E). E being an elliptic curve over a number field, returns a real number c such that the Birch and Swinnerton-Dyer conjecture predicts that $L_E^{(r)}(1)/r! = cRS$, where r is the rank, R the regulator and S the cardinal of the Tate-Shafarevich group.

```

? e = ellinit([0,-1,1,-10,-20]); \\ rank 0
? ellbsd(e)
%2 = 0.25384186085591068433775892335090946105
? lfun(e,1)
%3 = 0.25384186085591068433775892335090946104
? e = ellinit([0,0,1,-1,0]); \\ rank 1
? P = ellheegner(e);
? ellbsd(e)*ellheight(e,P)
%6 = 0.30599977383405230182048368332167647445
? lfun(e,1,1)
%7 = 0.30599977383405230182048368332167647445
? e = ellinit([1+a,0,1,0,0],nfinit(a^2+1)); \\ rank 0
? ellbsd(e)
%9 = 0.42521832235345764503001271536611593310
? lfun(e,1)

```

```
%10 = 0.42521832235345764503001271536611593309
```

The library syntax is `GEN ellbsd(GEN E, long prec)`.

3.15.12 ellcard($E, \{p\}$). Let E be an `ell` structure as output by `ellinit`, attached to an elliptic curve E/K . If $K = \mathbf{F}_q$ is finite, return the order of the group $E(\mathbf{F}_q)$.

```
? E = ellinit([-3,1], 5); ellcard(E)
%1 = 7
? t = ffgen([3,5], 't); E = ellinit([t,t^2+1]); ellcard(E)
%2 = 217
```

For other fields of definition and p defining a finite residue field $\mathbf{F}_q$, return the order of the reduction of E : the argument p is best left omitted if $K = \mathbf{Q}_\ell$ (else we must have $p = \ell$) and must be a prime number ($K = \mathbf{Q}$) or prime ideal (K a general number field) with residue field $\mathbf{F}_q$ otherwise. The equation need not be minimal or even integral at p ; of course, a minimal model will be more efficient. The function considers the group of nonsingular points of the reduction of a minimal model of the curve at p , so also makes sense when the curve has bad reduction.

```
? E = ellinit([-3,1]);
? factor(E.disc)
%2 =
[2 4]
[3 4]
? ellcard(E, 5) \\ as above !
%3 = 7
? ellcard(E, 2) \\ additive reduction
%4 = 2
```

When the characteristic of the finite field is large, the availability of the `seadata` package will speed the computation. See also `ellap` for the list of implemented algorithms.

The library syntax is `GEN ellcard(GEN E, GEN p = NULL)`. Also available is `GEN ellcard(GEN E, GEN p)` where p is not `NULL`.

3.15.13 ellchangecurve(E, v). Changes the data for the elliptic curve E by changing the coordinates using the vector $v=[u,r,s,t]$, i.e. if x' and y' are the new coordinates, then $x = u^2x' + r$, $y = u^3y' + su^2x' + t$. E must be an `ell` structure as output by `ellinit`. The special case $v = 1$ is also used instead of $[1,0,0,0]$ to denote the trivial coordinate change.

The library syntax is `GEN ellchangecurve(GEN E, GEN v)`.

3.15.14 ellchangept(x, v). Changes the coordinates of the point or vector of points x using the vector $v=[u,r,s,t]$, i.e. if x' and y' are the new coordinates, then $x = u^2x' + r$, $y = u^3y' + su^2x' + t$ (see also `ellchangecurve`).

```
? E0 = ellinit([1,1]); P0 = [0,1]; v = [1,2,3,4];
? E = ellchangecurve(E0, v);
? P = ellchangept(P0,v)
%3 = [-2, 3]
? ellisoncurve(E, P)
%4 = 1
? ellchangeptinv(P,v)
```

```
%5 = [0, 1]
```

The library syntax is `GEN ellchangepoint(GEN x, GEN v)`. The reciprocal function `GEN ellchangepointinv(GEN x, GEN ch)` inverts the coordinate change.

3.15.15 ellchangepointinv(x, v). Changes the coordinates of the point or vector of points x using the inverse of the isomorphism attached to $v=[u,r,s,t]$, i.e. if x' and y' are the old coordinates, then $x = u^2x' + r$, $y = u^3y' + su^2x' + t$ (inverse of `ellchangepoint`).

```
? E0 = ellinit([1,1]); P0 = [0,1]; v = [1,2,3,4];
? E = ellchangecurve(E0, v);
? P = ellchangepoint(P0,v)
%3 = [-2, 3]
? ellisoncurve(E, P)
%4 = 1
? ellchangepointinv(P,v)
%5 = [0, 1]  \\ we get back P0
```

The library syntax is `GEN ellchangepointinv(GEN x, GEN v)`.

3.15.16 ellconvertname($name$). Converts an elliptic curve name, as found in the `elldata` database, from a string to a triplet $[conductor, isogeny\ class, index]$. It will also convert a triplet back to a curve name. Examples:

```
? ellconvertname("123b1")
%1 = [123, 1, 1]
? ellconvertname(%)
%2 = "123b1"
```

The library syntax is `GEN ellconvertname(GEN name)`.

3.15.17 elldivpol($E, n, \{v = 'x\}$). n -division polynomial f_n for the curve E in the variable v . In standard notation, for any affine point $P = (X, Y)$ on the curve and any integer $n \geq 0$, we have

$$[n]P = (\phi_n(P)\psi_n(P) : \omega_n(P) : \psi_n(P)^3)$$

for some polynomials ϕ_n, ω_n, ψ_n in $\mathbf{Z}[a_1, a_2, a_3, a_4, a_6][X, Y]$. We have $f_n(X) = \psi_n(X)$ for n odd, and $f_n(X) = \psi_n(X, Y)(2Y + a_1X + a_3)$ for n even. We have

$$f_0 = 0, \quad f_1 = 1, \quad f_2 = 4X^3 + b_2X^2 + 2b_4X + b_6, \quad f_3 = 3X^4 + b_2X^3 + 3b_4X^2 + 3b_6X + b_8,$$

$$f_4 = f_2(2X^6 + b_2X^5 + 5b_4X^4 + 10b_6X^3 + 10b_8X^2 + (b_2b_8 - b_4b_6)X + (b_8b_4 - b_6^2)), \dots$$

When n is odd, the roots of f_n are the X -coordinates of the affine points in the n -torsion subgroup $E[n]$; when n is even, the roots of f_n are the X -coordinates of the affine points in $E[n] \setminus E[2]$ when $n > 2$, resp. in $E[2]$ when $n = 2$. For $n < 0$, we define $f_n := -f_{-n}$.

The library syntax is `GEN elldivpol(GEN E, long n, long v = -1)` where v is a variable number.

3.15.18 `elleisnum( $w, k, \{flag = 0\}$ )`. k being an even positive integer, computes the numerical value of the Eisenstein series of weight k at the lattice w , as given by `ellperiods`, namely

$$(2i\pi/\omega_2)^k \left(1 + 2/\zeta(1-k) \sum_{n>1} n^{k-1} q^n / (1-q^n) \right),$$

where $q = \exp(2i\pi\tau)$ and $\tau := \omega_1/\omega_2$ belongs to the complex upper half-plane. It is also possible to directly input $w = [\omega_1, \omega_2]$, or an elliptic curve E as given by `ellinit`.

```
? w = ellperiods([1,I]);  
? elleisnum(w, 4)  
%2 = 2268.8726415508062275167367584190557607  
? elleisnum(w, 6)  
%3 = -3.977978632282564763 E-33  
? E = ellinit([1, 0]);  
? elleisnum(E, 4)  
%5 = -48.0000000000000000000000000000000000
```

When *flag* is nonzero and $k = 4$ or 6 , returns the elliptic invariants g_2 or g_3 , such that

$$y^2 = 4x^3 - g_2x - g_3$$

is a Weierstrass equation for E .

[illegible]

The library syntax is `GEN elleisnum(GEN w, long k, long flag, long prec)`.

3.15.19 elleta(w). Returns the quasi-periods $[\eta_1, \eta_2]$ attached to the lattice basis $w = [\omega_1, \omega_2]$. Alternatively, w can be an elliptic curve E as output by **ellinit**, in which case, the quasi periods attached to the period lattice basis $E.\text{omega}$ (namely, $E.\text{eta}$) are returned.

```
? elleta([1, I])
%1 = [3.141592653589793238462643383, 9.424777960769379715387930149*I]
```

The library syntax is `GEN elleta(GEN w, long prec)`.

3.15.20 ellformaldifferential($E, \{n = \text{seriesprecision}\}, \{t = 'x'\}$). Let $\omega := dx/(2y + a_1x + a_3)$ be the invariant differential form attached to the model E of some elliptic curve (**ellinit** form), and $\eta := x(t)\omega$. Return n terms (**seriesprecision** by default) of $f(t), g(t)$ two power series in the formal parameter $t = -x/y$ such that $\omega = f(t)dt, \eta = g(t)dt$:

$$f(t) = 1 + a_1 t + (a_1^2 + a_2)t^2 + \dots, \quad g(t) = t^{-2} + \dots$$

```
? E = ellinit([-1,1/4]); [f,g] = ellformaldifferential(E,7,'t');
? f
%2 = 1 - 2*t^4 + 3/4*t^6 + 0(t^7)
? g
%3 = t^-2 - t^2 + 1/2*t^4 + 0(t^5)
```

The library syntax is `GEN ellformaldifferential(GEN E, long precdl, long n = -1)` where `n` is a variable number.

3.15.21 ellformalexp($E, \{n = \text{seriesprecision}\}, \{z = 'x'\}$). The elliptic formal exponential Exp attached to E is the isomorphism from the formal additive law to the formal group of E . It is normalized so as to be the inverse of the elliptic logarithm (see **ellformallog**): $\text{Exp} \circ L = \text{Id}$. Return n terms of this power series:

```
? E=ellinit([-1,1/4]); Exp = ellformalexp(E,10,'z')
%1 = z + 2/5*z^5 - 3/28*z^7 + 2/15*z^9 + 0(z^11)
? L = ellformallog(E,10,'t');
? subst(Exp,z,L)
%3 = t + 0(t^11)
```

The library syntax is **GEN ellformalexp**(GEN E , long precd1 , long $n = -1$) where n is a variable number.

3.15.22 ellformallog($E, \{n = \text{seriesprecision}\}, \{v = 'x'\}$). The formal elliptic logarithm is a series L in $tK[[t]]$ such that $dL = \omega = dx/(2y + a_1x + a_3)$, the canonical invariant differential attached to the model E . It gives an isomorphism from the formal group of E to the additive formal group.

```
? E = ellinit([-1,1/4]); L = ellformallog(E, 9, 't')
%1 = t - 2/5*t^5 + 3/28*t^7 + 2/3*t^9 + 0(t^10)
? [f,g] = ellformaldifferential(E,8,'t');
? L' - f
%3 = 0(t^8)
```

The library syntax is **GEN ellformallog**(GEN E , long precd1 , long $n = -1$) where n is a variable number.

3.15.23 ellformalpoint($E, \{n = \text{seriesprecision}\}, \{v = 'x'\}$). If E is an elliptic curve, return the coordinates $x(t), y(t)$ in the formal group of the elliptic curve E in the formal parameter $t = -x/y$ at ∞ :

$$x = t^{-2} - a_1 t^{-1} - a_2 - a_3 t + \dots$$

$$y = -t^{-3} - a_1 t^{-2} - a_2 t^{-1} - a_3 + \dots$$

Return n terms (**seriesprecision** by default) of these two power series, whose coefficients are in $\mathbf{Z}[a_1, a_2, a_3, a_4, a_6]$.

```
? E = ellinit([0,0,1,-1,0]); [x,y] = ellformalpoint(E,8,'t');
? x
%2 = t^-2 - t + t^2 - t^4 + 2*t^5 + 0(t^6)
? y
%3 = -t^-3 + 1 - t + t^3 - 2*t^4 + 0(t^5)
? E = ellinit([0,1/2]); ellformalpoint(E,7)
%4 = [x^-2 - 1/2*x^4 + 0(x^5), -x^-3 + 1/2*x^3 + 0(x^4)]
```

The library syntax is **GEN ellformalpoint**(GEN E , long precd1 , long $n = -1$) where n is a variable number.

3.15.24 ellformalw($E, \{n = \text{seriesprecision}\}, \{t = 'x'\}$). Return the formal power series w attached to the elliptic curve E , in the variable t :

$$w(t) = t^3(1 + a_1t + (a_2 + a_1^2)t^2 + \cdots + O(t^n)),$$

which is the formal expansion of $-1/y$ in the formal parameter $t := -x/y$ at ∞ (take $n = \text{seriesprecision}$ if n is omitted). The coefficients of w belong to $\mathbf{Z}[a_1, a_2, a_3, a_4, a_6]$.

```
? E=ellinit([3,2,-4,-2,5]); ellformalw(E, 5, 't)
%1 = t^3 + 3*t^4 + 11*t^5 + 35*t^6 + 101*t^7 + 0(t^8)
```

The library syntax is GEN `ellformalw(GEN E, long precdl, long n = -1)` where n is a variable number.

3.15.25 ellfromeqn(P). Given a genus 1 plane curve, defined by the affine equation $f(x, y) = 0$, return the coefficients $[a_1, a_2, a_3, a_4, a_6]$ of a Weierstrass equation for its Jacobian. This allows to recover a Weierstrass model for an elliptic curve given by a general plane cubic or by a binary quartic or biquadratic model. The function implements the $f \mapsto f^*$ formulae of Artin, Tate and Villegas (Advances in Math. 198 (2005), pp. 366–382).

In the example below, the function is used to convert between twisted Edwards coordinates and Weierstrass coordinates.

```
? e = ellfromeqn(a*x^2+y^2 - (1+d*x^2*y^2))
%1 = [0, -a - d, 0, -4*d*a, 4*d*a^2 + 4*d^2*a]
? E = ellinit(ellfromeqn(y^2-x^2 - 1 +(121665/121666*x^2*y^2)),2^255-19);
? isprime(ellcard(E) / 8)
%3 = 1
```

The elliptic curve attached to the sum of two cubes is given by

```
? ellfromeqn(x^3+y^3 - a)
%1 = [0, 0, -9*a, 0, -27*a^2]
```

Congruent number problem. Let n be an integer, if $a^2 + b^2 = c^2$ and $ab = 2n$, then by substituting b by $2n/a$ in the first equation, we get $((a^2 + (2n/a)^2) - c^2)a^2 = 0$. We set $x = a$, $y = ac$.

```
? En = ellfromeqn((x^2 + (2*n/x)^2 - (y/x)^2)*x^2)
%1 = [0, 0, 0, -16*n^2, 0]
```

For example 23 is congruent since the curve has a point of infinite order, namely:

```
? ellheegner( ellinit(subst(En, n, 23)) )
%2 = [168100/289, 68053440/4913]
```

The library syntax is GEN `ellfromeqn(GEN P)`.

3.15.26 ellfromj(*j*). Returns the coefficients $[a_1, a_2, a_3, a_4, a_6]$ of a fixed elliptic curve with *j*-invariant *j*. The given model is arbitrary; for instance, over the rationals, it is in general not minimal nor even integral.

```
? v = ellfromj(1/2)
%1 = [0, 0, 0, 10365/4, 11937025/4]
? E = ellminimalmodel(ellinit(v)); E[1..5]
%2 = [0, 0, 0, 41460, 190992400]
? F = ellminimalmodel(elltwtist(E, 24)); F[1..5]
%3 = [1, 0, 0, 72, 13822]
? [E.disc, F.disc]
%4 = [-15763098924417024000, -82484842750]
```

For rational *j*, the following program returns the integral curve of minimal discriminant and given *j* invariant:

```
ellfromjminimal(j)=
{ my(E = ellinit(ellfromj(j)));
  my(D = ellminimaltwist(E));
  ellminimalmodel(elltwtist(E,D));
}
? e = ellfromjminimal(1/2); e.disc
%1 = -82484842750
```

Using *flag* = 1 in `ellminimaltwist` would instead return the curve of minimal conductor. For instance, if *j* = 1728, this would return a different curve (of conductor 32 instead of 64).

The library syntax is GEN `ellfromj(GEN j)`.

3.15.27 ellgenerators(*E*). If *E* is an elliptic curve over the rationals, return a $\mathbf{Z}$ -basis of the free part of the Mordell-Weil group attached to *E*. This relies on the `elldata` database being installed and referencing the curve, and so is only available for curves over $\mathbf{Z}$ of small conductors. If *E* is an elliptic curve over a finite field $\mathbf{F}_q$ as output by `ellinit`, return a minimal set of generators for the group $E(\mathbf{F}_q)$.

Caution. When the group is not cyclic, of shape $\mathbf{Z}/d_1\mathbf{Z} \times \mathbf{Z}/d_2\mathbf{Z}$ with $d_2 \mid d_1$, the points $[P, Q]$ returned by `ellgenerators` need not have order d_1 and d_2 : it is true that *P* has order d_1 , but we only know that *Q* is a generator of $E(\mathbf{F}_q)/\langle P \rangle$ and that the Weil pairing $w(P, Q)$ has order d_2 , see `??ellgroup`. If you need generators $[P, R]$ with *R* of order d_2 , find *x* such that $R = Q - [x]P$ has order d_2 by solving the discrete logarithm problem $[d_2]Q = [x]([d_2]P)$ in a cyclic group of order d_1/d_2 . This will be very expensive if d_1/d_2 has a large prime factor.

The library syntax is GEN `ellgenerators(GEN E)`.

3.15.28 ellglobalred(E). Let E be an `ell` structure as output by `ellinit` attached to an elliptic curve defined over a number field. This function calculates the arithmetic conductor and the global Tamagawa number c . The result $[N, v, c, F, L]$ is slightly different if E is defined over $\mathbf{Q}$ (domain $D = 1$ in `ellinit`) or over a number field (domain D is a number field structure, including `nfinit(x)` representing $\mathbf{Q}$!):

- N is the arithmetic conductor of the curve,
- v is an obsolete field, left in place for backward compatibility. If E is defined over $\mathbf{Q}$, v gives the coordinate change for E to the standard minimal integral model (`ellminimalmodel` provides it in a cheaper way); if E is defined over another number field, v gives a coordinate change to an integral model (`ellintegralmodel` provides it in a cheaper way).
- c is the product of the local Tamagawa numbers c_p , a quantity which enters in the Birch and Swinnerton-Dyer conjecture,
- F is the factorization of N ,
- L is a vector, whose i -th entry contains the local data at the i -th prime ideal divisor of N , i.e. `L[i] = elllocalred(E, F[i, 1])`. If E is defined over $\mathbf{Q}$, the local coordinate change has been deleted and replaced by a 0; if E is defined over another number field the local coordinate change to a local minimal model is given relative to the integral model afforded by v (so either start from an integral model so that v be trivial, or apply v first).

The library syntax is `GEN ellglobalred(GEN E)`.

3.15.29 ellgroup($E, \{p\}, \{flag\}$). Let E be an `ell` structure as output by `ellinit`, attached to an elliptic curve E/K . We first describe the function when the field $K = \mathbf{F}_q$ is finite, it computes the structure of the finite abelian group $E(\mathbf{F}_q)$:

- if $flag = 0$, returns the structure `[]` (trivial group) or `[d1]` (nontrivial cyclic group) or `[d1, d2]` (noncyclic group) of $E(\mathbf{F}_q) \sim \mathbf{Z}/d_1\mathbf{Z} \times \mathbf{Z}/d_2\mathbf{Z}$, with $d_2 \mid d_1$.
- if $flag = 1$, returns a triple `[h, cyc, gen]`, where h is the curve cardinality, cyc gives the group structure as a product of cyclic groups (as per $flag = 0$). More precisely, if $d_2 > 1$, the output is `[d1d2, [d1, d2], [P, Q]]` where P is of order d_1 and $[P, Q]$ generates the curve.

Caution. It is not guaranteed that Q has order d_2 , which in the worst case requires an expensive discrete log computation. Only that `ellweilpairing(E, P, Q, d1)` has order d_2 .

For other fields of definition and p defining a finite residue field $\mathbf{F}_q$, returns the structure of the reduction of E : the argument p is best left omitted if $K = \mathbf{Q}_\ell$ (else we must have $p = \ell$) and must be a prime number ($K = \mathbf{Q}$) or prime ideal (K a general number field) with residue field $\mathbf{F}_q$ otherwise. The curve is allowed to have bad reduction at p and in this case we consider the (cyclic) group of nonsingular points for the reduction of a minimal model at p .

If $flag = 0$, the equation need not be minimal or even integral at p ; of course, a minimal model will be more efficient.

If $flag = 1$, the requested generators depend on the model, which must then be minimal at p , otherwise an exception is thrown. Use `ellintegralmodel` and/or `elllocalred` first to reduce to this case.

```
? E = ellinit([0,1]); \\ y^2 = x^3 + 0.x + 1, defined over Q
? ellgroup(E, 7)
```

```

%2 = [6, 2] \\ Z/6 x Z/2, noncyclic
? E = ellinit([0,1] * Mod(1,11)); \\ defined over F_11
? ellgroup(E) \\ no need to repeat 11
%4 = [12]
? ellgroup(E, 11) \\ ... but it also works
%5 = [12]
? ellgroup(E, 13) \\ ouch, inconsistent input!
*** at top-level: ellgroup(E,13)
*** ^-----
*** ellgroup: inconsistent moduli in Rg_to_Fp:
    11
    13
? ellgroup(E, 7, 1)
%6 = [12, [6, 2], [[Mod(2, 7), Mod(4, 7)], [Mod(4, 7), Mod(4, 7)]]]

```

Let us now consider curves of bad reduction, in this case we return the structure of the (cyclic) group of nonsingular points, satisfying $\#E_{ns}(\mathbf{F}_p) = p - a_p$:

```

? E = ellinit([0,5]);
? ellgroup(E, 5, 1)
%2 = [5, [5], [[Mod(4, 5), Mod(2, 5)]]]
? ellap(E, 5)
%3 = 0 \\ additive reduction at 5
? E = ellinit([0,-1,0,35,0]);
? ellgroup(E, 5, 1)
%5 = [4, [4], [[Mod(2, 5), Mod(2, 5)]]]
? ellap(E, 5)
%6 = 1 \\ split multiplicative reduction at 5
? ellgroup(E, 7, 1)
%7 = [8, [8], [[Mod(3, 7), Mod(5, 7)]]]
? ellap(E, 7)
%8 = -1 \\ nonsplit multiplicative reduction at 7

```

The library syntax is `GEN ellgroup0(GEN E, GEN p = NULL, long flag)`. Also available is `GEN ellgroup(GEN E, GEN p)`, corresponding to *flag* = 0.

3.15.30 ellheegner(*E*). Let *E* be an elliptic curve over the rationals, assumed to be of (analytic) rank 1. This returns a nontorsion rational point on the curve, whose canonical height is equal to the product of the elliptic regulator by the analytic Sha.

This uses the Heegner point method, described in Cohen GTM 239; the complexity is proportional to the product of the square root of the conductor and the height of the point (thus, it is preferable to apply it to strong Weil curves).

```

? E = ellinit([-157^2,0]);
? u = ellheegner(E); print(u[1], "\n", u[2])
69648970982596494254458225/166136231668185267540804
538962435089604615078004307258785218335/67716816556077455999228495435742408
? ellheegner(ellinit([0,1])) \\ E has rank 0 !
*** at top-level: ellheegner(E=ellinit
*** ^-----

```

***** ellheegner:** The curve has even analytic rank.

The library syntax is `GEN ellheegner(GEN E)`.

3.15.31 ellheight($E, \{P\}, \{Q\}$). Let E be an elliptic curve defined over $K = \mathbf{Q}$ or a number field, as output by `ellinit`; it needs not be given by a minimal model although the computation will be faster if it is.

- Without arguments P, Q , returns the Faltings height of the curve E using Deligne normalization. For a rational curve, the normalization is such that the function returns $-(1/2) \cdot \log(\text{ellminimalmodel}(E).\text{area})$.

- If the argument $P \in E(K)$ is present, returns the global Néron-Tate height $h(P)$ of the point, using the normalization in Cremona's *Algorithms for modular elliptic curves*.

- If the argument $Q \in E(K)$ is also present, computes the value of the bilinear form $(h(P + Q) - h(P - Q))/4$.

The canonical height is equal to the Néron-Tate height divided by the degree of the number field. For a curve over a number field, it is `ellheight(E, P{, Q})/#E.nf.zk`.

The library syntax is `GEN ellheight0(GEN E, GEN P = NULL, GEN Q = NULL, long prec)`. Also available is `GEN ellheight(GEN E, GEN P, long prec)` (Q omitted).

3.15.32 ellheightmatrix(E, x). x being a vector of points, this function outputs the Gram matrix of x with respect to the Néron-Tate height, in other words, the (i, j) component of the matrix is equal to `ellheight(E, x[i], x[j])`. The rank of this matrix, at least in some approximate sense, gives the rank of the set of points, and if x is a basis of the Mordell-Weil group of E , its determinant is equal to the regulator of E . Note our height normalization follows Cremona's *Algorithms for modular elliptic curves*: this matrix should be divided by 2 to be in accordance with, e.g., Silverman's normalizations.

The library syntax is `GEN ellheightmatrix(GEN E, GEN x, long prec)`.

3.15.33 ellidentify(E). Look up the elliptic curve E , defined by an arbitrary model over $\mathbf{Q}$, in the `elldata` database. Return `[[N, M, G], C]` where N is the curve name in Cremona's elliptic curve database, M is the minimal model, G is a $\mathbf{Z}$ -basis of the free part of the Mordell-Weil group $E(\mathbf{Q})$ and C is the change of coordinates from E to M , suitable for `ellchangecurve`.

The library syntax is `GEN ellidentify(GEN E)`.

3.15.34 ellinit($x, \{D = 1\}$). Initialize an `ell` structure, attached to the elliptic curve E . E is either

- a 5-component vector $[a_1, a_2, a_3, a_4, a_6]$ defining the elliptic curve with Weierstrass equation

$$Y^2 + a_1XY + a_3Y = X^3 + a_2X^2 + a_4X + a_6,$$

- a 2-component vector $[a_4, a_6]$ defining the elliptic curve with short Weierstrass equation

$$Y^2 = X^3 + a_4X + a_6,$$

- a single-component vector $[j]$ giving the j -invariant for the curve, with the same coefficients as given by `ellfromj`.

- a character string in Cremona's notation, e.g. "11a1", in which case the curve is retrieved from the `elldata` database if available.

The optional argument D describes the domain over which the curve is defined:

- the `t_INT` 1 (default): the field of rational numbers $\mathbf{Q}$.
- a `t_INT` p , where p is a prime number: the prime finite field $\mathbf{F}_p$.
- an `t_INTMOD` `Mod(a, p)`, where p is a prime number: the prime finite field $\mathbf{F}_p$.
- a `t_FFELT`, as returned by `ffgen`: the corresponding finite field $\mathbf{F}_q$.
- a `t_PADIC`, $O(p^n)$: the field $\mathbf{Q}_p$, where p -adic quantities will be computed to a relative accuracy of n digits. We advise to input a model defined over $\mathbf{Q}$ for such curves. In any case, if you input an approximate model with `t_PADIC` coefficients, it will be replaced by a lift to $\mathbf{Q}$ (an exact model "close" to the one that was input) and all quantities will then be computed in terms of this lifted model, at the given accuracy.

- a `t_REAL` x : the field $\mathbf{C}$ of complex numbers, where floating point quantities are by default computed to a relative accuracy of `precision(x)`. If no such argument is given, the value of `realprecision` at the time `ellinit` is called will be used.

- a number field K , given by a `nf` or `bnf` structure; a `bnf` is required for `ellminimalmodel`.
- a prime ideal $\mathfrak{p}$, given by a `prid` structure; valid if x is a curve defined over a number field K and the equation is integral and minimal at $\mathfrak{p}$.

This argument D is indicative: the curve coefficients are checked for compatibility, possibly changing D ; for instance if $D = 1$ and an `t_INTMOD` is found. If inconsistencies are detected, an error is raised:

```
? ellinit([1 + 0(5), 1], 0(7));
***   at top-level: ellinit([1+0(5),1],0
***                                   ^-----
*** ellinit: inconsistent moduli in ellinit: 7 != 5
```

If the curve coefficients are too general to fit any of the above domain categories, only basic operations, such as point addition, will be supported later.

If the curve (seen over the domain D) is singular, fail and return an empty vector `[]`.

```
? E = ellinit([0,0,0,0,1]); \\ y^2 = x^3 + 1, over Q
? E = ellinit([0,1]);      \\ the same curve, short form
? E = ellinit("36a1");     \\ still the same curve, Cremona's notations
? E = ellinit([0]);        \\ a curve of j-invariant 0
? E = ellinit([0,1], 2)    \\ over F2: singular curve
%4 = []
? E = ellinit(['a4,'a6] * Mod(1,5)); \\ over F_5[a4,a6], basic support !
```

Note that the given curve of j -invariant 0 happens to be 36a1 but a priori any model for an arbitrary twist could have been returned. See `ellfromj`.

The result of `ellinit` is an `ell` structure. It contains at least the following information in its components:

$$a_1, a_2, a_3, a_4, a_6, b_2, b_4, b_6, b_8, c_4, c_6, \Delta, j.$$

All are accessible via member functions. In particular, the discriminant is $E.\text{disc}$, and the j -invariant is $E.j$.

```
? E = ellinit([a4, a6]);
? E.disc
%2 = -64*a4^3 - 432*a6^2
? E.j
%3 = -6912*a4^3/(-4*a4^3 - 27*a6^2)
```

Further components contain domain-specific data, which are in general dynamic: only computed when needed, and then cached in the structure.

```
? E = ellinit([2,3], 10^60+7); \\ E over F_p, p large
? ellap(E)
time = 4,440 ms.
%2 = -1376268269510579884904540406082
? ellcard(E); \\ now instantaneous !
time = 0 ms.
? ellgenerators(E);
time = 5,965 ms.
? ellgenerators(E); \\ second time instantaneous
time = 0 ms.
```

See the description of member functions related to elliptic curves at the beginning of this section.

The library syntax is GEN `ellinit`(GEN x , GEN $D = \text{NULL}$, long prec).

3.15.35 `ellintegralmodel`($E, \{&v\}$). Let E be an `ell` structure over a number field K or $\mathbf{Q}_p$. This function returns an integral model. If v is present, sets $v = [u, 0, 0, 0]$ to the corresponding change of variable: the return value is identical to that of `ellchangecurve`(E, v).

```
? e = ellinit([1/17, 1/42]);
? e = ellintegralmodel(e, &v);
? e[1..5]
%3 = [0, 0, 0, 15287762448, 3154568630095008]
? v
%4 = [1/714, 0, 0, 0]
```

The library syntax is GEN `ellintegralmodel`(GEN E , GEN $*v = \text{NULL}$).

3.15.36 `elliscm`(E). Let E an elliptic curve over a number field. Return 0 if E is not CM, otherwise return the discriminant of its endomorphism ring.

```
? E = ellinit([0,0,-5,-750,7900]);
? D = elliscm(E)
%2 = -27
? w = quadgen(D, 'w);
? P = ellheegner(E)
%4 = [10,40]
? Q = ellmul(E,P,w)
%5 = [110/7-5/49*w, 85/49-225/343*w]
```

An example over a number field:

```

? nf=nfinit(a^2-5);
? E = ellinit([261526980*a-584793000,-3440201839360*a+7692525148000],nf);
? elliscm(E)
%3 = -20
? ellisomat(E)[2]
%4 = [1,2,5,10;2,1,10,5;5,10,1,2;10,5,2,1]

```

The library syntax is `long elliscm(GEN E)`.

3.15.37 ellisdivisible($E, P, n, \{&Q\}$). Given E/K a number field and P in $E(K)$ return 1 if $P = [n]R$ for some R in $E(K)$ and set Q to one such R ; and return 0 otherwise.

```

? K = nfinit(polcyclo(11,t));
? E = ellinit([0,-1,1,0,0], K);
? P = [0,0];
? ellorder(E,P)
%4 = 5
? ellisdivisible(E,P,5, &Q)
%5 = 1
? lift(Q)
%6 = [-t^7-t^6-t^5-t^4+1, -t^9-2*t^8-2*t^7-3*t^6-3*t^5-2*t^4-2*t^3-t^2-1]
? ellorder(E, Q)
%7 = 25

```

We use a fast multimodular algorithm over $\mathbf{Q}$ whose complexity is essentially independent of n (polynomial in $\log n$). Over number fields, we compute roots of division polynomials and the algebraic complexity of the underlying algorithm is in $O(p^4)$, where p is the largest prime divisor of n . The integer $n \geq 0$ may be given as `ellxn(E,n)`, if many points need to be tested; this provides a modest speedup over number fields but is likely to slow down the algorithm over $\mathbf{Q}$.

The library syntax is `long ellisdivisible(GEN E, GEN P, GEN n, GEN *Q = NULL)`.

3.15.38 ellisisom(E, F). Return 0 if the elliptic curves E and F defined over the same number field are not isomorphic, otherwise return `[u,r,s,t]` suitable for `ellchangecurve`, mapping E to F .

```

? E = ellinit([1,2]);
? ellisisom(E, ellinit([1,3]))
%2 = 0
? F = ellchangecurve(E, [-1,1,3,2]);
? ellisisom(E,F)
%4 = [1, 1, -3, -2]

? nf = nfinit(a^3-2); E = ellinit([a^2+1,2*a-5], nf);
? F = ellchangecurve(E,Mod([a, a+1, a^2, a^2+a-3], nf.pol));
? v = ellisisom(E,F)
%3 = [Mod(-a, a^3 - 2), Mod(a + 1, a^3 - 2), Mod(-a^2, a^3 - 2),
      Mod(-a^2 - a + 3, a^3 - 2)]
? ellchangecurve(E,v) == F
%4 = 1

```

The library syntax is `GEN ellisisom(GEN E, GEN F)`.

3.15.39 ellisogeny($E, G, \{only_image = 0\}, \{x = 'x\}, \{y = 'y\}$). Given an elliptic curve E , a finite subgroup G of E is given either as a generating point P (for a cyclic G) or as a polynomial whose roots vanish on the x -coordinates of the nonzero elements of G (general case and more efficient if available). This function returns the $[a_1, a_2, a_3, a_4, a_6]$ invariants of the quotient elliptic curve E/G and (if *only_image* is zero (the default)) a vector of rational functions $[f, g, h]$ such that the isogeny $E \rightarrow E/G$ is given by $(x, y) \mapsto (f(x)/h(x)^2, g(x, y)/h(x)^3)$.

```
? E = ellinit([0,1]);
? elltors(E)
%2 = [6, [6], [[2, 3]]]
? ellisogeny(E, [2,3], 1) \\ Weierstrass model for E/<P>
%3 = [0, 0, 0, -135, -594]
? ellisogeny(E, [-1,0])
%4 = [[0,0,0,-15,22], [x^3+2*x^2+4*x+3, y*x^3+3*y*x^2-2*y, x+1]]
```

The library syntax is GEN ellisogeny(GEN E, GEN G, long only_image, long x = -1, long y = -1) where x, y are variable numbers.

3.15.40 ellisogenyapply(f, g). Given an isogeny of elliptic curves $f : E' \rightarrow E$ (being the result of a call to **ellisogeny**), apply f to g :

- if g is a point P in the domain of f , return the image $f(P)$;
- if $g : E'' \rightarrow E'$ is a compatible isogeny, return the composite isogeny $f \circ g : E'' \rightarrow E$.

```
? one = ffgen(101, 't')^0;
? E = ellinit([6, 53, 85, 32, 34] * one);
? P = [84, 71] * one;
? ellorder(E, P)
%4 = 5
? [F, f] = ellisogeny(E, P); \\ f: E->F = E/<P>
? ellisogenyapply(f, P)
%6 = [0]
? F = ellinit(F);
? Q = [89, 44] * one;
? ellorder(F, Q)
%9 = 2
? [G, g] = ellisogeny(F, Q); \\ g: F->G = F/<Q>
? gof = ellisogenyapply(g, f); \\ gof: E -> G
```

The library syntax is GEN ellisogenyapply(GEN f, GEN g).

3.15.41 ellisomat($E, \{p = 0\}, \{flag = 0\}$). Given an elliptic curve E defined over a number field K , computes representatives of the set of isomorphism classes of elliptic curves defined over K and K -isogenous to E , assuming it is finite (see below). For any such curve E_i , let $f_i : E \rightarrow E_i$ be a rational isogeny of minimal degree and let $g_i : E_i \rightarrow E$ be the dual isogeny; and let M be the matrix such that $M_{i,j}$ is the minimal degree for an isogeny $E_i \rightarrow E_j$.

The function returns a vector $[L, M]$ where L is a list of triples $[E_i, f_i, g_i]$ ($flag = 0$), or simply the list of E_i ($flag = 1$, which saves time). The curves E_i are given in $[a_4, a_6]$ form and the first curve E_1 is isomorphic to E by f_1 .

The set of isomorphism classes is finite except when E has CM over a quadratic order contained in K . In that case the function only returns the discriminant of the quadratic order.

If p is set, it must be a prime number; in this case only isogenies of degree a power of p are considered.

Over a number field, the possible isogeny degrees are determined by Billerey's algorithm.

```
? E = ellinit("14a1");
? [L,M] = ellisomat(E);
? LE = apply(x->x[1], L) \\ list of curves
%3 = [[215/48,-5291/864],[-675/16,6831/32],[-8185/48,-742643/864],
      [-1705/48,-57707/864],[-13635/16,306207/32],[-131065/48,-47449331/864]]
? L[2][2] \\ isogeny f_2
%4 = [x^3+3/4*x^2+19/2*x-311/12,
      1/2*x^4+(y+1)*x^3+(y-4)*x^2+(-9*y+23)*x+(55*y+55/2),x+1/3]
? L[2][3] \\ dual isogeny g_2
%5 = [1/9*x^3-1/4*x^2-141/16*x+5613/64,
      -1/18*x^4+(1/27*y-1/3)*x^3+(-1/12*y+87/16)*x^2+(49/16*y-48)*x
      +(-3601/64*y+16947/512),x-3/4]
? apply(E->ellidentify(ellinit(E))[1][1], LE)
%6 = ["14a1","14a4","14a3","14a2","14a6","14a5"]
? M
%7 =
[1  3  3  2  6  6]
[3  1  9  6  2 18]
[3  9  1  6 18  2]
[2  6  6  1  3  3]
[6  2 18  3  1  9]
[6 18  2  3  9  1]
```

The library syntax is `GEN ellisomat(GEN E, long p, long flag)`.

3.15.42 ellisoncurve(E, z). Gives 1 (i.e. true) if the point z is on the elliptic curve E , 0 otherwise. If E or z have imprecise coefficients, an attempt is made to take this into account, i.e. an imprecise equality is checked, not a precise one. It is allowed for z to be a vector of points in which case a vector (of the same type) is returned.

The library syntax is `GEN ellisoncurve(GEN E, GEN z)`. Also available is `int oncurve(GEN E, GEN z)` which does not accept vectors of points.

3.15.43 ellisotree(E). Given an elliptic curve E defined over $\mathbf{Q}$ or a set of $\mathbf{Q}$ -isogenous curves as given by `ellisomat`, return a pair $[L, M]$ where

- L lists the minimal models of the isomorphism classes of elliptic curves $\mathbf{Q}$ -isogenous to E (or in the set of isogenous curves),

- M is the adjacency matrix of the prime degree isogenies tree: there is an edge from E_i to E_j if there is an isogeny $E_i \rightarrow E_j$ of prime degree such that the Néron differential forms are preserved.

```
? E = ellinit("14a1");
? [L,M] = ellisotree(E);
? M
%3 =
[0 0 3 2 0 0]
[3 0 0 0 2 0]
[0 0 0 0 0 2]
[0 0 0 0 0 3]
[0 0 0 3 0 0]
[0 0 0 0 0 0]
? [L2,M2] = ellisotree(ellisomat(E,2,1));
%4 =
[0 2]
[0 0]
? [L3,M3] = ellisotree(ellisomat(E,3,1));
? M3
%6 =
[0 0 3]
[3 0 0]
[0 0 0]
```

Compare with the result of `ellisomat`.

```
? [L,M]=ellisomat(E,,1);
? M
%7 =
[1 3 3 2 6 6]
[3 1 9 6 2 18]
[3 9 1 6 18 2]
[2 6 6 1 3 3]
[6 2 18 3 1 9]
[6 18 2 3 9 1]
```

The library syntax is `GEN ellisotree(GEN E)`.

3.15.44 ellissupersingular($E, \{p\}$). Return 1 if the elliptic curve E defined over a number field, $\mathbf{Q}_p$ or a finite field is supersingular at p , and 0 otherwise. If the curve is defined over $\mathbf{Q}$ or a number field, p must be explicitly given, and must be a prime number, resp. a maximal ideal; we return 1 if and only if E has supersingular good reduction at p .

Alternatively, E can be given by its j -invariant in a finite field. In this case p must be omitted.

```
? g = ffprimroot(ffgen([7,5]))
%1 = 4*x^4+5*x^3+6*x^2+5*x+6
? [g^n | n <- [1 .. 7^5 - 1], ellissupersingular(g^n)]
%2 = [6]
? j = ellsupersingularj(2^31-1)
%3 = 1618591527*w+1497042960
? ellissupersingular(j)
%4 = 1

? K = nfinit(y^3-2); P = idealprimedec(K, 2)[1];
? E = ellinit([y,1], K);
? ellissupersingular(E, P)
%7 = 1
? Q = idealprimedec(K,5)[1];
? ellissupersingular(E, Q)
%9 = 0
```

The library syntax is `int ellissupersingular(GEN E, GEN p = NULL)`. Also available is `int elljissupersingular(GEN j)` where j is a j -invariant of a curve over a finite field.

3.15.45 ellj(x). Elliptic j -invariant. x must be a complex number with positive imaginary part, or convertible into a power series or a p -adic number with positive valuation.

The library syntax is `GEN jell(GEN x, long prec)`.

3.15.46 elllocalred($E, \{p\}$). Calculates the Kodaira type of the local fiber of the elliptic curve E at p . E must be an `ell` structure as output by `ellinit`, over $\mathbf{Q}_\ell$ (p better left omitted, else equal to ℓ) over $\mathbf{Q}$ (p a rational prime) or a number field K (p a maximal ideal given by a `prid` structure). The result is a 4-component vector $[f, kod, v, c]$. Here f is the exponent of p in the arithmetic conductor of E , and kod is the Kodaira type which is coded as follows:

1 means good reduction (type I_0), 2, 3 and 4 mean types II, III and IV respectively, $4 + \nu$ with $\nu > 0$ means type I_ν ; finally the opposite values $-1, -2$, etc. refer to the starred types I_0^*, II^* , etc. The third component v is itself a vector $[u, r, s, t]$ giving the coordinate changes done during the local reduction; $u = 1$ if and only if the given equation was already minimal at p . Finally, the last component c is the local Tamagawa number c_p .

The library syntax is `GEN elllocalred(GEN E, GEN p = NULL)`.

3.15.47 elllog($E, P, G, \{o\}$). Given two points P and G on the elliptic curve $E/\mathbf{F}_q$, returns the discrete logarithm of P in base G , i.e. the smallest nonnegative integer n such that $P = [n]G$. See **znlog** for the limitations of the underlying discrete log algorithms. If present, o represents the order of G , see Section 3.8.2; the preferred format for this parameter is $[N, \text{factor}(N)]$, where N is the order of G .

If no o is given, assume that G generates the curve. The function also assumes that P is a multiple of G .

```
? a = ffgen(ffinit(2,8),'a');
? E = ellinit([a,1,0,0,1]); \\ over F_{2^8}
? x = a^3; y = ellordinate(E,x)[1];
? P = [x,y]; G = ellmul(E, P, 113);
? ord = [242, factor(242)]; \\ P generates a group of order 242. Initialize.
? ellorder(E, G, ord)
%4 = 242
? e = elllog(E, P, G, ord)
%5 = 15
? ellmul(E,G,e) == P
%6 = 1
```

The library syntax is `GEN elllog(GEN E, GEN P, GEN G, GEN o = NULL)`.

3.15.48 ellseries($E, s, \{A = 1\}$). This function is deprecated, use `lfun(E,s)` instead.

E being an elliptic curve, given by an arbitrary model over $\mathbf{Q}$ as output by `ellinit`, this function computes the value of the L -series of E at the (complex) point s . This function uses an $O(N^{1/2})$ algorithm, where N is the conductor.

The optional parameter A fixes a cutoff point for the integral and is best left omitted; the result must be independent of A , up to `realprecision`, so this allows to check the function's accuracy.

The library syntax is `GEN ellseries(GEN E, GEN s, GEN A = NULL, long prec)`.

3.15.49 ellmaninconstant(E). Let E be an elliptic curve over Q given by `ellinit` or a rational isogeny class given by `ellisomat`. Return the Manin constant of the curve, see `ellweilcurve`. The algorithm is slow but unconditional. The function also accepts the output of `ellisomat` and returns the list of Manin constants for all the isogeny class.

```
? E = ellinit("11a3");
? ellmaninconstant(E)
%2 = 5
? L=ellisomat(E,,1);
? ellmaninconstant(L)
%4 = [5,1,1]
```

The library syntax is `GEN ellmaninconstant(GEN E)`.

3.15.50 ellminimaldisc(E). E being an elliptic curve defined over a number field output by `ellinit`, return the minimal discriminant ideal of E .

The library syntax is `GEN ellminimaldisc(GEN E)`.

3.15.51 ellminimalmodel($E, \{&v\}$). Let E be an `ell` structure over a number field K . This function determines whether E admits a global minimal integral model. If so, it returns it and sets $v = [u, r, s, t]$ to the corresponding change of variable: the return value is identical to that of `ellchangecurve( $E, v$ )`.

Else return the (nonprincipal) Weierstrass class of E , i.e. the class of $\prod \mathfrak{p}^{(v_{\mathfrak{p}}\Delta - \delta_{\mathfrak{p}})/12}$ where $\Delta = E.\text{disc}$ is the model's discriminant and $\mathfrak{p}^{\delta_{\mathfrak{p}}}$ is the local minimal discriminant. This function requires either that E be defined over the rational field $\mathbf{Q}$ (with domain $D = 1$ in `ellinit`), in which case a global minimal model always exists, or over a number field given by a `bnf` structure. The Weierstrass class is given in `bnfisprincipal` format, i.e. in terms of the `K.gen` generators.

The resulting model has integral coefficients and is everywhere minimal, the coefficients a_1 and a_3 are reduced modulo 2 (in terms of the fixed integral basis `K.zk`) and a_2 is reduced modulo 3. Over $\mathbf{Q}$, we further require that a_1 and a_3 be 0 or 1, that a_2 be 0 or ± 1 and that $u > 0$ in the change of variable: both the model and the change of variable v are then unique.

```
? e = ellinit([6,6,12,55,233]);  \\ over Q
? E = ellminimalmodel(e, &v);
? E[1..5]
%3 = [0, 0, 0, 1, 1]
? v
%4 = [2, -5, -3, 9]

? K = bnfinit(a^2-65);  \\ over a nonprincipal number field
? K.cyc
%2 = [2]
? u = Mod(8+a, K.pol);
? E = ellinit([1,40*u+1,0,25*u^2,0], K);
? ellminimalmodel(E)  \\ no global minimal model exists over Z_K
%6 = [1]~
```

The library syntax is `GEN ellminimalmodel(GEN E, GEN *v = NULL)`.

3.15.52 ellminimaltwist($E, \{flag = 0\}$). Let E be an elliptic curve defined over $\mathbf{Q}$, return a discriminant D such that the twist of E by D is minimal among all possible quadratic twists, i.e. if $flag = 0$, its minimal model has minimal discriminant, or if $flag = 1$, it has minimal conductor.

In the example below, we find a curve with j -invariant 3 and minimal conductor.

```
? E = ellminimalmodel(ellinit(ellfromj(3)));
? ellglobalred(E)[1]
%2 = 357075
? D = ellminimaltwist(E,1)
%3 = -15
? E2 = ellminimalmodel(elltwtst(E,D));
? ellglobalred(E2)[1]
%5 = 14283
```

In the example below, $flag = 0$ and $flag = 1$ give different results.

```
? E = ellinit([1,0]);
? D0 = ellminimaltwist(E,0)
%7 = 1
```

```
? D1 = ellminimaltwist(E,1)
%8 = 8
? E0 = ellminimalmodel(elltwise(E,D0));
? [E0.disc, ellglobalred(E0)[1]]
%10 = [-64, 64]
? E1 = ellminimalmodel(elltwise(E,D1));
? [E1.disc, ellglobalred(E1)[1]]
%12 = [-4096, 32]
```

The library syntax is `GEN ellminimaltwist0(GEN E, long flag)`. Also available are `GEN ellminimaltwist(E)` for $flag = 0$, and `GEN ellminimaltwistcond(E)` for $flag = 1$.

3.15.53 ellmoddegree(e). e being an elliptic curve defined over $\mathbf{Q}$ output by `ellinit`, compute the modular degree of e divided by the square of the Manin constant c . It is conjectured that $c = 1$ for the strong Weil curve in the isogeny class (optimal quotient of $J_0(N)$) and this can be proven using `ellweilcurve` when the conductor N is moderate.

```
? E = ellinit("11a1"); \\ from Cremona table: strong Weil curve and c = 1
? [v,smith] = ellweilcurve(E); smith \\ proof of the above
%2 = [[1, 1], [5, 1], [1, 1/5]]
? ellmoddegree(E)
%3 = 1
? [ellidentify(e)[1][1] | e<-v]
%4 = ["11a1", "11a2", "11a3"]
? ellmoddegree(ellinit("11a2"))
%5 = 5
? ellmoddegree(ellinit("11a3"))
%6 = 1/5
```

The modular degree of `11a1` is 1 (because `ellweilcurve` or Cremona's table prove that the Manin constant is 1 for this curve); the output of `ellweilcurve` also proves that the Manin constants of `11a2` and `11a3` are 1 and 5 respectively, so the actual modular degree of both `11a2` and `11a3` is 5.

The library syntax is `GEN ellmoddegree(GEN e)`.

3.15.54 ellmodulareqn($N, \{x\}, \{y\}$). Given a prime $N < 500$, return a vector $[P, t]$ where $P(x, y)$ is a modular equation of level N , i.e. a bivariate polynomial with integer coefficients; t indicates the type of this equation: either *canonical* ($t = 0$) or *Atkin* ($t = 1$). This function requires the `seadata` package and its only use is to give access to the package contents. See `polmodular` for a more general and more flexible function.

Let j be the j -invariant function. The polynomial P satisfies the functional equation,

$$P(f, j) = P(f \mid W_N, j \mid W_N) = 0$$

for some modular function $f = f_N$ (hand-picked for each fixed N to minimize its size, see below), where $W_N(\tau) = -1/(N\tau)$ is the Atkin-Lehner involution. These two equations allow to compute the values of the classical modular polynomial Φ_N , such that $\Phi_N(j(\tau), j(N\tau)) = 0$, while being much smaller than the latter. More precisely, we have $j(W_N(\tau)) = j(N\tau)$; the function f is invariant under $\Gamma_0(N)$ and also satisfies

- for Atkin type: $f \mid W_N = f$;

- for canonical type: let $s = 12/\gcd(12, N-1)$, then $f \mid W_N = N^s/f$. In this case, f has a simple definition: $f(\tau) = N^s(\eta(N\tau)/\eta(\tau))^{2s}$, where η is Dedekind's eta function.

The following GP function returns values of the classical modular polynomial by eliminating $f_N(\tau)$ in the above functional equation, for $N \leq 31$ or $N \in \{41, 47, 59, 71\}$.

```
classicaleqn(N, X='X', Y='Y')=
{
  my([P,t] = ellmodulareqn(N), Q, d);
  if (poldegree(P,'y') > 2, error("level unavailable in classicaleqn"));
  if (t == 0, \\ Canonical
    my(s = 12/gcd(12,N-1));
    Q = 'x^(N+1) * substvec(P,['x','y'],[N^s/'x,Y]);
    d = N^(s*(2*N+1)) * (-1)^(N+1);
  , \\ Atkin
    Q = subst(P,'y,Y);
    d = (X-Y)^(N+1));
  polresultant(subst(P,'y,X), Q) / d;
}
```

The library syntax is GEN ellmodulareqn(long N, long x = -1, long y = -1) where x, y are variable numbers.

3.15.55 ellmul(E, z, n). Computes $[n]z$, where z is a point on the elliptic curve E . The exponent n is in $\mathbf{Z}$, or may be a complex quadratic integer if the curve E has complex multiplication by n (if not, an error message is issued).

```
? Ei = ellinit([1,0]); z = [0,0];
? ellmul(Ei, z, 10)
%2 = [0] \\ unsurprising: z has order 2
? ellmul(Ei, z, I)
%3 = [0, 0] \\ Ei has complex multiplication by Z[i]
? ellmul(Ei, z, quadgen(-4))
%4 = [0, 0] \\ an alternative syntax for the same query
? Ej = ellinit([0,1]); z = [-1,0];
? ellmul(Ej, z, I)
*** at top-level: ellmul(Ej,z,I)
*** ^-----
*** ellmul: not a complex multiplication in ellmul.
? ellmul(Ej, z, 1+quadgen(-3))
%6 = [1 - w, 0]
```

The simple-minded algorithm for the CM case assumes that we are in characteristic 0, and that the quadratic order to which n belongs has small discriminant.

The library syntax is GEN ellmul(GEN E, GEN z, GEN n).

3.15.56 ellneg(E, z). Opposite of the point z on elliptic curve E .

The library syntax is GEN ellneg(GEN E, GEN z).

3.15.57 `ellnonsingularmultiple(E, P)`. Given an elliptic curve $E/\mathbf{Q}$ (more precisely, a model defined over $\mathbf{Q}$ of a curve) and a rational point $P \in E(\mathbf{Q})$, returns the pair $[R, n]$, where n is the least positive integer such that $R := [n]P$ has good reduction at every prime. More precisely, its image in a minimal model is everywhere nonsingular.

```
? e = ellinit("57a1"); P = [2, -2];
? ellnonsingularmultiple(e, P)
%2 = [[1, -1], 2]
? e = ellinit("396b2"); P = [35, -198];
? [R,n] = ellnonsingularmultiple(e, P);
? n
%5 = 12
```

The library syntax is `GEN ellnonsingularmultiple(GEN E, GEN P)`.

3.15.58 `ellorder(E, z, {o})`. Gives the order of the point z on the elliptic curve E , defined over a finite field or a number field. Return (the impossible value) zero if the point has infinite order.

```
? E = ellinit([-157^2, 0]); \\ the "157-is-congruent" curve
? P = [0, 0]; ellorder(E, P)
%2 = 2
? P = ellheegner(E); ellorder(E, P) \\ infinite order
%3 = 0
? K = nfinit(polcyclo(11, t)); E = ellinit("11a3", K); T = elltors(E);
? ellorder(E, T.gen[1])
%5 = 25
? E = ellinit(ellfromj(ffgen(5^10)));
? ellcard(E)
%7 = 9767025
? P = random(E); ellorder(E, P)
%8 = 1953405
? p = 2^160+7; E = ellinit([1, 2], p);
? N = ellcard(E)
%9 = 1461501637330902918203686560289225285992592471152
? o = [N, factor(N)];
? for(i=1, 100, ellorder(E, random(E)))
time = 260 ms.
```

The parameter o , is now mostly useless, and kept for backward compatibility. If present, it represents a nonzero multiple of the order of z , see Section 3.8.2; the preferred format for this parameter is `[ord, factor(ord)]`, where `ord` is the cardinality of the curve. It is no longer needed since PARI is now able to compute it over large finite fields (was restricted to small prime fields at the time this feature was introduced), *and* caches the result in E so that it is computed and factored only once. Modifying the last example, we see that including this extra parameter provides no improvement:

```
? o = [N, factor(N)];
? for(i=1, 100, ellorder(E, random(E), o))
time = 260 ms.
```

The library syntax is `GEN ellorder(GEN E, GEN z, GEN o = NULL)`. The obsolete form `GEN orderell(GEN e, GEN z)` should no longer be used.

3.15.59 ellordinate(E, x). Gives a 0, 1 or 2-component vector containing the y -coordinates of the points of the curve E having x as x -coordinate.

The library syntax is `GEN ellordinate(GEN E, GEN x, long prec)`.

3.15.60 ellpadicL($E, p, n, \{s = 0\}, \{r = 0\}, \{D = 1\}$). Returns the value (or r -th derivative) on a character χ^s of $\mathbf{Z}_p^*$ of the p -adic L -function of the elliptic curve $E/\mathbf{Q}$, twisted by D , given modulo p^n .

Characters. The set of continuous characters of $\text{Gal}(\mathbf{Q}(\mu_{p^\infty})/\mathbf{Q})$ is identified to $\mathbf{Z}_p^*$ via the cyclotomic character χ with values in $\overline{\mathbf{Q}_p}^*$. Denote by $\tau : \mathbf{Z}_p^* \rightarrow \mathbf{Z}_p^*$ the Teichmüller character, with values in the $(p-1)$ -th roots of 1 for $p \neq 2$, and $\{-1, 1\}$ for $p = 2$; finally, let $\langle \chi \rangle = \chi\tau^{-1}$, with values in $1 + 2p\mathbf{Z}_p$. In GP, the continuous character of $\text{Gal}(\mathbf{Q}(\mu_{p^\infty})/\mathbf{Q})$ given by $\langle \chi \rangle^{s_1} \tau^{s_2}$ is represented by the pair of integers $s = (s_1, s_2)$, with $s_1 \in \mathbf{Z}_p$ and $s_2 \bmod p-1$ for $p > 2$, (resp. $\bmod 2$ for $p = 2$); s may be also an integer, representing (s, s) or χ^s .

The p -adic L function. The p -adic L function L_p is defined on the set of continuous characters of $\text{Gal}(\mathbf{Q}(\mu_{p^\infty})/\mathbf{Q})$, as $\int_{\mathbf{Z}_p^*} \chi^s d\mu$ for a certain p -adic distribution μ on $\mathbf{Z}_p^*$. The derivative is given by

$$L_p^{(r)}(E, \chi^s) = \int_{\mathbf{Z}_p^*} \log_p^r(a) \chi^s(a) d\mu(a).$$

More precisely:

- When E has good supersingular reduction, L_p takes its values in $D := H_{dR}^1(E/\mathbf{Q}) \otimes_{\mathbf{Q}} \mathbf{Q}_p$ and satisfies

$$(1 - p^{-1}F)^{-2} L_p(E, \chi^0) = (L(E, 1)/\Omega) \cdot \omega$$

where F is the Frobenius, $L(E, 1)$ is the value of the complex L function at 1, ω is the Néron differential and Ω the attached period on $E(\mathbf{R})$. Here, χ^0 represents the trivial character.

The function returns the components of $L_p^{(r)}(E, \chi^s)$ in the basis $(\omega, F\omega)$.

- When E has ordinary good reduction, this method only defines the projection of $L_p(E, \chi^s)$ on the α -eigenspace, where α is the unit eigenvalue for F . This is what the function returns. We have

$$(1 - \alpha^{-1})^{-2} L_{p,\alpha}(E, \chi^0) = L(E, 1)/\Omega.$$

Two supersingular examples:

```
? cxL(e) = bestappr( ellL1(e) / e.omega[1] );
? e = ellinit("17a1"); p=3; \\ supersingular, a3 = 0
? L = ellpadicL(e,p,4);
? F = [0,-p;1,ellap(e,p)]; \\ Frobenius matrix in the basis (omega,F(omega))
? (1-p^(-1)*F)^-2 * L / cxL(e)
%5 = [1 + 0(3^5), 0(3^5)]~ \\ [1,0]~

? e = ellinit("116a1"); p=3; \\ supersingular, a3 != 0~
? L = ellpadicL(e,p,4);
? F = [0,-p; 1,ellap(e,p)];
? (1-p^(-1)*F)^-2*L~ / cxL(e)
%9 = [1 + 0(3^4), 0(3^5)]~
```

Good ordinary reduction:

```
? e = ellinit("17a1"); p=5; ap = ellap(e,p)
%1 = -2 \\ ordinary
? L = ellpadicL(e,p,4)
%2 = 4 + 3*5 + 4*5^2 + 2*5^3 + 0(5^4)
? al = padicappr(x^2 - ap*x + p, ap + 0(p^7))[1];
? (1-al^(-1))^(-2) * L / cxL(e)
%4 = 1 + 0(5^4)
```

Twist and Teichmüller:

```
? e = ellinit("17a1"); p=5; \\ ordinary
\\ 2nd derivative at tau^1, twist by -7
? ellpadicL(e, p, 4, [0,1], 2, -7)
%2 = 2*5^2 + 5^3 + 0(5^4)
```

We give an example of non split multiplicative reduction (see `ellpadicbsd` for more examples).

```
? e=ellinit("15a1"); p=3; n=5;
? L = ellpadicL(e,p,n)
%2 = 2 + 3 + 3^2 + 3^3 + 3^4 + 0(3^5)
? (1 - ellap(e,p))^(-1) * L / cxL(e)
%3 = 1 + 0(3^5)
```

This function is a special case of `mspadicL` and it also appears as the first term of `mspadicseries`:

```
? e = ellinit("17a1"); p=5;
? L = ellpadicL(e,p,4)
%2 = 4 + 3*5 + 4*5^2 + 2*5^3 + 0(5^4)
? [M,phi] = msfromell(e, 1);
? Mp = mspadicinit(M, p, 4);
? mu = mspadicmoments(Mp, phi);
? mspadicL(mu)
%6 = 4 + 3*5 + 4*5^2 + 2*5^3 + 2*5^4 + 5^5 + 0(5^6)
? mspadicseries(mu)
%7 = (4 + 3*5 + 4*5^2 + 2*5^3 + 2*5^4 + 5^5 + 0(5^6))
      + (3 + 3*5 + 5^2 + 5^3 + 0(5^4))*x
      + (2 + 3*5 + 5^2 + 0(5^3))*x^2
      + (3 + 4*5 + 4*5^2 + 0(5^3))*x^3
      + (3 + 2*5 + 0(5^2))*x^4 + 0(x^5)
```

These are more cumbersome than `ellpadicL` but allow to compute at different characters, or successive derivatives, or to twist by a quadratic character essentially for the cost of a single call to `ellpadicL` due to precomputations.

The library syntax is `GEN ellpadicL(GEN E, GEN p, long n, GEN s = NULL, long r, GEN D = NULL)`.

3.15.61 ellpadicbsd($E, p, n, \{D = 1\}$). Given an elliptic curve E over $\mathbf{Q}$, its quadratic twist E_D and a prime number p , this function is a p -adic analog of the complex functions **ellanalyticrank** and **ellbsd**. It calls **ellpadicL** with initial accuracy p^n and may increase it internally; it returns a vector $[r, L_p]$ where

- L_p is a p -adic number (resp. a pair of p -adic numbers if E has good supersingular reduction) defined modulo p^N , conjecturally equal to $R_p S$, where R_p is the p -adic regulator as given by **ellpadicregulator** (in the basis $(\omega, F\omega)$) and S is the cardinal of the Tate-Shafarevich group for the quadratic twist E_D .

- r is an upper bound for the analytic rank of the p -adic L -function attached to E_D : we know for sure that the i -th derivative of $L_p(E_D, \cdot)$ at χ^0 is $O(p^N)$ for all $i < r$ and that its r -th derivative is nonzero; it is expected that the true analytic rank is equal to the rank of the Mordell-Weil group $E_D(\mathbf{Q})$, plus 1 if the reduction of E_D at p is split multiplicative; if $r = 0$, then both the analytic rank and the Mordell-Weil rank are unconditionnally 0.

Recall that the p -adic BSD conjecture (Mazur, Tate, Teitelbaum, Bernardi, Perrin-Riou) predicts an explicit link between $R_p S$ and

$$(1 - p^{-1}F)^{-2} \cdot L_p^{(r)}(E_D, \chi^0)/r!$$

where r is the analytic rank of the p -adic L -function attached to E_D and F is the Frobenius on H_{dR}^1 ; see **ellpadicL** for definitions.

```
? E = ellinit("11a1"); p = 7; n = 5; \\ good ordinary
? ellpadicbsd(E, 7, 5) \\ rank 0,
%2 = [0, 1 + 0(7^5)]

? E = ellinit("91a1"); p = 7; n = 5; \\ non split multiplicative
? [r,Lp] = ellpadicbsd(E, p, n)
%5 = [1, 2*7 + 6*7^2 + 3*7^3 + 7^4 + 0(7^5)]
? R = ellpadicregulator(E, p, n, E.gen)
%6 = 2*7 + 6*7^2 + 3*7^3 + 7^4 + 5*7^5 + 0(7^6)
? sha = Lp/R
%7 = 1 + 0(7^4)

? E = ellinit("91b1"); p = 7; n = 5; \\ split multiplicative
? [r,Lp] = ellpadicbsd(E, p, n)
%9 = [2, 2*7 + 7^2 + 5*7^3 + 0(7^4)]
? ellpadicregulator(E, p, n, E.gen)
%10 = 2*7 + 7^2 + 5*7^3 + 6*7^4 + 2*7^5 + 0(7^6)
? [rC, LC] = ellanalyticrank(E);
? [r, rC]
%12 = [2, 1] \\ r = rC+1 because of split multiplicative reduction

? E = ellinit("53a1"); p = 5; n = 5; \\ supersingular
? [r, Lp] = ellpadicbsd(E, p, n);
? r
%15 = 1
? Lp
%16 = [3*5 + 2*5^2 + 2*5^5 + 0(5^6), \
      5 + 3*5^2 + 4*5^3 + 2*5^4 + 5^5 + 0(5^6)]
? R = ellpadicregulator(E, p, n, E.gen)
```

```

%17 = [3*5 + 2*5^2 + 2*5^5 + 0(5^6), 5 + 3*5^2 + 4*5^3 + 2*5^4 + 0(5^5)]
\\ expect Lp = R*#Sha, hence (conjecturally) #Sha = 1
? E = ellinit("84a1"); p = 11; n = 6; D = -443;
? [r,Lp] = ellpadicbsd(E, 11, 6, D) \\ Mordell-Weil rank 0, no regulator
%19 = [0, 3 + 2*11 + 0(11^6)]
? lift(Lp) \\ expected cardinal for Sha is 5^2
%20 = 25
? ellpadicbsd(E, 3, 12, D) \\ at 3
%21 = [1, 1 + 2*3 + 2*3^2 + 0(3^8)]
? ellpadicbsd(E, 7, 8, D) \\ and at 7
%22 = [0, 4 + 3*7 + 0(7^8)]

```

The library syntax is GEN ellpadicbsd(GEN E, GEN p, long n, GEN D = NULL).

3.15.62 ellpadicfrobenius(E, p, n). If $p > 2$ is a prime and E is an elliptic curve on $\mathbf{Q}$ with good reduction at p , return the matrix of the Frobenius endomorphism φ on the crystalline module $D_p(E) = \mathbf{Q}_p \otimes H_{dR}^1(E/\mathbf{Q})$ with respect to the basis of the given model $(\omega, \eta = x\omega)$, where $\omega = dx/(2y + a_1x + a_3)$ is the invariant differential. The characteristic polynomial of φ is $x^2 - a_px + p$. The matrix is computed to absolute p -adic precision p^n .

```

? E = ellinit([1,-1,1,0,0]);
? F = ellpadicfrobenius(E,5,3);
? lift(F)
%3 =
[120 29]
[ 55  5]
? charpoly(F)
%4 = x^2 + 0(5^3)*x + (5 + 0(5^3))
? ellap(E, 5)
%5 = 0

```

The library syntax is GEN ellpadicfrobenius(GEN E, ulong p, long n).

3.15.63 ellpadicheight($E, p, n, P, \{Q\}$). Cyclotomic p -adic height of the rational point P on the elliptic curve E (defined over $\mathbf{Q}$), given to n p -adic digits. If the argument Q is present, computes the value of the bilinear form $(h(P + Q) - h(P - Q))/4$.

Let $D := H_{dR}^1(E) \otimes_{\mathbf{Q}} \mathbf{Q}_p$ be the $\mathbf{Q}_p$ vector space spanned by ω (invariant differential $dx/(2y + a_1x + a_3)$ related to the given model) and $\eta = x\omega$. Then the cyclotomic p -adic height h_E associates to $P \in E(\mathbf{Q})$ an element $f\omega + g\eta$ in D . This routine returns the vector $[f, g]$ to n p -adic digits. If $P \in E(\mathbf{Q})$ is in the kernel of reduction mod p and if its reduction at all finite places is non singular, then $g = -(\log_E P)^2$, where $\log_E$ is the logarithm for the formal group of E at p .

If furthermore the model is of the form $Y^2 = X^3 + aX + b$ and $P = (x, y)$, then

$$f = \log_p(\text{denominator}(x)) - 2\log_p(\sigma(P))$$

where $\sigma(P)$ is given by ellsigma(E, P).

Recall (*Advanced topics in the arithmetic of elliptic curves*, Theorem 3.2) that the local height function over the complex numbers is of the form

$$\lambda(z) = -\log(|E.\text{disc}|)/6 + \Re(z\eta(z)) - 2\log(\sigma(z)).$$

(N.B. our normalization for local and global heights is twice that of Silverman's).

```
? E = ellinit([1,-1,1,0,0]); P = [0,0];
? ellpadicheight(E,5,3, P)
%2 = [3*5 + 5^2 + 2*5^3 + 0(5^4), 5^2 + 4*5^4 + 0(5^5)]
? E = ellinit("11a1"); P = [5,5]; \\ torsion point
? ellpadicheight(E,19,6, P)
%4 = [0, 0]
? E = ellinit([0,0,1,-4,2]); P = [-2,1];
? ellpadicheight(E,3,3, P)
%6 = [2*3^2 + 2*3^3 + 3^4 + 0(3^5), 2*3^2 + 3^4 + 0(3^5)]
? ellpadicheight(E,3,5, P, elladd(E,P,P))
%7 = [3^2 + 2*3^3 + 0(3^7), 3^2 + 3^3 + 2*3^4 + 3^5 + 0(3^7)]
```

• When E has good ordinary reduction at p or non split multiplicative reduction, the “canonical” p -adic height is given by

```
s2 = ellpadics2(E,p,n);
ellpadicheight(E, p, n, P) * [1,-s2]~
```

Since s_2 does not depend on P , it is preferable to compute it only once:

```
? E = ellinit("5077a1"); p = 5; n = 7; \\ rank 3
? s2 = ellpadics2(E,p,n);
? M = ellpadicheightmatrix(E,p, n, E.gen) * [1,-s2]~;
? matdet(M) \\ p-adic regulator on the points in E.gen
%4 = 5 + 5^2 + 4*5^3 + 2*5^4 + 2*5^5 + 2*5^6 + 0(5^7)
```

• When E has split multiplicative reduction at p (Tate curve), the “canonical” p -adic height is given by

```
Ep = ellinit(E[1..5], 0(p^(n))); \\ E seen as a Tate curve over Qp
[u2,u,q] = Ep.tate;
ellpadicheight(E, p, n, P) * [1,-s2 + 1/log(q)/u2]]~
```

where s_2 is as above. For example,

```
? E = ellinit("91b1"); P = [-1, 3]; p = 7; n = 5;
? Ep = ellinit(E[1..5], 0(p^(n)));
? s2 = ellpadics2(E,p,n);
? [u2,u,q] = Ep.tate;
? H = ellpadicheight(E,p, n, P) * [1,-s2 + 1/log(q)/u2]~
%5 = 2*7 + 7^2 + 5*7^3 + 6*7^4 + 2*7^5 + 0(7^6)
```

These normalizations are chosen so that p -adic BSD conjectures are easy to state, see `ellpadicbsd`.

The library syntax is `GEN ellpadicheight0(GEN E, GEN p, long n, GEN P, GEN Q = NULL)`

3.15.64 ellpadicheightmatrix(E, p, n, Q). Q being a vector of points, this function returns the “Gram matrix” $[F, G]$ of the cyclotomic p -adic height h_E with respect to the basis (ω, η) of $D = H_{dR}^1(E) \otimes_{\mathbf{Q}} \mathbf{Q}_p$ given to n p -adic digits. In other words, if $\text{ellpadicheight}(E, p, n, Q[i], Q[j]) = [f, g]$, corresponding to $f\omega + g\eta$ in D , then $F[i, j] = f$ and $G[i, j] = g$.

```
? E = ellinit([0,0,1,-7,6]); Q = [[-2,3],[-1,3]]; p = 5; n = 5;
? [F,G] = ellpadicheightmatrix(E,p,n,Q);
? lift(F) \\ p-adic entries, integral approximation for readability
%3 =
[2364 3100]
[3100 3119]
? G
%4 =
[25225 46975]
[46975 61850]
? [F,G] * [1,-ellpadics2(E,p,n)]~
%5 =
[4 + 2*5 + 4*5^2 + 3*5^3 + 0(5^5)          4*5^2 + 4*5^3 + 5^4 + 0(5^5)]
[ 4*5^2 + 4*5^3 + 5^4 + 0(5^5) 4 + 3*5 + 4*5^2 + 4*5^3 + 5^4 + 0(5^5)]
```

The library syntax is GEN `ellpadicheightmatrix(GEN E, GEN p, long n, GEN Q)`.

3.15.65 ellpadiclamdamu($E, p, \{D = 1\}, \{i = 0\}$). Let p be a prime number and let $E/\mathbf{Q}$ be a rational elliptic curve with good or bad multiplicative reduction at p . Return the Iwasawa invariants λ and μ for the p -adic L function $L_p(E)$, twisted by $(D/\cdot)$ and the i -th power of the Teichmüller character τ , see `ellpadicL` for details about $L_p(E)$.

Let χ be the cyclotomic character and choose γ in $\text{Gal}(\mathbf{Q}_p(\mu_{p^\infty})/\mathbf{Q}_p)$ such that $\chi(\gamma) = 1 + 2p$. Let $\hat{L}^{(i),D} \in \mathbf{Q}_p[[X]] \otimes D_{\text{cris}}$ such that

$$(< \chi >^s \tau^i)(\hat{L}^{(i),D}(\gamma - 1)) = L_p(E, < \chi >^s \tau^i(D/\cdot)).$$

- When E has good ordinary or bad multiplicative reduction at p . By Weierstrass’s preparation theorem the series $\hat{L}^{(i),D}$ can be written $p^\mu(X^\lambda + pG(X))$ up to a p -adic unit, where $G(X) \in \mathbf{Z}_p[X]$. The function returns $[\lambda, \mu]$.

- When E has good supersingular reduction, we define a sequence of polynomials P_n in $\mathbf{Q}_p[X]$ of degree $< p^n$ (and bounded denominators), such that

$$\hat{L}^{(i),D} \equiv P_n \varphi^{n+1} \omega_E - \xi_n P_{n-1} \varphi^{n+2} \omega_E \pmod{((1+X)^{p^n} - 1) \mathbf{Q}_p[X] \otimes D_{\text{cris}}},$$

where $\xi_n = \text{polcyclo}(p^n, 1+X)$. Let λ_n, μ_n be the invariants of P_n . We find that

- μ_n is nonnegative and decreasing for n of given parity hence μ_{2n} tends to a limit μ^+ and μ_{2n+1} tends to a limit μ^- (both conjecturally 0).

- there exists integers λ^+, λ^- in $\mathbf{Z}$ (denoted with a $\sim$ in the reference below) such that

$$\lim_{n \rightarrow \infty} \lambda_{2n} + 1/(p+1) = \lambda^+ \quad \text{and} \quad \lim_{n \rightarrow \infty} \lambda_{2n+1} + p/(p+1) = \lambda^-.$$

The function returns $[[\lambda^+, \lambda^-], [\mu^+, \mu^-]]$.

Reference: B. Perrin-Riou, Arithmétique des courbes elliptiques à réduction supersingulière en p , *Experimental Mathematics*, **12**, 2003, pp. 155-186.

The library syntax is GEN `ellpadiclamdamu(GEN E, long p, long D, long i)`.

3.15.66 ellpadiclog(E, p, n, P). Given E defined over $K = \mathbf{Q}$ or $\mathbf{Q}_p$ and $P = [x, y]$ on $E(K)$ in the kernel of reduction mod p , let $t(P) = -x/y$ be the formal group parameter; this function returns $L(t)$ to relative p -adic precision p^n , where L denotes the formal logarithm (mapping the formal group of E to the additive formal group) attached to the canonical invariant differential: $dL = dx/(2y + a_1x + a_3)$.

```
? E = ellinit([0,0,1,-4,2]); P = [-2,1];
? ellpadiclog(E,2,10,P)
%2 = 2 + 2^3 + 2^8 + 2^9 + 2^10 + 0(2^11)
? E = ellinit([17,42]);
? p=3; Ep = ellinit(E,p); \\ E mod p
? P=[114,1218]; ellorder(Ep,P) \\ the order of P on (E mod p) is 2
%5 = 2
? Q = ellmul(E,P,2) \\ we need a point of the form 2*P
%6 = [200257/7056, 90637343/592704]
? ellpadiclog(E,3,10,Q)
%7 = 3 + 2*3^2 + 3^3 + 3^4 + 3^5 + 3^6 + 2*3^8 + 3^9 + 2*3^10 + 0(3^11)
```

The library syntax is GEN `ellpadiclog(GEN E, GEN p, long n, GEN P)`.

3.15.67 ellpadicregulator(E, p, n, S). Let $E/\mathbf{Q}$ be an elliptic curve. Return the determinant of the Gram matrix of the vector of points $S = (S_1, \dots, S_r)$ with respect to the “canonical” cyclotomic p -adic height on E , given to n (p -adic) digits.

When E has ordinary reduction at p , this is the expected Gram determinant in $\mathbf{Q}_p$.

In the case of supersingular reduction of E at p , the definition requires care: the regulator R is an element of $D := H_{dR}^1(E) \otimes_{\mathbf{Q}} \mathbf{Q}_p$, which is a two-dimensional $\mathbf{Q}_p$ -vector space spanned by ω and $\eta = x\omega$ (which are defined over $\mathbf{Q}$) or equivalently but now over $\mathbf{Q}_p$ by ω and $F\omega$ where F is the Frobenius endomorphism on D as defined in `ellpadicfrobenius`. On D we define the cyclotomic height $h_E = f\omega + g\eta$ (see `ellpadicheight`) and a canonical alternating bilinear form $[\cdot, \cdot]_D$ such that $[\omega, \eta]_D = 1$.

For any $\nu \in D$, we can define a height $h_\nu := [h_E, \nu]_D$ from $E(\mathbf{Q})$ to $\mathbf{Q}_p$ and $\langle \cdot, \cdot \rangle_\nu$ the attached bilinear form. In particular, if $h_E = f\omega + g\eta$, then $h_\eta = [h_E, \eta]_D = f$ and $h_\omega = [h_E, \omega]_D = -g$ hence $h_E = h_\eta\omega - h_\omega\eta$. Then, R is the unique element of D such that

$$[\omega, \nu]_D^{r-1} [R, \nu]_D = \det(\langle S_i, S_j \rangle_\nu)$$

for all $\nu \in D$ not in $\mathbf{Q}_p\omega$. The `ellpadicregulator` function returns R in the basis $(\omega, F\omega)$, which was chosen so that p -adic BSD conjectures are easy to state, see `ellpadicbsd`.

Note that by definition

$$[R, \eta]_D = \det(\langle S_i, S_j \rangle_\eta)$$

and

$$[R, \omega + \eta]_D = \det(\langle S_i, S_j \rangle_{\omega+\eta}).$$

The library syntax is GEN `ellpadicregulator(GEN E, GEN p, long n, GEN S)`.

3.15.68 ellpadics2(E, p, n). If $p > 2$ is a prime and $E/\mathbf{Q}$ is an elliptic curve with ordinary good reduction at p , returns the slope of the unit eigenvector of **ellpadicfrobenius**(E, p, n), i.e., the action of Frobenius φ on the crystalline module $D_p(E) = \mathbf{Q}_p \otimes H_{dR}^1(E/\mathbf{Q})$ in the basis of the given model $(\omega, \eta = x\omega)$, where ω is the invariant differential $dx/(2y + a_1x + a_3)$. In other words, $\eta + s_2\omega$ is an eigenvector for the unit eigenvalue of φ .

```
? e=ellinit([17,42]);
? ellpadics2(e,13,4)
%2 = 10 + 2*13 + 6*13^3 + 0(13^4)
```

This slope is the unique $c \in 3^{-1}\mathbf{Z}_p$ such that the odd solution $\sigma(t) = t + O(t^2)$ of

$$-d\left(\frac{1}{\sigma} \frac{d\sigma}{\omega}\right) = (x(t) + c)\omega$$

is in $t\mathbf{Z}_p[[t]]$.

It is equal to $b_2/12 - E_2/12$ where E_2 is the value of the Katz p -adic Eisenstein series of weight 2 on (E, ω) . This is used to construct a canonical p -adic height when E has good ordinary reduction at p as follows

```
s2 = ellpadics2(E,p,n);
h(E,p,n, P, s2) = ellpadicheight(E, [p,[1,-s2]],n, P);
```

Since s_2 does not depend on the point P , we compute it only once.

The library syntax is **GEN ellpadics2**(**GEN** E , **GEN** p , **long** n).

3.15.69 ellperiods($w, \{flag = 0\}$). Let w describe a complex period lattice ($w = [w_1, w_2]$ or an **ellinit** structure). Returns normalized periods $[W_1, W_2]$ generating the same lattice such that $\tau := W_1/W_2$ has positive imaginary part and lies in the standard fundamental domain for $\mathrm{SL}_2(\mathbf{Z})$.

If $flag = 1$, the function returns $[[W_1, W_2], [\eta_1, \eta_2]]$, where η_1 and η_2 are the quasi-periods attached to $[W_1, W_2]$, satisfying $\eta_2 W_1 - \eta_1 W_2 = 2i\pi$.

The output of this function is meant to be used as the first argument given to **ellwp**, **ellzeta**, **ellsigma** or **elleisnum**. Quasi-periods are needed by **ellzeta** and **ellsigma** only.

```
? L = ellperiods([1,I],1);
? [w1,w2] = L[1]; [e1,e2] = L[2];
? e2*w1 - e1*w2
%3 = 6.2831853071795864769252867665590057684*I
? ellzeta(L, 1/2 + 2*I)
%4 = 1.5707963... - 6.283185307...*I
? ellzeta([1,I], 1/2 + 2*I) \\ same but less efficient
%4 = 1.5707963... - 6.283185307...*I
```

The library syntax is **GEN ellperiods**(**GEN** w , **long** $flag$, **long** $prec$).

3.15.70 ellpointtoz(E, P). If $E/\mathbf{C} \simeq \mathbf{C}/\Lambda$ is a complex elliptic curve ($\Lambda = \mathbf{E}.\omega\mathbf{e}\mathbf{g}\mathbf{a}$), computes a complex number z , well-defined modulo the lattice Λ , corresponding to the point P ; i.e. such that $P = [\wp_\Lambda(z), \wp'_\Lambda(z)]$ satisfies the equation

$$y^2 = 4x^3 - g_2x - g_3,$$

where g_2, g_3 are the elliptic invariants.

If E is defined over $\mathbf{R}$ and $P \in E(\mathbf{R})$, we have more precisely, $0 \leq \Re(t) < w1$ and $0 \leq \Im(t) < \Im(w2)$, where $(w1, w2)$ are the real and complex periods of E .

[illegible]

If E is defined over a general number field, the function returns the values corresponding to the various complex embeddings of the curve and of the point, in the same order as `E.nf.roots`:

```
? E=ellinit([-22032-15552*x,0], nfini(x^2-2));
? P=[-72*x-108,0];
? ellisoncurve(E,P)
%3 = 1
? ellpointtoz(E,P)
%4 = [-0.52751724240790530394437835702346995884*I,
      -0.090507650025885335533571758708283389896*I]
? E.nf.roots
%5 = [-1.4142135623730950488016887242096980786, \\ x-> -sqrt(2)
      1.4142135623730950488016887242096980786] \\ x-> sqrt(2)
```

If $E/\mathbf{Q}_p$ has multiplicative reduction, then $E/\bar{\mathbf{Q}}_p$ is analytically isomorphic to $\bar{\mathbf{Q}}_p^*/q^{\mathbf{Z}}$ (Tate curve) for some p -adic integer q . The behavior is then as follows:

- If the reduction is split ($E.\text{tate}[2]$ is a `t_PADIC`), we have an isomorphism $\phi : E(\mathbf{Q}_p) \simeq \mathbf{Q}_p^*/q^{\mathbf{Z}}$ and the function returns $\phi(P) \in \mathbf{Q}_p$.

• If the reduction is *not* split ($E.\text{tate}[2]$ is a $\mathfrak{t_POLMOD}$), we only have an isomorphism $\phi : E(K) \simeq K^*/q^{\mathbf{Z}}$ over the unramified quadratic extension $K/\mathbf{Q}_p$. In this case, the output $\phi(P) \in K$ is a $\mathfrak{t_POLMOD}$; the function is not fully implemented in this case and may fail with a “ u not in $\mathbf{Q}_p$ ” exception:

```
? E = ellinit([0,-1,1,0,0], 0(11^5)); P = [0,0];
? [u2,u,q] = E.tate; type(u) \\ split multiplicative reduction
%2 = "t_PADIC"
? ellmul(E, P, 5) \\ P has order 5
%3 = [0]
? z = ellpointtoz(E, [0,0])
%4 = 3 + 11^2 + 2*11^3 + 3*11^4 + 6*11^5 + 10*11^6 + 8*11^7 + 0(11^8)
```

```

? z^5
%5 = 1 + 0(11^9)
? E = ellinit(ellfromj(1/4), 0(2^6)); x=1/2; y=ellordinate(E,x)[1];
? z = ellpointtoz(E,[x,y]); \\ t_POLMOD of t_POL with t_PADIC coeffs
? liftint(z) \\ lift all p-adics
%8 = Mod(8*u + 7, u^2 + 437)
? x=33/4; y=ellordinate(E,x)[1]; z = ellpointtoz(E,[x,y])
*** at top-level: ...;y=ellordinate(E,x)[1];z=ellpointtoz(E,[x,y])
***                                     ^-----
*** ellpointtoz: sorry, ellpointtoz when u not in Qp is not yet implemented.

```

The library syntax is GEN zell(GEN E, GEN P, long prec).

3.15.71 **ellpow**(E, z, n). Deprecated alias for **ellmul**.

The library syntax is GEN ellmul(GEN E, GEN z, GEN n).

3.15.72 **ellrank**($E, \{effort = 0\}, \{points\}$). If E is an elliptic curve over $\mathbf{Q}$, attempts to compute the Mordell-Weil group attached to the curve. The output is $[r_1, r_2, s, L]$, where $r_1 \leq \text{rank}(E) \leq r_2$, s gives informations on the Tate-Shafarevic group (see below), and L is a list of independent, non-torsion rational points on the curve. E can also be given as the output of **ellrankinit**(E).

If **points** is provided, it must be a vector of rational points on the curve, which are not computed again.

The parameter **effort** is a measure of the time employed to find rational points before giving up. If **effort** is not 0, the search is randomized, so rerunning the function might yield different or even a different number of rational points. Values up to 10 or so are reasonable but the parameter can be increased further, with running times increasing roughly like the *cube* of the **effort** value.

```

? E = ellinit([-127^2,0]);
? ellrank(E)
%2 = [1, 1, 0, []] \\ rank is 1 but no point has been found.
? ellrank(E,4) \\ with more effort we find a point.
%3 = [1, 1, 0, [[38902300445163190028032/305111826865145547009,
680061120400889506109527474197680/5329525731816164537079693913473]]]

```

In addition to the previous calls, the first argument E can be a pair $[e, f]$, where e is an elliptic curve given by **ellrankinit** and f is a quadratic twist of e . We then look for points on f . Note that the **ellrankinit** initialization is independent of f , so this can speed up computations significantly!

Technical explanation. The algorithm, which computes the 2-descent and the 2-part of the Cassels pairings has an intrinsic limitation: $r_1 = r_2$ never holds when the Tate-Shafarevic group G has 4-torsion. Thus, in this case we cannot determine the rank precisely. The algorithm computes unconditionally three quantities:

- the rank C of the 2-Selmer group.
- the rank T of the 2-torsion subgroup.
- the (even) rank s of $G[2]/2G[4]$; then r_2 is defined by $r_2 = C - T - s$.

The following quantities are also relevant:

- the rank R of the free part of $E(\mathbf{Q})$; it always holds that $r_1 \leq R \leq r_2$.
- the rank S of $G[2]$ (conjecturally even); it always holds that $s \leq S$ and that $C = T + R + S$.

Then $r_2 = C - T - s \geq R$.

When the conductor of E is small, the BSD conjecture can be used to (conditionally) find the true rank:

```
? E=ellinit([-113^2,0]);  
? ellrootno(E) \\ rank is even (parity conjecture)  
%2 = 1  
? ellrank(E)  
%3 = [0, 2, 0, []] \\ rank is either 0 or 2, $2$-rank of $G$ is  
? ellrank(E, 3) \\ try harder  
%4 = [0, 2, 0, []] \\ no luck  
? [r,L] = ellanalyticrank(E) \\ assume BSD  
%5 = [0, 3.9465...]  
? L / ellbsd(E) \\ analytic rank is 0, compute Sha  
%6 = 16.00000000000000000000000000000000000000000000000000000
```

We find that the rank is 0 and the cardinal of the Tate-Shafarevich group is 16 (assuming BSD!). Moreover, since $s = 0$, it is isomorphic to $(\mathbf{Z}/4\mathbf{Z})^2$.

When the rank is 1 and the conductor is small, `ellheegner` can be used to find a non-torsion point:

```
? E = ellinit([-157^2,0]);
? ellrank(E)
%2 = [1, 1, 0, []] \\ rank is 1, no point found
? ellrank(E, 5) \\ Try harder
time = 1,094 ms.
%3 = [1, 1, 0, []] \\ No luck
? ellheegner(E) \\ use analytic method
time = 492 ms.
%4 = [69648970982596494254458225/166136231668185267540804, ...]
```

In this last example, an **effort** about 10 would also (with probability about 80%) find a random point, not necessarily the Heegner point, in about 5 seconds.

The library syntax is `GEN ellrank(GEN E, long effort, GEN points = NULL, long prec)`

3.15.73 ellrankinit(E). If E is an elliptic curve over $\mathbf{Q}$, initialize data to speed up further calls to `ellrank`.

```
? E = ellinit([0,2429469980725060,0,275130703388172136833647756388,0]);
? rk = ellrankinit(E);
? [r, R, s, P] = ellrank(rk)
%3 = [12, 14, 0, [...]]
? [r, R, s, P] = ellrank(rk, 1, P) \\ more effort, using known points
%4 = [14, 14, 0, [...]] \\ this time all points are found
```

The library syntax is GEN `ellrankinit(GEN E, long prec)`.

3.15.74 ellratpoints($E, h, \{flag = 0\}$). E being an integral model of elliptic curve, return a vector containing the affine rational points on the curve of naive height less than h . If $flag = 1$, stop as soon as a point is found; return either an empty vector or a vector containing a single point. See `hyperellratpoints` for how h can be specified.

```
? E=ellinit([-25,1]);
? ellratpoints(E,10)
%2 = [[-5,1],[-5,-1],[-3,7],[-3,-7],[-1,5],[-1,-5],
      [0,1],[0,-1],[5,1],[5,-1],[7,13],[7,-13]]
? ellratpoints(E,10,1)
%3 = [[-5,1]]
```

The library syntax is GEN `ellratpoints(GEN E, GEN h, long flag)`.

3.15.75 ellrootno($E, \{p\}$). E being an `ell` structure over $\mathbf{Q}$ as output by `ellinit`, this function computes the local root number of its L -series at the place p (at the infinite place if $p = 0$). If p is omitted, return the global root number and in this case the curve can also be defined over a number field.

Note that the global root number is the sign of the functional equation and conjecturally is the parity of the rank of the Mordell-Weil group. The equation for E needs not be minimal at p , but if the model is already minimal the function will run faster.

The library syntax is long `ellrootno(GEN E, GEN p = NULL)`.

3.15.76 ellsaturation(E, V, B). Let E be an elliptic curve over $\mathbf{Q}$ and V be a set of independent non-torsion rational points on E of infinite order that generates a subgroup G of $E(\mathbf{Q})$ of finite index. Return a new set W of the same length that generates a subgroup H of $E(\mathbf{Q})$ containing G and such that $[E(\mathbf{Q}) : H]$ is not divisible by any prime number less than B . The running time is roughly quadratic in B .

```
? E = ellinit([0,0, 1, -7, 6]);
? [r,R,s,V] = ellrank(E)
%2 = [3, 3, 0, [[-1,3], [-3,0], [11,35]]]
? matdet(ellheightmatrix(E, V))
%3 = 3.7542920288254557283540759015628405708
? W = ellsaturation(E, V, 2) \\ index is now odd
time = 1 ms.
%4 = [[-1, 3], [-3, 0], [11, 35]]
? W = ellsaturation(E, W, 10) \\ index not divisible by p <= 10
```

```

%5 = [[1, -1], [2, -1], [0, -3]]
time = 2 ms.
? W = ellsaturation(E, V, 100) \\ looks OK now
time = 171 ms.
%6 = [[1, -1], [2, -1], [0, -3]]
? matdet(ellheightmatrix(E,V))
%7 = 0.41714355875838396981711954461809339675
? lfun(E,1,3)/3! / ellbsd(E) \\ conductor is small, check assuming BSD
%8 = 0.41714355875838396981711954461809339675

```

The library syntax is `GEN ellsaturation(GEN E, GEN V, long B, long prec)`.

3.15.77 ellsea($E, \{tors = 0\}$). Let E be an *ell* structure as output by `ellinit`, defined over a finite field $\mathbf{F}_q$. This low-level function computes the order of the group $E(\mathbf{F}_q)$ using the SEA algorithm; compared to the high-level function `ellcard`, which includes SEA among its choice of algorithms, the `tors` argument allows to speed up a search for curves having almost prime order and whose quadratic twist may also have almost prime order. When `tors` is set to a nonzero value, the function returns 0 as soon as it detects that the order has a small prime factor not dividing `tors`; SEA considers modular polynomials of increasing prime degree ℓ and we return 0 as soon as we hit an ℓ (coprime to `tors`) dividing $\#E(\mathbf{F}_q)$:

```

? ellsea(ellinit([1,1], 2^56+3477), 1)
%1 = 72057594135613381
? forprime(p=2^128,oo, q = ellcard(ellinit([1,1],p)); if(isprime(q),break))
time = 6,571 ms.
? forprime(p=2^128,oo, q = ellsea(ellinit([1,1],p),1);if(isprime(q),break))
time = 522 ms.

```

In particular, set `tors` to 1 if you want a curve with prime order, to 2 if you want to allow a cofactor which is a power of two (e.g. for Edwards's curves), etc. The early exit on bad curves yields a massive speedup compared to running the cardinal algorithm to completion.

When `tors` is negative, similar checks are performed for the quadratic twist of the curve.

The following function returns a curve of prime order over $\mathbf{F}_p$.

```

cryptocurve(p) =
{
  while(1,
    my(E, N, j = Mod(random(p), p));
    E = ellinit(ellfromj(j));
    N = ellsea(E, 1); if (!N, continue);
    if (isprime(N), return(E));
    \\ try the quadratic twist for free
    if (isprime(2*p+2 - N), return(elltwise(E)));
  );
}
? p = randomprime([2^255, 2^256]);
? E = cryptocurve(p); \\ insist on prime order
%2 = 47,447ms

```

The same example without early abort (using `ellcard(E)` instead of `ellsea(E, 1)`) runs for about 5 minutes before finding a suitable curve.

The availability of the `seadata` package will speed up the computation, and is strongly recommended. The generic function `ellcard` should be preferred when you only want to compute the cardinal of a given curve without caring about it having almost prime order:

- If the characteristic is too small ($p \leq 7$) or the field cardinality is tiny ($q \leq 523$) the generic algorithm `ellcard` is used instead and the `tors` argument is ignored. (The reason for this is that SEA is not implemented for $p \leq 7$ and that if $q \leq 523$ it is likely to run into an infinite loop.)
- If the field cardinality is smaller than about 2^{50} , the generic algorithm will be faster.
- Contrary to `ellcard`, `ellsea` does not store the computed cardinality in E .

The library syntax is `GEN ellsea(GEN E, long tors)`.

3.15.78 `ellsearch(N)`. This function finds all curves in the `elldata` database satisfying the constraint defined by the argument N :

- if N is a character string, it selects a given curve, e.g. `"11a1"`, or curves in the given isogeny class, e.g. `"11a"`, or curves with given conductor, e.g. `"11"`;
- if N is a vector of integers, it encodes the same constraints as the character string above, according to the `ellconvertname` correspondance, e.g. `[11,0,1]` for `"11a1"`, `[11,0]` for `"11a"` and `[11]` for `"11"`;
- if N is an integer, curves with conductor N are selected.

If N codes a full curve name, for instance `"11a1"` or `[11,0,1]`, the output format is $[N, [a_1, a_2, a_3, a_4, a_6], G]$ where $[a_1, a_2, a_3, a_4, a_6]$ are the coefficients of the Weierstrass equation of the curve and G is a $\mathbf{Z}$ -basis of the free part of the Mordell-Weil group attached to the curve.

```
? ellsearch("11a3")
%1 = ["11a3", [0, -1, 1, 0, 0], []]
? ellsearch([11,0,3])
%2 = ["11a3", [0, -1, 1, 0, 0], []]
```

If N is not a full curve name, then the output is a vector of all matching curves in the above format:

```
? ellsearch("11a")
%1 = ["11a1", [0, -1, 1, -10, -20], []],
      ["11a2", [0, -1, 1, -7820, -263580], []],
      ["11a3", [0, -1, 1, 0, 0], []]]
? ellsearch("11b")
%2 = []
```

The library syntax is `GEN ellsearch(GEN N)`. Also available is `GEN ellsearchcurve(GEN N)` that only accepts complete curve names (as `t_STR`).

3.15.79 ellsigma($L, \{z = 'x\}, \{flag = 0\}$). Computes the value at z of the Weierstrass σ function attached to the lattice L as given by `ellperiods`(, 1): including quasi-periods is useful, otherwise there are recomputed from scratch for each new z .

$$\sigma(z, L) = z \prod_{\omega \in L^*} \left(1 - \frac{z}{\omega}\right) e^{\frac{z}{\omega} + \frac{z^2}{2\omega^2}}.$$

It is also possible to directly input $L = [\omega_1, \omega_2]$, or an elliptic curve E as given by `ellinit` ($L = E.omega$).

```
? w = ellperiods([1,I], 1);
? ellsigma(w, 1/2)
%2 = 0.47494937998792065033250463632798296855
? E = ellinit([1,0]);
? ellsigma(E) \\ at 'x, implicitly at default seriesprecision
%4 = x + 1/60*x^5 - 1/10080*x^9 - 23/259459200*x^13 + 0(x^17)
```

If $flag = 1$, computes an arbitrary determination of $\log(\sigma(z))$.

The library syntax is GEN `ellsigma`(GEN L , GEN $z = \text{NULL}$, long $flag$, long $prec$).

3.15.80 ellsub($E, z1, z2$). Difference of the points $z1$ and $z2$ on the elliptic curve corresponding to E .

The library syntax is GEN `ellsub`(GEN E , GEN $z1$, GEN $z2$).

3.15.81 ellsupersingularj(p). Return a random supersingular j -invariant defined over $\mathbf{F}_p^2$ as a `t_FFELT` in the variable `w`, if p is a prime number, or over the field of definition of p if p is a `t_FFELT`. The field must be of even degree. The random distribution is close to uniform except when 0 or 1728 are supersingular j -invariants, in which case they are less likely to be returned. This bias becomes negligible as p grows.

```
? j = ellsupersingularj(1009)
%1 = 12*w+295
? ellissupersingular(j)
%2 = 1
? a = ffgen([1009,2], 'a);
? j = ellsupersingularj(a)
%4 = 867*a+721
? ellissupersingular(j)
%5 = 1
? E = ellinit([j]);
? F = elltwist(E);
? ellissupersingular(F)
%8 = 1
? ellap(E)
%9 = 2018
? ellap(F)
%10 = -2018
```

The library syntax is GEN `ellsupersingularj`(GEN p).

3.15.82 elltamagawa(E). The object E being an elliptic curve over a number field, returns the global Tamagawa number of the curve (including the factor at infinite places).

```
? e = ellinit([1, -1, 1, -3002, 63929]); \\ curve "90c6" from elldata
? elltamagawa(e)
%2 = 288
? [elllocalred(e,p)[4] | p<-[2,3,5]]
%3 = [6, 4, 6]
? vecprod(%) \\ since e.disc > 0 the factor at infinity is 2
%4 = 144
? ellglobalred(e)[4] \\ product without the factor at infinity
%5 = 144
```

The library syntax is GEN `elltamagawa(GEN E)`.

3.15.83 elltaniyama($E, \{n = \text{seriesprecision}\}$). Computes the modular parametrization of the elliptic curve $E/\mathbf{Q}$, where E is an `ell` structure as output by `ellinit`. This returns a two-component vector $[u, v]$ of power series, given to n significant terms (`seriesprecision` by default), characterized by the following two properties. First the point (u, v) satisfies the equation of the elliptic curve. Second, let N be the conductor of E and $\Phi : X_0(N) \rightarrow E$ be a modular parametrization; the pullback by Φ of the Néron differential $du/(2v + a_1u + a_3)$ is equal to $2i\pi f(z)dz$, a holomorphic differential form. The variable used in the power series for u and v is x , which is implicitly understood to be equal to $\exp(2i\pi z)$.

The algorithm assumes that E is a *strong* Weil curve and that the Manin constant is equal to 1: in fact, $f(x) = \sum_{n>0} \text{ellak}(E, n)x^n$.

The library syntax is GEN `elltaniyama(GEN E, long prec1)`.

3.15.84 elltatepairing(E, P, Q, m). Let E be an elliptic curve defined over a finite field k and $m \geq 1$ be an integer. This function computes the (nonreduced) Tate pairing of the points P and Q on E , where P is an m -torsion point. More precisely, let $f_{m,P}$ denote a Miller function with divisor $m[P] - m[O_E]$; the algorithm returns $f_{m,P}(Q) \in k^*/(k^*)^m$.

The library syntax is GEN `elltatepairing(GEN E, GEN P, GEN Q, GEN m)`.

3.15.85 elltors(E). If E is an elliptic curve defined over a number field or a finite field, outputs the torsion subgroup of E as a 3-component vector $[\mathbf{t}, \mathbf{v1}, \mathbf{v2}]$, where $\mathbf{t}$ is the order of the torsion group, $\mathbf{v1}$ gives the structure of the torsion group as a product of cyclic groups (sorted by decreasing order), and $\mathbf{v2}$ gives generators for these cyclic groups. E must be an `ell` structure as output by `ellinit`.

```
? E = ellinit([-1,0]);
? elltors(E)
%1 = [4, [2, 2], [[0, 0], [1, 0]]]
```

Here, the torsion subgroup is isomorphic to $\mathbf{Z}/2\mathbf{Z} \times \mathbf{Z}/2\mathbf{Z}$, with generators $[0, 0]$ and $[1, 0]$.

The library syntax is GEN `elltors(GEN E)`.

3.15.86 elltrace(E, P). Let E be an elliptic curve over a base field and a point P defined over an extension field using `t_POLMOD` constructs. Returns the sum of the Galois conjugates of P . The field over which P is defined must be specified, even in the (silly) case of a trivial extension:

```
? E = ellinit([1,15]); \\ y^2 = x^3 + x + 15, over Q
? P = Mod([a/8-1, 1/32*a^2-11/32*a-19/4], a^3-135*a-408);
? ellisoncurve(E,P) \\ P defined over a cubic extension
%3 = 1
? elltrace(E,P)
%4 = [2,-5]

? E = ellinit([-13^2, 0]);
? P = Mod([13,0], a^2-2); \\ defined over Q, seen over a quadratic extension
? elltrace(E,P) == ellmul(E,P,2)
%3 = 1
? elltrace(E,[13,0]) \\ number field of definition of the point unspecified!
*** at top-level: elltrace(E,[13,0])
*** ^-----
*** elltrace: incorrect type in elltrace (t_INT).
? elltrace(E,Mod([13,0],a)) \\ trivial extension
%5 = [Mod(13, a), Mod(0, a)]
? P = Mod([-10*x^3+10*x-13, -16*x^3+16*x-34], x^4-x^3+2*x-1);
? ellisoncurve(E,P)
%7 = 1
? Q = elltrace(E,P)
%8 = [11432100241 / 375584400, 1105240264347961 / 7278825672000]
? ellisoncurve(E,Q)
%9 = 1

? E = ellinit([2,3], 19); \\ over F_19
? T = a^5+a^4+15*a^3+16*a^2+3*a+1; \\ irreducible
? P = Mod([11*a^3+11*a^2+a+12,15*a^4+9*a^3+18*a^2+18*a+6], T);
? ellisoncurve(E, P)
%4 = 1
? Q = elltrace(E, P)
%5 = [Mod(1,19), Mod(14,19)]
? ellisoncurve(E, Q)
%6 = 1
```

The library syntax is `GEN elltrace(GEN E, GEN P)`.

3.15.87 elltwist($E, \{P\}$). Returns an `ell` structure (as given by `ellinit`) for the twist of the elliptic curve E by the quadratic extension of the coefficient ring defined by P (when P is a polynomial) or `quadpoly(P)` when P is an integer. If E is defined over a finite field, then P can be omitted, in which case a random model of the unique nontrivial twist is returned. If E is defined over a number field, the model should be replaced by a minimal model (if one exists).

The elliptic curve E can be given in some of the formats allowed by `ellinit`: an `ell` structure, a 5-component vector $[a_1, a_2, a_3, a_4, a_6]$ or a 2-component vector $[a_4, a_6]$.

Twist by discriminant -3 :

```
? elltwist([0,a2,0,a4,a6], -3)[1..5]
%1 = [0, -3*a2, 0, 9*a4, -27*a6]
? elltwist([a4,a6], -3)[1..5]
%2 = [0, 0, 0, 9*a4, -27*a6]
```

Twist by the Artin-Schreier extension given by $x^2 + x + T$ in characteristic 2:

```
? lift(elltwist([a1,a2,a3,a4,a6]*Mod(1,2), x^2+x+T)[1..5])
%1 = [a1, a2+a1^2*T, a3, a4, a6+a3^2*T]
```

Twist of an elliptic curve defined over a finite field:

```
? E = elltwist([1,7]*Mod(1,19)); lift([E.a4, E.a6])
%1 = [11, 12]
```

The library syntax is `GEN elltwist(GEN E, GEN P = NULL)`.

3.15.88 ellweilcurve($E, \{&ms\}$). If E' is an elliptic curve over $\mathbf{Q}$, let $L_{E'}$ be the sub- $\mathbf{Z}$ -module of $\text{Hom}_{\Gamma_0(N)}(\Delta_0, \mathbf{Q})$ attached to E' (It is given by $x[3]$ if $[M, x] = \text{msfromell}(E')$.)

On the other hand, if N is the conductor of E and f is the modular form for $\Gamma_0(N)$ attached to E , let L_f be the lattice of the f -component of $\text{Hom}_{\Gamma_0(N)}(\Delta_0, \mathbf{Q})$ given by the elements ϕ such that $\phi(\{0, \gamma^{-1}0\}) \in \mathbf{Z}$ for all $\gamma \in \Gamma_0(N)$ (see `mslattice`).

Let E' run through the isomorphism classes of elliptic curves isogenous to E as given by `ellisomat` (and in the same order). This function returns a pair `[vE, vS]` where `vE` contains minimal models for the E' and `vS` contains the list of Smith invariants for the lattices $L_{E'}$ in L_f . The function also accepts the output of `ellisomat`, i.e. the isogeny class. If the optional argument `ms` is present, it contains the output of `msfromell(vE, 0)`, i.e. the new modular symbol space M of level N and a vector of triples $[x^+, x^-, L]$ attached to each curve E' .

In particular, the strong Weil curve amongst the curves isogenous to E is the one whose Smith invariants are $[c, c]$, where c is the Manin constant, conjecturally equal to 1.

```
? E = ellinit("11a3");
? [vE, vS] = ellweilcurve(E);
? [n] = [ i | i<-[1..#vS], vS[i]==[1,1] ]  \\ lattice with invariant [1,1]
%3 = [2]
? ellidentify(vE[n])  \\ ... corresponds to strong Weil curve
%4 = [["11a1", [0, -1, 1, -10, -20], []], [1, 0, 0, 0]]
? [vE, vS] = ellweilcurve(E, &ms);  \\ vE, vS are as above
? [M, vx] = ms; msdim(M)  \\ ... but ms contains more information
%6 = 3
? #vx
%7 = 3
? vx[1]
%8 = [[1/25, -1/10, -1/10]~, [0, 1/2, -1/2]~, [1/25, 0; -3/5, 1; 2/5, -1]]
? forell(E, 11, 11, print(msfromell(ellinit(E[1]), 1)[2]))
[1/5, -1/2, -1/2]~
[1, -5/2, -5/2]~
[1/25, -1/10, -1/10]~
```

The last example prints the modular symbols x^+ in M^+ attached to the curves 11a1, 11a2 and 11a3.

The library syntax is `GEN ellweilcurve(GEN E, GEN *ms = NULL)`.

3.15.89 ellweilpairing(E, P, Q, m). Let E be an elliptic curve defined over a finite field and $m \geq 1$ be an integer. This function computes the Weil pairing of the two m -torsion points P and Q on E , which is an alternating bilinear map. More precisely, let $f_{m,R}$ denote a Miller function with divisor $m[R] - m[O_E]$; the algorithm returns the m -th root of unity

$$\varepsilon(P, Q)^m \cdot f_{m, P}(Q) / f_{m, Q}(P),$$

where $f(R)$ is the extended evaluation of f at the divisor $[R] - [O_E]$ and $\varepsilon(P, Q) \in \{\pm 1\}$ is given by Weil reciprocity: $\varepsilon(P, Q) = 1$ if and only if P, Q, O_E are not pairwise distinct.

The library syntax is `GEN ellweilpairing(GEN E, GEN P, GEN Q, GEN m)`.

3.15.90 ellwp($w, \{z = x\}, \{flag = 0\}$). Computes the value at z of the Weierstrass $\wp$ function attached to the lattice w as given by **ellperiods**. It is also possible to directly input $w = [\omega_1, \omega_2]$, or an elliptic curve E as given by **ellinit** ($w = E.\text{omega}$).

```
? w = ellperiods([1,I]);
? ellwp(w, 1/2)
%2 = 6.8751858180203728274900957798105571978
? E = ellinit([1,1]);
? ellwp(E, 1/2)
%4 = 3.94131112427016474646048282462709151389
```

One can also compute the series expansion around $z = 0$:

```
? E = ellinit([1,0]);
? ellwp(E) \\ 'x implicitly at default seriesprecision
%5 = x^-2 - 1/5*x^2 + 1/75*x^6 - 2/4875*x^10 + 0(x^14)
? ellwp(E, x + 0(x^12)) \\ explicit precision
%6 = x^-2 - 1/5*x^2 + 1/75*x^6 + 0(x^9)
```

Optional *flag* means 0 (default): compute only $\varphi(z)$, 1: compute $[\varphi(z), \varphi'(z)]$.

For instance, the Dickson elliptic functions sm and cm can be implemented as follows

```

smcm(z) =
{ my(a, b, E = ellinit([0,-1/(4*27)])); \\ ell. invariants (g2,g3)=(0,1/27)
  [a,b] = ellwp(E, z, 1);
  [6*a / (1-3*b), (3*b+1)/(3*b-1)];
}
? [s,c] = smcm(0.5);
? s
%2 = 0.4898258757782682170733218609
? c
%3 = 0.9591820206453842491187464098
? s^3+c^3
%4 = 1.000000000000000000000000000000
? smcm('x + O('x^11))
%5 = [x - 1/6*x^4 + 2/63*x^7 - 13/2268*x^10 + O(x^11),
      1 - 1/3*x^3 + 1/18*x^6 - 23/2268*x^9 + O(x^10)]

```

The library syntax is `GEN ellwp0(GEN w, GEN z = NULL, long flag, long prec)`. For `flag = 0`, we also have `GEN ellwp(GEN w, GEN z, long prec)`, and `GEN ellwpseries(GEN E, long v, long precdl)` for the power series in variable v .

3.15.91 ellxn($E, n, \{v = 'x\}$). For any affine point $P = (t, u)$ on the curve E , we have

$$[n]P = (\phi_n(P)\psi_n(P) : \omega_n(P) : \psi_n(P)^3)$$

for some ϕ_n, ω_n, ψ_n in $\mathbf{Z}[a_1, a_2, a_3, a_4, a_6][t, u]$ modulo the curve equation. This function returns a pair $[A, B]$ of polynomials in $\mathbf{Z}[a_1, a_2, a_3, a_4, a_6][v]$ such that $[A(t), B(t)] = [\phi_n(P), \psi_n(P)^2]$ in the function field of E , whose quotient give the abscissa of $[n]P$. If P is an n -torsion point, then $B(t) = 0$.

```
? E = ellinit([17,42]); [t,u] = [114,1218];
? T = ellxn(E, 2, 'X)
%2 = [X^4 - 34*X^2 - 336*X + 289, 4*X^3 + 68*X + 168]
? [a,b] = subst(T,'X,t);
%3 = [168416137, 5934096]
? a / b == ellmul(E, [t,u], 2)[1]
%4 = 1
```

The library syntax is `GEN ellxn(GEN E, long n, long v = -1)` where v is a variable number.

3.15.92 ellzeta($w, \{z = 'x\}$). Computes the value at z of the Weierstrass ζ function attached to the lattice w as given by `ellperiods(1)`: including quasi-periods is useful, otherwise there are recomputed from scratch for each new z .

$$\zeta(z, L) = \frac{1}{z} + z^2 \sum_{\omega \in L^*} \frac{1}{\omega^2(z - \omega)}.$$

It is also possible to directly input $w = [\omega_1, \omega_2]$, or an elliptic curve E as given by `ellinit` ($w = E.\text{omega}$). The quasi-periods of ζ , such that

$$\zeta(z + a\omega_1 + b\omega_2) = \zeta(z) + a\eta_1 + b\eta_2$$

for integers a and b are obtained as $\eta_i = 2\zeta(\omega_i/2)$. Or using directly `ellzeta`.

```
? w = ellperiods([1,I],1);
? ellzeta(w, 1/2)
%2 = 1.5707963267948966192313216916397514421
? E = ellinit([1,0]);
? ellzeta(E, E.omega[1]/2)
%4 = 0.84721308479397908660649912348219163647
```

One can also compute the series expansion around $z = 0$ (the quasi-periods are useless in this case):

```
? E = ellinit([0,1]);
? ellzeta(E) \\ at 'x, implicitly at default seriesprecision
%4 = x^-1 + 1/35*x^5 - 1/7007*x^11 + 0(x^15)
? ellzeta(E, x + 0(x^20)) \\ explicit precision
%5 = x^-1 + 1/35*x^5 - 1/7007*x^11 + 1/1440257*x^17 + 0(x^18)
```

The library syntax is `GEN ellzeta(GEN w, GEN z = NULL, long prec)`.


```

%4 = [Mod(0, a^2 - 3), Mod(9600*a, a^2 - 3), Mod(186624, a^2 - 3),
      Mod(-69120000, a^2 - 3), Mod(-241864704*a + 204800000, a^2 - 3)]
? a = ffgen([3,4], 'a'); \\ over  $\mathbf{F}_{3^4} = \mathbf{F}_3[a]$ 
? genus2igusa(x^6+a*x^5-a*x^4+2*x^3+a*x+a+1)
%6 = [2*a^2, a^3 + a^2 + a + 1, a^2 + a + 2, 2*a^3 + 2*a^2 + a + 1,
      2*a^2 + 2]
? a = ffgen([2,4], 'a'); \\  $\mathbf{F}_{2^4} = \mathbf{F}_2[a]$ 
? genus2igusa(x^6+a*x^5+a*x^4+a*x+a+1) \\ doesn't work in characteristic 2
*** at top-level: genus2igusa(x^6+a*x^5+a*x^4+a*x+a+1)
*** ^-----
*** genus2igusa: impossible inverse in FF_mul2n: 2.

```

The library syntax is GEN `genus2igusa(GEN PQ, long k)`.

3.15.95 `genus2red(PQ, {p})`. Let PQ be a polynomial P , resp. a vector $[P, Q]$ of polynomials, with rational coefficients. Determines the reduction at $p > 2$ of the (proper, smooth) genus 2 curve $C/\mathbf{Q}$, defined by the hyperelliptic equation $y^2 = P(x)$, resp. $y^2 + Q(x) * y = P(x)$. (The special fiber X_p of the minimal regular model X of C over $\mathbf{Z}$.)

If p is omitted, determines the reduction type for all (odd) prime divisors of the discriminant.

This function was rewritten from an implementation of Liu's algorithm by Cohen and Liu (1994), `genus2reduction-0.3`, see <https://www.math.u-bordeaux.fr/~liu/G2R/>.

CAVEAT. The function interface may change: for the time being, it returns $[N, FaN, [P_m, Q_m], V]$ where N is either the local conductor at p or the global conductor, FaN is its factorization, $y^2 + Q_my = P_m$ defines a minimal model over $\mathbf{Z}$ and V describes the reduction type at the various considered p . Unfortunately, the program is not complete for $p = 2$, and we may return the odd part of the conductor only: this is the case if the factorization includes the (impossible) term 2^{-1} ; if the factorization contains another power of 2, then this is the exact local conductor at 2 and N is the global conductor.

```

? default(debuglevel, 1);
? genus2red(x^6 + 3*x^3 + 63, 3)
(potential) stable reduction: [1, []]
reduction at p: [III{9}] page 184, [3, 3], f = 10
%1 = [59049, Mat([3, 10]), x^6 + 3*x^3 + 63, [3, [1, []],
      ["[III{9}] page 184", [3, 3]]]]
? [N, FaN, T, V] = genus2red(x^3-x^2-1, x^2-x); \\ X_1(13), global reduction
p = 13
(potential) stable reduction: [5, [Mod(0, 13), Mod(0, 13)]]
reduction at p: [I{0}-II-0] page 159, [], f = 2
? N
%3 = 169
? FaN
%4 = Mat([13, 2]) \\ in particular, good reduction at 2 !
? T
%5 = x^6 + 58*x^5 + 1401*x^4 + 18038*x^3 + 130546*x^2 + 503516*x + 808561
? V
%6 = [[13, [5, [Mod(0, 13), Mod(0, 13)]]], ["[I{0}-II-0] page 159", []]]

```

We now first describe the format of the vector $V = V_p$ in the case where p was specified (local reduction at p): it is a triple $[p, \text{stable}, \text{red}]$. The component $\text{stable} = [\text{type}, \text{vecj}]$ contains information about the stable reduction after a field extension; depending on types , the stable reduction is

- 1: smooth (i.e. the curve has potentially good reduction). The Jacobian $J(C)$ has potentially good reduction.
- 2: an elliptic curve E with an ordinary double point; vecj contains $j \bmod p$, the modular invariant of E . The (potential) semi-abelian reduction of $J(C)$ is the extension of an elliptic curve (with modular invariant $j \bmod p$) by a torus.
- 3: a projective line with two ordinary double points. The Jacobian $J(C)$ has potentially multiplicative reduction.
- 4: the union of two projective lines crossing transversally at three points. The Jacobian $J(C)$ has potentially multiplicative reduction.
- 5: the union of two elliptic curves E_1 and E_2 intersecting transversally at one point; vecj contains their modular invariants j_1 and j_2 , which may live in a quadratic extension of $\mathbf{F}_p$ and need not be distinct. The Jacobian $J(C)$ has potentially good reduction, isomorphic to the product of the reductions of E_1 and E_2 .
- 6: the union of an elliptic curve E and a projective line which has an ordinary double point, and these two components intersect transversally at one point; vecj contains $j \bmod p$, the modular invariant of E . The (potential) semi-abelian reduction of $J(C)$ is the extension of an elliptic curve (with modular invariant $j \bmod p$) by a torus.
- 7: as in type 6, but the two components are both singular. The Jacobian $J(C)$ has potentially multiplicative reduction.

The component $\text{red} = [N\text{Utype}, \text{neron}]$ contains two data concerning the reduction at p without any ramified field extension.

The $N\text{Utype}$ is a `t_STR` describing the reduction at p of C , following Namikawa-Ueno, *The complete classification of fibers in pencils of curves of genus two*, Manuscripta Math., vol. 9, (1973), pages 143-186. The reduction symbol is followed by the corresponding page number or page range in this article.

The second datum neron is the group of connected components (over an algebraic closure of $\mathbf{F}_p$) of the Néron model of $J(C)$, given as a finite abelian group (vector of elementary divisors).

If $p = 2$, the red component may be omitted altogether (and replaced by `[]`, in the case where the program could not compute it. When p was not specified, V is the vector of all V_p , for all considered p .

Notes about Namikawa-Ueno types.

- A lower index is denoted between braces: for instance, $[I\{2\}-II-5]$ means $[I_2-II-5]$.
- If K and K' are Kodaira symbols for singular fibers of elliptic curves, then $[K-K'-m]$ and $[K'-K-m]$ are the same.

We define a total ordering on Kodaira symbol by fixing $I < I^* < II < II^*, \dots$. If the reduction type is the same, we order by the number of components, e.g. $I_2 < I_4$, etc. Then we normalize our output so that $K \leq K'$.

- $[K-K'-1]$ is $[K-K'-\alpha]$ in the notation of Namikawa-Ueno.
- The figure $[2I_0-m]$ in Namikawa-Ueno, page 159, must be denoted by $[2I_0-(m+1)]$.

The library syntax is `GEN genus2red(GEN PQ, GEN p = NULL)`.

3.15.96 hyperellchangecurve(C, m). C being a nonsingular hyperelliptic model of a curve, apply the change of coordinate given by $m = [e, [a, b; c, d], H]$.

If (x, y) is a point on the new model, the corresponding point (X, Y) on C is given by

$$X = (a * x + b) / (c * x + d), \quad Y = e(y + H(x)) / (c * x + d)^{g+1}.$$

C can be given either by a squarefree polynomial P such that $C : y^2 = P(x)$ or by a vector $[P, Q]$ such that $C : y^2 + Q(x)y = P(x)$ and $Q^2 + 4P$ is squarefree.

The library syntax is `GEN hyperellchangecurve(GEN C, GEN m)`.

3.15.97 hyperellcharpoly(X). X being a nonsingular hyperelliptic curve defined over a finite field, return the characteristic polynomial of the Frobenius automorphism. X can be given either by a squarefree polynomial P such that $X : y^2 = P(x)$ or by a vector $[P, Q]$ such that $X : y^2 + Q(x)y = P(x)$ and $Q^2 + 4P$ is squarefree.

The library syntax is `GEN hyperellcharpoly(GEN X)`.

3.15.98 hyperelldisc(X). X being a nonsingular hyperelliptic model of a curve, defined over a field of characteristic distinct from 2, returns its discriminant. X can be given either by a squarefree polynomial P such that X has equation $y^2 = P(x)$ or by a vector $[P, Q]$ such that X has equation $y^2 + Q(x)y = P(x)$ and $Q^2 + 4P$ is squarefree.

```
? hyperelldisc([x^3,1])
%1 = -27
? hyperelldisc(x^5+1)
%2 = 800000
```

The library syntax is `GEN hyperelldisc(GEN X)`.

3.15.99 hyperellisoncurve(X, p). X being a nonsingular hyperelliptic model of a curve, test whether the point p is on the curve.

X can be given either by a squarefree polynomial P such that $X : y^2 = P(x)$ or by a vector $[P, Q]$ such that $X : y^2 + Q(x)y = P(x)$ and $Q^2 + 4P$ is squarefree.

```
? W = [2*x^6+3*x^5+x^4+x^3-x,x^3+1]; p = [px, py] = [1/3,-14/27];
? hyperellisoncurve(W, p)
%2 = 1
? [Px,Qx]=subst(W,x,px); py^2+py*Qx == Px
%3 = 1
```

The library syntax is `int hyperellisoncurve(GEN X, GEN p)`.

3.15.100 hyperellminimaldisc($C, \{pr\}$). C being a nonsingular integral hyperelliptic model of a curve, return the minimal discriminant of an integral model of C . If pr is given, it must be a list of primes and the discriminant is then only guaranteed minimal at the elements of pr . C can be given either by a squarefree polynomial P such that $C : y^2 = P(x)$ or by a vector $[P, Q]$ such that $C : y^2 + Q(x)y = P(x)$ and $Q^2 + 4P$ is squarefree.

```
? W = [x^6+216*x^3+324,0];
? D = hyperelldisc(W)
%2 = 1828422898924853919744000
? M = hyperellminimaldisc(W)
%4 = 29530050606000
```

The library syntax is `GEN hyperellminimaldisc(GEN C, GEN pr = NULL)`.

3.15.101 hyperellminimalmodel($C, \{&m\}, \{pr\}$). C being a nonsingular integral hyperelliptic model of a curve, return an integral model of C with minimal discriminant. If pr is given, it must be a list of primes and the model is then only guaranteed minimal at the elements of pr . If present, m is set to the mapping from the original model to the new one: a three-component vector $[e, [a, b; c, d], H]$ such that if (x, y) is a point on W , the corresponding point on C is given by

$$x_C = (a * x + b) / (c * x + d)$$

$$y_C = (e * y + H(x)) / (c * x + d)^{g+1}$$

where g is the genus. C can be given either by a squarefree polynomial P such that $C : y^2 = P(x)$ or by a vector $[P, Q]$ such that $C : y^2 + Q(x)y = P(x)$ and $Q^2 + 4P$ is squarefree.

```
? W = [x^6+216*x^3+324,0];
? D = hyperelldisc(W)
%2 = 1828422898924853919744000
? Wn = hyperellminimalmodel(W,&M)
%3 = [2*x^6+18*x^3+1,x^3];
? M
%4 = [18, [3, 0; 0, 1], 9*x^3]
? hyperelldisc(Wn)
%5 = 29530050606000
? hyperellchangecurve(W, M)
%6 = [2*x^6+18*x^3+1,x^3]
```

The library syntax is `GEN hyperellminimalmodel(GEN C, GEN *m = NULL, GEN pr = NULL)`

3.15.102 hyperellordinate(H, x). Gives a 0, 1 or 2-component vector containing the y -coordinates of the points of the curve H having x as x -coordinate.

The library syntax is GEN hyperellordinate(GEN H, GEN x).

3.15.103 hyperellpadicfrobenius(Q, q, n). Let X be the curve defined by $y^2 = Q(x)$, where Q is a polynomial of degree d over $\mathbf{Q}$ and $q \geq d$ is a prime such that X has good reduction at q . Return the matrix of the Frobenius endomorphism φ on the crystalline module $D_p(X) = \mathbf{Q}_p \otimes H_{dR}^1(X/\mathbf{Q})$ with respect to the basis of the given model $(\omega, x\omega, \dots, x^{g-1}\omega)$, where $\omega = dx/(2y)$ is the invariant differential, where g is the genus of X (either $d = 2g + 1$ or $d = 2g + 2$). The characteristic polynomial of φ is the numerator of the zeta-function of the reduction of the curve X modulo q . The matrix is computed to absolute q -adic precision q^n .

Alternatively, q may be of the form $[T, p]$ where p is a prime, T is a polynomial with integral coefficients whose projection to $\mathbf{F}_p[t]$ is irreducible, X is defined over $K = \mathbf{Q}[t]/(T)$ and has good reduction to the finite field $\mathbf{F}_q = \mathbf{F}_p[t]/(T)$. The matrix of φ on $D_q(X) = \mathbf{Q}_q \otimes H_{dR}^1(X/K)$ is computed to absolute p -adic precision p^n .

```
? M=hyperellpadicfrobenius(x^5+'a*x+1,['a^2+1,3],10);
? liftall(M)
[48107*a + 38874  9222*a + 54290  41941*a + 8931  39672*a + 28651]
[ 21458*a + 4763  3652*a + 22205  31111*a + 42559  39834*a + 40207]
[ 13329*a + 4140  45270*a + 25803  1377*a + 32931  55980*a + 21267]
[15086*a + 26714  33424*a + 4898  41830*a + 48013  5913*a + 24088]
? centerlift(simplify(liftpol(charpoly(M))))
%8 = x^4+4*x^2+81
? hyperellcharpoly((x^5+Mod(a,a^2+1)*x+1)*Mod(1,3))
%9 = x^4+4*x^2+81
```

The library syntax is GEN hyperellpadicfrobenius0(GEN Q, GEN q, long n). The functions GEN hyperellpadicfrobenius(GEN H, ulong p, long n) and GEN nfhyperellpadicfrobenius(GEN H, GEN T, ulong p, long n) are also available.

3.15.104 hyperellratpoints($X, h, \{flag = 0\}$). X being a nonsingular hyperelliptic curve given by an rational model, return a vector containing the affine rational points on the curve of naive height less than h . If $flag = 1$, stop as soon as a point is found; return either an empty vector or a vector containing a single point.

X is given either by a squarefree polynomial P such that $X : y^2 = P(x)$ or by a vector $[P, Q]$ such that $X : y^2 + Q(x)y = P(x)$ and $Q^2 + 4P$ is squarefree.

The parameter h can be

- an integer H : find the points $[n/d, y]$ whose abscissas $x = n/d$ have naive height ($= \max(|n|, d)$) less than H ;
- a vector $[N, D]$ with $D \leq N$: find the points $[n/d, y]$ with $|n| \leq N, d \leq D$.
- a vector $[N, [D_1, D_2]]$ with $D_1 < D_2 \leq N$ find the points $[n/d, y]$ with $|n| \leq N$ and $D_1 \leq d \leq D_2$.

The library syntax is GEN hyperellratpoints(GEN X, GEN h, long flag).

3.15.105 hyperellred($C, \{m\}$). Let C be a nonsingular integral hyperelliptic model of a curve of positive genus $g > 0$. Return an integral model of C with the same discriminant but small coefficients, using Cremona-Stoll reduction.

The optional argument m is set to the mapping from the original model to the new one, given by a three-component vector $[1, [a, b; c, d], H]$ such that $a * d - b * c = 1$ and if (x, y) is a point on W , the corresponding point (X, Y) on C is given by

$$X = (a * x + b) / (c * x + d), \quad Y = (y + H(x)) / (c * x + d)^{g+1}.$$

C can be given either by a squarefree polynomial P such that $C : y^2 = P(x)$ or by a vector $[P, Q]$ such that $C : y^2 + Q(x)y = P(x)$ and $Q^2 + 4P$ is squarefree.

```
? P = 1001*x^4 + 3704*x^3 + 5136*x^2 + 3163*x + 730;
? hyperellred(P, &m)
%2 = [x^3 + 1, 0]
? hyperellchangecurve(P, m)
%3 = [x^3 + 1, 0]
```

The library syntax is `GEN hyperellred(GEN C, GEN *m = NULL)`.

Also available is `GEN ZX_hyperellred(GEN P, GEN *M)` where $C : y^2 = P(x)$ and $*M$ is set to $[a, b; c, d]$

3.16 Hypergeometric Motives.

3.16.1 Templates.

A *hypergeometric template* is a pair of multisets (i.e., sets with possibly repeated elements) of rational numbers $(\alpha_1, \dots, \alpha_d)$ and $(\beta_1, \dots, \beta_d)$ having the same number of elements, and we set

$$A(x) = \prod_{1 \leq j \leq d} (x - e^{2\pi i \alpha_j}), \quad B(x) = \prod_{1 \leq k \leq d} (x - e^{2\pi i \beta_k}).$$

We make the following assumptions:

- $\alpha_j - \beta_k \notin \mathbf{Z}$ for all j and k , or equivalently $\gcd(A, B) = 1$.
- $\alpha_j \notin \mathbf{Z}$ for all j , or equivalently $A(1) \neq 0$.
- our template is *defined over* $\mathbf{Q}$, in other words $A, B \in \mathbf{Z}[x]$, or equivalently if some a/D with $\gcd(a, D) = 1$ occurs in the α_j or β_k , then all the b/D modulo 1 with $\gcd(b, D) = 1$ also occur.

The last assumption allows to abbreviate $[a_1/D, \dots, a_{\phi(D)}/D]$ (where the a_i range in $(\mathbf{Z}/D\mathbf{Z})^*$) to $[D]$. We thus have two possible ways of giving a hypergeometric template: either by the two vectors $[\alpha_1, \dots, \alpha_d]$ and $[\beta_1, \dots, \beta_d]$, or by their denominators $[D_1, \dots, D_m]$ and $[E_1, \dots, E_n]$, which are called the *cyclotomic parameters*; note that $\sum_j \phi(D_j) = \sum_k \phi(E_k) = d$. A third way is to give the *gamma vector* (γ_n) defined by $A(X)/B(X) = \prod_n (X^n - 1)^{\gamma_n}$, which satisfies $\sum_n n\gamma_n = 0$. To any such data we associate a hypergeometric template using the function `hgminit`; then the α_j and β_k are obtained using `hgmalph`, cyclotomic parameters using `hgmcyclo` and the gamma vectors using `hgmgamma`.

To such a hypergeometric template is associated a number of additional parameters, for which we do not give the definition but refer to the survey *Hypergeometric Motives* by Roberts and

Villegas, <https://arxiv.org/abs/2109.00027>: the degree d , the *weight* w , a *Hodge polynomial* P , a *Tate twist* T , and a normalizing M-factor $M = \prod_n n^{n\gamma_n}$. The `hgmparams` function returns

$$[d, w, [P, T], M] .$$

Example with cyclotomic parameters [5], [1, 1, 1, 1]:

```
? H = hgmininit([5]); \\ [1,1,1,1] can be omitted
? hgmparams(H)
%2 = [4, 3, [x^3+x^2+x+1,0], 3125]
? hgmalpha(H)
%3 = [[1/5, 2/5, 3/5, 4/5], [0, 0, 0, 0]]
? hgmcycolo(H)
%4 = [Vecsmall([5]), Vecsmall([1, 1, 1, 1])]
? hgmgamma(H)
%5 = Vecsmall([-5,0,0,0,1]) \\ A/B = (x^5-1) / (x-1)^5
```

3.16.2 Motives.

A *hypergeometric motive* (HGM for short) is a pair (H, t) , where H is a hypergeometric template and $t \in \mathbf{Q}^*$. To such a motive and a finite field $\mathbf{F}_q$ one can associate via an explicit but complicated formula an *integer* $N(H, t; q)$, see Beukers, Cohen and Mellit, *Finite hypergeometric functions* Pure and Applied Math Quarterly 11 (2015), pp 559 - 589, <https://arxiv.org/abs/1505.02900>.

Warning. Depending on the authors, t may have to be replaced with $1/t$. The Pari/GP convention is the same as the one in Magma, but is the inverse of the one in the last reference.

This formula does not make sense and is not valid for *bad primes* p : a *wild prime* is a prime which divides a denominator of the α_j or β_i . If a prime p is not wild, it can be *good* if $v_p(t) = v_p(t-1) = 0$, or *tame* otherwise. The *local Euler factor* P_p at a good prime p is then given by the usual formula

$$-\log P_p(T) = \sum_{f \geq 1} \frac{N(H, t; p^f) T^f}{f} ,$$

and in the case of HGM's P_p is always a polynomial (note that the Euler factor used in the global L -function is $1/P_p(p^{-s})$). At a tame prime p it is necessary to modify the above formula, and usually (but not always) the degree of the local Euler factor decreases. Wild primes are currently not implemented by a formula but can be guessed via the global functional equation (see the next section). Continuing the previous example, we find

```
? hgmeulerfactor(H, -1, 3) \\ good prime
%4 = 729*x^4 + 135*x^3 + 45*x^2 + 5*x + 1
? hgmeulerfactor(H, -1, 2) \\ tame prime
%5 = 16*x^3 + 6*x^2 + x + 1
? hgmeulerfactor(H, -1, 5) \\ wild primes not implemented
%6 = 0
```

To obtain the Euler factor at wild primes, use `lfuneuler` once the global L -function is computed.

3.16.3 The Global L -function.

A theorem of Katz tells us that if one suitably defines $P_p(T)$ for all primes p including the wild ones, then the L -function defined by $L(H, s) = \prod_p P_p(p^{-s})^{-1}$ is motivic, with analytic continuation and functional equation, as used in the L -function package of **Pari/GP**. The command `L = lfunhgm(H,t)` creates such an L -function. In particular it must guess the local Euler factors at wild primes, which can be very expensive when the conductor `lfunparams(L)[1]` is large.

In our example, `L = lfunhgm(H,1/64)` finishes in about 20 seconds (the conductor is only 525000); this L -function can then be used with all the functions of the `lfun` package. For instance we can now obtain the global conductor and check the Euler factors at all bad primes:

```
? [N] = lfunparams(L); N \\ the conductor
%7 = 525000
? factor(N)
%8 =
[2 3]
[3 1]
[5 5]
[7 1]
? lfuneuler(L,2)
%9 = 1/(-x + 1)
? lfuneuler(L,3)
%10 = 1/(81*x^3 + 6*x^2 - 4*x + 1)
? lfuneuler(L,5)
%11 = 1
? lfuneuler(L,7)
%12 = 1/(2401*x^3 + 301*x^2 + x + 1)
```

Two additional functions related to the global L -function are available which do *not* require its full initialization: `hgmcoefs(H,t,n)` computes the first n coefficients of the L -function by setting all wild Euler factors to 1; this will be identical to `lfunan(L,n)` when this is indeed the case (as in the above example: only 5 is wild), otherwise all coefficients divisible by a wild prime will be wrong.

The second is the function `hgmcoef(H,t,n)` which only computes the n th coefficient of the global L -function. It gives an error if n is divisible by a wild prime. Compare `hgmcoefs(H,1/64,7^6)[7^6]` which requires more than 1 minute (it computes more than 100000 coefficients), with `hgmcoef(H,1/64,7^6)` which outputs -25290600 instantaneously.

3.16.4 hgmalpha(H). Returns the alpha and beta parameters of the hypergeometric motive template H .

```
? H = hgminit([5]); \\ template given by cyclotomic parameters
? hgmalpha(H)
%2 = [[1/5, 2/5, 3/5, 4/5], [0, 0, 0, 0]]
```

The library syntax is `GEN hgmalpha(GEN H)`.

3.16.5 hgmbydegree(n). Outputs $[L(0), \dots, L(n-1)]$ where $L(w)$ is the list of cyclotomic parameters of all possible hypergeometric motive templates of degree n and weight w .

The library syntax is GEN `hgmbydegree(long n)`.

3.16.6 hgmcoef(H, t, n). (H, t) being a hypergeometric motive, returns the n -th coefficient of its L -function. This is not implemented for wild primes p and will raise an exception if such a p divides n .

The library syntax is GEN `hgmcoef(GEN H, GEN t, GEN n)`.

3.16.7 hgmcoefs(H, t, n). (H, t) being a hypergeometric motive, returns the first n coefficients of its L -function, where Euler factors at wild primes are set to 1. The argument t may be replaced by $[t, bad]$ where bad is a vector of pairs $[p, L_p]$, p being a prime and L_p being the corresponding local Euler factor, overriding the default.

If you hope that the wild Euler factors can be computed not too slowly from the functional equation, you can also set `L=lfunhgm(H,t)`, and then `lfunan(L,n)`, and then the Euler factors at wild primes should be correct.

The library syntax is GEN `hgmcoefs(GEN H, GEN t, long n)`.

3.16.8 hgmcyκλο(H). Returns the cyclotomic parameters (D, E) of the hypergeometric motive template H .

```
\\ template given by alpha (implied beta is [0,0,0,0])
? H = hgminit([1/5, 2/5, 3/5, 4/5]);
? hgmcyκλο(H)
%3 = [Vecsmall([5]), Vecsmall([1, 1, 1, 1])]
? apply(Vec, %) \\ for readability
%4 = [[5], [1, 1, 1, 1]]
```

The library syntax is GEN `hgmcyκλο(GEN H)`.

3.16.9 hgmeulerfactor($H, t, p, \{&e\}$). (H, t) being a hypergeometric motive, returns the inverse of its Euler factor at the prime p and the exponent e of the conductor at p . This is not implemented when p is a wild prime: the function returns 0 and sets e to -1 . Caveat: contrary to `lfuneuler`, this function returns the *inverse* of the Euler factor, given by a polynomial P_p such that the Euler factor is $1/P_p(p^{-s})$.

```
? H = hgminit([5]); \\ cyclotomic parameters [5] and [1,1,1,1]
? hgmeulerfactor(H, 1/2, 3)
%2 = 729*x^4 + 135*x^3 + 45*x^2 + 5*x + 1
? hgmeulerfactor(H, 1/2, 3, &e)
%3 = 729*x^4 + 135*x^3 + 45*x^2 + 5*x + 1
? e
%4 = 0
? hgmeulerfactor(H, 1/2, 2, &e)
%5 = -x + 1
? e
%6 = 3
? hgmeulerfactor(H, 1/2, 5)
%7 = 0 \\ 5 is wild
```

If the conductor is small, the wild Euler factors can be computed from the functional equation: set `L = lfunhgm(H,t)` (the complexity should be roughly proportional to the conductor) then the `lfuneuler` function should give you the correct Euler factors at all primes:

```
? L = lfunhgm(H, 1/2);
time = 790 ms. \\ fast in this case, only 5 is wild
? lfunparams(L) \\ ... and the conductor 5000 is small
%8 = [5000, 4, [-1, 0, 0, 1]]
? lfun Euler(L, 5)
%9 = 1 \\ trivial Euler factor

? L = lfunhgm(H, 1/64); lfunparams(L)
time = 20,122 ms. \\ slower: the conductor is larger
%10 = [525000, 4, [-1, 0, 0, 1]]

? L = lfunhgm(H, 1/128); lfunparams(L)
time = 2min, 16,205 ms. \\ even slower, etc.
%11 = [3175000, 4, [-1, 0, 0, 1]]
```

The library syntax is GEN `hgmeulerfactor(GEN H, GEN t, long p, GEN *e = NULL)`.

3.16.10 hgmgamma(*H*). Returns the gamma vector of the hypergeometric motive template *H*.

```
? H = hgmininit([5]);
? hgmgamma(H)
%2 = Vecsmall([-5, 0, 0, 0, 1])
```

The library syntax is GEN `hgmgamma(GEN H)`.

3.16.11 hgmininit(*a*, {*b*}). Create the template for the hypergeometric motive with parameters *a* and possibly *b*. The format of the parameters may be

- alpha: lists of rational numbers $a = (\alpha_j)$ and $b = (\beta_k)$ of the same length (and defined over $\mathbf{Q}$); if *b* is omitted, we take it to be $(0, \dots, 0)$.

- cyclo: lists $a = D$ and $b = E$ of positive integers corresponding to the denominators of the (α_i) and (β_i) ; if *b* is omitted we take it to be $(1, \dots, 1)$. This is the simplest and most compact input format.

- gamma: list of γ_n such that the $\prod_j (x - \exp(2\pi i \alpha_j)) / \prod_k (x - \exp(2\pi i \beta_k)) = \prod_n (x^n - 1)^{\gamma_n}$.

The hypergeometric motive itself is given by a pair (H, t) , where *H* is a template as above and $t \in \mathbf{Q}^*$. Note that the motives given by $(\alpha, \beta; t)$ and $(\beta, \alpha; 1/t)$ are identical.

```
? H = hgmininit([5]); \\ template given by cyclotomic parameters 5 and 1,1,1,1
? L = lfunhgm(H, 1); \\ global L-function attached to motive (H,1)
? lfunparams(L)
%3 = [25, 4, [0, 1]]

? hgmalpha(H)
%4 = [[1/5, 2/5, 3/5, 4/5], [0, 0, 0, 0]]
? hgmgamma(H)
%5 = Vecsmall([-5, 0, 0, 0, 1])
```

The library syntax is GEN `hgmininit(GEN a, GEN b = NULL)`.

3.16.12 hgmissymmetrical(H). Is the hypergeometric motive template H symmetrical at $t = 1$? This means that the α_j and β_k defining the template are obtained from one another by adding $1/2$ (modulo 1), see `hgmtwist`.

```
? H = hgminit([2,2]);
? hgmalpha(H)
%2 = [[1/2, 1/2], [0, 0]]
? hgmissymmetrical(H)
%3 = 1 \\ this template is symmetrical

? H = hgminit([5]);
? hgmalpha(H)
%5 = [[1/5, 2/5, 3/5, 4/5], [0, 0, 0, 0]]
? hgmissymmetrical(H)
%6 = 1 \\ this one is not
```

The library syntax is `long hgmissymmetrical(GEN H)`.

3.16.13 hgmparams(H). H being a hypergeometric motive template, returns $[d, w, [P, T], M]$, where d is the degree, w the weight, P the Hodge polynomial, and T the Tate twist number (so that the Hodge function itself is P/x^T); finally the normalizing factor M is the so-called M -value, $M = \prod_n n^{n\gamma_n}$.

The library syntax is `GEN hgmparams(GEN H)`.

3.16.14 hgmtwist(H). Twist by $1/2$ of alpha and beta of the hypergeometric motive template H .

```
? H = hgminit([5]);
? hgmalpha(H)
%2 = [[1/5, 2/5, 3/5, 4/5], [0, 0, 0, 0]]
? H2 = hgmtwist(H);
? hgmalpha(H2)
%4 = [[1/10, 3/10, 7/10, 9/10], [1/2, 1/2, 1/2, 1/2]]
```

The template is symmetrical (`hgmissymmetrical`) if it is equal to its twist.

The library syntax is `GEN hgmtwist(GEN H)`.

3.16.15 lfunhgm($H, t, \{hint\}$). (H, t) being a hypergeometric motive, returns the corresponding `lfuncreate` data for use with the L -function package. This function needs to guess local conductors and euler factors at wild primes and will be very costly if there are many such primes: the complexity is roughly proportional to the conductor. The optional parameter `hint` allows to speed up the function by making various assumptions:

- `hint = lim` a `t_INT`: assume that Euler factors at wild primes have degree less than `lim`, which may speed it up a little.
- `hint = [N]`: guess that the conductor is N .
- `hint = [N, lim]`: initial guess N for the conductor and limit degrees to `lim`.

If your guess for `lim` is wrong, the function will enter an infinite loop. If your guess for an initial N is wrong, the function silently restarts (it will not enter an infinite loop unless a simultaneous failed guess for `lim` is made).

```

? H = hgmininit([5]);
? L = lfunhgm(H, 1/64);
time = 23,113 ms.
? L=lfunhgm(H,1/64,0); \\ assume Euler factors at wild primes are trivial
time = 19,721 ms. \\ a little faster
? L=lfunhgm(H,1/64,[525000]); \\ initial guess N = 525000
time = 15,486 ms. \\ a little faster
? L=lfunhgm(H,1/64,[525000, 0]);
time = 15,293 ms. \\ marginally faster with both assumptions

```

The library syntax is GEN lfunhgm(GEN H, GEN t, GEN hint = NULL, long bitprec)

3.17 L -functions.

This section describes routines related to L -functions. We first introduce the basic concept and notations, then explain how to represent them in GP. Let $\Gamma_{\mathbf{R}}(s) = \pi^{-s/2}\Gamma(s/2)$, where Γ is Euler's gamma function. Given $d \geq 1$ and a d -tuple $A = [\alpha_1, \dots, \alpha_d]$ of complex numbers, we let $\gamma_A(s) = \prod_{\alpha \in A} \Gamma_{\mathbf{R}}(s + \alpha)$.

Given a sequence $a = (a_n)_{n \geq 1}$ of complex numbers (such that $a_1 = 1$), a positive *conductor* $N \in \mathbf{Z}$, and a *gamma factor* γ_A as above, we consider the Dirichlet series

$$L(a, s) = \sum_{n \geq 1} a_n n^{-s}$$

and the attached completed function

$$\Lambda(a, s) = N^{s/2} \gamma_A(s) \cdot L(a, s).$$

Such a datum defines an L -function if it satisfies the three following assumptions:

- [Convergence] The $a_n = O_{\epsilon}(n^{k_1+\epsilon})$ have polynomial growth, equivalently $L(s)$ converges absolutely in some right half-plane $\Re(s) > k_1 + 1$.
- [Analytic continuation] $L(s)$ has a meromorphic continuation to the whole complex plane with finitely many poles.
- [Functional equation] There exist an integer k , a complex number ϵ (usually of modulus 1), and an attached sequence a^* defining both an L -function $L(a^*, s)$ satisfying the above two assumptions and a completed function $\Lambda(a^*, s) = N^{s/2} \gamma_A(s) \cdot L(a^*, s)$, such that

$$\Lambda(a, k - s) = \epsilon \Lambda(a^*, s)$$

for all regular points.

More often than not in number theory we have $a^* = \bar{a}$ (which forces $|\epsilon| = 1$), but this needs not be the case. If a is a real sequence and $a = a^*$, we say that L is *self-dual*. We do not assume that the a_n are multiplicative, nor equivalently that $L(s)$ has an Euler product.

Remark. Of course, a determines the L -function, but the (redundant) datum $a, a^*, A, N, k, \epsilon$ describes the situation in a form more suitable for fast computations; knowing the polar part r of $\Lambda(s)$ (a rational function such that $\Lambda - r$ is holomorphic) is also useful. A subset of these, including only finitely many a_n -values will still completely determine L (in suitable families), and we provide routines to try and compute missing invariants from whatever information is available.

Important Caveat. The implementation assumes that the implied constants in the O_ϵ are small. In our generic framework, it is impossible to return proven results without more detailed information about the L function. The intended use of the L -function package is not to prove theorems, but to experiment and formulate conjectures, so all numerical results should be taken with a grain of salt. One can always increase `realbitprecision` and recompute: the difference estimates the actual absolute error in the original output.

Note. The requested precision has a major impact on runtimes. Because of this, most L -function routines, in particular `lfun` itself, specify the requested precision in *bits*, not in decimal digits. This is transparent for the user once `realprecision` or `realbitprecision` are set. We advise to manipulate precision via `realbitprecision` as it allows finer granularity: `realprecision` increases by increments of 64 bits, i.e. 19 decimal digits at a time.

3.17.1 Theta functions.

Given an L -function as above, we define an attached theta function via Mellin inversion: for any positive real $t > 0$, we let

$$\theta(a, t) := \frac{1}{2\pi i} \int_{\Re(s)=c} t^{-s} \Lambda(s) ds$$

where c is any positive real number $c > k_1 + 1$ such that $c + \Re(a) > 0$ for all $a \in A$. In fact, we have

$$\theta(a, t) = \sum_{n \geq 1} a_n K(nt/N^{1/2}) \quad \text{where} \quad K(t) := \frac{1}{2\pi i} \int_{\Re(s)=c} t^{-s} \gamma_A(s) ds.$$

Note that this function is analytic and actually makes sense for complex t , such that $\Re(t^{2/d}) > 0$, i.e. in a cone containing the positive real half-line. The functional equation for Λ translates into

$$\theta(a, 1/t) - \epsilon t^k \theta(a^*, t) = P_\Lambda(t),$$

where P_Λ is an explicit polynomial in t and $\log t$ given by the Taylor expansion of the polar part of Λ : there are no log's if all poles are simple, and $P = 0$ if Λ is entire. The values $\theta(t)$ are generally easier to compute than the $L(s)$, and this functional equation provides a fast way to guess possible values for missing invariants in the L -function definition.

3.17.2 Data structures describing L and theta functions.

We have 3 levels of description:

- an **Lmath** is an arbitrary description of the underlying mathematical situation (to which e.g., we associate the a_p as traces of Frobenius elements); this is done via constructors to be described in the subsections below.

- an **Ldata** is a computational description of situation, containing the complete datum $(a, a^*, A, k, N, \epsilon, r)$. Where a and a^* describe the coefficients (given n, b we must be able to compute $[a_1, \dots, a_n]$ with bit accuracy b), A describes the Euler factor, the (classical) weight is k , N is the conductor, and r describes the polar part of $L(s)$. This is obtained via the function `lfuncreate`. N.B. For motivic L -functions, the motivic weight w is $w = k - 1$; but we also support nonmotivic L -functions.

Technical note. When some components of an `Ldata` cannot be given exactly, usually r or ϵ , the `Ldata` may be given as a *closure*. When evaluated at a given precision, the closure must return all components as exact data or floating point numbers at the requested precision, see `??lfuncreate`. The reason for this technicality is that the accuracy to which we must compute is not bounded a priori and unknown at this stage: it depends on the domain where we evaluate the L -function.

- an `Linit` contains an `Ldata` and everything needed for fast *numerical* computations. It specifies the functions to be considered (either $L^{(j)}(s)$ or $\theta^{(j)}(t)$ for derivatives of order $j \leq m$, where m is now fixed) and specifies a *domain* which limits the range of arguments (t or s , respectively to certain cones and rectangular regions) and the output accuracy. This is obtained via the functions `lfuninit` or `lfunthetainit`.

All the functions which are specific to L or theta functions share the prefix `lfun`. They take as first argument either an `Lmath`, an `Ldata`, or an `Linit`. If a single value is to be computed, this makes no difference, but when many values are needed (e.g. for plots or when searching for zeros), one should first construct an `Linit` attached to the search range and use it in all subsequent calls. If you attempt to use an `Linit` outside the range for which it was initialized, a warning is issued, because the initialization is performed again, a major inefficiency:

```
? Z = lfuncreate(1); \\ Riemann zeta
? L = lfuninit(Z, [1/2, 0, 100]); \\ zeta(1/2+it), |t| < 100
? lfun(L, 1/2)    \\ OK, within domain
%3 = -1.4603545088095868128894991525152980125
? lfun(L, 0)      \\ not on critical line !
*** lfun: Warning: lfuninit: insufficient initialization.
%4 = -0.5000000000000000000000000000000000000000000000000000000
? lfun(L, 1/2, 1) \\ attempt first derivative !
*** lfun: Warning: lfuninit: insufficient initialization.
%5 = -3.9226461392091517274715314467145995137
```

For many L -functions, passing from `Lmath` to an `Ldata` is inexpensive: in that case one may use `lfuninit` directly from the `Lmath` even when evaluations in different domains are needed. The above example could equally have skipped the `lfuncreate`:

```
? L = lfuninit(1, [1/2, 0, 100]); \\ zeta(1/2+it), |t| < 100
```

In fact, when computing a single value, you can even skip `lfuninit`:

```
? L = lfun(1, 1/2, 1); \\ zeta'(1/2)
? L = lfun(1, 1+x+O(x^5)); \\ first 5 terms of Taylor expansion at 1
```

Both give the desired results with no warning.

Complexity. The implementation requires $O(N(|t| + 1))^{1/2}$ coefficients a_n to evaluate L of conductor N at $s = \sigma + it$.

We now describe the available high-level constructors, for built-in L functions.

3.17.3 Dirichlet L -functions.

Given a Dirichlet character $\chi : (\mathbf{Z}/N\mathbf{Z})^* \rightarrow \mathbf{C}$, we let

$$L(\chi, s) = \sum_{n \geq 1} \chi(n) n^{-s}.$$

Only primitive characters are supported. Given a nonzero integer D , the `t_INT` D encodes the function $L((D_0/.), s)$, for the quadratic Kronecker symbol attached to the fundamental discriminant $D_0 = \text{coredisc}(D)$. This includes Riemann ζ function via the special case $D = 1$.

More general characters can be represented in a variety of ways:

- via Conrey notation (see `znconreychar`): $\chi_N(m, \cdot)$ is given as the `t_INTMOD` `Mod(m, N)`.
- via a *znstar* structure describing the abelian group $(\mathbf{Z}/N\mathbf{Z})^*$, where the character is given in terms of the *znstar* generators:

```
? G = znstar(100, 1); \ (Z/100Z)^*
? G.cyc \ \ ~ Z/20 . g1 + Z/2 . g2 for some generators g1 and g2
%2 = [20, 2]
? G.gen
%3 = [77, 51]
? chi = [a, b] \ \ maps g1 to e(a/20) and g2 to e(b/2); e(x) = exp(2ipi x)
```

More generally, let $(\mathbf{Z}/N\mathbf{Z})^* = \oplus (\mathbf{Z}/d_j\mathbf{Z})g_j$ be given via a *znstar* structure G (`G.cyc` gives the d_j and `G.gen` the g_j). A character χ on G is given by a row vector $v = [a_1, \dots, a_n]$ such that $\chi(\prod_j g_j^{n_j}) = \exp(2\pi i \sum_j a_j n_j / d_j)$. The pair $[G, v]$ encodes the *primitive* character attached to χ .

- in fact, this construction $[G, m]$ describing a character is more general: m is also allowed to be a Conrey label as seen above, or a Conrey logarithm (see `znconreylog`), and the latter format is actually the fastest one. Instead of a single character as above, one may use the construction `lfuncreate([G, vchi])` where `vchi` is a nonempty vector of characters of the same conductor (more precisely, whose attached primitive characters have the same conductor) and same parity. The function is then vector-valued, where the values are ordered as the characters in `vchi`. Conrey labels cannot be used in this last format because of the need to distinguish a single character given by a row vector of integers and a vector of characters given by their labels: use `znconreylog(G, i)` first to convert a label to Conrey logarithm.

- it is also possible to view Dirichlet characters as Hecke characters over $K = \mathbf{Q}$ (see below), for a modulus $[N, [1]]$ but this is both more complicated and less efficient.

In all cases, a nonprimitive character is replaced by the attached primitive character.

3.17.4 Hecke L -functions of finite order characters.

The Dedekind zeta function of a number field $K = \mathbf{Q}[X]/(T)$ is encoded either by the defining polynomial T , or any absolute number fields structure (a *nf* is enough).

An alternative description for the Dedekind zeta function of an Abelian extension of a number field is to use class-field theory parameters `[bnr, subg]`, see `bnrinit`.

```
? bnf = bnfinit(a^2 - a - 9);
? bnr = bnrinit(bnf, [2, [0,0]]); subg = Mat(3);
? L = lfuncreate([bnr, subg]);
```

Let K be a number field given as a `bnfinit`. Given a finite order Hecke character $\chi : Cl_f(K) \rightarrow \mathbf{C}$, we let

$$L(\chi, s) = \sum_{A \subset O_K} \chi(A) (N_{K/\mathbf{Q}} A)^{-s}.$$

Let $Cl_f(K) = \oplus (\mathbf{Z}/d_j \mathbf{Z}) g_j$ given by a *bnr* structure with generators: the d_j are given by `K.cyc` and the g_j by `K.gen`. A character χ on the ray class group is given by a row vector $v = [a_1, \dots, a_n]$ such that $\chi(\prod_j g_j^{n_j}) = \exp(2\pi i \sum_j a_j n_j / d_j)$. The pair $[bnr, v]$ encodes the *primitive* character attached to χ .

```
? K = bnfinit(x^2-60);
? Cf = bnrinit(K, [7, [1,1]], 1); \ f = 7 oo_1 oo_2
? Cf.cyc
%3 = [6, 2, 2]
? Cf.gen
%4 = [[2, 1; 0, 1], [22, 9; 0, 1], [-6, 7]~]
? lfuncreate([Cf, [1,0,0]]); \ \chi(g_1) = \zeta_6, \chi(g_2) = \chi(g_3) = 1
```

Dirichlet characters on $(\mathbf{Z}/N\mathbf{Z})^*$ are a special case, where $K = \mathbf{Q}$:

```
? Q = bnfinit(x);
? Cf = bnrinit(Q, [100, [1]]); \ for odd characters on (Z/100Z)*
```

For even characters, replace by `bnrinit(K, N)`. Note that the simpler direct construction in the previous section will be more efficient. Instead of a single character as above, one may use the construction `lfuncreate([Cf, vchi])` where `vchi` is a nonempty vector of characters *of the same conductor* (more precisely, whose attached primitive characters have the same conductor). The function is then vector-valued, where the values are ordered as the characters in `vchi`.

3.17.5 General Hecke L -functions.

Given a Hecke *Grossencharacter* $\chi : \mathbf{A}^\times \rightarrow \mathbf{C}^\times$ of conductor $\mathfrak{f}$, we let

$$L(\chi, s) = \sum_{A \subset \mathbf{Z}_K, A + \mathfrak{f} = \mathbf{Z}_K} \chi(A) (N_{K/\mathbf{Q}} A)^{-s}.$$

Let $C_K(\mathfrak{m}) = \mathbf{A}^\times / (K^\times \cdot U(\mathfrak{m}))$ be an idèle class group of modulus $\mathfrak{m}$ given by a *gchar* structure *gc* (see `gcharinit` and Section 3.13.8). A Grossencharacter χ on $C_K(\mathfrak{m})$ is given by a row vector of size `#gc.cyc`.

```
? gc = gcharinit(x^3+4*x-1, [5, [1]]); \ mod = 5.oo
? gc.cyc
%3 = [4, 2, 0, 0]
? gcharlog(gc, idealprimedec(gc.bnf, 5)[1]) \ logarithm map C_K(m) -> R^n
? chi = [1, 0, 0, 1, 0]~;
? gcharduallog(gc, chi) \ row vector of coefficients in R^n
? L = lfuncreate([gc, chi]); \ non algebraic L-function
? lfunzeros(L, 1)
? lfuneuler(L, 2) \ Euler factor at 2
```

Finite order Hecke characters are a special case.

3.17.6 Artin L functions.

Given a Galois number field $N/\mathbf{Q}$ with group $G = \text{galoisinit}(N)$, a representation ρ of G over the cyclotomic field $\mathbf{Q}(\zeta_n)$ is specified by the matrices giving the images of $G.\text{gen}$ by ρ . The corresponding Artin L function is created using `lfunartin`.

```
P = quadhilbert(-47); \\ degree 5, Galois group D_5
N = nfinit(nfsplitting(P)); \\ Galois closure
G = galoisinit(N);
[s,t] = G.gen; \\ order 5 and 2
L = lfunartin(N,G, [[a,0;0,a^-1],[0,1;1,0]], 5); \\ irr. degree 2
```

In the above, the polynomial variable (here a) represents $\zeta_5 := \exp(2i\pi/5)$ and the two matrices give the images of s and t . Here, priority of a must be lower than the priority of x .

3.17.7 L -functions of algebraic varieties.

L -function of elliptic curves over number fields are supported.

```
? E = ellinit([1,1]);
? L = lfuncreate(E); \\ L-function of E/Q
? E2 = ellinit([1,a], nfinit(a^2-2));
? L2 = lfuncreate(E2); \\ L-function of E/Q(sqrt(2))
```

L -function of hyperelliptic genus-2 curve can be created with `lfungenus2`. To create the L function of the curve $y^2 + (x^3 + x^2 + 1)y = x^2 + x$:

```
? L = lfungenus2([x^2+x, x^3+x^2+1]);
```

Currently, the model needs to be minimal at 2, and if the conductor is even, its valuation at 2 might be incorrect (a warning is issued).

3.17.8 Eta quotients / Modular forms.

An eta quotient is created by applying `lfunetaquo` to a matrix with 2 columns $[m, r_m]$ representing

$$f(\tau) := \prod_m \eta(m\tau)^{r_m}.$$

It is currently assumed that f is a self-dual cuspidal form on $\Gamma_0(N)$ for some N . For instance, the L -function $\sum \tau(n)n^{-s}$ attached to Ramanujan's Δ function is encoded as follows

```
? L = lfunetaquo(Mat([1,24]));
? lfunan(L, 100) \\ first 100 values of tau(n)
```

More general modular forms defined by modular symbols will be added later.

3.17.9 Low-level `Ldata` format.

When no direct constructor is available, you can still input an L function directly by supplying $[a, a^*, A, k, N, \epsilon, r]$ to `lfuncreate` (see `??lfuncreate` for details).

It is *strongly* suggested to first check consistency of the created L -function:

```
? L = lfuncreate([a, as, A, k, N, eps, r]);
? lfuncheckfeq(L) \\ check functional equation
```

3.17.10 lfun($L, s, \{D = 0\}$). Compute the L -function value $L(s)$, or if D is set, the derivative of order D at s . The parameter L is either an `Lmath`, an `Ldata` (created by `lfuncreate`, or an `Linit` (created by `lfuninit`), preferably the latter if many values are to be computed.

The argument s is also allowed to be a power series; for instance, if $s = \alpha + x + O(x^n)$, the function returns the Taylor expansion of order n around α . The result is given with absolute error less than 2^{-B} , where $B = \text{realbitprecision}$.

Caveat. The requested precision has a major impact on runtimes. It is advised to manipulate precision via `realbitprecision` as explained above instead of `realprecision` as the latter allows less granularity: `realprecision` increases by increments of 64 bits, i.e. 19 decimal digits at a time.

```
? lfun(x^2+1, 2)  \\ Lmath: Dedekind zeta for Q(i) at 2
%1 = 1.5067030099229850308865650481820713960

? L = lfuncreate(ellinit("5077a1")); \\ Ldata: Hasse-Weil zeta function
? lfun(L, 1+x+O(x^4))  \\ zero of order 3 at the central point
%3 = 0.E-58 - 5.[...] E-40*x + 9.[...] E-40*x^2 + 1.7318[...] *x^3 + O(x^4)

\\ Linit: zeta(1/2+it), |t| < 100, and derivative
? L = lfunit(1, [100], 1);
? T = lfunitzeros(L, [1,25]);
%5 = [14.134725[...], 21.022039[...]]
? z = 1/2 + I*T[1];
? abs( lfun(L, z) )
%7 = 8.7066865533412207420780392991125136196 E-39
? abs( lfun(L, z, 1) )
%8 = 0.79316043335650611601389756527435211412  \\ simple zero
```

The library syntax is `GEN lfun0(GEN L, GEN s, long D, long bitprec)`.

3.17.11 lfunan(L, n). Compute the first n terms of the Dirichlet series attached to the L -function given by L (`Lmath`, `Ldata` or `Linit`).

```
? lfunan(1, 10)  \\ Riemann zeta
%1 = [1, 1, 1, 1, 1, 1, 1, 1, 1, 1]
? lfunan(5, 10)  \\ Dirichlet L-function for kronecker(5,.)
%2 = [1, -1, -1, 1, 0, 1, -1, -1, 1, 0]
```

The library syntax is `GEN lfunan(GEN L, long n, long prec)`.

3.17.12 lfunartin(nf, gal, rho, n). Returns the `Ldata` structure attached to the Artin L -function provided by the representation ρ of the Galois group of the extension $K/\mathbf{Q}$, defined over the cyclotomic field $\mathbf{Q}(\zeta_n)$, where nf is the `nfin` structure attached to K , gal is the `galoisinit` structure attached to $K/\mathbf{Q}$, and rho is given either

- by the values of its character on the conjugacy classes (see `galoisconjclasses` and `galois-chartable`)
- or by the matrices that are the images of the generators `gal.gen`.

Cyclotomic numbers in `rho` are represented by polynomials, whose variable is understood as the complex number $\exp(2i\pi/n)$.

In the following example we build the Artin L -functions attached to the two irreducible degree 2 representations of the dihedral group D_{10} defined over $\mathbf{Q}(\zeta_5)$, for the extension $H/\mathbf{Q}$ where H is the Hilbert class field of $\mathbf{Q}(\sqrt{-47})$. We show numerically some identities involving Dedekind ζ functions and Hecke L series.

```
? P = quadhilbert(-47)
%1 = x^5 + 2*x^4 + 2*x^3 + x^2 - 1
? N = nfinit(nfsplitting(P));
? G = galoisinit(N); \\ D_10
? [T,n] = galoischartable(G);
? T \\ columns give the irreducible characters
%5 =
[1  1          2          2]
[1 -1          0          0]
[1  1 -y^3 - y^2 - 1      y^3 + y^2]
[1  1      y^3 + y^2 -y^3 - y^2 - 1]
? n
%6 = 5
? L2 = lfunartin(N,G, T[,2], n);
? L3 = lfunartin(N,G, T[,3], n);
? L4 = lfunartin(N,G, T[,4], n);
? s = 1 + x + O(x^4);
? lfun(-47,s) - lfun(L2,s)
%11 ~ 0
? lfun(1,s)*lfun(-47,s)*lfun(L3,s)^2*lfun(L4,s)^2 - lfun(N,s)
%12 ~ 0
? lfun(1,s)*lfun(L3,s)*lfun(L4,s) - lfun(P,s)
%13 ~ 0
? bnr = bnrinit(bnfinit(x^2+47),1,1);
? bnr.cyc
%15 = [5] \\ Z/5Z: 4 nontrivial ray class characters
? lfun([bnr,[1]], s) - lfun(L3, s)
%16 ~ 0
? lfun([bnr,[2]], s) - lfun(L4, s)
%17 ~ 0
? lfun([bnr,[3]], s) - lfun(L3, s)
%18 ~ 0
? lfun([bnr,[4]], s) - lfun(L4, s)
%19 ~ 0
```

The first identity identifies the nontrivial abelian character with $(-47, \cdot)$; the second is the factorization of the regular representation of D_{10} ; the third is the factorization of the natural representation of $D_{10} \subset S_5$; and the final four are the expressions of the degree 2 representations as induced from degree 1 representations.

The library syntax is GEN lfunartin(GEN nf, GEN gal, GEN rho, long n, long bitprec)

3.17.13 lfuncheckfeq($L, \{t\}$). Given the data attached to an L -function (`Lmath`, `Ldata` or `Linit`), check whether the functional equation is satisfied. This is most useful for an `Ldata` constructed “by hand”, via `lfuncreate`, to detect mistakes.

If the function has poles, the polar part must be specified. The routine returns a bit accuracy b such that $|w - \hat{w}| < 2^b$, where w is the root number contained in `data`, and

$$\hat{w} = \theta(1/t)t^{-k}/\bar{\theta}(t)$$

is a computed value derived from the assumed functional equation. Of course, the expected result is a large negative value of the order of `realbitprecision`. But if $\bar{\theta}$ is very small at t , you should first increase `realbitprecision` by $-\log_2 |\bar{\theta}(t)|$, which is positive if θ is small, to get a meaningful result. Note that t should be close to the unit disc for efficiency and such that $\bar{\theta}(t) \neq 0$. If the parameter t is omitted, we check the functional equation at the “random” complex number $t = 335/339 + I/7$.

```
? \pb 128          \ 128 bits of accuracy
? default(realbitprecision)
%1 = 128
? L = lfuncreate(1); \ Riemann zeta
? lfuncheckfeq(L)
%3 = -124
```

i.e. the given data is consistent to within 4 bits for the particular check consisting of estimating the root number from all other given quantities. Checking away from the unit disc will either fail with a precision error, or give disappointing results (if $\theta(1/t)$ is large it will be computed with a large absolute error)

```
? lfuncheckfeq(L, 2+I)
%4 = -115
? lfuncheckfeq(L,10)
*** at top-level: lfuncheckfeq(L,10)
*** ^-----
*** lfuncheckfeq: precision too low in lfuncheckfeq.
```

The case of Dedekind zeta functions. Dedekind zeta function for a number field $K = \mathbf{Q}[X]/(T)$ is in general computed (assuming Artin conjecture) as $(\zeta_K/\zeta_k) \times \zeta_k$, where k is a maximal subfield, applied recursively if possible. When $K/\mathbf{Q}$ is Galois, the zeta function is directly decomposed as a product of Artin L -functions.

These decompositions are computed when `lfuninit` is called. The behavior of `lfuncheckfeq` is then different depending of its argument

- the artificial query `lfuncheckfeq( $T$ )` serves little purpose since we already know that the technical parameters are theoretically correct; we just obtain an estimate on the accuracy they allow. This is computed directly, without using the above decomposition. And is likely to be very costly if the degree of T is large, possibly overflowing the possibilities of the implementation.

- a query `L = lfuninit( $T, \dots$ ); lfuncheckfeq( $L$ )` on the other hand returns the maximum of the `lfuncheckfeq` values for all involved L -functions, giving a general consistency check and again an estimate for the accuracy of computed values.

At the default accuracy of 128 bits:

```

? T = polcyclo(43);
? lfuncheckfeq(T);
*** at top-level: lfuncheckfeq(T)
*** ^-----
*** lfuncheckfeq: overflow in lfunthetacost.
? lfuncheckfeq(lfuninit(T, [2]))
time = 107 ms.
%2 = -122

```

The library syntax is `long lfuncheckfeq(GEN L, GEN t = NULL, long bitprec)`.

3.17.14 lfunconductor($L, \{setN = 10000\}, \{flag = 0\}$). Computes the conductor of the given L -function (if the structure contains a conductor, it is ignored). Two methods are available, depending on what we know about the conductor, encoded in the `setN` parameter:

- `setN` is a scalar: we know nothing but expect that the conductor lies in the interval $[1, setN]$.

If *flag* is 0 (default), gives either the conductor found as an integer, or a vector (possibly empty) of conductors found. If *flag* is 1, same but gives the computed floating point approximations to the conductors found, without rounding to integers. If *flag* is 2, gives all the conductors found, even those far from integers.

Caveat. This is a heuristic program and the result is not proven in any way:

```

? L = lfuncreate(857); \\ Dirichlet L function for kronecker(857,.)
? \p19
  realprecision = 19 significant digits
? lfunconductor(L)
%2 = [17, 857]
? lfunconductor(L,,1) \\ don't round
%3 = [16.999999999999999, 857.0000000000000000]
? \p38
  realprecision = 38 significant digits
? lfunconductor(L)
%4 = 857

```

Increasing `setN` or increasing `realbitprecision` slows down the program but gives better accuracy for the result. This algorithm should only be used if the primes dividing the conductor are unknown, which is uncommon.

- `setN` is a vector of possible conductors; for instance of the form $D_1 * \text{divisors}(D_2)$, where D_1 is the known part of the conductor and D_2 is a multiple of the contribution of the bad primes.

In that case, *flag* is ignored and the routine uses `lfuncheckfeq`. It returns $[N, e]$ where N is the best conductor in the list and e is the value of `lfuncheckfeq` for that N . When no suitable conductor exist or there is a tie among best potential conductors, return the empty vector `[]`.

```

? E = ellinit([0,0,0,4,0]); /* Elliptic curve y^2 = x^3+4x */
? E.disc \\ |disc E| = 2^12
%2 = -4096
\\ create Ldata by hand. Guess that root number is 1 and conductor N
? L(N) = lfuncreate([n->ellan(E,n), 0, [0,1], 2, N, 1]);
\\ lfunconductor ignores conductor = 1 in Ldata !

```

```

? lfunconductor(L(1), divisors(E.disc))
%5 = [32, -127]
? fordiv(E.disc, d, print(d,": ",lfuncheckfeq(L(d)))) \\ direct check
1: 0
2: 0
4: -1
8: -2
16: -3
32: -127
64: -3
128: -2
256: -2
512: -1
1024: -1
2048: 0
4096: 0

```

The above code assumed that root number was 1; had we set it to -1 , none of the `lfuncheckfeq` values would have been acceptable:

```

? L2 = lfuncreate([n->ellan(E,n), 0, [0,1], 2, 0, -1]);
? lfunconductor(L2, divisors(E.disc))
%7 = []

```

The library syntax is `GEN lfunconductor(GEN L, GEN setN = NULL, long flag, long bitprec)`.

3.17.15 lfuncost($L, \{sdom\}, \{der = 0\}$). Estimate the cost of running `lfuninit(L,sdom,der)` at current bit precision, given by a vector $[t, b]$.

- If L contains the root number, indicate that t coefficients a_n will be computed, as well as t values of `gammamellininv`, all at bit accuracy b . A subsequent call to `lfun` at s evaluates a polynomial of degree t at $\exp(hs)$ for some real parameter h , at the same bit accuracy b .

- If the root number is *not* known, then more values of a_n may be needed in order to compute it, and the returned value of t takes this into account (it may not be the exact value in this case but is always an upper bound). Fewer than t `gammamellininv` will be needed, and a call to `lfun` evaluates a polynomial of degree less than t , still at bit accuracy b .

If L is already an `Linit`, then $sdom$ and der are ignored and are best left omitted; the bit accuracy is also inferred from L : in short we get an estimate of the cost of using that particular `Linit`. Note that in this case, the root number is always already known and you get the right value of t (corresponding to the number of past calls to `gammamellininv` and the actual degree of the evaluated polynomial).

```

? \pb 128
? lfuncost(1, [100]) \\ for zeta(1/2+I*t), |t| < 100
%1 = [7, 242] \\ 7 coefficients, 242 bits
? lfuncost(1, [1/2, 100]) \\ for zeta(s) in the critical strip, |Im s| < 100
%2 = [7, 246] \\ now 246 bits
? lfuncost(1, [100], 10) \\ for zeta(1/2+I*t), |t| < 100
%3 = [8, 263] \\ 10th derivative increases the cost by a small amount

```

```

? lfuncost(1, [10^5])
%3 = [158, 113438] \\ larger imaginary part: huge accuracy increase
? L = lfuncreate(polcyclo(5)); \\ Dedekind zeta for Q(zeta_5)
? lfuncost(L, [100]) \\ at s = 1/2+I*t, |t| < 100
%5 = [11457, 582]
? lfuncost(L, [200]) \\ twice higher
%6 = [36294, 1035]
? lfuncost(L, [10^4]) \\ much higher: very costly !
%7 = [70256473, 45452]
? \pb 256
? lfuncost(L, [100]); \\ doubling bit accuracy is cheaper
%8 = [17080, 710]
? \p38
? K = bnfinit(y^2 - 4493); [P] = idealprimedec(K,1123); f = [P,[1,1]];
? R = bnrinit(K, f); R.cyc
%10 = [1122]
? L = lfuncreate([R, [7]]); \\ Hecke L-function
? L[6]
%12 = 0 \\ unknown root number
? \pb 3000
? lfuncost(L, [0], 1)
%13 = [1171561, 3339]
? L = lfuninit(L, [0], 1);
time = 1min, 56,426 ms.
? lfuncost(L)
%14 = [826966, 3339]

```

In the final example, the root number was unknown and extra coefficients a_n were needed to compute it (1171561). Once the initialization is performed we obtain the lower value $t = 826966$, which corresponds to the number of `gammamellinv` computed and the actual degree of the polynomial to be evaluated to compute a value within the prescribed domain.

Finally, some L functions can be factorized algebraically by the `lfuninit` call, e.g. the Dedekind zeta function of abelian fields, leading to much faster evaluations than the above upper bounds. In that case, the function returns a vector of costs as above for each individual function in the product actually evaluated:

```

? L = lfuncreate(polcyclo(5)); \\ Dedekind zeta for Q(zeta_5)
? lfuncost(L, [100]) \\ a priori cost
%2 = [11457, 582]
? L = lfuninit(L, [100]); \\ actually perform all initializations
? lfuncost(L)
%4 = [[16, 242], [16, 242], [7, 242]]

```

The Dedekind function of this abelian quartic field is the product of four Dirichlet L -functions attached to the trivial character, a nontrivial real character and two complex conjugate characters. The nontrivial characters happen to have the same conductor (hence same evaluation costs), and correspond to two evaluations only since the two conjugate characters are evaluated simultaneously. For a total of three L -functions evaluations, which explains the three components above. Note that the actual cost is much lower than the a priori cost in this case.

The library syntax is `GEN lfuncost0(GEN L, GEN sdom = NULL, long der, long bitprec)`. Also available is `GEN lfuncost(GEN L, GEN dom, long der, long bitprec)` when L is *not* an `Linit`; the return value is a `t_VECSMALL` in this case.

3.17.16 lfuncreate(*obj*). This low-level routine creates `Ldata` structures, needed by *lfun* functions, describing an L -function and its functional equation. We advise using a high-level constructor when one is available, see `??lfun`, and this function accepts them:

```
? L = lfuncreate(1); \\ Riemann zeta
? L = lfuncreate(5); \\ Dirichlet L-function for quadratic character (5/.)
? L = lfuncreate(x^2+1); \\ Dedekind zeta for Q(i)
? L = lfuncreate(ellinit([0,1])); \\ L-function of E/Q: y^2=x^3+1
```

One can then use, e.g., `lfun(L,s)` to directly evaluate the respective L -functions at s , or `lfuninit(L, [c,w,h])` to initialize computations in the rectangular box $\Re(s - c) \leq w$, $\Im(s) \leq h$.

We now describe the low-level interface, used to input nonbuiltin L -functions. The input is now a 6 or 7 component vector $V = [a, astar, Vga, k, N, eps, poles]$, whose components are as follows:

- $V[1]=a$ encodes the Dirichlet series coefficients (a_n) . The preferred format is a closure of arity 1: `n->vector(n,i,a(i))` giving the vector of the first n coefficients. The closure is allowed to return a vector of more than n coefficients (only the first n will be considered) or even less than n , in which case loss of accuracy will occur and a warning that `#an` is less than expected is issued. This allows to precompute and store a fixed large number of Dirichlet coefficients in a vector v and use the closure `n->v`, which does not depend on n . As a shorthand for this latter case, you can input the vector v itself instead of the closure.

```
? z = lfuncreate([n->vector(n,i,1), 1, [0], 1, 1, 1, 1]); \\ Riemann zeta
? lfun(z,2) - Pi^2/6
%2 = -5.877471754111437540 E-39
```

A second format is limited to L -functions affording an Euler product. It is a closure of arity 2 `(p,d)->F(p)` giving the local factor $L_p(X)$ at p as a rational function, to be evaluated at p^{-s} as in `direuler`; d is set to `logint(n,p) + 1`, where n is the total number of Dirichlet coefficients $(a_1, \dots, a_n)$ that will be computed. In other words, the smallest integer d such that $p^d > n$. This parameter d allows to compute only part of L_p when p is large and L_p expensive to compute: any polynomial (or `t_SER`) congruent to L_p modulo X^d is acceptable since only the coefficients of $X^0, \dots, X^{d-1}$ are needed to expand the Dirichlet series. The closure can of course ignore this parameter:

```
? z = lfuncreate([(p,d)->1/(1-x), 1, [0], 1, 1, 1, 1]); \\ Riemann zeta
? lfun(z,2) - Pi^2/6
%4 = -5.877471754111437540 E-39
```

One can describe separately the generic local factors coefficients and the bad local factors by setting `dir = [F, Lbad]`, where $L_{bad} = [[p_1, L_{p_1}], \dots, [p_k, L_{p_k}]]$, where F describes the generic local factors as above, except that when $p = p_i$ for some $i \leq k$, the coefficient a_p is directly set to L_{p_i} instead of calling F .

```
N = 15;
E = ellinit([1, 1, 1, -10, -10]); \\ = "15a1"
F(p,d) = 1 / (1 - ellap(E,p)*'x + p*'x^2);
Lbad = [[3, 1/(1+'x)], [5, 1/(1-'x)]];

```

```
L = lfuncreate([F,Lbad], 0, [0,1], 2, N, ellrootno(E));
```

Of course, in this case, `lfuncreate(E)` is preferable!

- **V[2]=astar** is the Dirichlet series coefficients of the dual function, encoded as **a** above. The sentinel values 0 and 1 may be used for the special cases where $a = a^*$ and $a = \overline{a^*}$, respectively.

- **V[3]=Vga** is the vector of α_j such that the gamma factor of the L -function is equal to

$$\gamma_A(s) = \prod_{1 \leq j \leq d} \Gamma_{\mathbf{R}}(s + \alpha_j),$$

where $\Gamma_{\mathbf{R}}(s) = \pi^{-s/2} \Gamma(s/2)$. This same syntax is used in the `gammamellininv` functions. In particular the length d of **Vga** is the degree of the L -function. In the present implementation, the α_j are assumed to be exact rational numbers. However when calling theta functions with *complex* (as opposed to real) arguments, determination problems occur which may give wrong results when the α_j are not integral.

- **V[4]=k** is a positive integer k . The functional equation relates values at s and $k - s$. For instance, for an Artin L -series such as a Dedekind zeta function we have $k = 1$, for an elliptic curve $k = 2$, and for a modular form, k is its weight. For motivic L -functions, the *motivic* weight w is $w = k - 1$.

By default we assume that $a_n = O_\epsilon(n^{k_1+\epsilon})$, where $k_1 = w$ and even $k_1 = w/2$ when the L function has no pole (Ramanujan-Petersson). If this is not the case, you can replace the k argument by a vector $[k, k_1]$, where k_1 is the upper bound you can assume.

- **V[5]=N** is the conductor, an integer $N \geq 1$, such that $\Lambda(s) = N^{s/2} \gamma_A(s) L(s)$ with $\gamma_A(s)$ as above.

- **V[6]=eps** is the root number ε , i.e., the complex number (usually of modulus 1) such that $\Lambda(a, k - s) = \varepsilon \Lambda(a^*, s)$.

- The last optional component **V[7]=poles** encodes the poles of the L or Λ -functions, and is omitted if they have no poles. A polar part is given by a list of 2-component vectors $[\beta, P_\beta(x)]$, where β is a pole and the power series $P_\beta(x)$ describes the attached polar part, such that $L(s) - P_\beta(s - \beta)$ is holomorphic in a neighbourhood of β . For instance $P_\beta = r/x + O(1)$ for a simple pole at β or $r_1/x^2 + r_2/x + O(1)$ for a double pole. The type of the list describing the polar part allows to distinguish between L and Λ : a **t_VEC** is attached to L , and a **t_COL** is attached to Λ . Unless $a = \overline{a^*}$ (coded by **astar** equal to 0 or 1), it is mandatory to specify the polar part of Λ rather than those of L since the poles of L^* cannot be inferred from the latter ! Whereas the functional equation allows to deduce the polar part of Λ^* from the polar part of Λ .

Finally, if $a = \overline{a^*}$, we allow a shortcut to describe the frequent situation where L has at most simple pole, at $s = k$, with residue r a complex scalar: you may then input **poles** = r . This value r can be set to 0 if unknown and it will be computed.

When one component is not exact. Alternatively, `obj` can be a closure of arity 0 returning the above vector to the current real precision. This is needed if some components are not available exactly but only through floating point approximations. The closure allows algorithms to recompute them to higher accuracy when needed. Compare

```
? Ld1() = [n->lfunan(Mod(2,7),n),1,[0],1,7,((-13-3*sqrt(-3))/14)^(1/6)];
? Ld2 = [n->lfunan(Mod(2,7),n),1,[0],1,7,((-13-3*sqrt(-3))/14)^(1/6)];
? L1 = lfuncreate(Ld1);
? L2 = lfuncreate(Ld2);
? lfun(L1,1/2+I*200) \\ OK
%5 = 0.55943925130316677665287870224047183265 -
      0.42492662223174071305478563967365980756*I
? lfun(L2,1/2+I*200) \\ all accuracy lost
%6 = 0.E-38 + 0.E-38*I
```

The accuracy lost in `Ld2` is due to the root number being given to an insufficient precision. To see what happens try

```
? Ld3() = printf("prec needed: %ld bits",getlocalbitprec());Ld1()
? L3 = lfuncreate(Ld3);
prec needed: 64 bits
? z3 = lfun(L3,1/2+I*200)
prec needed: 384 bits
%16 = 0.55943925130316677665287870224047183265 -
      0.42492662223174071305478563967365980756*I
```

The library syntax is `GEN lfuncreate(GEN obj)`.

3.17.17 `lfundiv(L1,L2)`. Creates the `Ldata` structure (without initialization) corresponding to the quotient of the Dirichlet series L_1 and L_2 given by `L1` and `L2`. Assume that $v_z(L_1) \geq v_z(L_2)$ at all complex numbers z : the construction may not create new poles, nor increase the order of existing ones.

The library syntax is `GEN lfundiv(GEN L1, GEN L2, long bitprec)`.

3.17.18 `lfundual(L)`. Creates the `Ldata` structure (without initialization) corresponding to the dual L-function $\hat{L}$ of L . If k and ε are respectively the weight and root number of L , then the following formula holds outside poles, up to numerical errors:

$$\Lambda(L, s) = \varepsilon \Lambda(\hat{L}, k - s).$$

```
? L = lfunqf(matdiagonal([1,2,3,4]));
? eps = lfunrootres(L)[3]; k = L[4];
? M = lfundual(L); lfuncheckfeq(M)
%3 = -127
? s= 1+Pi*I;
? a = lfunlambda(L,s);
? b = eps * lfunlambda(M,k-s);
? exponent(a - b)
%7 = -130
```

The library syntax is `GEN lfundual(GEN L, long bitprec)`.

3.17.19 lfunetaquo(M). Returns the `Ldata` structure attached to the L function attached to the modular form $z \mapsto \prod_{i=1}^n \eta(M_{i,1}z)^{M_{i,2}}$. It is currently assumed that f is a self-dual cuspidal form on $\Gamma_0(N)$ for some N . For instance, the L -function $\sum \tau(n)n^{-s}$ attached to Ramanujan's Δ function is encoded as follows

```
? L = lfunetaquo(Mat([1,24]));
? lfunan(L, 100) \\ first 100 values of tau(n)
```

For convenience, a `t_VEC` is also accepted instead of a factorization matrix with a single row:

```
? L = lfunetaquo([1,24]); \\ same as above
```

The library syntax is `GEN lfunetaquo(GEN M)`.

3.17.20 lfuneuler(L, p). Return the Euler factor at p of the L -function given by `L` (`Lmath`, `Ldata` or `Linit`), assuming the L -function admits an Euler product factorization and that it can be determined.

```
? E=ellinit([1,3]);
? lfuneuler(E,7)
%2 = 1/(7*x^2-2*x+1)
? L=lfunsympow(E,2);
? lfuneuler(L,11)
%4 = 1/(-1331*x^3+275*x^2-25*x+1)
```

The library syntax is `GEN lfuneuler(GEN L, GEN p, long prec)`.

3.17.21 lfungenus2(F). Returns the `Ldata` structure attached to the L function attached to the genus-2 curve defined by $y^2 = F(x)$ or $y^2 + Q(x)y = P(x)$ if $F = [P, Q]$. Currently, if the conductor is even, its valuation at 2 might be incorrect (a warning is issued).

The library syntax is `GEN lfungenus2(GEN F)`.

3.17.22 lfunhardy(L, t). Variant of the Hardy Z -function given by `L`, used for plotting or locating zeros of $L(k/2 + it)$ on the critical line. The precise definition is as follows: let $k/2$ be the center of the critical strip, d be the degree, $\mathbf{vga} = (\alpha_j)_{j \leq d}$ given the gamma factors, and ε be the root number; we set $s = k/2 + it = \rho e^{i\theta}$ and $2E = d(k/2 - 1) + \Re(\sum_{1 \leq j \leq d} \alpha_j)$. Assume first that Λ is self-dual, then the computed function at t is equal to

$$Z(t) = \varepsilon^{-1/2} \Lambda(s) \cdot \rho^{-E} e^{dt\theta/2},$$

which is a real function of t vanishing exactly when $L(k/2 + it)$ does on the critical line. The normalizing factor $|s|^{-E} e^{dt\theta/2}$ compensates the exponential decrease of $\gamma_A(s)$ as $t \rightarrow \infty$ so that $Z(t) \approx 1$. For non-self-dual Λ , the definition is the same except we drop the $\varepsilon^{-1/2}$ term (which is not well defined since it depends on the chosen dual sequence $a^*(n)$): $Z(t)$ is still of the order of 1 and still vanishes where $L(k/2 + it)$ does, but it needs no longer be real-valued.

```
? T = 100; \\ maximal height
? L = lfunit(1, [T]); \\ initialize for zeta(1/2+it), |t|<T
? \p19 \\ no need for large accuracy
? plot(t = 0, T, lfunhardy(L,t))
```

Using `lfuninit` is critical for this particular applications since thousands of values are computed. Make sure to initialize up to the maximal t needed: otherwise expect to see many warnings for insufficient initialization and suffer major slowdowns.

The library syntax is `GEN lfunhardy(GEN L, GEN t, long bitprec)`.

3.17.23 lfuninit($L, sdom, \{der = 0\}$). Initialization function for all functions linked to the computation of the L -function $L(s)$ encoded by L , where s belongs to the rectangular domain $sdom = [center, w, h]$ centered on the real axis, $|\Re(s) - center| \leq w$, $|\Im(s)| \leq h$, where all three components of $sdom$ are real and w, h are nonnegative. der is the maximum order of derivation that will be used. The subdomain $[k/2, 0, h]$ on the critical line (up to height h) can be encoded as $[h]$ for brevity. The subdomain $[k/2, w, h]$ centered on the critical line can be encoded as $[w, h]$ for brevity.

The argument L is an `Lmath`, an `Ldata` or an `Linit`. See `??Ldata` and `??lfuncreate` for how to create it.

The height h of the domain is a *crucial* parameter: if you only need $L(s)$ for real s , set h to 0. The running time is roughly proportional to

$$(B/d + \pi h/4)^{d/2+3} N^{1/2},$$

where B is the default bit accuracy, d is the degree of the L -function, and N is the conductor (the exponent $d/2 + 3$ is reduced to $d/2 + 2$ when $d = 1$ and $d = 2$). There is also a dependency on w , which is less crucial, but make sure to use the smallest rectangular domain that you need.

```
? L0 = lfuncreate(1); \\ Riemann zeta
? L = lfuninit(L0, [1/2, 0, 100]); \\ for zeta(1/2+it), |t| < 100
? lfun(L, 1/2 + I)
? L = lfuninit(L0, [100]); \\ same as above !
```

Riemann-Siegel formula. If L is a function of degree $d = 1$, then a completely different algorithm is implemented which can compute with complexity $N\sqrt{h}$ (for fixed accuracy B). So it handles larger imaginary parts than the default implementation. But this variant is less efficient when the imaginary part of s is tiny and the dependency in B is still in $O(B^{2+1/2})$.

For such functions, you can use $sdom = []$ to indicate that you are only interested in relatively high imaginary parts and do not want to perform any initialization:

```
? L = lfuninit(1, []); \\ Riemann zeta
? #lfunzeros(L, [10^12, 10^12+1])
time = 1min, 31,496 ms.
%2 = 4
```

If you ask instead for `lfuninit(1, [10^12+1])`, the initialization is restricted by some cutoff value (depending on the conductor, but less than 10^4 in any case): up to that point, the standard algorithm is used (and uses the initialization); and above the cutoff, we switch to Riemann-Siegel. Note that this is quite wasteful if only values with imaginary parts larger than 10^4 are needed.

The library syntax is `GEN lfuninit0(GEN L, GEN sdom, long der, long bitprec)`.

3.17.24 lfunlambda($L, s, \{D = 0\}$). Compute the completed L -function $\Lambda(s) = N^{s/2}\gamma(s)L(s)$, or if D is set, the derivative of order D at s . The parameter L is either an **Lmath**, an **Ldata** (created by **lfuncreate**, or an **Linit** (created by **lfuninit**), preferably the latter if many values are to be computed.

The result is given with absolute error less than $2^{-B}|\gamma(s)N^{s/2}|$, where $B = \text{realbitprecision}$.

The library syntax is `GEN lfunlambda0(GEN L, GEN s, long D, long bitprec)`.

3.17.25 lfunmfspec(L). Let L be the L -function attached to a modular eigenform f of weight k , as given by **lfunmf**. In even weight, returns $[\text{vo}, \text{ve}, \text{om}, \text{op}]$, where vo (resp., ve) is the vector of odd (resp., even) periods of f and om and op the corresponding real numbers ω^- and ω^+ normalized in a noncanonical way. In odd weight om is the same as op and we return $[\text{v}, \text{op}]$ where v is the vector of all periods.

```
? D = mfDelta(); mf = mfinit(D); L = lfunmf(mf, D);
? [vo, ve, om, op] = lfunmfspec(L)
%2 = [[1, 25/48, 5/12, 25/48, 1], [1620/691, 1, 9/14, 9/14, 1, 1620/691], \
      0.0074154209298961305890064277459002287248, \
      0.0050835121083932868604942901374387473226]
? DS = mfsymbol(mf, D); bestappr(om*op / mfpetersson(DS), 10^8)
%3 = 8192/225
? mf = mfinit([4, 9, -4], 0);
? F = mfeigenbasis(mf)[1]; L = lfunmf(mf, F);
? [v, om] = lfunmfspec(L)
%6 = [[1, 10/21, 5/18, 5/24, 5/24, 5/18, 10/21, 1], \
      1.1302643192034974852387822584241400608]
? FS = mfsymbol(mf, F); bestappr(om^2 / mfpetersson(FS), 10^8)
%7 = 113246208/325
```

The library syntax is `GEN lfunmfspec(GEN L, long bitprec)`.

3.17.26 lfunmul($L1, L2$). Creates the **Ldata** structure (without initialization) corresponding to the product of the Dirichlet series given by $L1$ and $L2$.

The library syntax is `GEN lfunmul(GEN L1, GEN L2, long bitprec)`.

3.17.27 lfunorderzero($L, \{m = -1\}$). Computes the order of the possible zero of the L -function at the center $k/2$ of the critical strip; return 0 if $L(k/2)$ does not vanish.

If m is given and has a nonnegative value, assumes the order is at most m . Otherwise, the algorithm chooses a sensible default:

- if the L argument is an **Linit**, assume that a multiple zero at $s = k/2$ has order less than or equal to the maximal allowed derivation order.
- else assume the order is less than 4.

You may explicitly increase this value using optional argument m ; this overrides the default value above. (Possibly forcing a recomputation of the **Linit**.)

The library syntax is `long lfunorderzero(GEN L, long m, long bitprec)`.

3.17.28 lfunparams(*ldata*). Returns the parameters $[N, k, Vga]$ of the L -function defined by *ldata*, corresponding respectively to the conductor, the functional equation relating values at s and $k - s$, and the gamma shifts of the L -function (see `lfuncreate`). The gamma shifts are returned to the current precision.

```
? L = lfuncreate(1); /* Riemann zeta function */
? lfunparams(L)
%2 = [1, 1, [0]]
```

The library syntax is `GEN lfunparams(GEN ldata, long prec)`.

3.17.29 lfunqf(*Q*). Returns the *Ldata* structure attached to the Θ function of the lattice attached to the primitive form proportional to the definite positive quadratic form Q .

```
? L = lfunqf(matid(2));
? lfunqf(L, 2)
%2 = 6.0268120396919401235462601927282855839
? lfun(x^2+1, 2)*4
%3 = 6.0268120396919401235462601927282855839
```

The following computes the Madelung constant:

```
? L1=lfunqf(matdiagonal([1,1,1]));
? L2=lfunqf(matdiagonal([4,1,1]));
? L3=lfunqf(matdiagonal([4,4,1]));
? F(s)=6*lfun(L2,s)-12*lfun(L3,s)-lfun(L1,s)*(1-8/4^s);
? F(1/2)
%5 = -1.7475645946331821906362120355443974035
```

The library syntax is `GEN lfunqf(GEN Q, long prec)`.

3.17.30 lfunrootres(*data*). Given the *Ldata* attached to an L -function (or the output of `lfun-thetainit`), compute the root number and the residues.

The output is a 3-component vector $[[[a_1, r_1], \dots, [a_n, r_n], [[b_1, R_1], \dots, [b_m, R_m]], w]$, where r_i is the polar part of $L(s)$ at a_i , R_i is the polar part of $\Lambda(s)$ at b_i or $[0, 0, r]$ if there is no pole, and w is the root number. In the present implementation,

- either the polar part must be completely known (and is then arbitrary): the function determines the root number,

```
? L = lfunmul(1,1); \\ zeta^2
? [r,R,w] = lfunrootres(L);
? r \\ single pole at 1, double
%3 = [[1, 1.[...]*x^-2 + 1.1544[...]*x^-1 + O(x^0)]]
? w
%4 = 1
? R \\ double pole at 0 and 1
%5 = [[1,[...]], [0,[...]]]~
```

- or at most a single pole is allowed: the function computes both the root number and the residue (0 if no pole).

The library syntax is `GEN lfunrootres(GEN data, long bitprec)`.


```

? lfunthetacost(L, 1 + I/2) \\ for theta(1+I/2).
%2 = 23
? lfunthetacost(L, 1 + I/2, 10) \\ for theta^((10))(1+I/2).
%3 = 24
? lfunthetacost(L, [2, 1/10]) \\ cost for theta(t), |t| >= 2, |arg(t)| < 1/10
%4 = 8
? L = lfuncreate( ellinit([1,1]) );
? lfunthetacost(L) \\ for t >= 1
%6 = 2471

```

The library syntax is `long lfunthetacost0(GEN L, GEN tdom = NULL, long m, long bitprec)`.

3.17.35 lfunthetainit($L, \{tdom\}, \{m = 0\}$). Initialization function for evaluating the m -th derivative of theta functions with argument t in domain $tdom$. By default ($tdom$ omitted), t is real, $t \geq 1$. Otherwise, $tdom$ may be

- a positive real scalar ρ : t is real, $t \geq \rho$.
- a nonreal complex number: compute at this particular t ; this allows to compute $\theta(z)$ for any complex z satisfying $|z| \geq |t|$ and $|\arg z| \leq |\arg t|$; we must have $|2 \arg z/d| < \pi/2$, where d is the degree of the Γ factor.
- a pair $[\rho, \alpha]$: assume that $|t| \geq \rho$ and $|\arg t| \leq \alpha$; we must have $|2\alpha/d| < \pi/2$, where d is the degree of the Γ factor.

```

? \p500
? L = lfuncreate(1); \\ Riemann zeta
? t = 1+I/2;
? lfuntheta(L, t); \\ direct computation
time = 30 ms.
? T = lfunthetainit(L, 1+I/2);
time = 30 ms.
? lfuntheta(T, t); \\ instantaneous

```

The T structure would allow to quickly compute $\theta(z)$ for any z in the cone delimited by t as explained above. On the other hand

```

? lfuntheta(T,I)
*** at top-level: lfuntheta(T,I)
*** ~-----
*** lfuntheta: domain error in lfunthetaneed: arg t > 0.785398163397448

```

The initialization is equivalent to

```

? lfunthetainit(L, [abs(t), arg(t)])

```

The library syntax is `GEN lfunthetainit(GEN L, GEN tdom = NULL, long m, long bitprec)`.

3.17.36 lfuntwist(L, χ). Creates the L data structure (without initialization) corresponding to the twist of L by the primitive character attached to the Dirichlet character χ . The conductor of the character must be coprime to the conductor of the L -function L .

The library syntax is `GEN lfuntwist(GEN L, GEN chi, long bitprec)`.

3.17.37 lfunzeros($L, \text{lim}, \{\text{divz} = 8\}$). lim being either a positive upper limit or a nonempty real interval, computes an ordered list of zeros of $L(s)$ on the critical line up to the given upper limit or in the given interval. Use a naive algorithm which may miss some zeros: it assumes that two consecutive zeros at height $T \geq 1$ differ at least by $2\pi/\omega$, where

$$\omega := \text{divz} \cdot (d \log(T/2\pi) + d + 2 \log(N/(\pi/2)^d)).$$

To use a finer search mesh, set divz to some integral value larger than the default ($= 8$).

```
? lfunzeros(1, 30) \\ zeros of Rieman zeta up to height 30
%1 = [14.134[...], 21.022[...], 25.010[...]]
? #lfunzeros(1, [100,110]) \\ count zeros with 100 <= Im(s) <= 110
%2 = 4
```

The algorithm also assumes that all zeros are simple except possibly on the real axis at $s = k/2$ and that there are no poles in the search interval. (The possible zero at $s = k/2$ is repeated according to its multiplicity.)

If you pass an **Linit** to the function, the algorithm assumes that a multiple zero at $s = k/2$ has order less than or equal to the maximal derivation order allowed by the **Linit**. You may increase that value in the **Linit** but this is costly: only do it for zeros of low height or in **lfunorderzero** instead.

The library syntax is **GEN lfunzeros**(**GEN L**, **GEN lim**, **long divz**, **long bitprec**).

3.18 Modular forms.

This section describes routines for working with modular forms and modular form spaces.

3.18.1 Modular form spaces.

These structures are initialized by the **mfini**t command; supported modular form *spaces* with corresponding flags are the following:

- The full modular form space $M_k(\Gamma_0(N), \chi)$, where k is an integer or a half-integer and χ a Dirichlet character modulo N (*flag* = 4, the default).
- The cuspidal space $S_k(\Gamma_0(N), \chi)$ (*flag* = 1).
- The Eisenstein space $\mathcal{E}_k(\Gamma_0(N), \chi)$ (*flag* = 3), so that $M_k = \mathcal{E}_k \oplus S_k$.
- The new space $S_k^{\text{new}}(\Gamma_0(N), \chi)$ (*flag* = 0).
- The old space $S_k^{\text{old}}(\Gamma_0(N), \chi)$ (*flag* = 2), so that $S_k = S_k^{\text{new}} \oplus S_k^{\text{old}}$.

These resulting **mf** structure contains a basis of modular forms, which is accessed by the function **mfbasis**; the elements of this basis have Fourier coefficients in the cyclotomic field $\mathbf{Q}(\chi)$. These coefficients are given algebraically, as rational numbers or **t_POLMODs**. The member function **mf.mod** recovers the modulus used to define $\mathbf{Q}(\chi)$, which is a cyclotomic polynomial $\Phi_n(t)$. When needed, the elements of $\mathbf{Q}(\chi)$ are considered to be canonically embedded into $\mathbf{C}$ via $\text{Mod}(t, \Phi_n(t)) \mapsto \exp(2i\pi/n)$.

The basis of eigenforms for the new space is obtained by the function **mfeigenbasis**: the elements of this basis now have Fourier coefficients in a relative field extension of $\mathbf{Q}(\chi)$. Note that if the space is larger than the new space (i.e. is the cuspidal or full space) we nevertheless obtain only the eigenbasis for the new space.

3.18.2 Generalized modular forms.

A modular form is represented in a special internal format giving the possibility to compute an arbitrary number of terms of its Fourier coefficients at infinity $[a(0), a(1), \dots, a(n)]$ using the function `mfcoefs`. These coefficients are given algebraically, as rational numbers or `t_POLMODs`. The member function `f.mod` recovers the modulus used in the coefficients of f , which will be the same as for $k = \mathbf{Q}(\chi)$ (a cyclotomic polynomial), or define a number field extension K/k .

Modular forms are obtained either directly from other mathematical objects, e.g., elliptic curves, or by a specific formula, e.g., Eisenstein series or Ramanujan's Delta function, or by applying standard operators to existing forms (Hecke operators, Rankin–Cohen brackets, ...). A function `mfparams` is provided so that one can recover the level, weight, character and field of definition corresponding to a given modular form.

A number of creation functions and operations are provided. It is however important to note that strictly speaking some of these operations create objects which are *not* modular forms: typical examples are derivation or integration of modular forms, the Eisenstein series E_2 , eta quotients, or quotients of modular forms. These objects are nonetheless very important in the theory, so are not considered as errors; however the user must be aware that no attempt is made to check that the objects that he handles are really modular. When the documentation of a function does not state that it applies to generalized modular forms, then the output is undefined if the input is not a true modular form.

3.18.3 `lfunmf(mf, {F})`. If F is a modular form in `mf`, output the L -functions corresponding to its $[\mathbf{Q}(F) : \mathbf{Q}(\chi)]$ complex embeddings, ready for use with the `lfun` package. If F is omitted, output the L -functions attached to all eigenforms in the new space; the result is a vector whose length is the number of Galois orbits of newforms. Each entry contains the vector of L -functions corresponding to the d complex embeddings of an orbit of dimension d over $\mathbf{Q}(\chi)$.

```
? mf = mfinit([35,2],0);mffields(mf)
%1 = [y, y^2 - y - 4]
? f = mfeigenbasis(mf)[2]; mfparams(f) \\ orbit of dimension two
%2 = [35, 2, 1, y^2 - y - 4, t - 1]
? [L1,L2] = lfunmf(mf, f); \\ Two L-functions
? lfun(L1,1)
%4 = 0.81018461849460161754947375433874745585
? lfun(L2,1)
%5 = 0.46007635204895314548435893464149369804
? [ lfun(L,1) | L <- concat(lfunmf(mf)) ]
%6 = [0.70291..., 0.81018..., 0.46007...]
```

The `concat` instruction concatenates the vectors corresponding to the various (here two) orbits, so that we obtain the vector of all the L -functions attached to eigenforms.

The library syntax is `GEN lfunmf(GEN mf, GEN F = NULL, long bitprec)`.

3.18.4 `mfDelta()`. `Mf` structure corresponding to the Ramanujan Delta function Δ .

```
? mfcoefs(mfDelta(),4)
%1 = [0, 1, -24, 252, -1472]
```

The library syntax is `GEN mfDelta()`.

3.18.5 mfEH(k). k being in $1/2 + \mathbf{Z}_{\geq 0}$, return the mf structure corresponding to the Cohen-Eisenstein series H_k of weight k on $\Gamma_0(4)$.

```
? H = mfEH(13/2); mfcoefs(H,4)
%1 = [691/32760, -1/252, 0, 0, -2017/252]
```

The coefficients of H are given by the Cohen-Hurwitz function $H(k - 1/2, N)$ and can be obtained for moderately large values of N (the algorithm uses $\tilde{O}(N)$ time):

```
? mfcoef(H, 10^5+1)
time = 55 ms.
%2 = -12514802881532791504208348
? mfcoef(H, 10^7+1)
time = 6,044 ms.
%3 = -1251433416009877455212672599325104476
```

The library syntax is GEN mfEH(GEN k).

3.18.6 mfEk(k). K being an even nonnegative integer, return the mf structure corresponding to the standard Eisenstein series E_k .

```
? mfcoefs(mfEk(8), 4)
%1 = [1, 480, 61920, 1050240, 7926240]
```

The library syntax is GEN mfEk(long k).

3.18.7 mfTheta($\{psi = 1\}$). The unary theta function corresponding to the primitive Dirichlet character ψ . Its level is $4F(\psi)^2$ and its weight is $1 - \psi(-1)/2$.

```
? Ser(mfcoefs(mfTheta(), 30))
%1 = 1 + 2*x + 2*x^4 + 2*x^9 + 2*x^16 + 2*x^25 + 0(x^31)

? f = mfTheta(8); Ser(mfcoefs(f, 30))
%2 = 2*x - 2*x^9 - 2*x^25 + 0(x^31)
? mfparams(f)
%3 = [256, 1/2, 8, y, t + 1]

? g = mfTheta(-8); Ser(mfcoefs(g, 30))
%4 = 2*x + 6*x^9 - 10*x^25 + 0(x^31)
? mfparams(g)
%5 = [256, 3/2, 8, y, t + 1]

? h = mfTheta(Mod(2,5)); mfparams(h)
%6 = [100, 3/2, Mod(7, 20), y, t^2 + 1]
```

The library syntax is GEN mfTheta(GEN psi = NULL).

3.18.8 mfatkin(*mfatk*, *f*). Given a *mfatk* output by *mfatk* = *mfatkininit*(*mf*, *Q*) and a modular form *f* belonging to the space *mf*, returns the modular form $g = C \times f|W_Q$, where $C = \text{mfatk}[3]$ is a normalizing constant such that *g* has the same field of coefficients as *f*; *mfatk*[3] gives the constant *C*, and *mfatk*[1] gives the modular form space to which *g* belongs (or is set to 0 if it is *mf*).

```
? mf = mfinit([35,2],0); [f] = mfbasis(mf);
? mfcoefs(f, 4)
%2 = [0, 3, -1, 0, 3]
? mfatk = mfatkininit(mf,7);
? g = mfatkin(mfatk, f); mfcoefs(g, 4)
%4 = [0, 1, -1, -2, 7]
? mfatk = mfatkininit(mf,35);
? g = mfatkin(mfatk, f); mfcoefs(g, 4)
%6 = [0, -3, 1, 0, -3]
```

The library syntax is GEN *mfatkin*(GEN *mfatk*, GEN *f*).

3.18.9 mfatkineigenvalues(*mf*, *Q*). Given a modular form space *mf* of integral weight *k* and a primitive divisor *Q* of the level *N* of *mf*, outputs the Atkin–Lehner eigenvalues of w_Q on the new space, grouped by orbit. If the Nebentypus χ of *mf* is a (trivial or) quadratic character defined modulo N/Q , the result is rounded and the eigenvalues are $\pm i^k$.

```
? mf = mfinit([35,2],0); mffields(mf)
%1 = [y, y^2 - y - 4] \\ two orbits, dimension 1 and 2
? mfatkineigenvalues(mf,5)
%2 = [[1], [-1, -1]]
? mf = mfinit([12,7,Mod(3,4)],0);
? mfatkineigenvalues(mf,3)
%4 = [[I, -I, -I, I, I, -I]] \\ one orbit
```

To obtain the eigenvalues on a larger space than the new space, e.g., the full space, you can directly call [*mfB*, *M*, *C*] = *mfatkininit* and compute the eigenvalues as the roots of the characteristic polynomial of M/C , by dividing the roots of *charpoly*(*M*) by *C*. Note that the characteristic polynomial is computed exactly since *M* has coefficients in $\mathbf{Q}(\chi)$, whereas *C* may be given by a complex number. If the coefficients of the characteristic polynomial are polmods modulo *T* they must be embedded to $\mathbf{C}$ first using *subst*(*lift*(), *t*, *exp*(2*I*Pi/*n*)), when *T* is *poliscyclo*(*n*); note that $T = \text{mf.mod}$.

The library syntax is GEN *mfatkineigenvalues*(GEN *mf*, long *Q*, long *prec*).

3.18.10 mfatkininit(*mf*, *Q*). Given a modular form space with parameters *N*, *k*, χ and a primitive divisor *Q* of the level *N*, initializes data necessary for working with the Atkin–Lehner operator W_Q , for now only the function *mfatkin*. We write $\chi \sim \chi_Q \chi_{N/Q}$ where the two characters are primitive with (coprime) conductors dividing *Q* and N/Q respectively. For $F \in M_k(\Gamma_0(N), \chi)$, the form $F|W_Q$ still has level *N* and weight *k* but its Nebentypus may no longer be χ : it becomes $\overline{\chi_Q} \chi_{N/Q}$ if *k* is integral and $\overline{\chi_Q} \chi_{N/Q}(4Q/\cdot)$ if not.

The result is a technical 4-component vector [*mfB*, *MC*, *C*, *mf*], where

- *mfB* encodes the modular form space to which $F|W_Q$ belongs when $F \in M_k(\Gamma_0(N), \chi)$: an *mfinit* corresponding to a new Nebentypus or the integer 0 when the character does not change. This does not depend on *F*.

- **MC** is the matrix of W_Q on the bases of **mf** and **mfB** multiplied by a normalizing constant $C(k, \chi, Q)$. This matrix has polmod coefficients in $\mathbf{Q}(\chi)$.

- **C** is the complex constant $C(k, \chi, Q)$. For k integral, let $A(k, \chi, Q) = Q^\varepsilon / g(\chi_Q)$, where $\varepsilon = 0$ for k even and $1/2$ for k odd and where $g(\chi_Q)$ is the Gauss sum attached to χ_Q . (A similar, more complicated, definition holds in half-integral weight depending on the parity of $k - 1/2$.) Then if M denotes the matrix of W_Q on the bases of **mf** and **mfB**, $A \cdot M$ has coefficients in $\mathbf{Q}(\chi)$. If A is rational, we let $C = 1$ and $C = A$ as a floating point complex number otherwise, and finally $\text{MC} := M \cdot C$.

```
? mf=mfinit([32,4],0); [mfB,MC,C]=mfatkininit(mf,32); MC
%1 =
[5/16 11/2 55/8]
[ 1/8    0 -5/4]
[1/32 -1/4 11/16]
? C
%2 = 1
? mf=mfinit([32,4,8],0); [mfB,MC,C]=mfatkininit(mf,32); MC
%3 =
[ 1/8 -7/4]
[-1/16 -1/8]
? C
%4 = 0.35355339059327376220042218105242451964
? algdep(C,2)  \\ C = 1/sqrt(8)
%5 = 8*x^2 - 1
```

The library syntax is GEN `mfatkininit(GEN mf, long Q, long prec)`.

3.18.11 mfbasis($NK, \{space = 4\}$). If $NK = [N, k, CHI]$ as in `mfinit`, gives a basis of the corresponding subspace of $M_k(\Gamma_0(N), \chi)$. NK can also be the output of `mfinit`, in which case **space** can be omitted. To obtain the eigenforms, use `mfeigenbasis`.

If **space** is a full space M_k , the output is the union of first, a basis of the space of Eisenstein series, and second, a basis of the cuspidal space.

```
? see(L) = apply(f->mfcoefs(f,3), L);
? mf = mfinit([35,2],0);
? see( mfbasis(mf) )
%2 = [[0, 3, -1, 0], [0, -1, 9, -8], [0, 0, -8, 10]]
? see( mfeigenbasis(mf) )
%3 = [[0, 1, 0, 1], [Mod(0, z^2 - z - 4), Mod(1, z^2 - z - 4), \
      Mod(-z, z^2 - z - 4), Mod(z - 1, z^2 - z - 4)]]
? mf = mfinit([35,2]);
? see( mfbasis(mf) )
%5 = [[1/6, 1, 3, 4], [1/4, 1, 3, 4], [17/12, 1, 3, 4], \
      [0, 3, -1, 0], [0, -1, 9, -8], [0, 0, -8, 10]]
? see( mfbasis([48,4],0) )
%6 = [[0, 3, 0, -3], [0, -3, 0, 27], [0, 2, 0, 30]]
```

The library syntax is GEN `mfbasis(GEN NK, long space)`.

3.18.12 mfbd(F, d). F being a generalized modular form, return $B(d)(F)$, where $B(d)$ is the expanding operator $\tau \mapsto d\tau$.

```
? D2=mfbd(mfDelta(),2); mfcoefs(D2, 6)
%1 = [0, 0, 1, 0, -24, 0, 252]
```

The library syntax is GEN mfbd(GEN F, long d).

3.18.13 mfbracket($F, G, \{m = 0\}$). Compute the m -th Rankin–Cohen bracket of the generalized modular forms F and G .

```
? E4 = mfEk(4); E6 = mfEk(6);
? D1 = mfbracket(E4,E4,2); mfcoefs(D1,5)/4800
%2 = [0, 1, -24, 252, -1472, 4830]
? D2 = mfbracket(E4,E6,1); mfcoefs(D2,10)/(-3456)
%3 = [0, 1, -24, 252, -1472, 4830]
```

The library syntax is GEN mfbracket(GEN F, GEN G, long m).

3.18.14 mfcoef(F, n). Compute the n -th Fourier coefficient $a(n)$ of the generalized modular form F . Note that this is the $n + 1$ -st component of the vector **mfcoefs**(F, n) as well as the second component of **mfcoefs**($F, 1, n$).

```
? mfcoef(mfDelta(),10)
%1 = -115920
```

The library syntax is GEN mfcoef(GEN F, long n).

3.18.15 mfcoefs($F, n, \{d = 1\}$). Compute the vector of Fourier coefficients $[a[0], a[d], \dots, a[nd]]$ of the generalized modular form F ; d must be positive and $d = 1$ by default.

```
? D = mfDelta();
? mfcoefs(D,10)
%2 = [0, 1, -24, 252, -1472, 4830, -6048, -16744, 84480, -113643, -115920]
? mfcoefs(D,5,2)
%3 = [0, -24, -1472, -6048, 84480, -115920]
? mfcoef(D,10)
%4 = -115920
```

This function also applies when F is a modular form space as output by **mfinit**; it then returns the matrix whose columns give the Fourier expansions of the elements of **mfbasis**(F):

```
? mf = mfinit([1,12]);
? mfcoefs(mf,5)
%2 =
[691/65520    0]
[      1     1]
[   2049   -24]
[  177148   252]
[ 4196353 -1472]
[48828126  4830]
```

The library syntax is GEN mfcoefs(GEN F, long n, long d).

3.18.16 mfconductor(mf, F). mf being output by `mfini`t for the cuspidal space and F a modular form, gives the smallest level at which F is defined. In particular, if F is cuspidal and we write $F = \sum_j B(d_j) f_j$ for new forms f_j of level N_j (see `mftnew`), then its conductor is the least common multiple of the $d_j N_j$.

```
? mf=mfinit([96,6],1); vF = mfbasis(mf); mfdim(mf)
%1 = 72
? vector(10,i, mfconductor(mf, vF[i]))
%2 = [3, 6, 12, 24, 48, 96, 4, 8, 12, 16]
```

The library syntax is `long mfconductor(GEN mf, GEN F)`.

3.18.17 mfcosets(N). Let N be a positive integer. Return the list of right cosets of $\Gamma_0(N) \backslash \Gamma$, i.e., matrices $\gamma_j \in \Gamma$ such that $\Gamma = \bigsqcup_j \Gamma_0(N) \gamma_j$. The γ_j are chosen in the form $[a, b; c, d]$ with $c \mid N$.

```
? mfcosets(4)
%1 = [[0, -1; 1, 0], [1, 0; 1, 1], [0, -1; 1, 2], [0, -1; 1, 3], \
      [1, 0; 2, 1], [1, 0; 4, 1]]
```

We also allow the argument N to be a modular form space, in which case it is replaced by the level of the space:

```
? M = mfinit([4, 12, 1], 0); mfcosets(M)
%2 = [[0, -1; 1, 0], [1, 0; 1, 1], [0, -1; 1, 2], [0, -1; 1, 3], \
      [1, 0; 2, 1], [1, 0; 4, 1]]
```

Warning. In the present implementation, the trivial coset is represented by $[1, 0; N, 1]$ and is the last in the list.

The library syntax is `GEN mfcosets(GEN N)`.

3.18.18 mfcuspisregular($NK, cusp$). In the space defined by $NK = [N, k, CHI]$ or $NK = mf$, determine if $cusp$ in canonical format (oo or denominator dividing N) is regular or not.

```
? mfcuspisregular([4,3,-4],1/2)
%1 = 0
```

The library syntax is `long mfcuspisregular(GEN NK, GEN cusp)`.

3.18.19 mfcusps(N). Let N be a positive integer. Return the list of cusps of $\Gamma_0(N)$ in the form a/b with $b \mid N$.

```
? mfcusps(24)
%1 = [0, 1/2, 1/3, 1/4, 1/6, 1/8, 1/12, 1/24]
```

We also allow the argument N to be a modular form space, in which case it is replaced by the level of the space:

```
? M = mfinit([4, 12, 1], 0); mfcusps(M)
%2 = [0, 1/2, 1/4]
```

The library syntax is `GEN mfcusps(GEN N)`.

3.18.20 mfcuspval(*mf*, *F*, *cusp*). Valuation of modular form *F* in the space *mf* at *cusp*, which can be either ∞ or any rational number. The result is either a rational number or ∞ if *F* is zero. Let χ be the Nebentypus of the space *mf*; if $\mathbf{Q}(F) \neq \mathbf{Q}(\chi)$, return the vector of valuations attached to the $[\mathbf{Q}(F) : \mathbf{Q}(\chi)]$ complex embeddings of *F*.

```
? T=mfTheta(); mf=mfinit([12,1/2]); mfcusps(12)
%1 = [0, 1/2, 1/3, 1/4, 1/6, 1/12]
? apply(x->mfcuspval(mf,T,x), %1)
%2 = [0, 1/4, 0, 0, 1/4, 0]
? mf=mfinit([12,6,12],1); F=mbasis(mf)[5];
? apply(x->mfcuspval(mf,F,x),%1)
%4 = [1/12, 1/6, 1/2, 2/3, 1/2, 2]
? mf=mfinit([12,3,-4],1); F=mbasis(mf)[1];
? apply(x->mfcuspval(mf,F,x),%1)
%6 = [1/12, 1/6, 1/4, 2/3, 1/2, 1]

? mf = mfinit([625,2],0); [F] = mfeigenbasis(mf); mfparams(F)
%7 = [625, 2, 1, y^2 - y - 1, t - 1] \\ [Q(F):Q(chi)] = 2
? mfcuspval(mf, F, 1/25)
%8 = [1, 2] \\ one conjugate has valuation 1, and the other is 2
? mfcuspval(mf, F, 1/5)
%9 = [1/25, 1/25]
```

The library syntax is GEN mfcuspval(GEN mf, GEN F, GEN cusp, long bitprec).

3.18.21 mfcuspwidth(*N*, *cusp*). Width of *cusp* in $\Gamma_0(N)$.

```
? mfcusps(12)
%1 = [0, 1/2, 1/3, 1/4, 1/6, 1/12]
? [mfcuspwidth(12,c) | c <- mfcusps(12)]
%2 = [12, 3, 4, 3, 1, 1]
? mfcuspwidth(12, oo)
%3 = 1
```

We also allow the argument *N* to be a modular form space, in which case it is replaced by the level of the space:

```
? M = mfinit([4, 12, 1], 0); mfcuspwidth(M, 1/2)
%4 = 1
```

The library syntax is long mfcuspwidth(GEN N, GEN cusp).

3.18.22 mfderiv(*F*, {*m* = 1}). *m*-th formal derivative of the power series corresponding to the generalized modular form *F*, with respect to the differential operator *qd/dq* (default *m* = 1).

```
? D=mfDelta();
? mfcoefs(D, 4)
%2 = [0, 1, -24, 252, -1472]
? mfcoefs(mfderiv(D), 4)
%3 = [0, 1, -48, 756, -5888]
```

The library syntax is GEN mfderiv(GEN F, long m).

3.18.23 mfderivE2($F, \{m = 1\}$). Compute the Serre derivative $(qd/dq)F - kE_2F/12$ of the generalized modular form F , which has weight $k + 2$; if F is a true modular form, then its Serre derivative is also modular. If $m > 1$, compute the m -th iterate, of weight $k + 2m$.

```
? mfcoefs(mfderivE2(mfEk(4)),5)*(-3)
%1 = [1, -504, -16632, -122976, -532728]
? mfcoefs(mfEk(6),5)
%2 = [1, -504, -16632, -122976, -532728]
```

The library syntax is GEN mfderivE2(GEN F, long m).

3.18.24 mfdescribe($F, \{&G\}$). Gives a human-readable description of F , which is either a modular form space or a generalized modular form. If the address of G is given, puts into G the vector of parameters of the outermost operator defining F ; this vector is empty if F is a leaf (an atomic object such as mfDelta(), not defined in terms of other forms) or a modular form space.

```
? E1 = mfeisenstein(4,-3,-4); mfdescribe(E1)
%1 = "F_4(-3, -4)"
? E2 = mfeisenstein(3,5,-7); mfdescribe(E2)
%2 = "F_3(5, -7)"
? E3 = mfderivE2(mfmul(E1,E2), 3); mfdescribe(E3,&G)
%3 = "DERE2^3(MUL(F_4(-3, -4), F_3(5, -7)))"
? mfdescribe(G[1][1])
%4 = "MUL(F_4(-3, -4), F_3(5, -7))"
? G[2]
%5 = 3
? for (i = 0, 4, mf = mfini([37,4],i); print(mfdescribe(mf)));
S_4^new(G_0(37, 1))
S_4(G_0(37, 1))
S_4^old(G_0(37, 1))
E_4(G_0(37, 1))
M_4(G_0(37, 1))
```

The library syntax is GEN mfdescribe(GEN F, GEN *G = NULL).

3.18.25 mfdim($NK, \{space = 4\}$). If $NK = [N, k, CHI]$ as in mfini, gives the dimension of the corresponding subspace of $M_k(\Gamma_0(N), \chi)$. NK can also be the output of mfini, in which case space must be omitted.

The subspace is described by the small integer **space**: 0 for the newspace $S_k^{\text{new}}(\Gamma_0(N), \chi)$, 1 for the cuspidal space S_k , 2 for the oldspace S_k^{old} , 3 for the space of Eisenstein series E_k and 4 for the full space M_k .

Wildcards. As in `mfin`, *CHI* may be the wildcard 0 (all Galois orbits of characters); in this case, the output is a vector of $[order, conrey, dim, dimdih]$ corresponding to the nontrivial spaces, where

- *order* is the order of the character,
- *conrey* its Conrey label from which the character may be recovered via `znchar(conrey)`,
- *dim* the dimension of the corresponding space,
- *dimdih* the dimension of the subspace of dihedral forms corresponding to Hecke characters if $k = 1$ (this is not implemented for the old space and set to -1 for the time being) and 0 otherwise.

The spaces are sorted by increasing order of the character; the characters are taken up to Galois conjugation and the Conrey number is the minimal one among Galois conjugates. In weight 1, this is only implemented when the space is 0 (newspace), 1 (cusp space), 2(old space) or 3(Eisenstein series).

Wildcards for sets of characters. *CHI* may be a set of characters, and we return the set of $[dim, dimdih]$.

Wildcard for $M_k(\Gamma_1(N))$. Additionally, the wildcard *CHI* = -1 is available in which case we output the total dimension of the corresponding subspace of $M_k(\Gamma_1(N))$. In weight 1, this is not implemented when the space is 4 (fullspace).

```
? mfdim([23,2], 0) \\ new space
%1 = 2
? mfdim([96,6], 0)
%2 = 10
? mfdim([10^9,4], 3) \\ Eisenstein space
%1 = 40000
? mfdim([10^9+7,4], 3)
%2 = 2
? mfdim([68,1,-1],0)
%3 = 3
? mfdim([68,1,0],0)
%4 = [[2, Mod(67, 68), 1, 1], [4, Mod(47, 68), 1, 1]]
? mfdim([124,1,0],0)
%5 = [[6, Mod(67, 124), 2, 0]]
```

This last example shows that there exists a nondihedral form of weight 1 in level 124.

The library syntax is `GEN mfdim(GEN NK, long space)`.

3.18.26 mfddiv(F, G). Given two generalized modular forms F and G , compute F/G assuming that the quotient will not have poles at infinity. If this is the case, use `mfshift` before doing the division.

```
? D = mfDelta(); \\ Delta
? H = mfpow(mfEk(4), 3);
? J = mfddiv(H, D)
*** at top-level: J=mfddiv(H,mfdeltac
*** ^-----
*** mfddiv: domain error in mfddiv: ord(G) > ord(F)
? J = mfddiv(H, mfshift(D,1));
? mfcoefs(J, 4)
%4 = [1, 744, 196884, 21493760, 864299970]
```

The library syntax is `GEN mfddiv(GEN F, GEN G)`.

3.18.27 mfeigenbasis(mf). Vector of the eigenforms for the space `mf`. The initial basis of forms computed by `mfinit` before splitting is also available via `mfbasis`.

```
? mf = mfinit([26,2],0);
? see(L) = for(i=1,#L,print(mfcoefs(L[i],6)));
? see( mfeigenbasis(mf) )
[0, 1, -1, 1, 1, -3, -1]
[0, 1, 1, -3, 1, -1, -3]
? see( mfbasis(mf) )
[0, 2, 0, -2, 2, -4, -4]
[0, -2, -4, 10, -2, 0, 8]
```

The eigenforms are internally expressed as (algebraic) linear combinations of `mfbasis(mf)` and it is very inefficient to compute many coefficients of those forms individually: you should rather use `mfcoefs(mf)` to expand the basis once and for all, then multiply by `mftobasis(mf,f)` for the forms you're interested in:

```
? mf = mfinit([96,6],0); B = mfeigenbasis(mf); #B
%1 = 8;
? vector(#B, i, mfcoefs(B[i],1000)); \\ expanded individually: slow
time = 7,881 ms.
? M = mfcoefs(mf, 1000); \\ initialize once
time = 982 ms.
? vector(#B, i, M * mftobasis(mf,B[i])); \\ then expand: much faster
time = 623 ms.
```

When the eigenforms are defined over an extension field of $\mathbf{Q}(\chi)$ for a nonrational character, their coefficients are hard to read and you may want to lift them or to express them in an absolute number field. In the construction below T defines $\mathbf{Q}(f)$ over $\mathbf{Q}$, a is the image of the generator $\text{Mod}(t, t^2 + t + 1)$ of $\mathbf{Q}(\chi)$ in $\mathbf{Q}(f)$ and $y - ka$ is the image of the root y of `f.mod`:

```
? mf = mfinit([31, 2, Mod(25,31)], 0); [f] = mfeigenbasis(mf);
? f.mod
%2 = Mod(1, t^2 + t + 1)*y^2 + Mod(2*t + 2, t^2 + t + 1)
? v = liftpol(mfcoefs(f,5))
%3 = [0, 1, (-t - 1)*y - 1, t*y + (t + 1), (2*t + 2)*y + 1, t]
```

```

? [T,a,k] = rnfequation(mf.mod, f.mod, 1)
%4 = [y^4 + 2*y^2 + 4, Mod(-1/2*y^2 - 1, y^4 + 2*y^2 + 4), 0]
? liftpol(substvec(v, [t,y], [a, y-k*a]))
%5 = [0, 1, 1/2*y^3 - 1, -1/2*y^3 - 1/2*y^2 - y, -y^3 + 1, -1/2*y^2 - 1]

```

Beware that the meaning of y has changed in the last line is different: it now represents of root of T , no longer of $f.mod$ (the notions coincide if $k = 0$ as here but it will not always be the case). This can be avoided with an extra variable substitution, for instance

```

? [T,a,k] = rnfequation(mf.mod, subst(f.mod,'y','x'), 1)
%6 = [x^4 + 2*x^2 + 4, Mod(-1/2*x^2 - 1, x^4 + 2*x^2 + 4), 0]
? liftpol(substvec(v, [t,y], [a, x-k*a]))
%7 = [0, 1, 1/2*x^3 - 1, -1/2*x^3 - 1/2*x^2 - x, -x^3 + 1, -1/2*x^2 - 1]

```

The library syntax is `GEN mfeigenbasis(GEN mf)`.

3.18.28 mfeigensearch($NK, \{AP\}$). Search for a normalized rational eigen cuspform with quadratic character given restrictions on a few initial coefficients. The meaning of the parameters is as follows:

- NK governs the limits of the search: it is of the form $[N, k]$: search for given level N , weight k and quadratic character; note that the character $(D/.)$ is uniquely determined by (N, k) . The level N can be replaced by a vector of allowed levels.

- AP is the search criterion, which can be omitted: a list of pairs $[\dots, [p, a_p], \dots]$, where p is a prime number and a_p is either a `t_INT` (the p -th Fourier coefficient must match a_p exactly) or a `t_INTMOD Mod(a, b)` (the p -th coefficient must be congruent to a modulo b).

The result is a vector of newforms f matching the search criteria, sorted by increasing level then increasing $|D|$.

```

? #mfeigensearch([[1..80],2], [[2,2],[3,-1]])
%1 = 1
? #mfeigensearch([[1..80],2], [[2,2],[5,2]])
%2 = 1
? v = mfeigensearch([[1..20],2], [[3,Mod(2,3)],[7,Mod(5,7)]]); #v
%3 = 1
? F=v[1]; [mfparams(F)[1], mfcoefs(F,15)]
%4 = [11, [0, 1, -2, -1, 2, 1, 2, -2, 0, -2, -2, 1, -2, 4, 4, -1]]

```

The library syntax is `GEN mfeigensearch(GEN NK, GEN AP = NULL)`.

3.18.29 mfeisenstein($k, \{CHI1\}, \{CHI2\}$). Create the Eisenstein series $E_k(\chi_1, \chi_2)$, where $k \geq 1$, χ_i are Dirichlet characters and an omitted character is considered as trivial. This form belongs to $\mathcal{E}_k(\Gamma_0(N), \chi)$ with $\chi = \chi_1 \chi_2$ and N is the product of the conductors of χ_1 and χ_2 .

```

? CHI = Mod(3,4);
? E = mfeisenstein(3, CHI);
? mfcoefs(E, 6)
%2 = [-1/4, 1, 1, -8, 1, 26, -8]
? CHI2 = Mod(4,5);
? mfcoefs(mfeisenstein(3,CHI,CHI2), 6)
%3 = [0, 1, -1, -10, 1, 25, 10]

```

```
? mfcoefs(mfeisenstein(4,CHI,CHI), 6)
%4 = [0, 1, 0, -28, 0, 126, 0]
? mfcoefs(mfeisenstein(4), 6)
%5 = [1/240, 1, 9, 28, 73, 126, 252]
```

Note that $\text{mfeisenstein}(k)$ is 0 for k odd and $-B_k/(2k) \cdot E_k$ for k even, where

$$E_k(q) = 1 - (2k/B_k) \sum_{n \geq 1} \sigma_{k-1}(n) q^n$$

is the standard Eisenstein series. In other words it is normalized so that its linear coefficient is 1.

Important note. This function is currently implemented only when $\mathbf{Q}(\chi)$ is the field of definition of $E_k(\chi_1, \chi_2)$. If it is a strict subfield, an error is raised:

```
? mfeisenstein(6, Mod(7,9), Mod(4,9));
*** at top-level: mfeisenstein(6,Mod(7,9),Mod(4,9))
*** ^-----
*** mfeisenstein: sorry, mfeisenstein for these characters is not
*** yet implemented.
```

The reason for this is that each modular form is attached to a modular form space $M_k(\Gamma_0(N), \chi)$. This is a $\mathbf{C}$ -vector space but it allows a basis of forms defined over $\mathbf{Q}(\chi)$ and is only implemented as a $\mathbf{Q}(\chi)$ -vector space: there is in general no mechanism to take linear combinations of forms in the space with coefficients belonging to a larger field. (Due to their importance, eigenforms are the single exception to this restriction; for an eigenform F , $\mathbf{Q}(F)$ is built on top of $\mathbf{Q}(\chi)$.) When the property $\mathbf{Q}(\chi) = \mathbf{Q}(E_k(\chi_1, \chi_2))$ does not hold, we cannot express E as a $\mathbf{Q}(\chi)$ -linear combination of the basis forms and many operations will fail. For this reason, the construction is currently disabled.

The library syntax is `GEN mfeisenstein(long k, GEN CHI1 = NULL, GEN CHI2 = NULL)`.

3.18.30 mfembed($f, \{v\}$). Let f be a generalized modular form with parameters $[N, k, \chi, P]$ (see `mfparams`, we denote $\mathbf{Q}(\chi)$ the subfield of $\mathbf{C}$ generated by the values of χ and $\mathbf{Q}(f)$ the field of definition of f . In this context $\mathbf{Q}(\chi)$ has a single canonical complex embedding given by $s : \text{Mod}(\mathfrak{t}, \text{polcyclo}(\mathfrak{n}, \mathfrak{t})) \mapsto \exp(2i\pi/n)$ and the number field $\mathbf{Q}(f)$ has $[\mathbf{Q}(f) : \mathbf{Q}(\chi)]$ induced embeddings attached to the complex roots of the polynomial $s(P)$. If $\mathbf{Q}(f)$ is strictly larger than $\mathbf{Q}(\chi)$ we only allow an f which is an eigenform, produced by `mfeigenbasis`.

This function is meant to create embeddings of $\mathbf{Q}(f)$ and/or apply them to the object v , typically a vector of Fourier coefficients of f from `mfcoefs`.

- If v is omitted and f is a modular form as above, we return the embedding of $\mathbf{Q}(\chi)$ if $\mathbf{Q}(\chi) = \mathbf{Q}(f)$ and a vector containing $[\mathbf{Q}(f) : \mathbf{Q}(\chi)]$ embeddings of $\mathbf{Q}(f)$ otherwise.

- If v is given, it must be a scalar in $\mathbf{Q}(f)$, or a vector/matrix of such, we apply the embeddings coefficientwise and return either a single result if $\mathbf{Q}(f) = \mathbf{Q}(\chi)$ and a vector of $[\mathbf{Q}(f) : \mathbf{Q}(\chi)]$ results otherwise.

- Finally f can be replaced by a single embedding produced by `mfembed(f)` (v was omitted) and we apply that particular embedding to v .

```
? mf = mfinit([35,2,Mod(11,35)], 0);
```

```

? [f] = mfbasis(mf);
? f.mod \\  $\mathbf{Q}(\chi) = \mathbf{Q}(\zeta_3)$ 
%3 = t^2 + t + 1
? v = mfcoefs(f,5); lift(v) \\ coefficients in  $\mathbf{Q}(\chi)$ 
%4 = [0, 2, -2*t - 2, 2*t, 2*t, -2*t - 2]
? mfembed(f, v) \\ single embedding
%5 = [0, 2, -1 - 1.7320...*I, -1 + 1.73205...*I, -1 + 1.7320...*I, ...]

? [F] = mfeigenbasis(mf);
? mffieldds(mf)
%7 = [y^2 + Mod(-2*t, t^2 + t + 1)] \\ [ $\mathbf{Q}(f) : \mathbf{Q}(\chi)$ ] = 2
? V = liftpol( mfcoefs(F,5) );
%8 = [0, 1, y + (-t - 1), (t + 1)*y + t, (-2*t - 2)*y + t, -t - 1]
? vall = mfembed(F, V); #vall
%9 = 2 \\ 2 embeddings, both applied to V
? vall[1] \\ the first
%10 = [0, 1, -1.2071... - 2.0907...*I, 0.2071... - 0.3587...*I, ...]
? vall[2] \\ and the second one
%11 = [0, 1, 0.2071... + 0.3587...*I, -1.2071... + 2.0907...*I, ...]

? vE = mfembed(F); #vE \\ same 2 embeddings
%12 = 2
? mfembed(vE[1], V) \\ apply first embedding to V
%13 = [0, 1, -1.2071... - 2.0907...*I, 0.2071... - 0.3587...*I, ...]

```

For convenience, we also allow a modular form space from `mfininit` instead of `f`, corresponding to the single embedding of $\mathbf{Q}(\chi)$.

```

? [mfB,MC,C] = mfatkininit(mf,7); MC \\ coefs in  $\mathbf{Q}(\chi)$ 
%13 =
[      Mod(2/7*t, t^2 + t + 1) Mod(-1/7*t - 2/7, t^2 + t + 1)]
[Mod(-1/7*t - 2/7, t^2 + t + 1)      Mod(2/7*t, t^2 + t + 1)]

? C \\ normalizing constant
%14 = 0.33863... - 0.16787*I
? M = mfembed(mf, MC) / C \\ the true matrix for the action of  $w_7$ 
[-0.6294... + 0.4186...*I -0.3625... - 0.5450...*I]
[-0.3625... - 0.5450...*I -0.6294... + 0.4186...*I]

? exponent(M*conj(M) - 1) \\ M * conj(M) is close to 1
%16 = -126

```

The library syntax is `GEN mfembed0(GEN f, GEN v = NULL, long prec)`.

3.18.31 mfeval(*mf*, *F*, *vtau*). Computes the numerical value of the modular form *F*, belonging to *mf*, at the complex number *vtau* or the vector *vtau* of complex numbers in the completed upper-half plane. The result is given with absolute error less than 2^{-B} , where $B = \text{realbitprecision}$.

If the field of definition $\mathbf{Q}(F)$ is larger than $\mathbf{Q}(\chi)$ then *F* may be embedded into $\mathbf{C}$ in $d = [\mathbf{Q}(F) : \mathbf{Q}(\chi)]$ ways, in which case a vector of the *d* results is returned.

```
? mf = mfinit([11,2],0); F = mfbasis(mf)[1]; mfparams(F)
%1 = [11, 2, 1, y, t-1] \\ Q(F) = Q(chi) = Q
? mfeval(mf,F,I/2)
%2 = 0.039405471130100890402470386372028382117
? mf = mfinit([35,2],0); F = mfeigenbasis(mf)[2]; mfparams(F)
%3 = [35, 2, 1, y^2 - y - 4, t - 1] \\ [Q(F) : Q(chi)] = 2
? mfeval(mf,F,I/2)
%4 = [0.045..., 0.0385...] \\ sigma_1(F) and sigma_2(F) at I/2
? mf = mfinit([12,4],1); F = mfbasis(mf)[1];
? mfeval(mf, F, 0.318+10^(-7)*I)
%6 = 3.379... E-21 + 6.531... E-21*I \\ instantaneous !
```

In order to maximize the imaginary part of the argument, the function computes $(f|_k \gamma)(\gamma^{-1} \cdot \tau)$ for a suitable γ not necessarily in $\Gamma_0(N)$ (in which case $f|_k \gamma$ is evaluated using `mfslashexpansion`).

```
? T = mfTheta(); mf = mfinit(T); mfeval(mf,T,[0,1/2,1,oo])
%1 = [1/2 - 1/2*I, 0, 1/2 - 1/2*I, 1]
```

The library syntax is `GEN mfeval(GEN mf, GEN F, GEN vtau, long bitprec)`.

3.18.32 mffields(*mf*). Given *mf* as output by `mfinit` with parameters (N, k, χ) , returns the vector of polynomials defining each Galois orbit of newforms over $\mathbf{Q}(\chi)$.

```
? mf = mfinit([35,2],0); mffields(mf)
%1 = [y, y^2 - y - 4]
```

Here the character is trivial so $\mathbf{Q}(\chi) = \mathbf{Q}$ and there are 3 newforms: one is rational (corresponding to *y*), the other two are conjugate and defined over the quadratic field $\mathbf{Q}[y]/(y^2 - y - 4)$.

```
? [G,chi] = znchar(Mod(3,35));
? zncharconductor(G,chi)
%2 = 35
? charorder(G,chi)
%3 = 12
? mf = mfinit([35, 2, [G,chi]],0); mffields(mf)
%4 = [y, y]
```

Here the character is primitive of order 12 and the two newforms are defined over $\mathbf{Q}(\chi) = \mathbf{Q}(\zeta_{12})$.

```
? mf = mfinit([35, 2, Mod(13,35)],0); mffields(mf)
%3 = [y^2 + Mod(5*t, t^2 + 1)]
```

This time the character has order 4 and there are two conjugate newforms over $\mathbf{Q}(\chi) = \mathbf{Q}(i)$.

The library syntax is `GEN mffields(GEN mf)`.

3.18.33 mffromell(E). E being an elliptic curve defined over Q given by an integral model in `ellinit` format, computes a 3-component vector `[mf,F,v]`, where F is the newform corresponding to E by modularity, `mf` is the newspace to which F belongs, and `v` gives the coefficients of F on `mfbasis(mf)`.

```
? E = ellinit("26a1");
? [mf,F,co] = mffromell(E);
? co
%2 = [3/4, 1/4]~
? mfcoefs(F, 5)
%3 = [0, 1, -1, 1, 1, -3]
? ellan(E, 5)
%4 = [1, -1, 1, 1, -3]
```

The library syntax is `GEN mffromell(GEN E)`.

3.18.34 mffrometaquo(η , { $flag = 0$ }). Modular form corresponding to the eta quotient matrix `eta`. If the valuation v at infinity is fractional, returns 0. If the eta quotient is not holomorphic but simply meromorphic, returns 0 if $flag = 0$; returns the eta quotient (divided by q to the power $-v$ if $v < 0$, i.e., with valuation 0) if $flag$ is set.

```
? mffrometaquo(Mat([1,1]),1)
%1 = 0
? mfcoefs(mffrometaquo(Mat([1,24])),6)
%2 = [0, 1, -24, 252, -1472, 4830, -6048]
? mfcoefs(mffrometaquo([1,1;23,1]),10)
%3 = [0, 1, -1, -1, 0, 0, 1, 0, 1, 0, 0]
? F = mffrometaquo([1,2;2,-1]); mfparams(F)
%4 = [16, 1/2, 1, y, t - 1]
? mfcoefs(F,10)
%5 = [1, -2, 0, 0, 2, 0, 0, 0, 0, -2, 0]
? mffrometaquo(Mat([1,-24]))
%6 = 0
? f = mffrometaquo(Mat([1,-24]),1); mfcoefs(f,6)
%7 = [1, 24, 324, 3200, 25650, 176256, 1073720]
```

For convenience, a `t_VEC` is also accepted instead of a factorization matrix with a single row:

```
? f = mffrometaquo([1,24]); \\ also valid
```

The library syntax is `GEN mffrometaquo(GEN eta, long flag)`.

3.18.35 mffromlfun(L). Let L being an L -function in any of the `lfun` formats representing a self-dual modular form (for instance an eigenform). Return `[NK,space,v]` when `mf = mfini(NK,space)` is the modular form space containing the form and `mftobasis(mf, v)` will represent it on the space basis. If L has rational coefficients, this will be enough to recognize the modular form in mf :

```
? L = lfuncreate(x^2+1);
? lfunan(L,10)
%2 = [1, 1, 0, 1, 2, 0, 0, 1, 1, 2]
? [NK,space,v] = mffromlfun(L); NK
%4 = [4, 1, -4]
? mf=mfinit(NK,space); w = mftobasis(mf,v)
%5 = [1.0000000000000000000000000000000000000000000000000000000]~
? [f] = mfbasis(mf); mfcoefs(f,10)    \\ includes a_0 !
%6 = [1/4, 1, 1, 0, 1, 2, 0, 0, 1, 1, 2]
```

If L has inexact complex coefficients, one can for instance compute an eigenbasis for mf and check whether one of the attached L -function is reasonably close to L . In the example, we cheat by producing the L function from an eigenform in a known space, but the function does not use this information:

```
? mf = mfini([32,6,Mod(5,32)],0);
? [poldegree(K) | K<-mffields(mf)]
%2 = [19] \\ one orbit, [Q(F) : Q(chi)] = 19
? L = lfunmf(mf)[1][1]; \\ one of the 19 L-functions attached to F
? lfunan(L,3)
%4 = [1, 5.654... - 0.1812...*I, -7.876... - 19.02...*I]
? [NK,space,v] = mfffromlfun(L); NK
%5 = [32, 6, Mod(5, 32)]
? vL = concat(lfunmf(mf)); \\ L functions for all cuspidal eigenforms
? an = lfunan(L,10);
? for (i = 1, #vL, if (normlp(lfunan(vL[i],10) - an, oo) < 1e-10, print(i)));
1
```

The library syntax is `GEN mffromlfun(GEN L, long prec)`.

3.18.36 mffromqf($Q, \{P\}$). Q being an even integral positive definite quadratic form and P a homogeneous spherical polynomial for Q , computes a 3-component vector $[mf, F, v]$, where F is the theta function corresponding to (Q, P) , mf is the corresponding space of modular forms (from **mfninit**), and v gives the coefficients of F on **mbasis(mf)**.

```
? [mf,F,v] = mffromqf(2*matid(10)); v
%1 = [64/5, 4/5, 32/5]~
? mfcoefs(F, 5)
%2 = [1, 20, 180, 960, 3380, 8424]
? mfcoef(F, 10000) \\ number of ways of writing 10000 as sum of 10 squares
%3 = 128205250571893636
? mfcoefs(F, 10000); \\ fast !
time = 220ms
? [mf,F,v] = mffromqf([2,0;0,2],x^4-6*x^2*y^2+y^4);
? mfcoefs(F,10)
```

```
%6 = [0, 4, -16, 0, 64, -56, 0, 0, -256, 324, 224]
? mfcoef(F,100000)  \\ instantaneous
%7 = 41304367104
```

Odd dimensions are supported, corresponding to forms of half-integral weight:

```
? [mf,F,v] = mffromqf(2*matid(3));
? mfisequal(F, mfpow(mfTheta(),3))
%2 = 1
? mfcoefs(F, 32)  \\ illustrate Legendre's 3-square theorem
%3 = [ 1,
      6, 12,  8, 6, 24, 24, 0, 12,
      30, 24, 24, 8, 24, 48, 0, 6,
      48, 36, 24,24, 48, 24, 0, 24,
      30, 72, 32, 0, 72, 48, 0, 12]
```

The library syntax is GEN mffromqf(GEN Q, GEN P = NULL).

3.18.37 mfgaloisprojrep(*mf*, *F*). *mf* being an mf output by mfinit in weight 1, return a polynomial defining the field fixed by the kernel of the projective Artin representation attached to *F* (by Deligne–Serre). Currently only implemented for projective images A_4 , A_5 and S_4 . The type A_5 requires the nflistdata package to be installed.

```
\\ A4 example
? mf = mfinit([4*31,1,Mod(87,124)],0);
? F = mfeigenbasis(mf)[1];
? mfgaloistype(mf,F)
%3 = -12
? pol = mfgaloisprojrep(mf,F)
%4 = x^12 + 68*x^10 + 4808*x^8 + ... + 4096
? G = galoisinit(pol); galoisidentify(G)
%5 = [12,3]  \\A4
? pol4 = polredbest(galoisfixedfield(G,G.gen[3], 1))
%6 = x^4 + 7*x^2 - 2*x + 14
? polgalois(pol4)
%7 = [12, 1, 1, "A4"]
? factor(nfdisc(pol4))
%8 =
[ 2 4]
[31 2]

\\ S4 example
? mf = mfinit([4*37,1,Mod(105,148)],0);
? F = mfeigenbasis(mf)[1];
? mfgaloistype(mf,F)
%11 = -24
? pol = mfgaloisprojrep(mf,F)
%12 = x^24 + 24*x^22 + 256*x^20 + ... + 255488256
? G = galoisinit(pol); galoisidentify(G)
%13 = [24, 12]  \\S4
? pol4 = polredbest(galoisfixedfield(G,G.gen[3..4], 1))
```

```

%14 = x^4 - x^3 + 5*x^2 - 7*x + 12
? polgalois(pol4)
%15 = [24, -1, 1, "S4"]
? factor(nfdisc(pol4))
%16 =
[ 2 2]
[37 3]

```

The library syntax is `GEN mfgaloisprojrep(GEN mf, GEN F, long prec)`.

3.18.38 mfgaloistype($NK, \{F\}$). NK being either $[N, 1, \text{CHI}]$ or an `mf` output by `mfininit` in weight 1, gives the vector of types of Galois representations attached to each cuspidal eigenform, unless the modular form F is specified, in which case only for F (note that it is not tested whether F belongs to the correct modular form space, nor whether it is a cuspidal eigenform). Types A_4 , S_4 , A_5 are represented by minus their cardinality -12 , -24 , or -60 , and type D_n is represented by its cardinality, the integer $2n$:

```

? mfgaloistype([124,1, Mod(67,124)]) \\ A4
%1 = [-12]
? mfgaloistype([148,1, Mod(105,148)]) \\ S4
%2 = [-24]
? mfgaloistype([633,1, Mod(71,633)]) \\ D10, A5
%3 = [10, -60]
? mfgaloistype([239,1, -239]) \\ D6, D10, D30
%4 = [6, 10, 30]
? mfgaloistype([71,1, -71])
%5 = [14]
? mf = mfininit([239,1, -239],0); F = mfeigenbasis(mf)[2];
? mfgaloistype(mf, F)
%7 = 10

```

The function may also return 0 as a type when it failed to determine it; in this case the correct type is either -12 or -60 , and most likely -12 .

The library syntax is `GEN mfgaloistype(GEN NK, GEN F = NULL)`.

3.18.39 mfhecke(mf, F, n). F being a modular form in modular form space mf , returns $T(n)F$, where $T(n)$ is the n -th Hecke operator.

Warning. If F is of level $M < N$, then $T(n)F$ is in general not the same in $M_k(\Gamma_0(M), \chi)$ and in $M_k(\Gamma_0(N), \chi)$. We take $T(n)$ at the same level as the one used in `mf`.

```
? mf = mfinit([26,2],0); F = mfbasis(mf)[1]; mftobasis(mf,F)
%1 = [1, 0]~
? G2 = mfhecke(mf,F,2); mftobasis(mf,G2)
%2 = [0, 1]~
? G5 = mfhecke(mf,F,5); mftobasis(mf,G5)
%3 = [-2, 1]~
```

Modular forms of half-integral weight are supported, in which case n must be a perfect square, else T_n will act as 0 (the operator T_p for $p \mid N$ is not supported yet):

```
? F = mfpow(mfTheta(),3); mf = mfinit(F);
? mfisequal(mfhecke(mf,F,9), mflinear([F],[4]))
%2 = 1
```

(F is an eigenvector of all T_{p^2} , with eigenvalue $p + 1$ for odd p .)

Warning. When n is a large composite, resp. the square of a large composite in half-integral weight, it is in general more efficient to use `mfheckemat` on the `mftobasis` coefficients:

```
? mfcoefs(mfhecke(mf,F,3^10), 10)
time = 917 ms.
%3 = [324, 1944, 3888, 2592, 1944, 7776, 7776, 0, 3888, 9720, 7776]
? M = mfheckemat(mf,3^10) \\ instantaneous
%4 =
[324]
? G = mflinear(mf, M*mftobasis(mf,F));
? mfcoefs(G, 10) \\ instantaneous
%6 = [324, 1944, 3888, 2592, 1944, 7776, 7776, 0, 3888, 9720, 7776]
```

The library syntax is `GEN mfhecke(GEN mf, GEN F, long n)`.

3.18.40 mfheckemat($mf, vecn$). If `vecn` is an integer, matrix of the Hecke operator $T(n)$ on the basis formed by `mfbasis(mf)`. If it is a vector, vector of such matrices, usually faster than calling each one individually.

```
? mf=mfinit([32,4],0); mfheckemat(mf,3)
%1 =
[0 44 0]
[1 0 -10]
[0 -2 0]
? mfheckemat(mf,[5,7])
%2 = [[0, 0, 220; 0, -10, 0; 1, 0, 12], [0, 88, 0; 2, 0, -20; 0, -4, 0]]
```

The library syntax is `GEN mfheckemat(GEN mf, GEN vecn)`.

3.18.41 mfininit($NK, \{space = 4\}$). Create the space of modular forms corresponding to the data contained in NK and **space**. NK is a vector which can be either $[N, k]$ (N level, k weight) corresponding to a subspace of $M_k(\Gamma_0(N))$, or $[N, k, CHI]$ (CHI a character) corresponding to a subspace of $M_k(\Gamma_0(N), \chi)$. Alternatively, it can be a modular form F or modular form space, in which case we use **mfparams** to define the space parameters.

The subspace is described by the small integer **space**: 0 for the newspace $S_k^{\text{new}}(\Gamma_0(N), \chi)$, 1 for the cuspidal space S_k , 2 for the oldspace S_k^{old} , 3 for the space of Eisenstein series E_k and 4 for the full space M_k .

Wildcards. For given level and weight, it is advantageous to compute simultaneously spaces attached to different Galois orbits of characters, especially in weight 1. The parameter CHI may be set to 0 (wildcard), in which case we return a vector of all **mfininit**(s) of non trivial spaces in $S_k(\Gamma_1(N))$, one for each Galois orbit (see **znchargalois**). One may also set CHI to a vector of characters and we return a vector of all mfinits of subspaces of $M_k(G_0(N), \chi)$ for χ in the list, in the same order. In weight 1, only S_1^{new} , S_1 and E_1 support wildcards.

The output is a technical structure S , or a vector of structures if CHI was a wildcard, which contains the following information: $[N, k, \chi]$ is given by **mfparams**(S), the space dimension is **mfdim**(S) and a **C**-basis for the space is **mfbasis**(S). The structure is entirely algebraic and does not depend on the current **realbitprecision**.

```
? S = mfininit([36,2], 0); \\ new space
? mfdim(S)
%2 = 1
? mfparams
%3 = [36, 2, 1, y] \\ trivial character
? f = mfbasis(S)[1]; mfcoefs(f,10)
%4 = [0, 1, 0, 0, 0, 0, 0, 0, -4, 0, 0, 0]
? vS = mfininit([36,2,0],0); \\ with wildcard
? #vS
%6 = 4 \\ 4 non trivial spaces (mod Galois action)
? apply(mfdim,vS)
%7 = [1, 2, 1, 4]
? mfdim([36,2,0], 0)
%8 = [[1, Mod(1, 36), 1, 0], [2, Mod(35, 36), 2, 0], [3, Mod(13, 36), 1, 0],
      [6, Mod(11, 36), 4, 0]]
```

The library syntax is **GEN mfininit**(**GEN** NK , **long** $space$).

3.18.42 mfisCM(F). Tests whether the eigenform F is a CM form. The answer is 0 if it is not, and if it is, either the unique negative discriminant of the CM field, or the pair of two negative discriminants of CM fields, this latter case occurring only in weight 1 when the projective image is $D_2 = C_2 \times C_2$, i.e., coded 4 by **mfgaloistype**.

```
? F = mffromell(ellinit([0,1]))[2]; mfisCM(F)
%1 = -3
? mf = mfininit([39,1,-39],0); F=mfeigenbasis(mf)[1]; mfisCM(F)
%2 = Vecsmall([-3, -39])
? mfgaloistype(mf)
%3 = [4]
```

The library syntax is **GEN mfisCM**(**GEN** F).

3.18.43 mfisequal($F, G, \{lim = 0\}$). Checks whether the modular forms F and G are equal. If lim is nonzero, only check equality of the first $lim + 1$ Fourier coefficients and the function then also applies to generalized modular forms.

```
? D = mfDelta(); F = mfderiv(D);
? G = mfmul(mfEk(2), D);
? mfisequal(F, G)
%2 = 1
```

The library syntax is `long mfisequal(GEN F, GEN G, long lim)`.

3.18.44 mfisetaquo($f, \{flag = 0\}$). If the generalized modular form f is a holomorphic eta quotient, return the eta quotient matrix, else return 0. If *flag* is set, also accept meromorphic eta quotients: check whether $f = q^{-v(g)}g(q)$ for some eta quotient g ; if so, return the eta quotient matrix attached to g , else return 0. See `mffrometaquo`.

```
? mfisetaquo(mfDelta())
%1 =
[1 24]
? f = mffrometaquo([1,1;23,1]);
? mfisetaquo(f)
%3 =
[ 1 1]
[23 1]
? f = mffrometaquo([1,-24], 1);
? mfisetaquo(f) \\ nonholomorphic
%5 = 0
? mfisetaquo(f,1)
%6 =
[1 -24]
```

The library syntax is `GEN mfisetaquo(GEN f, long flag)`.

3.18.45 mfkohnenbasis(mf). mf being a cuspidal space of half-integral weight $k \geq 3/2$ with level N and character χ , gives a basis B of the Kohnen $+$ -space of mf as a matrix whose columns are the coefficients of B on the basis of mf . The conductor of either χ or $\chi \cdot (-4/.)$ must divide $N/4$.

```
? mf = mfinit([36,5/2],1); K = mfkohnenbasis(mf); K~
%1 =
[-1 0 0 2 0 0]
[ 0 0 0 0 1 0]
? (mfcoefs(mf,20) * K)~
%4 =
[0 -1 0 0 2 0 0 0 0 0 0 0 0 -6 0 0 8 0 0 0 0]
[0 0 0 0 0 1 0 0 -2 0 0 0 0 0 0 0 0 1 0 0 2]
? mf = mfinit([40,3/2,8],1); mfkohnenbasis(mf)
*** at top-level: mfkohnenbasis(mf)
*** ^-----
*** mfkohnenbasis: incorrect type in mfkohnenbasis [incorrect CHI] (t_VEC).
```

In the final example both $\chi = (8/.)$ and $\chi \cdot (-4/.)$ have conductor 8, which does not divide $N/4 = 10$.

The library syntax is `GEN mfkohnenbasis(GEN mf)`.

3.18.46 mfkohnenbijection(mf). Let `mf` be a cuspidal space of half-integral weight and weight $4N$, with N squarefree and let $S_k^+(\Gamma_0(4N), \chi)$ be the Kohnen $+$ -space. Returns `[mf2,M,K,shi]`, where

- `mf2` gives the cuspidal space $S_{2k-1}(\Gamma_0(N), \chi^2)$;
- `M` is a matrix giving a Hecke-module isomorphism from that space to the Kohnen $+$ -space $S_k^+(\Gamma_0(4N), \chi)$;
- `K` represents a basis B of the Kohnen $+$ -space as a matrix whose columns are the coefficients of B on the basis of `mf`;
- `shi` is a vector of pairs (t_i, n_i) gives the linear combination of Shimura lifts giving M^{-1} : t_i is a squarefree positive integer and n_i is a small nonzero integer.

```
? mf=mfinit([60,5/2],1); [mf2,M,K,shi]=mfkohnenbijection(mf); M
%2 =
[-3      0 5/2 7/2]
[ 1 -1/2  -7  -7]
[ 1  1/2   0  -3]
[ 0      0 5/2 5/2]
? shi
%2 = [[1, 1], [2, 1]]
```

This last command shows that the map giving the bijection is the sum of the Shimura lift with $t = 1$ and the one with $t = 2$.

Since it gives a bijection of Hecke modules, this matrix can be used to transport modular form data from the easily computed space of level N and weight $2k - 1$ to the more difficult space of level $4N$ and weight k : matrices of Hecke operators, new space, splitting into eigenspaces and eigenforms. Examples:

```
? K^(-1)*mfheckemat(mf,121)*K /* matrix of T_11^2 on K. Slowish. */
time = 1,280 ms.
%1 =
[ 48  24  24  24]
[  0  32   0 -20]
[-48 -72 -40 -72]
[  0   0   0  52]
? M*mfheckemat(mf2,11)*M^(-1) /* instantaneous via T_11 on S_{2k-1} */
time = 0 ms.
%2 =
[ 48  24  24  24]
[  0  32   0 -20]
[-48 -72 -40 -72]
```

```

[ 0 0 0 52]
? mf20=mfinit(mf2,0); [mftobasis(mf2,b) | b<-mfbasis(mf20)]
%3 = [[0, 0, 1, 0]~, [0, 0, 0, 1]~]
? F1=M*[0,0,1,0]~
%4 = [1/2, 1/2, -3/2, -1/2]~
? F2=M*[0,0,0,1]~
%5 = [3/2, 1/2, -9/2, -1/2]
? K*F1
%6 = [1, 0, 0, 1, 1, 0, 0, 1, -3, 0, 0, -3, 0, 0]~
? K*F2
%7 = [3, 0, 0, 3, 1, 0, 0, 1, -9, 0, 0, -3, 0, 0]~

```

This gives a basis of the new space of $S_{5/2}^+(\Gamma_0(60))$ expressed on the initial basis of $S_{5/2}(\Gamma_0(60))$. To obtain the eigenforms, we write instead:

```

? BE=mfeigenbasis(mf20); [E1,E2]=apply(x->K*M*mftobasis(mf2,x),BE)
%1 = [[1, 0, 0, 1, 0, 0, 0, 0, -3, 0, 0, 0, 0, 0]~, \
      [0, 0, 0, 0, 1, 0, 0, 1, 0, 0, 0, -3, 0, 0]~]
? EI1 = mflinear(mf, E1); EI2=mflinear(mf, E2);

```

These are the two eigenfunctions in the space **mf**, the first (resp., second) will have Shimura image a multiple of $BE[1]$ (resp., $BE[2]$). The function **mfkohneneigenbasis** does this directly.

The library syntax is **GEN mfkohnenbijection(GEN mf)**.

3.18.47 mfkohneneigenbasis(mf, bij). **mf** being a cuspidal space of half-integral weight $k \geq 3/2$ and **bij** being the output of **mfkohnenbijection(mf)**, outputs a 3-component vector **[mf0,BNEW,BEIGEN]**, where **BNEW** and **BEIGEN** are two matrices whose columns are the coefficients of a basis of the Kohnen new space and of the eigenforms on the basis of **mf** respectively, and **mf0** is the corresponding new space of integral weight $2k - 1$.

```

? mf=mfinit([44,5/2],1);bij=mfkohnenbijection(mf);
? [mf0,BN,BE]=mfkohneneigenbasis(mf,bij);
? BN~
%2 =
[2 0 0 -2 2 0 -8]
[2 0 0 4 14 0 -32]
? BE~
%3 = [1 0 0 Mod(y-1, y^2-3) Mod(2*y+1, y^2-3) 0 Mod(-4*y-4, y^2-3)]
? lift(mfcoefs(mf,20)*BE[,1])
%4 = [0, 1, 0, 0, y - 1, 2*y + 1, 0, 0, 0, -4*y - 4, 0, 0, \
      -5*y + 3, 0, 0, 0, -6, 0, 0, 0, 7*y + 9]~

```

The library syntax is **GEN mfkohneneigenbasis(GEN mf, GEN bij)**.

3.18.48 mflinear(vF, v). **vF** being a vector of generalized modular forms and **v** a vector of coefficients of same length, compute the linear combination of the entries of **vF** with coefficients **v**.

Note. Use this in particular to subtract two forms F and G (with $vF = [F, G]$ and $v = [1, -1]$), or to multiply an form by a scalar λ (with $vF = [F]$ and $v = [\lambda]$).

```
? D = mfDelta(); G = mflinear([D], [-3]);
? mfcoefs(G, 4)
%2 = [0, -3, 72, -756, 4416]
```

For user convenience, we allow

- a modular form space $\mathbf{mf}$ as a $\mathbf{vF}$ argument, which is understood as $\mathbf{mfbasis}(\mathbf{mf})$;
- in this case, we also allow a modular form f as v , which is understood as $\mathbf{mftobasis}(\mathbf{mf}, f)$.

```
? T = mfpow(mfTheta(), 7); F = mfShimura(T, -3); \\ Shimura lift for D=-3
? mfcoefs(F, 8)
%2 = [-5/9, 280, 9240, 68320, 295960, 875280, 2254560, 4706240, 9471000]
? mf = mfinit(F); G = mflinear(mf, F);
? mfcoefs(G, 8)
%4 = [-5/9, 280, 9240, 68320, 295960, 875280, 2254560, 4706240, 9471000]
```

This last construction allows to replace a general modular form by a simpler linear combination of basis functions, which is often more efficient:

```
? T10=mfpow(mfTheta(), 10); mfcoef(T10, 10^4) \\ direct evaluation
time = 399 ms.
%5 = 128205250571893636
? mf=mfinit(T10); F=mflinear(mf, T10); \\ instantaneous
? mfcoef(F, 10^4) \\ after linearization
time = 67 ms.
%7 = 128205250571893636
```

The library syntax is `GEN mflinear(GEN vF, GEN v)`.

3.18.49 mfmanin(FS). Given the modular symbol FS associated to an eigenform F by `mfsymbol(mf, F)`, computes the odd and even special polynomials as well as the odd and even periods ω^- and ω^+ as a vector $[[P^-, P^+], [\omega^-, \omega^+, r]]$, where $r = \Im(\omega^+ \overline{\omega^-}) / \langle F, F \rangle$. If F has several embeddings into $\mathbf{C}$, give the vector of results corresponding to each embedding.

```
? D=mfDelta(); mf=mfinit(D); DS=mfsymbol(mf, D);
? [pols, oms]=mfmanin(DS); pols
%2 = [[4*x^9 - 25*x^7 + 42*x^5 - 25*x^3 + 4*x], \
      [-36*x^10 + 691*x^8 - 2073*x^6 + 2073*x^4 - 691*x^2 + 36]]
? oms
%3 = [0.018538552324740326472516069364750571812, \
      -0.00033105361053212432521308691198949874026*I, 4096/691]
? mf=mfinit([11, 2], 0); F=mfeigenbasis(mf)[1]; FS=mfsymbol(mf, F);
? [pols, oms]=mfmanin(FS); pols
%5 = [[0, 0, 0, 1, 1, 0, 0, -1, -1, 0, 0, 0], \
      [2, 0, 10, 5, -5, -10, -10, -5, 5, 10, 0, -2]]
? oms[3]
%6 = 24/5
```

The library syntax is `GEN mfmanin(GEN FS, long bitprec)`.

3.18.50 mfmul(F, G). Multiply the two generalized modular forms F and G .

```
? E4 = mfeK(4); G = mfmul(mfmul(E4,E4),E4);
? mfcoefs(G, 4)
%2 = [1, 720, 179280, 16954560, 396974160]
? mfcoefs(mfpow(E4,3), 4)
%3 = [1, 720, 179280, 16954560, 396974160]
```

The library syntax is GEN mfmul(GEN F, GEN G).

3.18.51 mfnumcusps(N). Number of cusps of $\Gamma_0(N)$

```
? mfnumcusps(24)
%1 = 8
? mfcusps(24)
%1 = [0, 1/2, 1/3, 1/4, 1/6, 1/8, 1/12, 1/24]
```

The library syntax is GEN mfnumcusps(GEN N).

3.18.52 mfparams(F). If F is a modular form space, returns $[N, k, \text{CHI}, \text{space}, \Phi]$, level, weight, character χ , and space code; where Φ is the cyclotomic polynomial defining the field of values of CHI. If F is a generalized modular form, returns $[N, k, \text{CHI}, P, \Phi]$, where P is the (polynomial giving the) field of definition of F as a relative extension of the cyclotomic field $\mathbf{Q}(\chi) = \mathbf{Q}[t]/(\Phi)$: in that case the level N may be a multiple of the level of F and the polynomial P may define a larger field than $\mathbf{Q}(F)$. If you want the true level of F from this result, use `mfconductor(mfinit(F), F)`. The polynomial P defines an extension of $\mathbf{Q}(\chi) = \mathbf{Q}[t]/(\Phi(t))$; it has coefficients in that number field (polmods in t).

In contrast with `mfparams(F)[4]` which always gives the polynomial P defining the relative extension $\mathbf{Q}(F)/\mathbf{Q}(\chi)$, the member function `F.mod` returns the polynomial used to define $\mathbf{Q}(F)$ over $\mathbf{Q}$ (either a cyclotomic polynomial or a polynomial with cyclotomic coefficients).

```
? E1 = mfeisenstein(4,-3,-4); E2 = mfeisenstein(3,5,-7); E3 = mfmul(E1,E2);
? apply(mfparams, [E1,E2,E3])
%2 = [[12, 4, 12, y, t-1], [35, 3, -35, y, t-1], [420, 7, -420, y, t-1]]
? mf = mfinit([36,2,Mod(13,36)],0); [f] = mfeigenbasis(mf); mfparams(mf)
%3 = [36, 2, Mod(13, 36), 0, t^2 + t + 1]
? mfparams(f)
%4 = [36, 2, Mod(13, 36), y, t^2 + t + 1]
? f.mod
%5 = t^2 + t + 1
? mf = mfinit([36,4,Mod(13,36)],0); [f] = mfeigenbasis(mf);
? lift(mfparams(f))
%7 = [36, 4, 13, y^3 + (2*t-2)*y^2 + (-4*t+6)*y + (10*t-1), t^2+t+1]
```

The library syntax is GEN mfparams(GEN F).

3.18.53 mfperiodpol(*mf*, *f*, {*flag* = 0}). Period polynomial of the cuspidal part of the form *f*, in other words $\int_0^\infty (X - \tau)^{k-2} f(\tau) d\tau$. If *flag* = 0, ordinary period polynomial. If it is 1 or -1, even or odd part of that polynomial. *f* can also be the modular symbol output by **mfsymbol**(*mf*, *f*).

```
? D = mfDelta(); mf = mfinit(D,0);
? PP = mfperiodpol(mf, D, -1); PP/=polcoef(PP, 1); bestappr(PP)
%1 = x^9 - 25/4*x^7 + 21/2*x^5 - 25/4*x^3 + x
? PM = mfperiodpol(mf, D, 1); PM/=polcoef(PM, 0); bestappr(PM)
%2 = -x^10 + 691/36*x^8 - 691/12*x^6 + 691/12*x^4 - 691/36*x^2 + 1
```

The library syntax is GEN mfperiodpol(GEN mf, GEN f, long flag, long bitprec).

3.18.54 mfperiodpolbasis(*k*, {*flag* = 0}). Basis of period polynomials for weight *k*. If *flag* = 1 or -1, basis of odd or even period polynomials.

```
? mfperiodpolbasis(12,1)
%1 = [x^8 - 3*x^6 + 3*x^4 - x^2, x^10 - 1]
? mfperiodpolbasis(12,-1)
%2 = [4*x^9 - 25*x^7 + 42*x^5 - 25*x^3 + 4*x]
```

The library syntax is GEN mfperiodpolbasis(long k, long flag).

3.18.55 mfpetersson(*fs*, {*gs*}). Petersson scalar product of the modular forms *f* and *g* belonging to the same modular form space *mf*, given by the corresponding “modular symbols” *fs* and *gs* output by **mfsymbol** (also in weight 1 and half-integral weight, where symbols do not exist). If *gs* is omitted it is understood to be equal to *fs*. The scalar product is normalized by the factor $1/[\Gamma : \Gamma_0(N)]$. Note that *f* and *g* can both be noncuspidal, in which case the program returns an error if the product is divergent. If the fields of definition $\mathbf{Q}(f)$ and $\mathbf{Q}(g)$ are equal to $\mathbf{Q}(\chi)$ the result is a scalar. If $[\mathbf{Q}(f) : \mathbf{Q}(\chi)] = d > 1$ and $[\mathbf{Q}(g) : \mathbf{Q}(\chi)] = e > 1$ the result is a $d \times e$ matrix corresponding to all the embeddings of *f* and *g*. In the intermediate cases $d = 1$ or $e = 1$ the result is a row or column vector.

```
? D=mfDelta(); mf=mfinit(D); DS=mfsymbol(mf,D); mfpetersson(DS)
%1 = 1.0353620568043209223478168122251645932 E-6
? mf=mfinit([11,6],0); B=mfeigenbasis(mf); BS=vector(#B,i,mfsymbol(mf,B[i]));
? mfpetersson(BS[1])
%3 = 1.6190120685220988139111708455305245466 E-5
? mfpetersson(BS[1],BS[2])
%4 = [-3.826479006582967148 E-42 - 2.801547395385577002 E-41*I,\
      1.6661127341163336125 E-41 + 1.1734725972345985061 E-41*I,\
      0.E-42 - 6.352626992842664490 E-41*I]~
? mfpetersson(BS[2])
%5 =
[ 2.7576133733... E-5 2.0... E-42 6.3... E-43 ]
[ -4.1... E-42 6.77837030070... E-5 3.3...E-42 ]
[ -6.32...E-43 3.6... E-42 2.27268958069... E-5]
? mf=mfinit([23,2],0); F=mfeigenbasis(mf)[1]; FS=mfsymbol(mf,F);
? mfpetersson(FS)
%5 =
[0.0039488965740025031688548076498662860143 -3.56 ... E-40]
```

```
[ -3.5... E-40  0.0056442542987647835101583821368582485396]
```

Noncuspidal example:

```
? E1=mfeisenstein(5,1,-3);E2=mfeisenstein(5,-3,1);
? mf=mfinit([12,5,-3]); cusps=mfcusps(12);
? apply(x->mfcuspval(mf,E1,x),cusps)
%3 = [0, 0, 1, 0, 1, 1]
? apply(x->mfcuspval(mf,E2,x),cusps)
%4 = [1/3, 1/3, 0, 1/3, 0, 0]
? E1S=mfsymbol(mf,E1);E2S=mfsymbol(mf,E2);
? mfpetersson(E1S,E2S)
%6 = -1.884821671646... E-5 - 1.9... E-43*I
```

Weight 1 and 1/2-integral weight example:

```
? mf=mfinit([23,1,-23],1);F=mfbasis(mf)[1];FS=mfsymbol(mf,F);
? mfpetersson(mf,FS)
%2 = 0.035149946790370230814006345508484787443
? mf=mfinit([4,9/2],1);F=mfbasis(mf)[1];FS=mfsymbol(mf,F);
? mfpetersson(FS)
%4 = 0.00015577084407139192774373662467908966030
```

The library syntax is GEN mfpetersson(GEN fs, GEN gs = NULL).

3.18.56 mfpow(F, n). Compute F^n , where n is an integer and F is a generalized modular form:

```
? G = mfpow(mfEk(4), 3); \\ E4^3
? mfcoefs(G, 4)
%2 = [1, 720, 179280, 16954560, 396974160]
```

The library syntax is GEN mfpow(GEN F, long n).

3.18.57 mfsearch($NK, V, \{space\}$). NK being of the form $[N, k]$ with k possibly half-integral, search for a modular form with rational coefficients, of weight k and level N , whose initial coefficients $a(0), \dots$ are equal to V ; $space$ specifies the modular form spaces in which to search, in `mfinit` or `mfdim` notation. The output is a list of matching forms with that given level and weight. Note that the character is of the form $(D/.)$, where D is a (positive or negative) fundamental discriminant dividing N . The forms are sorted by increasing $|D|$.

The parameter N can be replaced by a vector of allowed levels, in which case the list of forms is sorted by increasing level, then increasing $|D|$. If a form is found at level N , any multiple of N with the same D is not considered. Some useful possibilities are

- $[N_1..N_2]$: all levels between N_1 and N_2 , endpoints included;
- $F * [N_1..N_2]$: same but levels divisible by F ;
- `divisors( $N_0$ )`: all levels dividing N_0 .

Note that this is different from `mfeigensearch`, which only searches for rational eigenforms.

```
? F = mfsearch([1..40], 2, [0,1,2,3,4], 1); #F
%1 = 3
? [ mfparams(f)[1..3] | f <- F ]
%2 = [[38, 2, 1], [40, 2, 8], [40, 2, 40]]
```

```
? mfcoefs(F[1],10)
%3 = [0, 1, 2, 3, 4, -5, -8, 1, -7, -5, 7]
```

The library syntax is GEN mfsearch(GEN NK, GEN V, long space).

3.18.58 mfshift(F, s). Divide the generalized modular form F by q^s , omitting the remainder if there is one. One can have $s < 0$.

```
? D=mfDelta(); mfcoefs(mfshift(D,1), 4)
%1 = [1, -24, 252, -1472, 4830]
? mfcoefs(mfshift(D,2), 4)
%2 = [-24, 252, -1472, 4830, -6048]
? mfcoefs(mfshift(D,-1), 4)
%3 = [0, 0, 1, -24, 252]
```

The library syntax is GEN mfshift(GEN F, long s).

3.18.59 mfshimura($mf, F, \{D = 1\}$). F being a modular form of half-integral weight $k \geq 3/2$ and D a positive squarefree integer, returns the Shimura lift G of weight $2k - 1$ corresponding to D . This function returns $[mf2, G, v]$ where $mf2$ is a modular form space containing G and v expresses G in terms of $mfbasis(mf2)$; so that G is $mflinear(mf2, v)$.

```
? F = mfpow(mfTheta(), 7); mf = mfininit(F);
? [mf2, G, v] = mfshimura(mf, F, 3); mfcoefs(G,5)
%2 = [-5/9, 280, 9240, 68320, 295960, 875280]
? mfparams(G) \\ the level may be lower than expected
%3 = [1, 6, 1, y, t - 1]
? mfparams(mf2)
%4 = [2, 6, 1, 4, t - 1]
? v
%5 = [280, 0]~
? mfcoefs(mf2, 5)
%6 =
[-1/504 -1/504]
[      1      0]
[     33      1]
[    244      0]
[   1057     33]
[   3126      0]
? mf = mfininit([60,5/2],1); F = mflinear(mf,mfkohnenbasis(mf)[,1]);
? mfparams(mfshimura(mf,F)[2])
%8 = [15, 4, 1, y, t - 1]
? mfparams(mfshimura(mf,F,6)[2])
%9 = [15, 4, 1, y, t - 1]
```

The library syntax is GEN mfshimura(GEN mf, GEN F, long D).

3.18.60 mflslashexpansion(*mf*, *f*, *g*, *n*, *flrat*, {¶ms}). Let *mf* be a modular form space in level *N*, *f* a modular form belonging to *mf* and let *g* be in $M_2^+(Q)$. This function computes the Fourier expansion of $f|_k g$ to *n* terms. We first describe the behaviour when *flrat* is 0: the result is a vector *v* of floating point complex numbers such that

$$f|_k g(\tau) = q^\alpha \sum_{m \geq 0} v[m+1] q^{m/w},$$

where $q = e(\tau)$, *w* is the width of the cusp $g(i\infty)$ (namely $(N/(c^2, N))$ if *g* is integral) and α is a rational number. If *params* is given, it is set to the parameters $[\alpha, w, \text{matid}(2)]$.

If *flrat* is 1, the program tries to rationalize the expression, i.e., to express the coefficients as rational numbers or polmods. We write $g = \lambda \cdot M \cdot A$ where $\lambda \in \mathbf{Q}^*$, $M \in \text{SL}_2(\mathbf{Z})$ and $A = \begin{bmatrix} a & b \\ 0 & d \end{bmatrix}$ is upper triangular, integral and primitive with $a > 0$, $d > 0$ and $0 \leq b < d$. Let α and *w* be the parameters attached to the expansion of $F := f|_k M$ as above, i.e.

$$F(\tau) = q^\alpha \sum_{m \geq 0} v[m+1] q^{m/w}.$$

The function returns the expansion *v* of $F = f|_k M$ and sets the parameters to $[\alpha, w, A]$. Finally, the desired expansion is $(a/d)^{k/2} F(\tau + b/d)$. The latter is identical to the returned expansion when *A* is the identity, i.e. when $g \in \text{PSL}_2(\mathbf{Z})$. If this is not the case, the expansion differs from *v* by the multiplicative constant $(a/d)^{k/2} e(\alpha b/(dw))$ and a twist by a root of unity $q^{1/w} \rightarrow e(b/(dw)) q^{1/w}$. The complications introduced by this extra matrix *A* allow to recognize the coefficients in a much smaller cyclotomic field, hence to obtain a simpler description overall. (Note that this rationalization step may result in an error if the program cannot perform it.)

```
? mf = mfinit([32,4],0); f = mfbasis(mf)[1];
? mfcoefs(f, 10)
%2 = [0, 3, 0, 0, 0, 2, 0, 0, 0, 47, 0]
? mfatkin = mfatkininit(mf,32); mfcoefs(mfatkin(mfatkin,f),10) / mfatkin[3]
%3 = [0, 1, 0, 16, 0, 22, 0, 32, 0, -27, 0]
? mfatkin[3] \\ here normalizing constant C = 1, but need in general
%4 = 1
? mflslashexpansion(mf,f,[0,-1;1,0],10,1,&params) * 32^(4/2)
%5 = [0, 1, 0, 16, 0, 22, 0, 32, 0, -27, 0]
? params
%6 = [0, 32, [1, 0; 0, 1]]

? mf = mfinit([12,8],0); f = mfbasis(mf)[1];
? mflslashexpansion(mf,f,[1,0;2,1],7,0)
%7 = [0, 0, 0, 0.6666666... + 0.E-38*I, 0, -3.999999... + 6.92820...*I, 0, \
      -11.99999999... - 20.78460969...*I]
? mflslashexpansion(mf,f,[1,0;2,1],7,1, &params)
%8 = [0, 0, 0, 2/3, 0, Mod(8*t, t^2+t+1), 0, Mod(-24*t-24, t^2+t+1)]
? params
%9 = [0, 3, [1, 0; 0, 1]]
```

If $[\mathbf{Q}(f) : \mathbf{Q}(\chi)] > 1$, the coefficients may be polynomials in *y*, where *y* is any root of the polynomial giving the field of definition of *f* (*f.mod* or *mfparams*(*f*)[4]).

```
? mf=mfinit([23,2],0);f=mfeigenbasis(mf)[1];
```

```

? mfcoefs(f,5)
%1 = [Mod(0, y^2 - y - 1), Mod(1, y^2 - y - 1), Mod(-y, y^2 - y - 1),\
      Mod(2*y - 1, y^2 - y - 1), Mod(y - 1, y^2 - y - 1), Mod(-2*y, y^2 - y - 1)]
? mfslashexpansion(mf,f,[1,0;0,1],5,1)
%2 = [0, 1, -y, 2*y - 1, y - 1, -2*y]
? mfslashexpansion(mf,f,[0,-1;1,0],5,1)
%3 = [0, -1/23, 1/23*y, -2/23*y + 1/23, -1/23*y + 1/23, 2/23*y]

```

Caveat. In half-integral weight, we *define* the “slash” operation as

$$(f|_k g)(\tau) := ((c\tau + d)^{-1/2})^{2k} f(g \cdot \tau),$$

with the principal determination of the square root. In particular, the standard cocycle condition is no longer satisfied and we only have $f|(gg') = \pm(f|g)|g'$.

The library syntax is `GEN mfslashexpansion(GEN mf, GEN f, GEN g, long n, long flrat, GEN *params = NULL, long prec)`.

3.18.61 mfspace(mf, {f}). Identify the modular space mf , resp. the modular form f in mf if present, as the flag given to `mfini`. Returns 0 (newspace), 1 (cuspidal space), 2 (old space), 3 (Eisenstein space) or 4 (full space).

```

? mf = mfini([1,12],1); mfspace(mf)
%1 = 1
? mfspace(mf, mfDelta())
%2 = 0 \\ new space

```

This function returns -1 when the form f is modular but does not belong to the space.

```

? mf = mfini([1,2]; mfspace(mf, mfEk(2))
%3 = -1

```

When f is not modular and is for instance only quasi-modular, the function returns nonsense:

```

? M6 = mfini([1,6]);
? dE4 = mfderiv(mfEk(4)); \\ not modular !
? mfspace(M6,dE4) \\ asserts (wrongly) that E4' belongs to new space
%3 = 0

```

The library syntax is `long mfspace(GEN mf, GEN f = NULL)`.

3.18.62 mfsplit(mf, {dimlim = 0}, {flag = 0}). `mf` from `mfini` with integral weight containing the new space (either the new space itself or the cuspidal space or the full space), and preferably the newspace itself for efficiency, split the space into Galois orbits of eigenforms of the newspace, satisfying various restrictions.

The functions returns $[vF, vK]$, where vF gives (Galois orbit of) eigenforms and vK is a list of polynomials defining each Galois orbit. The eigenforms are given in `mftobasis` format, i.e. as a matrix whose columns give the forms with respect to `mfbasis(mf)`.

If `dimlim` is set, only the Galois orbits of dimension $\leq \text{dimlim}$ are computed (i.e. the rational eigenforms if `dimlim = 1` and the character is real). This can considerably speed up the function when a Galois orbit is defined over a large field.

flag speeds up computations when the dimension is large: if $flag = d > 0$, when the dimension of the eigenspace is $> d$, only the Galois polynomial is computed.

Note that the function `mfeigenbasis` returns all eigenforms in an easier to use format (as modular forms which can be input as is in other functions); `mfsplit` is only useful when you can restrict to orbits of small dimensions, e.g. rational eigenforms.

```
? mf=mfinit([11,2],0); f=mfeigenbasis(mf)[1]; mfcoefs(f,16)
%1 = [0, 1, -2, -1, ...]
? mf=mfinit([23,2],0); f=mfeigenbasis(mf)[1]; mfcoefs(f,16)
%2 = [Mod(0, z^2 - z - 1), Mod(1, z^2 - z - 1), Mod(-z, z^2 - z - 1), ...]
? mf=mfinit([179,2],0); apply(poldegree, mffields(mf))
%3 = [1, 3, 11]
? mf=mfinit([719,2],0);
? [vF,vK] = mfsplit(mf, 5); \\ fast when restricting to small orbits
time = 192 ms.
? #vF \\ a single orbit
%5 = 1
? poldegree(vK[1]) \\ of dimension 5
%6 = 5
? [vF,vK] = mfsplit(mf); \\ general case is slow
time = 2,104 ms.
? apply(poldegree,vK)
%8 = [5, 10, 45] \\ because degree 45 is large...
```

The library syntax is `GEN mfsplit(GEN mf, long dimlim, long flag)`.

3.18.63 mfsturm(NK). Gives the Sturm bound for modular forms on $\Gamma_0(N)$ and weight k , i.e., an upper bound for the order of the zero at infinity of a nonzero form. NK is either

- a pair $[N, k]$, in which case the bound is the floor of $(kN/12) \cdot \prod_{p|N} (1 + 1/p)$;
- or the output of `mfinit` in which case the exact upper bound is returned.

```
? NK = [96,6]; mfsturm(NK)
%1 = 97
? mf=mfinit(NK,1); mfsturm(mf)
%2 = 76
? mfdim(NK,0) \\ new space
%3 = 72
```

The library syntax is `long mfsturm(GEN NK)`.

3.18.64 mfsymbol(mf, f). Initialize data for working with all period polynomials of the modular form f : this is essential for efficiency for functions such as `mfsymboleval`, `mfmanin`, and `mfpetersson`. An `mfsymbol` contains an `mf` structure and can always be used whenever an `mf` would be needed.

```
? mf=mfinit([23,2],0);F=mfeigenbasis(mf)[1];
? FS=mfsymbol(mf,F);
? mfsymboleval(FS,[0,oo])
%3 = [8.762565143790690142 E-39 + 0.0877907874...*I,
      -5.617375463602574564 E-39 + 0.0716801031...*I]
? mfpetersson(FS)
%4 =
[0.0039488965740025031688548076498662860143 1.2789721111175127425 E-40]
[1.2630501762985554269 E-40 0.0056442542987647835101583821368582485396]
```

By abuse of language, initialize data for working with `mfpetersson` in weight 1 and half-integral weight (where no symbol exist); the `mf` argument may be an `mfsymbol` attached to a form on the space, which avoids recomputing data independent of the form.

```
? mf=mfinit([12,9/2],1); F=mbasis(mf);
? fs=mfsymbol(mf,F[1]);
time = 476 ms
? mfpetersson(fs)
%2 = 1.9722437519492014682047692073275406145 E-5
? f2s = mfsymbol(mf,F[2]);
time = 484 ms.
? mfpetersson(f2s)
%4 = 1.2142222531326333658647877864573002476 E-5
? gs = mfsymbol(fs,F[2]); \\ re-use existing symbol, a little faster
time = 430 ms.
? mfpetersson(gs) == %4 \\ same value
%6 = 1
```

For simplicity, we also allow `mfsymbol(f)` instead of `mfsymbol(mfinit(f), f)`:

The library syntax is `GEN mfsymbol(GEN mf, GEN f = NULL, long bitprec).`

3.18.65 mfsymboleval(fs, path, {ga = id}). Evaluation of the modular symbol fs (corresponding to the modular form f) output by `mfsymbol` on the given path `path`, where `path` is either a vector $[s_1, s_2]$ or an integral matrix $[a, b; c, d]$ representing the path $[a/c, b/d]$. In both cases s_1 or s_2 (or a/c or b/d) can also be elements of the upper half-plane. To avoid possibly lengthy `mfsymbol` computations, the program also accepts fs of the form `[mf, F]`, but in that case s_1 and s_2 are limited to ∞ and elements of the upper half-plane. The result is the polynomial equal to $\int_{s_1}^{s_2} (X - \tau)^{k-2} F(\tau) d\tau$, the integral being computed along a geodesic joining s_1 and s_2 . If ga in $GL_2^+(\mathbf{Q})$ is given, replace F by $F|_k \gamma$. Note that if the integral diverges, the result will be a rational function. If the field of definition $\mathbf{Q}(f)$ is larger than $\mathbf{Q}(\chi)$ then f can be embedded into $\mathbf{C}$ in $d = [\mathbf{Q}(f) : \mathbf{Q}(\chi)]$ ways, in which case a vector of the d results is returned.

```
? mf=mfinit([35,2],1);f=mbasis(mf)[1];fs=mfsymbol(mf,f);
? mfsymboleval(fs,[0,oo])
%1 = 0.31404011074188471664161704390256378537*I
```

```

? mfsymboleval(fs,[1,3;2,5])
%2 = -0.1429696291... - 0.2619975641...*I
? mfsymboleval(fs,[I,2*I])
%3 = 0.00088969563028739893631700037491116258378*I
? E2=mfEk(2);E22=mflinear([E2,mfbd(E2,2)],[1,-2]);mf=mfinit(E22);
? E2S = mfsymbol(mf,E22);
? mfsymboleval(E2S,[0,1])
%6 = (-1.00000...*x^2 + 1.00000...*x - 0.50000...)/(x^2 - x)

```

The rational function which is given in case the integral diverges is easy to interpret. For instance:

```

? E4=mfEk(4);mf=mfinit(E4);ES=mfsymbol(mf,E4);
? mfsymboleval(ES,[I,oo])
%2 = 1/3*x^3 - 0.928067...*I*x^2 - 0.833333...*x + 0.234978...*I
? mfsymboleval(ES,[0,I])
%3 = (-0.234978...*I*x^3 - 0.833333...*x^2 + 0.928067...*I*x + 0.333333...)/x

```

`mfsymboleval(ES,[a,oo])` is the limit as $T \rightarrow \infty$ of

$$\int_a^{iT} (X - \tau)^{k-2} F(\tau) d\tau + a(0)(X - iT)^{k-1}/(k-1),$$

where $a(0)$ is the 0th coefficient of F at infinity. Similarly, `mfsymboleval(ES,[0,a])` is the limit as $T \rightarrow \infty$ of

$$\int_{i/T}^a (X - \tau)^{k-2} F(\tau) d\tau + b(0)(1 + iTX)^{k-1}/(k-1),$$

where $b(0)$ is the 0th coefficient of $F|_k S$ at infinity.

The library syntax is GEN `mfsymboleval(GEN fs, GEN path, GEN ga = NULL, long bit-prec)`.

3.18.66 mftaylor($F, n, \{flreal = 0\}$). F being a form in $M_k(SL_2(\mathbf{Z}))$, computes the first $n + 1$ canonical Taylor expansion of F around $\tau = I$. If `flreal=0`, computes only an algebraic equivalence class. If `flreal` is set, compute p_n such that for τ close enough to I we have

$$f(\tau) = (2I/(\tau + I))^k \sum_{n \geq 0} p_n ((\tau - I)/(\tau + I))^n.$$

```

? D=mfDelta();
? mftaylor(D,8)
%2 = [1/1728, 0, -1/20736, 0, 1/165888, 0, 1/497664, 0, -11/3981312]

```

The library syntax is GEN `mftaylor(GEN F, long n, long flreal, long prec)`.

3.18.67 mftobasis(*mf*, *F*, {*flag* = 0}). Coefficients of the form *F* on the basis given by **mfbasis**(*mf*). A *q*-expansion or vector of coefficients can also be given instead of *F*, but in this case an error message may occur if the expansion is too short. An error message is also given if *F* does not belong to the modular form space. If *flag* is set, instead of error messages the output is an affine space of solutions if a *q*-expansion or vector of coefficients is given, or the empty column otherwise.

```
? mf = mfininit([26,2],0); mfdim(mf)
%1 = 2
? F = mflinear(mf,[a,b]); mftobasis(mf,F)
%2 = [a, b]~
```

A *q*-expansion or vector of coefficients can also be given instead of *F*.

```
? Th = 1 + 2*sum(n=1, 8, q^(n^2), 0(q^80));
? mf = mfininit([4,5,Mod(3,4)]);
? mftobasis(mf, Th^10)
%3 = [64/5, 4/5, 32/5]~
```

If *F* does not belong to the corresponding space, the result is incorrect and simply matches the coefficients of *F* up to some bound, and the function may either return an empty column or an error message. If *flag* is set, there are no error messages, and the result is an empty column if *F* is a modular form; if *F* is supplied via a series or vector of coefficients which does not contain enough information to force a unique (potential) solution, the function returns [*v*, *K*] where *v* is a solution and *K* is a matrix of maximal rank describing the affine space of potential solutions $v + K \cdot x$.

```
? mf = mfininit([4,12],1);
? mftobasis(mf, q-24*q^2+0(q^3), 1)
%2 = [[43/64, -63/8, 800, 21/64]~, [1, 0; 24, 0; 2048, 768; -1, 0]]
? mftobasis(mf, [0,1,-24,252], 1)
%3 = [[1, 0, 1472, 0]~, [0; 0; 768; 0]]
? mftobasis(mf, [0,1,-24,252,-1472], 1)
%4 = [1, 0, 0, 0]~ \\ now uniquely determined
? mftobasis(mf, [0,1,-24,252,-1472,0], 1)
%5 = [1, 0, 0, 0]~ \\ wrong result: no such form exists
? mfcoefs(mflinear(mf,%), 5) \\ double check
%6 = [0, 1, -24, 252, -1472, 4830]
? mftobasis(mf, [0,1,-24,252,-1472,0])
*** at top-level: mftobasis(mf,[0,1,
*** ^-----
*** mftobasis: domain error in mftobasis: form does not belong to space
? mftobasis(mf, mfEk(10))
*** at top-level: mftobasis(mf,mfEk(
*** ^-----
*** mftobasis: domain error in mftobasis: form does not belong to space
? mftobasis(mf, mfEk(10), 1)
%7 = []~
```

The library syntax is GEN **mftobasis**(GEN *mf*, GEN *F*, long *flag*).

3.18.68 mftocoset($N, M, Lcosets$). M being a matrix in $SL_2(Z)$ and $Lcosets$ being `mfcosets(N)`, a list of right cosets of $\Gamma_0(N)$, find the coset to which M belongs. The output is a pair $[\gamma, i]$ such that $M = \gamma Lcosets[i]$, $\gamma \in \Gamma_0(N)$.

```
? N = 4; L = mfcosets(N);
? mftocoset(N, [1,1;2,3], L)
%2 = [[-1, 1; -4, 3], 5]
```

The library syntax is `GEN mftocoset(ulong N, GEN M, GEN Lcosets)`.

3.18.69 mftonew(mf, F). mf being being a full or cuspidal space with parameters $[N, k, \chi]$ and F a cusp form in that space, returns a vector of 3-component vectors $[M, d, G]$, where $f(\chi) \mid M \mid N$, $d \mid N/M$, and G is a form in $S_k^{\text{new}}(\Gamma_0(M), \chi)$ such that F is equal to the sum of the $B(d)(G)$ over all these 3-component vectors.

```
? mf = mfini([96,6],1); F = mfbasis(mf)[60]; s = mftonew(mf,F); #s
%1 = 1
? [M,d,G] = s[1]; [M,d]
%2 = [48, 2]
? mfcoefs(F,10)
%3 = [0, 0, -160, 0, 0, 0, 0, 0, 0, 0, -14400]
? mfcoefs(G,5)
%4 = [0, -160, 0, 0, 0, -14400]
```

The library syntax is `GEN mftonew(GEN mf, GEN F)`.

3.18.70 mftraceform($NK, \{space = 0\}$). If $NK = [N, k, CHI, .]$ as in `mfini` with k integral, gives the trace form in the corresponding subspace of $S_k(\Gamma_0(N), \chi)$. The supported values for `space` are 0: the newspace (default), 1: the full cuspidal space.

```
? F = mftraceform([23,2]); mfcoefs(F,16)
%1 = [0, 2, -1, 0, -1, -2, -5, 2, 0, 4, 6, -6, 5, 6, 4, -10, -3]
? F = mftraceform([23,1,-23]); mfcoefs(F,16)
%2 = [0, 1, -1, -1, 0, 0, 1, 0, 1, 0, 0, 0, 0, -1, 0, 0, -1]
```

The library syntax is `GEN mftraceform(GEN NK, long space)`.

3.18.71 mftwist(F, D). F being a generalized modular form, returns the twist of F by the integer D , i.e., the form G such that `mfcoef(G,n)=(D/n)mfcoef(F,n)`, where (D/n) is the Kronecker symbol.

```
? mf = mfini([11,2],0); F = mfbasis(mf)[1]; mfcoefs(F, 5)
%1 = [0, 1, -2, -1, 2, 1]
? G = mftwist(F,-3); mfcoefs(G, 5)
%2 = [0, 1, 2, 0, 2, -1]
? mf2 = mfini([99,2],0); mftobasis(mf2, G)
%3 = [1/3, 0, 1/3, 0]~
```

Note that twisting multiplies the level by D^2 . In particular it is not an involution:

```
? H = mftwist(G,-3); mfcoefs(H, 5)
%4 = [0, 1, -2, 0, 2, 1]
? mfparams(G)
%5 = [99, 2, 1, y, t - 1]
```

The library syntax is `GEN mftwist(GEN F, GEN D)`.

3.19 Modular symbols.

Let $\Delta_0 := \text{Div}^0(\mathbf{P}^1(\mathbf{Q}))$ be the abelian group of divisors of degree 0 on the rational projective line. The standard $\text{GL}(2, \mathbf{Q})$ action on $\mathbf{P}^1(\mathbf{Q})$ via homographies naturally extends to Δ_0 . Given

- G a finite index subgroup of $\text{SL}(2, \mathbf{Z})$,
- a field F and a finite dimensional representation V/F of $\text{GL}(2, \mathbf{Q})$,

we consider the space of *modular symbols* $M := \text{Hom}_G(\Delta_0, V)$. This finite dimensional F -vector space is a G -module, canonically isomorphic to $H_c^1(X(G), V)$, and allows to compute modular forms for G .

Currently, we only support the groups $\Gamma_0(N)$ ($N > 0$ an integer) and the representations $V_k = \mathbf{Q}[X, Y]_{k-2}$ ($k \geq 2$ an integer) over $\mathbf{Q}$. We represent a space of modular symbols by an *ms* structure, created by the function `msinit`. It encodes basic data attached to the space: chosen $\mathbf{Z}[G]$ -generators (g_i) for Δ_0 (and relations among those) and an F -basis of M . A modular symbol s is thus given either in terms of this fixed basis, or as a collection of values $s(g_i)$ satisfying certain relations.

A subspace of M (e.g. the cuspidal or Eisenstein subspaces, the new or old modular symbols, etc.) is given by a structure allowing quick projection and restriction of linear operators; its first component is a matrix whose columns form an F -basis of the subspace.

3.19.1 msatkinlehner($M, Q, \{H\}$). Let M be a full modular symbol space of level N , as given by `msinit`, let $Q \mid N$, $(Q, N/Q) = 1$, and let H be a subspace stable under the Atkin-Lehner involution w_Q . Return the matrix of w_Q acting on H (M if omitted).

```
? M = msinit(36,2); \\ M_2(Gamma_0(36))
? w = msatkinlehner(M,4); w^2 == 1
%2 = 1
? #w    \\ involution acts on a 13-dimensional space
%3 = 13
? M = msinit(36,2, -1); \\ M_2(Gamma_0(36))^-
? w = msatkinlehner(M,4); w^2 == 1
%5 = 1
? #w
%6 = 4
```

The library syntax is `GEN msatkinlehner(GEN M, long Q, GEN H = NULL)`.

3.19.2 mscosets(gen, inH). **gen** being a system of generators for a group G and H being a subgroup of finite index in G , return a list of right cosets of $H \backslash G$ and the right action of G on $H \backslash G$. The subgroup H is given by a criterion **inH** (closure) deciding whether an element of G belongs to H . The group G is restricted to types handled by generic multiplication $(*)$ and inversion (g^{-1}) , such as matrix groups or permutation groups.

Let **gens** = $[g_1, \dots, g_r]$. The function returns $[C, M]$ where C lists the $h = [G : H]$ representatives $[\gamma_1, \dots, \gamma_h]$ for the right cosets $H\gamma_1, \dots, H\gamma_h$; γ_1 is always the neutral element in G . For all $i \leq h$, $j \leq r$, if $M[i][j] = k$ then $H\gamma_i g_j = H\gamma_k$.

```
? PSL2 = [[0,1;-1,0], [1,1;0,1]]; \\ S and T
\\ G = PSL2, H = Gamma0(2)
? [C, M] = mscosets(PSL2, g->g[2,1] % 2 == 0);
```

```
? C \\ three cosets
%3 = [[1, 0; 0, 1], [0, 1; -1, 0], [0, 1; -1, -1]]
? M
%4 = [Vecsmall([2, 1]), Vecsmall([1, 3]), Vecsmall([3, 2])]
```

Looking at $M[1]$ we see that S belongs to the second coset and T to the first (trivial) coset.

The library syntax is `GEN mscosets0(GEN gen, GEN inH)`. Also available is the function `GEN mscosets(GEN G, void *E, long (*inH)(void *, GEN))`

3.19.3 mscuspidal($M, \{flag = 0\}$). M being a full modular symbol space, as given by `msinit`, return its cuspidal part S . If $flag = 1$, return $[S, E]$ its decomposition into cuspidal and Eisenstein parts.

A subspace is given by a structure allowing quick projection and restriction of linear operators; its first component is a matrix with integer coefficients whose columns form a $\mathbf{Q}$ -basis of the subspace.

```
? M = msinit(2,8, 1); \\ M_8(Gamma_0(2))^+
? [S,E] = mscuspidal(M, 1);
? E[1] \\ 2-dimensional
%3 =
[0 -10]
[0 -15]
[0 -3]
[1 0]
? S[1] \\ 1-dimensional
%4 =
[ 3]
[30]
[ 6]
[-8]
```

The library syntax is `GEN mscuspidal(GEN M, long flag)`.

3.19.4 msdim(M). M being a full modular symbol space or subspace, for instance as given by `msinit` or `mscuspidal`, return its dimension as a $\mathbf{Q}$ -vector space.

```
? M = msinit(11,4); msdim(M)
%1 = 6
? M = msinit(11,4,1); msdim(M)
%2 = 4 \\ dimension of the '+' part
? [S,E] = mscuspidal(M,1);
? [msdim(S), msdim(E)]
%4 = [2, 2]
```

Note that `mfdim([N,k])` is going to be much faster if you only need the dimension of the space and not really to work with it. This function is only useful to quickly check the dimension of an existing space.

The library syntax is `long msdim(GEN M)`.

3.19.5 mseisenstein(M). M being a full modular symbol space, as given by `msinit`, return its Eisenstein subspace. A subspace is given by a structure allowing quick projection and restriction of linear operators; its first component is a matrix with integer coefficients whose columns form a $\mathbf{Q}$ -basis of the subspace. This is the same basis as given by the second component of `mscuspidal( $M, 1$ )`.

```
? M = msinit(2,8, 1); \\ M_8(Gamma_0(2))^+
? E = mseisenstein(M);
? E[1] \\ 2-dimensional
%3 =
[0 -10]
[0 -15]
[0 -3]
[1 0]
? E == mscuspidal(M,1)[2]
%4 = 1
```

The library syntax is `GEN mseisenstein(GEN M)`.

3.19.6 mseval($M, s, \{p\}$). Let $\Delta_0 := \text{Div}^0(\mathbf{P}^1(\mathbf{Q}))$. Let M be a full modular symbol space, as given by `msinit`, let s be a modular symbol from M , i.e. an element of $\text{Hom}_G(\Delta_0, V)$, and let $p = [a, b] \in \Delta_0$ be a path between two elements in $\mathbf{P}^1(\mathbf{Q})$, return $s(p) \in V$. The path extremities a and b may be given as `t_INT`, `t_FRAC` or `oo` $= (1 : 0)$; it is also possible to describe the path by a 2×2 integral matrix whose columns give the two cusps. The symbol s is either

- a `t_COL` coding a modular symbol in terms of the fixed basis of $\text{Hom}_G(\Delta_0, V)$ chosen in M ; if M was initialized with a nonzero *sign* (+ or −), then either the basis for the full symbol space or the $\pm$ -part can be used (the dimension being used to distinguish the two).

- a `t_MAT` whose columns encode modular symbols as above. This is much faster than evaluating individual symbols on the same path p independently.

- a `t_VEC` (v_i) of elements of V , where the $v_i = s(g_i)$ give the image of the generators g_i of Δ_0 , see `mspathgens`. We assume that s is a proper symbol, i.e. that the v_i satisfy the `mspathgens` relations.

If p is omitted, convert a single symbol s to the second form: a vector of the $s(g_i)$. A `t_MAT` is converted to a vector of such.

```
? M = msinit(2,8,1); \\ M_8(Gamma_0(2))^+
? g = mspathgens(M)[1]
%2 = [[+oo, 0], [0, 1]]
? N = msnew(M)[1]; #N \\ Q-basis of new subspace, dimension 1
%3 = 1
? s = N[,1] \\ t_COL representation
%4 = [-3, 6, -8]~
? S = mseval(M, s) \\ t_VEC representation
%5 = [64*x^6-272*x^4+136*x^2-8, 384*x^5+960*x^4+192*x^3-672*x^2-432*x-72]
? mseval(M,s, g[1])
%6 = 64*x^6 - 272*x^4 + 136*x^2 - 8
? mseval(M,S, g[1])
%7 = 64*x^6 - 272*x^4 + 136*x^2 - 8
```

Note that the symbol should have values in $V = \mathbf{Q}[x, y]_{k-2}$, we return the de-homogenized values corresponding to $y = 1$ instead.

The library syntax is `GEN mseval(GEN M, GEN s, GEN p = NULL)`.

3.19.7 msfarey($F, inH, \{\&CM\}$). F being a Farey symbol attached to a group G contained in $\mathrm{PSL}_2(\mathbf{Z})$ and H a subgroup of G , return a Farey symbol attached to H . The subgroup H is given by a criterion `inH` (closure) deciding whether an element of G belongs to H . The symbol F can be created using

- `mspolygon`: $G = \Gamma_0(N)$, which runs in time $\tilde{O}(N)$;
- or `msfarey` itself, which runs in time $O([G : H]^2)$.

If present, the argument `CM` is set to `mscosets(F[3])`, giving the right cosets of $H \backslash G$ and the action of G by right multiplication. Since `msfarey`'s algorithm is quadratic in the index $[G : H]$, it is advisable to construct subgroups by a chain of inclusions if possible.

```
\\ Gamma_0(N)
G0(N) = mspolygon(N);

\\ Gamma_1(N): direct construction, slow
G1(N) = msfarey(mspolygon(1), g -> my(a = g[1,1]%N, c = g[2,1]%N);\
                                     c == 0 && (a == 1 || a == N-1));

\\ Gamma_1(N) via Gamma_0(N): much faster
G1(N) = msfarey(G0(N), g -> my(a=g[1,1]%N); a==1 || a==N-1);
```

Note that the simpler criterion `g[1,1]%N == 1` would not be correct since it must apply to elements of $\mathrm{PSL}_2(\mathbf{Z})$ hence be invariant under $g \mapsto -g$. Here are other examples:

```
\\ Gamma(N)
G(N) = msfarey(G1(N), g -> g[1,2]%N==0);

G_00(N) = msfarey(G0(N), x -> x[1,2]%N==0);
G1_0(N1,N2) = msfarey(G0(1), x -> x[2,1]%N1==0 && x[1,2]%N2==0);

\\ Gamma_0(91) has 4 elliptic points of order 3, Gamma_1(91) has none
D0 = mspolygon(G0(91), 2)[4];
D1 = mspolygon(G1(91), 2)[4];
write("F.tex", "\\documentclass{article}\\usepackage{tikz}\\begin{document}", \
      D0, "\\n", D1, "\\end{document}");
```

The library syntax is `GEN msfarey0(GEN F, GEN inH, GEN *CM = NULL)`. Also available is `GEN msfarey(GEN F, void *E, long (*inH)(void *, GEN), GEN *pCM)`.

3.19.8 msfromcusp(M, c). Returns the modular symbol attached to the cusp c , where M is a modular symbol space of level N , attached to $G = \Gamma_0(N)$. The cusp c in $\mathbf{P}^1(\mathbf{Q})/G$ is given either as ∞ ($= (1 : 0)$) or as a rational number a/b ($= (a : b)$). The attached symbol maps the path $[b] - [a] \in \text{Div}^0(\mathbf{P}^1(\mathbf{Q}))$ to $E_c(b) - E_c(a)$, where $E_c(r)$ is 0 when $r \neq c$ and $X^{k-2} \mid \gamma_r$ otherwise, where $\gamma_r \cdot r = (1 : 0)$. These symbols span the Eisenstein subspace of M .

```
? M = msinit(2,8);  \\ M_8(Gamma_0(2))
? E = mseisenstein(M);
? E[1]  \\ two-dimensional
%3 =
[0 -10]
[0 -15]
[0 -3]
[1  0]
? s = msfromcusp(M,oo)
%4 = [0, 0, 0, 1]~
? mseval(M, s)
%5 = [1, 0]
? s = msfromcusp(M,1)
%6 = [-5/16, -15/32, -3/32, 0]~
? mseval(M,s)
%7 = [-x^6, -6*x^5 - 15*x^4 - 20*x^3 - 15*x^2 - 6*x - 1]
```

In case M was initialized with a nonzero *sign*, the symbol is given in terms of the fixed basis of the whole symbol space, not the $+$ or $-$ part (to which it need not belong).

```
? M = msinit(2,8, 1);  \\ M_8(Gamma_0(2))^+
? E = mseisenstein(M);
? E[1]  \\ still two-dimensional, in a smaller space
%3 =
[ 0 -10]
[ 0  3]
[-1  0]
? s = msfromcusp(M,oo)  \\ in terms of the basis for M_8(Gamma_0(2)) !
%4 = [0, 0, 0, 1]~
? mseval(M, s)  \\ same symbol as before
%5 = [1, 0]
```

The library syntax is `GEN msfromcusp(GEN M, GEN c)`.

3.19.9 msfromell($E, \{sign = 0\}$). Let $E/\mathbf{Q}$ be an elliptic curve of conductor N . For $\varepsilon = \pm 1$, we define the (cuspidal, new) modular symbol x^ε in $H_c^1(X_0(N), \mathbf{Q})^\varepsilon$ attached to E . For all primes p not dividing N we have $T_p(x^\varepsilon) = a_p x^\varepsilon$, where $a_p = p + 1 - \#E(\mathbf{F}_p)$.

Let $\Omega^+ = \mathbf{E.omega}[1]$ be the real period of E (integration of the Néron differential $dx/(2y + a_1x + a_3)$ on the connected component of $E(\mathbf{R})$, i.e. the generator of $H_1(E, \mathbf{Z})^+$) normalized by $\Omega^+ > 0$. Let $i\Omega^-$ the integral on a generator of $H_1(E, \mathbf{Z})^-$ with $\Omega^- \in \mathbf{R}_{>0}$. If c_∞ is the number of connected components of $E(\mathbf{R})$, Ω^- is equal to $(-2/c_\infty) \times \text{imag}(\mathbf{E.omega}[2])$. The complex modular symbol is defined by

$$F : \delta \rightarrow 2i\pi \int_{\delta} f(z) dz$$

The modular symbols x^ε are normalized so that $F = x^+ \Omega^+ + x^- i\Omega^-$. In particular, we have

$$x^+([0] - [\infty]) = L(E, 1)/\Omega^+,$$

which defines $x^\pm$ unless $L(E, 1) = 0$. Furthermore, for all fundamental discriminants D such that $\varepsilon \cdot D > 0$, we also have

$$\sum_{0 \leq a < |D|} (D|a) x^\varepsilon([a/|D|] - [\infty]) = L(E, (D|\cdot), 1)/\Omega^\varepsilon,$$

where $(D|\cdot)$ is the Kronecker symbol. The period Ω^- is also $2/c_\infty \times$ the real period of the twist $E^{(-4)} = \mathbf{elltwist}(\mathbf{E}, -4)$.

This function returns the pair $[M, x]$, where M is $\mathbf{msinit}(N, 2)$ and x is x^{sign} as above when $sign = \pm 1$, and $x = [x^+, x^-, L_E]$ when $sign$ is 0, where L_E is a matrix giving the canonical $\mathbf{Z}$ -lattice attached to E in the sense of $\mathbf{mslattice}$ applied to $\mathbf{Q}x^+ + \mathbf{Q}x^-$. Explicitly, it is generated by (x^+, x^-) when $E(\mathbf{R})$ has two connected components and by $(x^+ - x^-, 2x^-)$ otherwise.

The modular symbols $x^\pm$ are given as a $\mathbf{t_COL}$ (in terms of the fixed basis of $\text{Hom}_G(\Delta_0, \mathbf{Q})$ chosen in M).

```
? E=ellinit([0,-1,1,-10,-20]);  \\ X_0(11)
? [M,xp]= msfromell(E,1);
? xp
%3 = [1/5, -1/2, -1/2]~
? [M,x]= msfromell(E);
? x      \\ x^+, x^- and L_E
%5 = [[1/5, -1/2, -1/2]~, [0, 1/2, -1/2]~, [1/5, 0; -1, 1; 0, -1]]
? p = 23; (mshecke(M,p) - ellap(E,p))*x[1]
%6 = [0, 0, 0]~ \\ true at all primes, including p = 11; same for x[2]
? (mshecke(M,p) - ellap(E,p))*x[3] == 0
%7 = 1
```

Instead of a single curve E , one may use instead a vector of *isogenous* curves. The function then returns M and the vector of attached modular symbols.

The library syntax is `GEN msfromell(GEN E, long sign)`.

3.19.10 msfromhecke($M, v, \{H\}$). Given a msinit M and a vector v of pairs $[p, P]$ (where p is prime and P is a polynomial with integer coefficients), return a basis of all modular symbols such that $P(T_p)(s) = 0$. If H is present, it must be a Hecke-stable subspace and we restrict to $s \in H$. When T_p has a rational eigenvalue and $P(x) = x - a_p$ has degree 1, we also accept the integer a_p instead of P .

```
? E = ellinit([0,-1,1,-10,-20]) \\11a1
? ellap(E,2)
%2 = -2
? ellap(E,3)
%3 = -1
? M = msinit(11,2);
? S = msfromhecke(M, [[2,-2],[3,-1]])
%5 =
[ 1  1]
[-5  0]
[ 0 -5]
? mshecke(M, 2, S)
%6 =
[-2  0]
[ 0 -2]
? M = msinit(23,4);
? S = msfromhecke(M, [[5, x^4-14*x^3-244*x^2+4832*x-19904]]);
? factor( charpoly(mshecke(M,5,S)) )
%9 =
[x^4 - 14*x^3 - 244*x^2 + 4832*x - 19904 2]
```

The library syntax is GEN msfromhecke(GEN M, GEN v, GEN H = NULL).

3.19.11 msgetlevel(M). M being a full modular symbol space, as given by msinit, return its level N .

The library syntax is long msgetlevel(GEN M).

3.19.12 msgetsign(M). M being a full modular symbol space, as given by msinit, return its sign: ± 1 or 0 (unset).

```
? M = msinit(11,4, 1);
? msgetsign(M)
%2 = 1
? M = msinit(11,4);
? msgetsign(M)
%4 = 0
```

The library syntax is long msgetsign(GEN M).

3.19.13 msgetweight(M). M being a full modular symbol space, as given by `msinit`, return its weight k .

```
? M = msinit(11,4);
? msgetweight(M)
%2 = 4
```

The library syntax is `long msgetweight(GEN M)`.

3.19.14 mshecke($M, p, \{H\}$). M being a full modular symbol space, as given by `msinit`, p being a prime number, and H being a Hecke-stable subspace (M if omitted), return the matrix of T_p acting on H (U_p if p divides N). Result is undefined if H is not stable by T_p (resp. U_p).

```
? M = msinit(11,2); \\ M_2(Gamma_0(11))
? T2 = mshecke(M,2)
%2 =
[3  0  0]
[1 -2  0]
[1  0 -2]
? M = msinit(11,2, 1); \\ M_2(Gamma_0(11))^+
? T2 = mshecke(M,2)
%4 =
[ 3  0]
[-1 -2]
? N = msnew(M)[1] \\ Q-basis of new cuspidal subspace
%5 =
[-2]
[-5]
? p = 1009; mshecke(M, p, N) \\ action of T_1009 on N
%6 =
[-10]
? ellap(ellinit("11a1"), p)
%7 = -10
```

The library syntax is `GEN mshecke(GEN M, long p, GEN H = NULL)`.

3.19.15 msinit($G, V, \{sign = 0\}$). Given G a finite index subgroup of $SL(2, \mathbf{Z})$ and a finite dimensional representation V of $GL(2, \mathbf{Q})$, creates a space of modular symbols, the G -module $\text{Hom}_G(\text{Div}^0(\mathbf{P}^1(\mathbf{Q})), V)$. This is canonically isomorphic to $H_c^1(X(G), V)$, and allows to compute modular forms for G . If $sign$ is present and nonzero, it must be ± 1 and we consider the subspace defined by $\text{Ker}(\sigma - sign)$, where σ is induced by $[-1, 0; 0, 1]$. Currently the only supported groups are the $\Gamma_0(N)$, coded by the integer $N > 0$. The only supported representation is $V_k = \mathbf{Q}[X, Y]_{k-2}$, coded by the integer $k \geq 2$.

```
? M = msinit(11,2); msdim(M) \\ Gamma0(11), weight 2
%1 = 3
? mshecke(M,2) \\ T_2 acting on M
%2 =
[3  1  1]
```

```

[0 -2  0]
[0  0 -2]
? msstar(M) \\ * involution
%3 =
[1 0 0]
[0 0 1]
[0 1 0]
? Mp = msinit(11,2, 1); msdim(Mp) \\ + part
%4 = 2
? mshecke(Mp,2) \\ T_2 action on M^+
%5 =
[3  2]
[0 -2]
? msstar(Mp)
%6 =
[1 0]
[0 1]

```

The library syntax is GEN msinit(GEN G, GEN V, long sign).

3.19.16 msissymbol(M, s). M being a full modular symbol space, as given by `msinit`, check whether s is a modular symbol attached to M . If A is a matrix, check whether its columns represent modular symbols and return a 0 – 1 vector.

```

? M = msinit(7,8, 1); \\ M_8(Gamma_0(7))^+
? A = msnew(M)[1];
? s = A[,1];
? msissymbol(M, s)
%4 = 1
? msissymbol(M, A)
%5 = [1, 1, 1]
? S = mseval(M,s);
? msissymbol(M, S)
%7 = 1
? [g,R] = mspathgens(M); g
%8 = [[+oo, 0], [0, 1/2], [1/2, 1]]
? #R \\ 3 relations among the generators g_i
%9 = 3
? T = S; T[3]++; \\ randomly perturb S(g_3)
? msissymbol(M, T)
%11 = 0 \\ no longer satisfies the relations

```

The library syntax is GEN msissymbol(GEN M, GEN s).

3.19.17 mslattice($M, \{H\}$). Let $\Delta_0 := \text{Div}^0(\mathbf{P}^1(\mathbf{Q}))$ and $V_k = \mathbf{Q}[x, y]_{k-2}$. Let M be a full modular symbol space, as given by `msinit` and let H be a subspace, e.g. as given by `mscuspidal`. This function returns a canonical $\mathbf{Z}$ -structure on H defined as follows. Consider the map $c : M = \text{Hom}_{\Gamma_0(N)}(\Delta_0, V_k) \rightarrow H^1(\Gamma_0(N), V_k)$ given by $\phi \mapsto \text{class}(\gamma \rightarrow \phi(\{0, \gamma^{-1}0\}))$. Let $L_k = \mathbf{Z}[x, y]_{k-2}$ be the natural $\mathbf{Z}$ -structure of V_k . The result of `mslattice` is a $\mathbf{Z}$ -basis of the inverse image by c of $H^1(\Gamma_0(N), L_k)$ in the space of modular symbols generated by H .

For user convenience, H can be defined by a matrix representing the $\mathbf{Q}$ -basis of H (in terms of the canonical $\mathbf{Q}$ -basis of M fixed by `msinit` and used to represent modular symbols).

If omitted, H is the cuspidal part of M as given by `mscuspidal`. The Eisenstein part $\text{Hom}_{\Gamma_0(N)}(\text{Div}(\mathbf{P}^1(\mathbf{Q})), V_k)$ is in the kernel of c , so the result has no meaning for the Eisenstein part H .

```
? M=msinit(11,2);
? [S,E] = mscuspidal(M,1); S[1] \\ a primitive Q-basis of S
%2 =
[ 1  1]
[-5  0]
[ 0 -5]
? mslattice(M,S)
%3 =
[-1/5 -1/5]
[  1    0]
[  0    1]
? mslattice(M,E)
%4 =
[1]
[0]
[0]
? M=msinit(5,4);
? S=mscuspidal(M); S[1]
%6 =
[ 7  20]
[ 3   3]
[-10 -23]
[-30 -30]
? mslattice(M,S)
%7 =
[-1/10 -11/130]
[  0    -1/130]
[ 1/10    6/65]
[  0     1/13]
```

The library syntax is `GEN mslattice(GEN M, GEN H = NULL)`.

3.19.18 msnew(M). M being a full modular symbol space, as given by `msinit`, return the *new* part of its cuspidal subspace. A subspace is given by a structure allowing quick projection and restriction of linear operators; its first component is a matrix with integer coefficients whose columns form a $\mathbf{Q}$ -basis of the subspace.

```
? M = msinit(11,8, 1); \\ M_8(Gamma_0(11))^+
? N = msnew(M);
? #N[1] \\ 6-dimensional
%3 = 6
```

The library syntax is `GEN msnew(GEN M)`.

3.19.19 msomseval($Mp, PHI, path$). Return the vectors of moments of the p -adic distribution attached to the path `path` by the overconvergent modular symbol `PHI`.

```
? M = msinit(3,6,1);
? Mp= mspadicinit(M,5,10);
? phi = [5,-3,-1]~;
? msissymbol(M,phi)
%4 = 1
? PHI = mstooms(Mp,phi);
? ME = msomseval(Mp,PHI,[oo, 0]);
```

The library syntax is `GEN msomseval(GEN Mp, GEN PHI, GEN path)`.

3.19.20 mspadicL($\mu, \{s = 0\}, \{r = 0\}$). Returns the value (or r -th derivative) on a character χ^s of $\mathbf{Z}_p^*$ of the p -adic L -function attached to μ .

Let Φ be the p -adic distribution-valued overconvergent symbol attached to a modular symbol ϕ for $\Gamma_0(N)$ (eigenvector for $T_N(p)$ for the eigenvalue a_p). Then $L_p(\Phi, \chi^s) = L_p(\mu, s)$ is the p -adic L function defined by

$$L_p(\Phi, \chi^s) = \int_{\mathbf{Z}_p^*} \chi^s(z) d\mu(z)$$

where μ is the distribution on $\mathbf{Z}_p^*$ defined by the restriction of $\Phi([\infty] - [0])$ to $\mathbf{Z}_p^*$. The r -th derivative is taken in direction $\langle \chi \rangle$:

$$L_p^{(r)}(\Phi, \chi^s) = \int_{\mathbf{Z}_p^*} \chi^s(z) (\log z)^r d\mu(z).$$

In the argument list,

- `mu` is as returned by `mspadicmoments` (distributions attached to Φ by restriction to discs $a + p^\nu \mathbf{Z}_p$, $(a, p) = 1$).

- $s = [s_1, s_2]$ with $s_1 \in \mathbf{Z} \subset \mathbf{Z}_p$ and $s_2 \bmod p - 1$ or $s_2 \bmod 2$ for $p = 2$, encoding the p -adic character $\chi^s := \langle \chi \rangle^{s_1} \tau^{s_2}$; here χ is the cyclotomic character from $\text{Gal}(\mathbf{Q}_p(\mu_{p^\infty})/\mathbf{Q}_p)$ to $\mathbf{Z}_p^*$, and τ is the Teichmüller character (for $p > 2$ and the character of order 2 on $(\mathbf{Z}/4\mathbf{Z})^*$ if $p = 2$); for convenience, the character $[s, s]$ can also be represented by the integer s .

When a_p is a p -adic unit, L_p takes its values in $\mathbf{Q}_p$. When a_p is not a unit, it takes its values in the two-dimensional $\mathbf{Q}_p$ -vector space $D_{\text{cris}}(M(\phi))$ where $M(\phi)$ is the “motive” attached to ϕ , and we return the two p -adic components with respect to some fixed $\mathbf{Q}_p$ -basis.

```
? M = msinit(3,6,1); phi=[5, -3, -1]~;
```

```

? msissymbol(M,phi)
%2 = 1
? Mp = mspadicinit(M, 5, 4);
? mu = mspadicmoments(Mp, phi); \\ no twist
\\ End of initializations

? mspadicL(mu,0) \\ L_p(chi^0)
%5 = 5 + 2*5^2 + 2*5^3 + 2*5^4 + ...
? mspadicL(mu,1) \\ L_p(chi), zero for parity reasons
%6 = [0(5^13)]~
? mspadicL(mu,2) \\ L_p(chi^2)
%7 = 3 + 4*5 + 4*5^2 + 3*5^5 + ...
? mspadicL(mu,[0,2]) \\ L_p(tau^2)
%8 = 3 + 5 + 2*5^2 + 2*5^3 + ...
? mspadicL(mu, [1,0]) \\ L_p(<chi>)
%9 = 3*5 + 2*5^2 + 5^3 + 2*5^7 + 5^8 + 5^10 + 2*5^11 + 0(5^13)
? mspadicL(mu,0,1) \\ L_p'(chi^0)
%10 = 2*5 + 4*5^2 + 3*5^3 + ...
? mspadicL(mu, 2, 1) \\ L_p'(chi^2)
%11 = 4*5 + 3*5^2 + 5^3 + 5^4 + ...

```

Now several quadratic twists: `mstooms` is indicated.

```

? PHI = mstooms(Mp,phi);
? mu = mspadicmoments(Mp, PHI, 12); \\ twist by 12
? mspadicL(mu)
%14 = 5 + 5^2 + 5^3 + 2*5^4 + ...
? mu = mspadicmoments(Mp, PHI, 8); \\ twist by 8
? mspadicL(mu)
%16 = 2 + 3*5 + 3*5^2 + 2*5^4 + ...
? mu = mspadicmoments(Mp, PHI, -3); \\ twist by -3 < 0
? mspadicL(mu)
%18 = 0(5^13) \\ always 0, phi is in the + part and D < 0

```

One can locate interesting symbols of level N and weight k with `msnew` and `mssplit`. Note that instead of a symbol, one can input a 1-dimensional Hecke-subspace from `mssplit`: the function will automatically use the underlying basis vector.

```

? M=msinit(5,4,1); \\ M_4(Gamma_0(5))^+
? L = mssplit(M, msnew(M)); \\ list of irreducible Hecke-subspaces
? phi = L[1]; \\ one Galois orbit of newforms
? #phi[1] \\... this one is rational
%4 = 1
? Mp = mspadicinit(M, 3, 4);
? mu = mspadicmoments(Mp, phi);
? mspadicL(mu)
%7 = 1 + 3 + 3^3 + 3^4 + 2*3^5 + 3^6 + 0(3^9)

? M = msinit(11,8, 1); \\ M_8(Gamma_0(11))^+
? Mp = mspadicinit(M, 3, 4);
? L = mssplit(M, msnew(M));
? phi = L[1]; #phi[1] \\ ... this one is two-dimensional

```

```

%11 = 2
? mu = mspadicmoments(Mp, phi);
***   at top-level: mu=mspadicmoments(Mp,ph
***                                     ^-----
*** mspadicmoments: incorrect type in mstooms [dim_Q (eigenspace) > 1]

```

The library syntax is GEN mspadicL(GEN mu, GEN s = NULL, long r).

3.19.21 mspadicinit($M, p, n, \{flag\}$). M being a full modular symbol space, as given by `msinit`, and p a prime, initialize technical data needed to compute with overconvergent modular symbols, modulo p^n . If *flag* is unset, allow all symbols; else initialize only for a restricted range of symbols depending on *flag*: if *flag* = 0 restrict to ordinary symbols, else restrict to symbols ϕ such that $T_p(\phi) = a_p\phi$, with $v_p(a_p) \geq flag$, which is faster as *flag* increases. (The fastest initialization is obtained for *flag* = 0 where we only allow ordinary symbols.) For supersingular eigensymbols, such that $p \mid a_p$, we must further assume that p does not divide the level.

```

? E = ellinit("11a1");
? [M,phi] = msfromell(E,1);
? ellap(E,3)
%3 = -1
? Mp = mspadicinit(M, 3, 10, 0); \\ commit to ordinary symbols
? PHI = mstooms(Mp,phi);

```

If we restrict the range of allowed symbols with *flag* (for faster initialization), exceptions will occur if $v_p(a_p)$ violates this bound:

```

? E = ellinit("15a1");
? [M,phi] = msfromell(E,1);
? ellap(E,7)
%3 = 0
? Mp = mspadicinit(M,7,5,0); \\ restrict to ordinary symbols
? PHI = mstooms(Mp,phi)
***   at top-level: PHI=mstooms(Mp,phi)
***                                     ^-----
*** mstooms: incorrect type in mstooms [v_p(ap) > mspadicinit flag] (t_VEC).
? Mp = mspadicinit(M,7,5); \\ no restriction
? PHI = mstooms(Mp,phi);

```

This function uses $O(N^2(n+k)^2p)$ memory, where N is the level of M .

The library syntax is GEN mspadicinit(GEN M, long p, long n, long flag).

3.19.22 mspadicmoments($Mp, PHI, \{D = 1\}$). Given Mp from `mspadicinit`, an overconvergent eigensymbol PHI from `mstooms` and a fundamental discriminant D coprime to p , let PHI^D denote the twisted symbol. This function computes the distribution $\mu = PHI^D([0] - [\infty]) \mid \mathbf{Z}_p^*$ restricted to $\mathbf{Z}_p^*$. More precisely, it returns the moments of the $p-1$ distributions $PHI^D([0] - [\infty]) \mid (a + p\mathbf{Z}_p)$, $0 < a < p$. We also allow PHI to be given as a classical symbol, which is then lifted to an overconvergent symbol by `mstooms`; but this is wasteful if more than one twist is later needed.

The returned data μ (p -adic distributions attached to PHI) can then be used in `mspadicL` or `mspadicseries`. This precomputation allows to quickly compute derivatives of different orders or values at different characters.

```

? M = msinit(3,6, 1);
? phi = [5,-3,-1]~;
? msissymbol(M, phi)
%3 = 1
? p = 5; mshecke(M,p) * phi \\ eigenvector of T_5, a_5 = 6
%4 = [30, -18, -6]~
? Mp = mspadicinit(M, p, 10, 0); \\ restrict to ordinary symbols, mod p^10
? PHI = mstooms(Mp, phi);
? mu = mspadicmoments(Mp, PHI);
? mspadicL(mu)
%8 = 5 + 2*5^2 + 2*5^3 + ...
? mu = mspadicmoments(Mp, PHI, 12); \\ twist by 12
? mspadicL(mu)
%10 = 5 + 5^2 + 5^3 + 2*5^4 + ...

```

The library syntax is GEN mspadicmoments(GEN Mp, GEN PHI, long D).

3.19.23 mspadicseries($\mu, \{i = 0\}$). Let Φ be the p -adic distribution-valued overconvergent symbol attached to a modular symbol ϕ for $\Gamma_0(N)$ (eigenvector for $T_N(p)$ for the eigenvalue a_p). If μ is the distribution on $\mathbf{Z}_p^*$ defined by the restriction of $\Phi([\infty] - [0])$ to $\mathbf{Z}_p^*$, let

$$\hat{L}_p(\mu, \tau^i)(x) = \int_{\mathbf{Z}_p^*} \tau^i(t)(1+x)^{\log_p(t)/\log_p(u)} d\mu(t)$$

Here, τ is the Teichmüller character and u is a specific multiplicative generator of $1+2p\mathbf{Z}_p$, namely $1+p$ if $p > 2$ or 5 if $p = 2$. To explain the formula, let $G_\infty := \text{Gal}(\mathbf{Q}(\mu_{p^\infty})/\mathbf{Q})$, let $\chi : G_\infty \rightarrow \mathbf{Z}_p^*$ be the cyclotomic character (isomorphism) and γ the element of G_∞ such that $\chi(\gamma) = u$; then $\chi(\gamma)^{\log_p(t)/\log_p(u)} = \langle t \rangle$.

The p -adic precision of individual terms is maximal given the precision of the overconvergent symbol μ .

```

? [M,phi] = msfromell(ellinit("17a1"),1);
? Mp = mspadicinit(M, 5,7);
? mu = mspadicmoments(Mp, phi,1); \\ overconvergent symbol
? mspadicseries(mu)
%4 = (4 + 3*5 + 4*5^2 + 2*5^3 + 2*5^4 + 5^5 + 4*5^6 + 3*5^7 + 0(5^9)) \
+ (3 + 3*5 + 5^2 + 5^3 + 2*5^4 + 5^6 + 0(5^7))*x \
+ (2 + 3*5 + 5^2 + 4*5^3 + 2*5^4 + 0(5^5))*x^2 \
+ (3 + 4*5 + 4*5^2 + 0(5^3))*x^3 \
+ (3 + 0(5))*x^4 + 0(x^5)

```

An example with nonzero Teichmüller:

```

? [M,phi] = msfromell(ellinit("11a1"),1);
? Mp = mspadicinit(M, 3,10);
? mu = mspadicmoments(Mp, phi,1);
? mspadicseries(mu, 2)
%4 = (2 + 3 + 3^2 + 2*3^3 + 2*3^5 + 3^6 + 3^7 + 3^10 + 3^11 + 0(3^12)) \
+ (1 + 3 + 2*3^2 + 3^3 + 3^5 + 2*3^6 + 2*3^8 + 0(3^9))*x \
+ (1 + 2*3 + 3^4 + 2*3^5 + 0(3^6))*x^2 \

```

```
+ (3 + 0(3^2))*x^3 + 0(x^4)
```

Supersingular example (not checked)

```
? E = ellinit("17a1"); ellap(E,3)
%1 = 0
? [M,phi] = msfromell(E,1);
? Mp = mspadicinit(M, 3,7);
? mu = mspadicmoments(Mp, phi,1);
? mspadicseries(mu)
%5 = [(2*3^-1 + 1 + 3 + 3^2 + 3^3 + 3^4 + 3^5 + 3^6 + 0(3^7)) \
+ (2 + 3^3 + 0(3^5))*x \
+ (1 + 2*3 + 0(3^2))*x^2 + 0(x^3),\
(3^-1 + 1 + 3 + 3^2 + 3^3 + 3^4 + 3^5 + 3^6 + 0(3^7)) \
+ (1 + 2*3 + 2*3^2 + 3^3 + 2*3^4 + 0(3^5))*x \
+ (3^-2 + 3^-1 + 0(3^2))*x^2 + 0(3^-2)*x^3 + 0(x^4)]
```

Example with a twist:

```
? E = ellinit("11a1");
? [M,phi] = msfromell(E,1);
? Mp = mspadicinit(M, 3,10);
? mu = mspadicmoments(Mp, phi,5); \\ twist by 5
? L = mspadicseries(mu)
%5 = (2*3^2 + 2*3^4 + 3^5 + 3^6 + 2*3^7 + 2*3^10 + 0(3^12)) \
+ (2*3^2 + 2*3^6 + 3^7 + 3^8 + 0(3^9))*x \
+ (3^3 + 0(3^6))*x^2 + 0(3^2)*x^3 + 0(x^4)
? mspadicL(mu)
%6 = [2*3^2 + 2*3^4 + 3^5 + 3^6 + 2*3^7 + 2*3^10 + 0(3^12)]~
? ellpadicL(E,3,10,,5)
%7 = 2 + 2*3^2 + 3^3 + 2*3^4 + 2*3^5 + 3^6 + 2*3^7 + 0(3^10)
? mspadicseries(mu,1) \\ must be 0
%8 = 0(3^12) + 0(3^9)*x + 0(3^6)*x^2 + 0(3^2)*x^3 + 0(x^4)
```

The library syntax is GEN mspadicseries(GEN mu, long i).

3.19.24 mspathgens(M). Let $\Delta_0 := \text{Div}^0(\mathbf{P}^1(\mathbf{Q}))$. Let M being a full modular symbol space, as given by `msinit`, return a set of $\mathbf{Z}[G]$ -generators for Δ_0 . The output is $[g, R]$, where g is a minimal system of generators and R the vector of $\mathbf{Z}[G]$ -relations between the given generators. A relation is coded by a vector of pairs $[a_i, i]$ with $a_i \in \mathbf{Z}[G]$ and i the index of a generator, so that $\sum_i a_i g[i] = 0$.

An element $[v] - [u]$ in Δ_0 is coded by the “path” $[u, v]$, where ∞ denotes the point at infinity $(1 : 0)$ on the projective line. An element of $\mathbf{Z}[G]$ is either an integer n ($= n[\text{id}_2]$) or a “factorization matrix”: the first column contains distinct elements g_i of G and the second integers n_i and the matrix codes $\sum_i n_i [g_i]$:

```
? M = msinit(11,8); \\ M_8(Gamma_0(11))
? [g,R] = mspathgens(M);
? g
%3 = [[+oo, 0], [0, 1/3], [1/3, 1/2]] \\ 3 paths
? #R \\ a single relation
%4 = 1
```

```

? r = R[1]; #r \\ ...involving all 3 generators
%5 = 3
? r[1]
%6 = [[1, 1; [1, 1; 0, 1], -1], 1]
? r[2]
%7 = [[1, 1; [7, -2; 11, -3], -1], 2]
? r[3]
%8 = [[1, 1; [8, -3; 11, -4], -1], 3]

```

The given relation is of the form $\sum_i (1 - \gamma_i)g_i = 0$, with $\gamma_i \in \Gamma_0(11)$. There will always be a single relation involving all generators (corresponding to a round trip along all cusps), then relations involving a single generator (corresponding to 2 and 3-torsion elements in the group):

```

? M = msinit(2,8); \\ M_8(Gamma_0(2))
? [g,R] = mspathgens(M);
? g
%3 = [[+oo, 0], [0, 1]]

```

Note that the output depends only on the group G , not on the representation V .

The library syntax is `GEN mspathgens(GEN M)`.

3.19.25 mspathlog(M, p). Let $\Delta_0 := \text{Div}^0(\mathbf{P}^1(\mathbf{Q}))$. Let M being a full modular symbol space, as given by `msinit`, encoding fixed $\mathbf{Z}[G]$ -generators (g_i) of Δ_0 (see `mspathgens`). A path $p = [a, b]$ between two elements in $\mathbf{P}^1(\mathbf{Q})$ corresponds to $[b] - [a] \in \Delta_0$. The path extremities a and b may be given as `t_INT`, `t_FRAC` or `oo` $= (1 : 0)$. Finally, we also allow to input a path as a 2×2 integer matrix, whose first and second column give a and b respectively, with the convention $[x, y]^\sim = (x : y)$ in $\mathbf{P}^1(\mathbf{Q})$.

Returns (p_i) in $\mathbf{Z}[G]$ such that $p = \sum_i p_i g_i$.

```

? M = msinit(2,8); \\ M_8(Gamma_0(2))
? [g,R] = mspathgens(M);
? g
%3 = [[+oo, 0], [0, 1]]
? p = mspathlog(M, [1/2, 2/3]);
? p[1]
%5 =
[[1, 0; 2, 1] 1]
? p[2]
%6 =
[[1, 0; 0, 1] 1]
[[3, -1; 4, -1] 1]
? mspathlog(M, [1,2;2,3]) == p \\ give path via a 2x2 matrix
%7 = 1

```

Note that the output depends only on the group G , not on the representation V .

The library syntax is `GEN mspathlog(GEN M, GEN p)`.

3.19.26 mspetersson($M, \{F\}, \{G = F\}$). M being a full modular symbol space for $\Gamma = \Gamma_0(N)$, as given by **msinit**, calculate the intersection product $\{F, G\}$ of modular symbols F and G on $M = \text{Hom}_\Gamma(\Delta_0, V_k)$ extended to an hermitian bilinear form on $M \otimes \mathbf{C}$ whose radical is the Eisenstein subspace of M .

Suppose that f_1 and f_2 are two parabolic forms. Let F_1 and F_2 be the attached modular symbols

$$F_i(\delta) = \int_{\delta} f_i(z) \cdot (zX + Y)^{k-2} dz$$

and let $F_1^{\mathbf{R}}, F_2^{\mathbf{R}}$ be the attached real modular symbols

$$F_i^{\mathbf{R}}(\delta) = \int_{\delta} \Re(f_i(z) \cdot (zX + Y)^{k-2} dz)$$

Then we have

$$\{F_1^{\mathbf{R}}, F_2^{\mathbf{R}}\} = -2(2i)^{k-2} \cdot \Im(\langle f_1, f_2 \rangle_{Petersson})$$

and

$$\{F_1, \bar{F}_2\} = (2i)^{k-2} \langle f_1, f_2 \rangle_{Petersson}$$

In weight 2, the intersection product $\{F, G\}$ has integer values on the $\mathbf{Z}$ -structure on M given by **mslattice** and defines a Riemann form on $H_{par}^1(\Gamma, \mathbf{R})$.

For user convenience, we allow F and G to be matrices and return the attached Gram matrix. If F is omitted: treat it as the full modular space attached to M ; if G is omitted, take it equal to F .

```
? M = msinit(37,2);
? C = mscuspidal(M)[1];
? mspetersson(M, C)
%3 =
[ 0 -17 -8 -17]
[17  0 -8 -25]
[ 8  8  0 -17]
[17 25 17  0]
? mspetersson(M, mslattice(M,C))
%4 =
[0 -1 0 -1]
[1  0 0 -1]
[0  0 0 -1]
[1  1 1  0]
? E = ellinit("33a1");
? [M,xpm] = msfromell(E); [xp,xm,L] = xpm;
? mspetersson(M, mslattice(M,L))
%7 =
[0 -3]
[3  0]
? ellmoddegree(E)
%8 = [3, -126]
```

The coefficient 3 in the matrix is the degree of the modular parametrization.

The library syntax is **GEN mspetersson(GEN M, GEN F = NULL, GEN G = NULL)**.

3.19.27 mspolygon($M, \{flag = 0\}$). M describes a subgroup G of finite index in the modular group $PSL_2(\mathbf{Z})$, as given by **msinit** or a positive integer N (encoding the group $G = \Gamma_0(N)$), or by **msfarey** (arbitrary subgroup). Return an hyperbolic polygon (Farey symbol) attached to G . More precisely:

- Its vertices are an ordered list in $\mathbf{P}^1(\mathbf{Q})$ and contain a representatives of all cusps.
- Its edges are hyperbolic arcs joining two consecutive vertices; each edge e is labelled by an integer $\mu(e) \in \{\infty, 2, 3\}$.
- Given a path (a, b) between two elements of $\mathbf{P}^1(\mathbf{Q})$, let $\overline{(a, b)} = (b, a)$ be the opposite path. There is an involution $e \rightarrow e^*$ on the edges. We have $\mu(e) = \infty$ if and only if $e \neq e^*$; when $\mu(e) \neq 3$, e is G -equivalent to $\overline{e^*}$, i.e. there exists $\gamma_e \in G$ such that $e = \gamma_e \overline{e^*}$; if $\mu(e) = 3$ there exists $\gamma_e \in G$ of order 3 such that the hyperbolic triangle $(e, \gamma_e e, \gamma_e^2 e)$ is invariant by γ_e . In all cases, to each edge we have attached $\gamma_e \in G$ of order $\mu(e)$.

The polygon is given by a triple $[E, A, g]$

- The list E of its consecutive edges as matrices in $M_2(\mathbf{Z})$.
- The permutation A attached to the involution: if $e = E[i]$ is the i -th edge, then $A[i]$ is the index of e^* in E .
- The list g of pairing matrices γ_e . Remark that $\gamma_{e^*} = \gamma_e^{-1}$ if $\mu(e) \neq 3$, i.e., $g[i]^{-1} = g[A[i]]$ whenever $i \neq A[i]$ ($\mu(g[i]) = 1$) or $\mu(g[i]) = 2$ ($g[i]^2 = 1$). Modulo these trivial relations, the pairing matrices form a system of independant generators of G . Note that γ_e is elliptic if and only if $e^* = e$.

The above data yields a fundamental domain for G acting on Poincaré's half-plane: take the convex hull of the polygon defined by

- The edges in E such that $e \neq e^*$ or $e^* = e$, where the pairing matrix γ_e has order 2;
- The edges (r, t) and (t, s) where the edge $e = (r, s) \in E$ is such that $e = e^*$ and γ_e has order 3 and the triangle (r, t, s) is the image of $(0, \exp(2i\pi/3), \infty)$ by some element of $PSL_2(\mathbf{Q})$ formed around the edge.

Binary digits of flag mean:

1: return a normalized hyperbolic polygon if set, else a polygon with unimodular edges (matrices of determinant 1). A polygon is normalized in the sense of compact orientable surfaces if the distance $d(a, a^*)$ between an edge a and its image by the involution a^* is less than 2, with equality if and only if a is *linked* with another edge b (a, b, a^* et b^* appear consecutively in E up to cyclic permutation). In particular, the vertices of all edges such that that $d(a, a^*) \neq 1$ (distance is 0 or 2) are all equivalent to 0 modulo G . The external vertices of aa^* such that $d(a, a^*) = 1$ are also equivalent to 0; the internal vertices $a \cap a^*$ (a single point), together with 0, form a system of representatives of the cusps of $G \backslash \mathbf{P}^1(\mathbf{Q})$. This is useful to compute the homology group $H_1(G, \mathbf{Z})$ as it gives a symplectic basis for the intersection pairing. In this case, the number of parabolic matrices (trace 2) in the system of generators G is $2(t-1)$, where t is the number of non equivalent cusps for G . This is currently only implemented for $G = \Gamma_0(N)$.

2: add graphical representations (in LaTeX form) for the hyperbolic polygon in Poincaré's half-space and the involution $a \rightarrow a^*$ of the Farey symbol. The corresponding character strings can be included in a LaTeX document provided the preamble contains `\usepackage{tikz}`.

```
? [V,A,g] = mspolygon(3);
? V
```

```

%2 = [[-1, 1; -1, 0], [1, 0; 0, 1], [0, 1; -1, 1]]
? A
%3 = Vecsmall([2, 1, 3])
? g
%4 = [[-1, -1; 0, -1], [1, -1; 0, 1], [1, -1; 3, -2]]
? [V,A,g, D1,D2] = mspolygon(11,2); \\ D1 and D2 contains pictures
? {write("F.tex",
      "\\documentclass{article}\\usepackage{tikz}\\begin{document}"
      D1, "\\n", D2,
      "\\end{document}");}
? [V1,A1] = mspolygon(6,1); \\ normalized
? V1
%8 = [[-1, 1; -1, 0], [1, 0; 0, 1], [0, 1; -1, 3],
      [1, -2; 3, -5], [-2, 1; -5, 2], [1, -1; 2, -1]]
? A1
%9 = Vecsmall([2, 1, 4, 3, 6, 5])
? [V0,A0] = mspolygon(6); \\ not normalized V[3]^* = V[6], d(V[3],V[6]) = 3
? A0
%11 = Vecsmall([2, 1, 6, 5, 4, 3])
? [V,A] = mspolygon(14, 1);
? A
%13 = Vecsmall([2, 1, 4, 3, 6, 5, 9, 10, 7, 8])

```

One can see from this last example that the (normalized) polygon has the form

$$(a_1, a_1^*, a_2, a_2^*, a_3, a_3^*, a_4, a_5, a_4^*, a_5^*),$$

that $X_0(14)$ is of genus 1 (in general the genus is the number of blocks of the form aba^*b^*), has no elliptic points (A has no fixed point) and 4 cusps (number of blocks of the form aa^* plus 1). The vertices of edges a_4 and a_5 all project to 0 in $X_0(14)$: the paths a_4 and a_5 project as loops in $X_0(14)$ and give a symplectic basis of the homology $H_1(X_0(14), \mathbf{Z})$.

```

? [V,A] = mspolygon(15);
? apply(matdet, V) \\ all unimodular
%2 = [1, 1, 1, 1, 1, 1, 1, 1, 1, 1, 1]
? [V,A] = mspolygon(15,1);
? apply(matdet, V) \\ normalized polygon but no longer unimodular edges
%4 = [1, 1, 1, 1, 2, 2, 47, 11, 47, 11]

```

The library syntax is GEN mspolygon(GEN M, long flag).

3.19.28 msqexpansion($M, \text{proj}H, \{B = \text{seriesprecision}\}$). M being a full modular symbol space, as given by `msinit`, and $\text{proj}H$ being a projector on a Hecke-simple subspace (as given by `mssplit`), return the Fourier coefficients $a_n, n \leq B$ of the corresponding normalized newform. If B is omitted, use `seriesprecision`.

This function uses a naive $O(B^2 d^3)$ algorithm, where $d = O(kN)$ is the dimension of $M_k(\Gamma_0(N))$.

```
? M = msinit(11,2, 1); \\ M_2(Gamma_0(11))^+
? L = mssplit(M, msnew(M));
? msqexpansion(M,L[1], 20)
%3 = [1, -2, -1, 2, 1, 2, -2, 0, -2, -2, 1, -2, 4, 4, -1, -4, -2, 4, 0, 2]
? ellan(ellinit("11a1"), 20)
%4 = [1, -2, -1, 2, 1, 2, -2, 0, -2, -2, 1, -2, 4, 4, -1, -4, -2, 4, 0, 2]
```

The shortcut `msqexpansion(M, s, B)` is available for a symbol s , provided it is a Hecke eigenvector:

```
? E = ellinit("11a1");
? [M,S] = msfromell(E); [sp,sm] = S;
? msqexpansion(M,sp,10) \\ in the + eigenspace
%3 = [1, -2, -1, 2, 1, 2, -2, 0, -2, -2]
? msqexpansion(M,sm,10) \\ in the - eigenspace
%4 = [1, -2, -1, 2, 1, 2, -2, 0, -2, -2]
? ellan(E, 10)
%5 = [1, -2, -1, 2, 1, 2, -2, 0, -2, -2]
```

The library syntax is `GEN msqexpansion(GEN M, GEN projH, long precdl)`.

3.19.29 mssplit($M, \{H\}, \{\text{dimlim}\}$). Let M denote a full modular symbol space, as given by `msinit`($N, k, 1$) or `msinit`($N, k, -1$) and let H be a Hecke-stable subspace of `msnew`(M) (the full new subspace if H is omitted). This function splits H into Hecke-simple subspaces. If `dimlim` is present and positive, restrict to subspaces of dimension $\leq \text{dimlim}$. A subspace is given by a structure allowing quick projection and restriction of linear operators; its first component is a matrix with integer coefficients whose columns form a $\mathbf{Q}$ -basis of the subspace.

```
? M = msinit(11,8, 1); \\ M_8(Gamma_0(11))^+
? L = mssplit(M); \\ split msnew(M)
? #L
%3 = 2
? f = msqexpansion(M,L[1],5); f[1].mod
%4 = x^2 + 8*x - 44
? lift(f)
%5 = [1, x, -6*x - 27, -8*x - 84, 20*x - 155]
? g = msqexpansion(M,L[2],5); g[1].mod
%6 = x^4 - 558*x^2 + 140*x + 51744
```

To a Hecke-simple subspace corresponds an orbit of (normalized) newforms, defined over a number field. In the above example, we printed the polynomials defining the said fields, as well as the first 5 Fourier coefficients (at the infinite cusp) of one such form.

The library syntax is `GEN mssplit(GEN M, GEN H = NULL, long dimlim)`.

3.19.30 msstar($M, \{H\}$). M being a full modular symbol space, as given by `msinit`, return the matrix of the $*$ involution, induced by complex conjugation, acting on the (stable) subspace H (M if omitted).

```
? M = msinit(11,2); \\ M_2(Gamma_0(11))
? w = msstar(M);
? w^2 == 1
%3 = 1
```

The library syntax is `GEN msstar(GEN M, GEN H = NULL)`.

3.19.31 mstooms(Mp, ϕ). Given Mp from `mspadicinit`, lift the (classical) eigen symbol ϕ to a p -adic distribution-valued overconvergent symbol in the sense of Pollack and Stevens. More precisely, let ϕ belong to the space W of modular symbols of level N , $v_p(N) \leq 1$, and weight k which is an eigenvector for the Hecke operator $T_N(p)$ for a nonzero eigenvalue a_p and let $N_0 = \text{lcm}(N, p)$.

Under the action of $T_{N_0}(p)$, ϕ generates a subspace W_ϕ of dimension 1 (if $p \mid N$) or 2 (if p does not divide N) in the space of modular symbols of level N_0 .

Let $V_p = [p, 0; 0, 1]$ and $C_p = [a_p, p^{k-1}; -1, 0]$. When p does not divide N and a_p is divisible by p , `mstooms` returns the lift Φ of $(\phi, \phi|_k V_p)$ such that

$$T_{N_0}(p)\Phi = C_p\Phi$$

When p does not divide N and a_p is not divisible by p , `mstooms` returns the lift Φ of $\phi - \alpha^{-1}\phi|_k V_p$ which is an eigenvector of $T_{N_0}(p)$ for the unit eigenvalue where $\alpha^2 - a_p\alpha + p^{k-1} = 0$.

The resulting overconvergent eigensymbol can then be used in `mspadicmoments`, then `mspadicL` or `mspadicseries`.

```
? M = msinit(3,6, 1); p = 5;
? Tp = mshecke(M, p); factor(charpoly(Tp))
%2 =
[x - 3126 2]
[ x - 6 1]
? phi = matker(Tp - 6)[,1] \\ generator of p-Eigenspace, a_p = 6
%3 = [5, -3, -1]~
? Mp = mspadicinit(M, p, 10, 0); \\ restrict to ordinary symbols, mod p^10
? PHI = mstooms(Mp, phi);
? mu = mspadicmoments(Mp, PHI);
? mspadicL(mu)
%7 = 5 + 2*5^2 + 2*5^3 + ...
```

A non ordinary symbol.

```
? M = msinit(4,6,1); p = 3;
? Tp = mshecke(M, p); factor(charpoly(Tp))
%2 =
[x - 244 3]
[ x + 12 1]
? phi = matker(Tp + 12)[,1] \\ a_p = -12 is divisible by p = 3
%3 = [-1/32, -1/4, -1/32, 1]~
```

```

? msissymbol(M,phi)
%4 = 1
? Mp = mspadicinit(M,3,5,0);
? PHI = mstooms(Mp,phi);
***   at top-level: PHI=mstooms(Mp,phi)
***   ^-----
*** mstooms: incorrect type in mstooms [v_p(ap) > mspadicinit flag] (t_VEC).
? Mp = mspadicinit(M,3,5,1);
? PHI = mstooms(Mp,phi);

```

The library syntax is GEN mstooms(GEN Mp, GEN phi).

3.20 Plotting functions.

Although plotting is not even a side purpose of PARI, a number of plotting functions are provided. There are three types of graphic functions.

3.20.1 High-level plotting functions. (all the functions starting with `plot`) in which the user has little to do but explain what type of plot he wants, and whose syntax is similar to the one used in the preceding section.

3.20.2 Low-level plotting functions. (called *rectplot* functions, sharing the prefix `plot`), where every drawing primitive (point, line, box, etc.) is specified by the user. These low-level functions work as follows. You have at your disposal 16 virtual windows which are filled independently, and can then be physically ORed on a single window at user-defined positions. These windows are numbered from 0 to 15, and must be initialized before being used by the function `plotinit`, which specifies the height and width of the virtual window (called a *rectwindow* in the sequel). At all times, a virtual cursor (initialized at [0,0]) is attached to the window, and its current value can be obtained using the function `plotcursor`.

A number of primitive graphic objects (called *rect* objects) can then be drawn in these windows, using a default color attached to that window (which can be changed using the `plotcolor` function) and only the part of the object which is inside the window will be drawn, with the exception of polygons and strings which are drawn entirely. The ones sharing the prefix `plotr` draw relatively to the current position of the virtual cursor, the others use absolute coordinates. Those having the prefix `plotrecth` put in the rectwindow a large batch of rect objects corresponding to the output of the related `plot` function.

Finally, the actual physical drawing is done using `plotdraw`. The rectwindows are preserved so that further drawings using the same windows at different positions or different windows can be done without extra work. To erase a window, use `plotkill`. It is not possible to partially erase a window: erase it completely, initialize it again, then fill it with the graphic objects that you want to keep.

In addition to initializing the window, you may use a scaled window to avoid unnecessary conversions. For this, use `plotscale`. As long as this function is not called, the scaling is simply the number of pixels, the origin being at the upper left and the *y*-coordinates going downwards.

Plotting functions are platform independent, but a number of graphical drivers are available for screen output: X11-windows (including Openwindows and Motif), Windows's Graphical Device Interface, and the FLTK graphical libraries and one may even write the graphical objects to a

PostScript or SVG file and use an external viewer to open it. The physical window opened by `plotdraw` or any of the `plot*` functions is completely separated from `gp` (technically, a `fork` is done, and all memory unrelated to the graphics engine is immediately freed in the child process), which means you can go on working in the current `gp` session, without having to kill the window first. This window can be closed, enlarged or reduced using the standard window manager functions. No zooming procedure is implemented though.

3.20.3 Functions for PostScript or SVG output. in the same way that `printtex` allows you to have a $\text{\TeX}$ output corresponding to printed results, the functions `plotexport`, `plotexport` and `plotdrawexport` convert a plot to a character string in either PostScript or Scalable Vector Graphics format. This string can then be written to a file in the customary way, using `write`. These export routines are available even if no Graphic Library is.

3.20.4 `parplot`($X = a, b, \textit{expr}, \{\textit{flags} = 0\}, \{n = 0\}$). Parallel version of `plot`. High precision plot of the function $y = f(x)$ represented by the expression *expr*, x going from a to b . This opens a specific window (which is killed whenever you click on it), and returns a four-component vector giving the coordinates of the bounding box in the form $[xmin, xmax, ymin, ymax]$.

Important note. `parplot` may evaluate `expr` thousands of times; given the relatively low resolution of plotting devices, few significant digits of the result will be meaningful. Hence you should keep the current precision to a minimum (e.g. 9) before calling this function.

The parameter n specifies the number of reference point on the graph, where a value of 0 means we use the hardwired default values; the binary digits of *flag* have the same meaning as in `plot`: 1 = Parametric; 2 = Recursive; 4 = no_Rescale; 8 = no_X_axis; 16 = no_Y_axis; 32 = no_Frame; 64 = no_Lines; 128 = Points_too; 256 = Splines; 512 = no_X_ticks; 1024 = no_Y_ticks; 2048 = Same_ticks; 4096 = Complex.

For instance:

```
\\ circle
parplot(X=0,2*Pi,[sin(X),cos(X)], "Parametric")
\\ two entwined sinusoidal curves
parplot(X=0,2*Pi,[sin(X),cos(X)])
\\ circle cut by the line y = x
parplot(X=0,2*Pi,[X,X,sin(X),cos(X)], "Parametric")
\\ circle
parplot(X=0,2*Pi,exp(I*X), "Complex")
\\ circle cut by the line y = x
parplot(X=0,2*Pi,[(1+I)*X,exp(I*X)], "Complex")
```

The library syntax is `parplot`(GEN a, GEN b, GEN code, long flag, long n, long prec).

3.20.5 parplothexport(*fmt*, $X = a, b$, *expr*, {*flags* = 0}, {*n* = 0}). Parallel version of **plothexport**. Plot of expression *expr*, X goes from a to b in high resolution, returning the resulting picture as a character string which can then be written to a file.

The format **fmt** is either "ps" (PostScript output) or "svg" (Scalable Vector Graphics). All other parameters and flags are as in **ploth**.

```
? s = parplothexport("svg", x=1,10, x^2+3);
? write("graph.svg", s);
```

The above only works if **graph.svg** does not already exist, otherwise **write** will append to the existing file and produce an invalid **svg**. Here is a version that truncates an existing file (beware!):

```
? n = fopen("graph.svg", "w");
? fwrite(n, s);
? fclose(n);
```

This is intentionally more complicated.

The library syntax is **parplothexport**(GEN *fmt*, GEN *a*, GEN *b*, GEN *code*, long *flags*, long *n*, long *prec*),

3.20.6 plot($X = a, b$, *expr*, {*Ymin*}, {*Ymax*}). Crude ASCII plot of the function represented by expression *expr* from a to b , with Y ranging from *Ymin* to *Ymax*. If *Ymin* (resp. *Ymax*) is not given, the minimum (resp. the maximum) of the computed values of the expression is used instead.

The library syntax is **pariplot**(void *E, GEN (*eval)(void*, GEN), GEN *a*, GEN *b*, GEN *ymin*, GEN *ymax*, long *prec*)

3.20.7 plotarc($w, x2, y2$, {*filled* = 0}). Let $(x1, y1)$ be the current position of the virtual cursor. Draws in the rectwindow w the outline of the ellipse that fits inside the box such that the points $(x1, y1)$ and $(x2, y2)$ are opposite corners. The virtual cursor does *not* move. If *filled* = 1, fills the ellipse.

```
? plotinit(1);plotmove(1,0,0);
? plotarc(1,50,50); plotdraw([1,100,100]);
```

The library syntax is void **plotarc**(long *w*, GEN *x2*, GEN *y2*, long *filled*).

3.20.8 plotbox($w, x2, y2$, {*filled* = 0}). Let $(x1, y1)$ be the current position of the virtual cursor. Draw in the rectwindow w the outline of the rectangle which is such that the points $(x1, y1)$ and $(x2, y2)$ are opposite corners. Only the part of the rectangle which is in w is drawn. The virtual cursor does *not* move. If *filled* = 1, fill the box.

The library syntax is void **plotbox**(long *w*, GEN *x2*, GEN *y2*, long *filled*).

3.20.9 plotclip(w). 'clips' the content of rectwindow w , i.e remove all parts of the drawing that would not be visible on the screen. Together with **plotcopy** this function enables you to draw on a scratchpad before committing the part you're interested in to the final picture.

The library syntax is void **plotclip**(long *w*).

3.20.10 plotcolor(*w*, *c*). Set default color to *c* in rectwindow *w*. Return [R,G,B] value attached to color. Possible values for *c* are

- a `t_VEC` or `t_VECSMALL` [*R*, *G*, *B*] giving the color RGB value (all 3 values are between 0 and 255), e.g. [250,235,215] or equivalently [0xfa, 0xeb, 0xd7] for `antiquewhite`;

- a `t_STR` giving a valid colour name (see the `rgb.txt` file in X11 distributions), e.g. "`antiquewhite`" or an RGV value given by a # followed by 6 hexadecimal digits, e.g. "`#faebd7`" for `antiquewhite`;

- a `t_INT`, an index in the `graphcolormap` default, factory setting are

0=white, 1=black, 2=blue, 3=violetred, 4=red, 5=green, 6=grey, 7=gainsborough

and the color index is a non-negative integer in [0, 7]. But this can be changed (see `??graphcolormap`); note that for historical reasons, `graphcolormap` is 0-based, so the color *c* is a non-negative integer, strictly less than the length of the colormap.

```
? plotinit(0,100,100);
? plotcolor(0, "turquoise")
%2 = [64, 224, 208]
? plotbox(0, 50,50,1);
? plotmove(0, 50,50);
? plotcolor(0, 2) \\ blue
%4 = [0, 0, 255]
? plotbox(0, 50,50,1);
? plotdraw(0);
```

The library syntax is `GEN plotcolor(long w, GEN c)`.

3.20.11 plotcopy(*sourcew*, *destw*, *dx*, *dy*, {*flag* = 0}). Copy the contents of rectwindow *sourcew* to rectwindow *destw* with offset (*dx*,*dy*). If *flag*'s bit 1 is set, *dx* and *dy* express fractions of the size of the current output device, otherwise *dx* and *dy* are in pixels. *dx* and *dy* are relative positions of northwest corners if other bits of *flag* vanish, otherwise of: 2: southwest, 4: southeast, 6: northeast corners.

The library syntax is `void plotcopy(long sourcew, long destw, GEN dx, GEN dy, long flag)`.

3.20.12 plotcursor(*w*). Give as a 2-component vector the current (scaled) position of the virtual cursor corresponding to the rectwindow *w*.

The library syntax is `GEN plotcursor(long w)`.

3.20.13 plotdraw(*w*, {*flag* = 0}). Physically draw the rectwindow *w*. More generally, *w* can be of the form [*w*₁, *x*₁, *y*₁, *w*₂, *x*₂, *y*₂, ...] (number of components must be divisible by 3; the windows *w*₁, *w*₂, etc. are physically placed with their upper left corner at physical position (*x*₁, *y*₁), (*x*₂, *y*₂), ... respectively, and are then drawn together. Overlapping regions will thus be drawn twice, and the windows are considered transparent. Then display the whole drawing in a window on your screen. If *flag* ≠ 0, *x*₁, *y*₁ etc. express fractions of the size of the current output device

The library syntax is `void plotdraw(GEN w, long flag)`.

3.20.14 plotexport(*fmt*, *list*, {*flag* = 0}). Draw list of rectwindows as in `plotdraw(list, flag)`, returning the resulting picture as a character string which can then be written to a file. The format *fmt* is either "ps" (PostScript output) or "svg" (Scalable Vector Graphics).

```
? plotinit(0, 100, 100);
? plotbox(0, 50, 50);
? plotcolor(0, 2);
? plotbox(0, 30, 30);
? plotdraw(0); \\ watch result on screen
? s = plotexport("svg", 0);
? write("graph.svg", s); \\ dump result to file
```

The library syntax is GEN `plotexport(GEN fmt, GEN list, long flag)`.

3.20.15 ploth($X = a, b, \text{expr}, \{\text{flag} = 0\}, \{n = 0\}$). High precision plot of the function $y = f(x)$ represented by the expression *expr*, x going from a to b . This opens a specific window (which is killed whenever you click on it), and returns a four-component vector giving the coordinates of the bounding box in the form $[xmin, xmax, ymin, ymax]$.

Important note. `ploth` may evaluate *expr* thousands of times; given the relatively low resolution of plotting devices, few significant digits of the result will be meaningful. Hence you should keep the current precision to a minimum (e.g. 9) before calling this function.

n specifies the number of reference point on the graph, where a value of 0 means we use the hardwired default values (1000 for general plot, 1500 for parametric plot, and 8 for recursive plot).

If no *flag* is given, *expr* is either a scalar expression $f(X)$, in which case the plane curve $y = f(X)$ will be drawn, or a vector $[f_1(X), \dots, f_k(X)]$, and then all the curves $y = f_i(X)$ will be drawn in the same window.

The binary digits of *flag* mean:

- **1 = Parametric:** *parametric plot*. Here *expr* must be a vector with an even number of components. Successive pairs are then understood as the parametric coordinates of a plane curve. Each of these are then drawn.

For instance:

```
ploth(X=0,2*Pi,[sin(X),cos(X)], "Parametric")
ploth(X=0,2*Pi,[sin(X),cos(X)])
ploth(X=0,2*Pi,[X,X,sin(X),cos(X)], "Parametric")
```

draw successively a circle, two entwined sinusoidal curves and a circle cut by the line $y = x$.

- **2 = Recursive:** *recursive plot*. If this is set, only *one* curve can be drawn at a time, i.e. *expr* must be either a two-component vector (for a single parametric curve, and the parametric flag *has* to be set), or a scalar function. The idea is to choose pairs of successive reference points, and if their middle point is not too far away from the segment joining them, draw this as a local approximation to the curve. Otherwise, add the middle point to the reference points. This is fast, and usually more precise than usual plot. Compare the results of

```
\pb 32
ploth(X=-1,1, sin(1/X))
ploth(X=-1,1, sin(1/X), "Recursive")
```

for instance. Note that this example is pathological as it is impossible to evaluate $\sin(1/X)$ close to 0. It is better to avoid the singularity as follows.

```
plotth(X=1e-10,1, sin(1/X), "Recursive")
```

Beware that if you are extremely unlucky, or choose too few reference points, you may draw some nice polygon bearing little resemblance to the original curve. For instance you should *never* plot recursively an odd function in a symmetric interval around 0. Try

```
plotth(x = -20, 20, sin(x), "Recursive")
```

to see why. Hence, it's usually a good idea to try and plot the same curve with slightly different parameters.

The other values toggle various display options:

- 4 = `no_Rescale`: do not rescale plot according to the computed extrema. This is used in conjunction with `plotscale` when graphing multiple functions on a rectwindow (as a `plotrecth` call):

```
s = plothsizes();
plotinit(0, s[2]-1, s[2]-1);
plotscale(0, -1,1, -1,1);
plotrecth(0, t=0,2*Pi, [cos(t),sin(t)], "Parametric|no_Rescale")
plotdraw([0, -1,1]);
```

This way we get a proper circle instead of the distorted ellipse produced by

```
plotth(t=0,2*Pi, [cos(t),sin(t)], "Parametric")
```

- 8 = `no_X_axis`: do not print the x -axis.
- 16 = `no_Y_axis`: do not print the y -axis.
- 32 = `no_Frame`: do not print frame.
- 64 = `no_Lines`: only plot reference points, do not join them.
- 128 = `Points_too`: plot both lines and points.
- 256 = `Splines`: use splines to interpolate the points.
- 512 = `no_X_ticks`: plot no x -ticks.
- 1024 = `no_Y_ticks`: plot no y -ticks.
- 2048 = `Same_ticks`: plot all ticks with the same length.
- 4096 = `Complex`: is a parametric plot but where each member of `expr` is considered a complex number encoding the two coordinates of a point. For instance:

```
plotth(X=0,2*Pi,exp(I*X), "Complex")
plotth(X=0,2*Pi,[(1+I)*X,exp(I*X)], "Complex")
```

will draw respectively a circle and a circle cut by the line $y = x$.

- 8192 = `no_MinMax`: do not print the boundary numbers (in both directions).

The library syntax is `plotth(void *E, GEN (*eval)(void*, GEN), GEN a, GEN b, long flag, long n, long prec)`,

3.20.16 plotlexport(*fmt*, $X = a, b$, *expr*, {*flags* = 0}, {*n* = 0}). Plot of expression *expr*, X goes from a to b in high resolution, returning the resulting picture as a character string which can then be written to a file.

The format *fmt* is either "ps" (PostScript output) or "svg" (Scalable Vector Graphics). All other parameters and flags are as in **plot**.

```
? s = plotlexport("svg", x=1,10, x^2+3);
? write("graph.svg", s);
```

The library syntax is **plotlexport**(GEN *fmt*, void **E*, GEN (**eval*)(void*, GEN), GEN *a*, GEN *b*, long *flags*, long *n*, long *prec*),

3.20.17 plotdraw(X, Y , {*flag* = 0}). Given X and Y two vectors of equal length, plots (in high precision) the points whose (x, y) -coordinates are given in X and Y . Automatic positioning and scaling is done, but with the same scaling factor on x and y . If *flag* is 1, join points, other nonzero flags toggle display options and should be combinations of bits 2^k , $k \geq 3$ as in **plot**.

The library syntax is GEN **plotdraw**(GEN X , GEN Y , long *flag*).

3.20.18 plotdrawexport(*fmt*, X, Y , {*flag* = 0}). Given X and Y two vectors of equal length, plots (in high precision) the points whose (x, y) -coordinates are given in X and Y , returning the resulting picture as a character string which can then be written to a file. The format *fmt* is either "ps" (PostScript output) or "svg" (Scalable Vector Graphics).

Automatic positioning and scaling is done, but with the same scaling factor on x and y . If *flag* is 1, join points, other nonzero flags toggle display options and should be combinations of bits 2^k , $k \geq 3$ as in **plot**.

The library syntax is GEN **plotdrawexport**(GEN *fmt*, GEN X , GEN Y , long *flag*).

3.20.19 plotsizes({*flag* = 0}). Return data corresponding to the output window in the form of a 8-component vector: window width and height, sizes for ticks in horizontal and vertical directions (this is intended for the **gnuplot** interface and is currently not significant), width and height of characters, width and height of display, if applicable. If display has no sense, e.g. for svg plots or postscript plots, then width and height of display are set to 0.

If *flag* = 0, sizes of ticks and characters are in pixels, otherwise are fractions of the screen size

The library syntax is GEN **plotsizes**(long *flag*).

3.20.20 plotinit($w, \{x\}, \{y\}, \{flag = 0\}$). Initialize the rectwindow w , destroying any rect objects you may have already drawn in w . The virtual cursor is set to (0,0). The rectwindow size is set to width x and height y ; omitting either x or y means we use the full size of the device in that direction. If *flag* = 0, x and y represent pixel units. Otherwise, x and y are understood as fractions of the size of the current output device (hence must be between 0 and 1) and internally converted to pixels.

The plotting device imposes an upper bound for x and y , for instance the number of pixels for screen output. These bounds are available through the **plotsizes** function. The following sequence initializes in a portable way (i.e independent of the output device) a window of maximal size, accessed through coordinates in the $[0, 1000] \times [0, 1000]$ range:

```
s = plotsizes();
```

```
plotinit(0, s[1]-1, s[2]-1);
plotscale(0, 0,1000, 0,1000);
```

The library syntax is `void plotinit(long w, GEN x = NULL, GEN y = NULL, long flag)`

3.20.21 plotkill(*w*). Erase rectwindow *w* and free the corresponding memory. Note that if you want to use the rectwindow *w* again, you have to use `plotinit` first to specify the new size. So it's better in this case to use `plotinit` directly as this throws away any previous work in the given rectwindow.

The library syntax is `void plotkill(long w).`

3.20.22 plotlines(*w*, *X*, *Y*, {*flag* = 0}). Draw on the rectwindow *w* the polygon such that the (*x*,*y*)-coordinates of the vertices are in the vectors of equal length *X* and *Y*. For simplicity, the whole polygon is drawn, not only the part of the polygon which is inside the rectwindow. If *flag* is nonzero, close the polygon. In any case, the virtual cursor does not move.

X and *Y* are allowed to be scalars (in this case, both have to). There, a single segment will be drawn, between the virtual cursor current position and the point (*X*,*Y*). And only the part thereof which actually lies within the boundary of *w*. Then *move* the virtual cursor to (*X*,*Y*), even if it is outside the window. If you want to draw a line from (*x*₁,*y*₁) to (*x*₂,*y*₂) where (*x*₁,*y*₁) is not necessarily the position of the virtual cursor, use `plotmove(w,x1,y1)` before using this function.

The library syntax is `void plotlines(long w, GEN X, GEN Y, long flag).`

3.20.23 plotlinetype(*w*, *type*). This function is obsolete and currently a no-op.

Change the type of lines subsequently plotted in rectwindow *w*. *type* -2 corresponds to frames, -1 to axes, larger values may correspond to something else. *w* = -1 changes highlevel plotting.

The library syntax is `void plotlinetype(long w, long type).`

3.20.24 plotmove(*w*, *x*, *y*). Move the virtual cursor of the rectwindow *w* to position (*x*,*y*).

The library syntax is `void plotmove(long w, GEN x, GEN y).`

3.20.25 plotpoints(*w*, *X*, *Y*). Draw on the rectwindow *w* the points whose (*x*,*y*)-coordinates are in the vectors of equal length *X* and *Y* and which are inside *w*. The virtual cursor does *not* move. This is basically the same function as `plothraw`, but either with no scaling factor or with a scale chosen using the function `plotscale`.

As was the case with the `plotlines` function, *X* and *Y* are allowed to be (simultaneously) scalar. In this case, draw the single point (*X*,*Y*) on the rectwindow *w* (if it is actually inside *w*), and in any case *move* the virtual cursor to position (*x*,*y*).

If you draw few points in the rectwindow, they will be hard to see; in this case, you can use filled boxes instead. Compare:

```
? plotinit(0, 100,100); plotpoints(0, 50,50);
? plotdraw(0)
? plotinit(1, 100,100); plotmove(1,48,48); plotrbox(1, 4,4, 1);
? plotdraw(1)
```

The library syntax is `void plotpoints(long w, GEN X, GEN Y).`

3.20.26 plotpointsize($w, size$). This function is obsolete. It is currently a no-op.

Changes the “size” of following points in rectwindow w . If $w = -1$, change it in all rectwindows.

The library syntax is `void plotpointsize(long w, GEN size)`.

3.20.27 plotpointtype($w, type$). This function is obsolete and currently a no-op.

change the type of points subsequently plotted in rectwindow w . $type = -1$ corresponds to a dot, larger values may correspond to something else. $w = -1$ changes highlevel plotting.

The library syntax is `void plotpointtype(long w, long type)`.

3.20.28 plotrbox($w, dx, dy, \{filled\}$). Draw in the rectwindow w the outline of the rectangle which is such that the points $(x1, y1)$ and $(x1 + dx, y1 + dy)$ are opposite corners, where $(x1, y1)$ is the current position of the cursor. Only the part of the rectangle which is in w is drawn. The virtual cursor does *not* move. If $filled = 1$, fill the box.

The library syntax is `void plotrbox(long w, GEN dx, GEN dy, long filled)`.

3.20.29 plotrecth($w, X = a, b, expr, \{flag = 0\}, \{n = 0\}$). Writes to rectwindow w the curve output of `ploth`($w, X = a, b, expr, flag, n$). Returns a vector for the bounding box.

3.20.30 plotrecthraw($w, data, \{flags = 0\}$). Plot graph(s) for $data$ in rectwindow w ; $flag$ has the same meaning here as in `ploth`, though recursive plot is no longer significant.

The argument $data$ is a vector of vectors, each corresponding to a list a coordinates. If parametric plot is set, there must be an even number of vectors, each successive pair corresponding to a curve. Otherwise, the first one contains the x coordinates, and the other ones contain the y -coordinates of curves to plot.

The library syntax is `GEN plotrecthraw(long w, GEN data, long flags)`.

3.20.31 plotrline(w, dx, dy). Draw in the rectwindow w the part of the segment $(x1, y1) - (x1 + dx, y1 + dy)$ which is inside w , where $(x1, y1)$ is the current position of the virtual cursor, and move the virtual cursor to $(x1 + dx, y1 + dy)$ (even if it is outside the window).

The library syntax is `void plotrline(long w, GEN dx, GEN dy)`.

3.20.32 plotrmove(w, dx, dy). Move the virtual cursor of the rectwindow w to position $(x1 + dx, y1 + dy)$, where $(x1, y1)$ is the initial position of the cursor (i.e. to position (dx, dy) relative to the initial cursor).

The library syntax is `void plotrmove(long w, GEN dx, GEN dy)`.

3.20.33 plotrpoint(w, dx, dy). Draw the point $(x1 + dx, y1 + dy)$ on the rectwindow w (if it is inside w), where $(x1, y1)$ is the current position of the cursor, and in any case move the virtual cursor to position $(x1 + dx, y1 + dy)$.

If you draw few points in the rectwindow, they will be hard to see; in this case, you can use filled boxes instead. Compare:

```
? plotinit(0, 100,100); plotrpoint(0, 50,50); plotrpoint(0, 10,10);
? plotdraw(0)

? thickpoint(w,x,y)= plotmove(w,x-2,y-2); plotrbox(w,4,4,1);
? plotinit(1, 100,100); thickpoint(1, 50,50); thickpoint(1, 60,60);
? plotdraw(1)
```

The library syntax is `void plotrpoint(long w, GEN dx, GEN dy)`.

3.20.34 plotscale($w, x1, x2, y1, y2$). Scale the local coordinates of the rectwindow w so that x goes from $x1$ to $x2$ and y goes from $y1$ to $y2$ ($x2 < x1$ and $y2 < y1$ being allowed). Initially, after the initialization of the rectwindow w using the function `plotinit`, the default scaling is the graphic pixel count, and in particular the y axis is oriented downwards since the origin is at the upper left. The function `plotscale` allows to change all these defaults and should be used whenever functions are graphed.

The library syntax is `void plotscale(long w, GEN x1, GEN x2, GEN y1, GEN y2)`.

3.20.35 plotstring($w, x, \{flags = 0\}$). Draw on the rectwindow w the String x (see Section 2.9.1), at the current position of the cursor.

flag is used for justification: bits 1 and 2 regulate horizontal alignment: left if 0, right if 2, center if 1. Bits 4 and 8 regulate vertical alignment: bottom if 0, top if 8, v-center if 4. Can insert additional small gap between point and string: horizontal if bit 16 is set, vertical if bit 32 is set (see the tutorial for an example).

The library syntax is `void plotstring(long w, const char *x, long flags)`.

3.20.36 psdraw($list, \{flag = 0\}$). This function is obsolete, use `plotexport` and write the result to file.

The library syntax is `void psdraw(GEN list, long flag)`.

3.20.37 psplot($X = a, b, expr, \{flags = 0\}, \{n = 0\}$). This function is obsolete, use `plotexport` and write the result to file.

The library syntax is `GEN psplot0(GEN X, GEN b, GEN expr, long flags, long prec)`.

3.20.38 psplotdraw($listx, listy, \{flag = 0\}$). This function is obsolete, use `plotdrawexport` and write the result to file.

The library syntax is `GEN psplotdraw(GEN listx, GEN listy, long flag)`.

Index

SomeWord refers to PARI-GP concepts.

SomeWord is a PARI-GP keyword.

SomeWord is a generic index entry.

+

+oo 160, 163

-

-LONG_MAX 253

A

a1 522

a2 522

a3 522

a4 522

a6 522

Abelian extension 469, 480

abs 315

accuracy 9

acos 316

acosh 316

addhelp 50, 86, 87

addprimes 125, 179, 199, 436, 448, 465

adj 278

adjoint matrix 278

adjsafe 278

agm 316

airy 316

akell 527

alarm 85, 87

algadd 493

algalgtobasis 493, 494

algaut 494

algb 494

algbasis 494

algbasistoalg 493, 494, 495

algcenter 495

algcentralproj 495

algchar 496

algcharpoly 496

algdegree 496

algdep 273, 274

algdep0 274

algdim 496, 497

algdisc 497

algdivl 497

algdivr 497

algebraic dependence 267, 273

algebraic number 375, 377

alggroup 497, 498

alggroupcenter 498

alghasse 498, 499

alghassef 499

alghassei 499

algindex 499, 500

alginit 493, 494, 495, 496, 497, 498, 499, 500,
502, 503, 504, 505, 506, 510,
511, 512, 514, 516, 518, 520

alginv 502

alginvbasis 502, 503

algisassociative 503

algiscommutative 503

algisdivision 503, 504

algisdivl 504

alginv 504

algisramified 504, 505

algissemisimple 505

algissimple 505, 506

algissplit 506

alglatadd 506

alglatcontains 506, 507

alglatelement 507

alglathnf 507

alglatindex 507

alglatinter 508

alglatlefttransporter 508

alglatmul 508, 509

alglatrighttransporter 509

alglatsubset 509

algmakeintegral 510

algmul 510

algmutable 510, 511

algneg 511

algnorm 511, 512

algpoleval 512

algpow 512, 513

algprimesubalg 513

algquotient 513

algradical 513, 514

algramifiedplaces 514

algrandom 514

algrelmutable 514, 515

algsimpledec 515

algsplit 515, 516

algsplittingdata 516

algsplittingfield 516, 517

algsqr 517

algsub	517	bernreal	170
algsuvalg	517, 518	bernvec	169, 170, 171
alhtableinit	495, 496, 503, 505, 510, 511, 512, 513, 515, 517, 518, 519, 520	besselh1	317
algtensor	519	besselh2	317
algtobasis	433	besseli	317
algtomatrix	519, 520	besselj	317
algtrace	520	besseljh	317
algtype	520, 521	besseljzero	317, 318
alg_centralproj	496	besselk	318
alg_quotient	513	besseln	318
alias	50, 88	bessely	318
alias0	89	besselyzero	318
allocatemem	89, 111	bestappr	153, 179, 180
alternating series	360	bestapprnf	274
and	135	bestapprPade	180, 181
and	151	bestapprPade0	181
apply	10, 90, 91	Bezout relation	209
area	522, 524	bezout	181
arg	316	bezoutres	245
arity	91	<i>bid</i>	47, 378
arity0	91	bid	379
Artin L-function	400	bigomega	181
Artin root number	400	bilhell	529
asin	316	binaire	151
asinh	316	binary file	119
asypnum	337, 338	binary file	61, 111
asypnum0	338	binary flag	68
asypnumraw	338, 339	binary quadratic form	23, 147
asypnumraw0	339	binary	151
atan	317	binomial coefficient	171
atanh	317	binomial	171
automatic simplification	133	binomial0	172
available commands	60	Birch and Swinnerton-Dyer conjecture	536
B			
b2	522	bitand	151
b4	522	bitneg	151
b6	522	bitnegimply	152
b8	522	bitor	152
backslash character	16	bitprecision	152
basistoalg	436	bitprecision00	153
bernfrac	169, 170	bittest	153
Bernoulli numbers	169, 170, 332	bitwise and	151
Bernoulli polynomial	170	bitwise exclusive or	153
bernpol	170	bitwise inclusive or	152
bernpol_eval	170	bitwise negation	151
		bitxor	153
		<i>bnf</i>	47, 375
		bnf	379, 524
		bnfcertify	382
		bnfcertify0	382

diffop 246, 247
diffop0 247
digits 155, 156
dilog 319
dirdiv 190
direuler 190, 191
Dirichlet series 190, 191, 401
dirmul 191
dirpowers 276, 277
dirpowersssum 191, 192
dirpowersssumfun 192
dirzetak 401
disc 379, 522
divisors 72, 192, 193
divisors0 193
divisorslenstra 193
divisors_factored 193
divrem 35, 142
dvi 65
dynamic scoping 36

E

echo 60, 125
ECM 177, 199
ecpp0 219
editing characters 16
Egyptian fraction 189
eigen 282
eint1 319
elementary divisors 292
ell 47, 521, 539
ell 538
ell2cover 525
elladd 526
ellak 526
ellan 527
ellanalyticrank 526, 527, 552
ellanQ_zv 527
ellap 527, 529, 530
ellbil 529
ellbsd 529, 552
ellcard 530
ellchangecurve 530
ellchangept 530
ellchangeptinv 530, 531
ellconvertname 531, 564
elldata 73, 524, 531, 535, 538, 564
elldivpol 531

ellE 319
elleisnum 531, 532
elleta 532, 570
ellformaldifferential 532
ellformalexp 532, 533
ellformallog 532, 533
ellformalpoint 533
ellformalw 533, 534
ellfromeqn 534
ellfromj 534, 535
ellgenerators 524, 535
ellglobalred 535, 536
ellgroup 536, 537
ellgroup0 537
ellheegner 537
ellheight 538
ellheight0 538
ellheightmatrix 538
ellidentify 524, 538
ellinit 521, 524, 535, 538, 539, 540
ellintegralmodel 535, 540
elliscm 540, 541
ellisdivisible 541
ellisism 541
ellisogeny 541, 542
ellisogenyapply 542
ellisomat 542, 543
ellisoncurve 543
ellisotree 543, 544
ellissupersingular 544, 545
ellj 545
elljissupersingular 545
ellK 319, 320
ellL1 525, 526
elllocalred 545
elllog 545, 546
elllseries 546
ellmaninconstant 546
ellminimaldisc 546
ellminimalmodel 535, 546, 547
ellminimaltwist 547, 548
ellminimaltwist0 548
ellminimaltwistcond 548
ellmoddegree 548
ellmodulareqn 548, 549
ellmul 549, 560
ellneg 549
ellnonsingularmultiple 549, 550
ellorder 550

galoisidentify	406	gchar_identify	414
galoisinit	398, 404, 405, 406, 408, 445	gconcat	276
galoisisabelian	408	gconcat1	276
galoisisnormal	408	gconj	154
galoisnbpol	405, 406	gcos	319
galoispermtopol	408	gcosh	319
galoissplittinginit	408	gcotan	319
galoissubcyclo	79, 264, 400, 408, 409	gcotanh	319
galoissubfields	404, 409, 461	gcvtoi	165
galoissubgroups	409, 410	gdeflate	269
gamma	321	gdiv	137
gamma-function	321	gdivent	137
gammah	321	gdiventres	142
gammamellininv	321, 322, 593	gdivround	138
gammamellininvasymp	322	gen (member function)	379
gammamellininvinit	321, 322, 323	GEN	8
garg	316	gen	523, 524
gasin	316	genapply	91
gasinh	317	generic matrix	50
gatan	317	genfold	97
gatanh	317	genindexselect	113
gauss	293	genrand	162
gaussmodulo	294	genselect	113
gaussmodulo2	294	GENtostr	149
gbitand	151	genus2igusa	571, 572
gbitneg	151	genus2red	572, 573
gbitnegimply	152	gen_I	315
gbitor	152	gerfc	320
gbittest	153	getabstime	97, 98, 99
gbitxor	153	getcache	98
gboundcf	189	getenv	98
gcd	208	getheap	98, 99
gcdext	208, 209	getlocalbitprec	99
gcdext0	181, 210	getlocalprec	99
gceil	153	getrand	99, 162
gcf	189	getstack	99
gcf2	189	gettime	99
<i>gchar</i>	381	getwalltime	98, 99
gcharalgebraic	410, 412	geval	248
gcharconductor	412	gexp	320
gcharduallog	412, 413	gexpm1	321
gchareval	413	gexpo	157
gcharidentify	413	gfloor	157
gcharinit	414, 415, 587	gfrac	157
gcharisalgebraic	415, 416	ggamma	321
gcharlocal	416, 417	ggammah	321
gcharlog	417	ggcd	209
gcharnewprec	417	ggcd0	209
gchar_conductor	412	ggrando	245

ghalfgcd	210	gpsi_der	328
gideallist	425	gpsystem	116
gimag	157	gpvaluation	165
gisanypower	211	gpwritebin	119
gisprime	212	gp_alarm	88
gispseudoprime	212	gp_allocatemem	90
gissquare	213	gp_fileclose	94
gissquareall	213	gp_fileextern	95
glambertW	325	gp_fileflush	95
glcm0	215	gp_fileflush0	95
glength	158	gp_fileopen	95
glngamma	326	gp_fileread	96
global	99	gp_filereadstr	96
glog	326	gp_filewrite	97
glog1p	327	gp_filewrite1	97
gmax	143	gp_getenv	98
gmin	143	gp_input	99
gmod	138	gp_readvec_file	112
gmodulo	146	gp_readvec_stream	112
gmul	137	gp_read_file	111
gmul2n	143	Graeffe	255
gneg	136	graphcolormap	127, 664
gnorm	159	graphcolors	127
gnorml2	295	greal	163
gnormlp	295	GRH	272, 380, 382, 384, 385, 480
gp	5	grndtoi	163
GP	5	grootsof1	328
gp	13	<i>Grossencharacter</i>	587
gp2c	5	ground	163
gpexponent	157	group	523
gpextern	93	gshift	143
gphelp	60	gsigne	143
gpidealfactor	421	gsin	329
gpidealval	432	gsinc	329
gpinstall	101	gsinh	329
gpnfvalrem	443	gsizebyte	164
gpolve	167	gsizeword	164
gpolylog	327	gsqr	137, 329
gpow	141	gsqrt	329
gpowers	296	gsqrtn	330
gpowers0	296	gsub	137
gppadicprec	161	gsubst	269
gppoldegree	253	gsubstpol	269
gprc	13, 34, 62	gsubstvec	269
GPRC	63	gtan	330
gprc	129, 130	gtanh	330
gprec	161	gtocol	144
gpserprec	164	gtocol0	144
gpsi	328	gtocolrev	144

gtocolrev0	144
gtolist	145
gtomap	145
gtomat	146
gtopoly	147
gtopolyrev	147
gtoset	149
gtovec	150
gtovec0	150
gtovecrev	150
gtovecrev0	150
gtovecsmall	151
gtovecsmall0	151
gtrace	309
gtrans	294
gtranslength	158
gtrunc	165
gvaluation	165
gvar	167
gzeta	332
gzip	61, 119

H

Hadamard product	267
halfgcd	210
hammingweight	173, 232
harmonic	173, 174
harmonic0	174
hbessel1	317
hbessel2	317
hclassno	223
<i>heap</i>	61
help	127
Hermite normal form	282, 284, 378, 422, 446
Hermite	255
hess	282
hgmalpha	579
hgmbdegree	579
hgmcoef	579, 580
hgmcoefs	580
hgmcy clo	580
hgmeulerfactor	580, 581
hgmgamma	581
hgminit	581
hgmissymmetrical	581, 582
hgmparams	582
hgmtwist	582
Hilbert class field	226

Hilbert matrix	282
Hilbert symbol	210, 446
hilbert	210
histfile	128
histsize	16, 128
hnf	284
hnfall	284
hnfmod	284
hnfmodid	285
Householder transform	285, 289
Hurwitz class number	223
hyperellchange curve	574
hyperellcharpoly	574
hyperelldisc	574
hyperellisoncurve	574
hyperellminimaldisc	575
hyperellminimalmodel	575
hyperellordinate	575, 576
hyperellpadicfrobenius	576
hyperellpadicfrobenius0	576
hyperellratpoints	576
hyperellred	576, 577
hypergeom	323, 324
hyperu	324

I

I	20, 315
ibessel	317
ibitand	151
ibitnegimply	152
ibitor	152
ibitxor	153
ideal (extended)	376, 424, 426, 427
<i>ideal list</i>	377
<i>ideal</i>	376
idealadd	417, 418
idealaddtoone	418
idealaddtoone0	418
idealappr	418, 419
idealappr0	419
idealchinese	419, 420
idealchineseinit	420
idealcoprime	420
idealdiv	420
idealdiv0	420
idealdivexact	420
idealdivdown	420
idealfactor	397, 414, 420, 421, 431

idealfactorback	421, 422	index	379
idealfactor_limit	421	indexrank	287
idealfrobenius	422	infinite product	359
idealhnf	422, 423	infinity	354
idealhnf0	423	inline	99
idealintersect	287, 423, 424	input	99
idealinv	424, 447	install	50, 56, 99, 133
idealismaximal	424	intcirc	340, 341
idealispower	424, 425	integ	250
ideallist	425	integer	7, 8, 17
ideallist0	425	integral basis	434
ideallistarch	425, 426	<i>integral pseudo-matrix</i>	377
ideallog	241, 376, 426, 431	internal longword format	61
ideallogmod	426	internal representation	62
idealmin	426	interpolating polynomial	256
idealmul	426, 427	intersect	287
idealmul0	427	intformal	249
idealmulred	427	intfuncinit	341, 342, 350
idealnrm	427	<i>intmod</i>	7
idealnumden	427	intmod	8, 18
idealpow	427, 429	intnum	103, 104, 336, 342, 348, 353, 368, 369
idealpow0	427	intnumgauss	336, 348, 349
idealpowred	427	intnumgauss0	348
idealpows	427	intnumgaussinit	336, 348, 349
idealprimedec	427, 428, 444	intnuminit	336, 343, 349, 350
idealprimedec_limit_f	428	intnumosc	336, 350
idealprincipalunits	428	intnumosc0	353
idealramgroups	428, 429	intnumromb	336, 353, 354
idealred	376, 422, 429	int_bit	153
idealred0	430	inverse	140
idealredmodpower	430	inverseimage	287
idealstar	397, 414, 428, 431	isdiagonal	288
Idealstar	431	isfundamental	210
idealstarmod	431	isideal	450
Idealstarmod	431	ispolygonal	210
idealtwoelt	431, 432	ispower	211
idealtwoelt0	432	ispowerful	211
idealtwoelt2	432	isprime	211, 212
idealval	432	isprimepower	212
if	82	isprincipalray	399
iferr	28, 53, 70, 82, 116, 150	ispseudoprime	199, 211, 212, 216, 230
imag	157	ispseudoprimepower	212
image	286	issquare	211, 212, 213
imagecompl	286	issquareall	213
incgam	324, 325	issquarefree	177, 213, 214
incgam0	325	istotient	214
incgamc	325		
inclusive or	135		
index	379		

J

listsort	103, 308
LLL	240, 282, 288, 301, 429
lll	302
lllgram	302
lllgramint	302
lllgramkerim	302
lllint	302
lllkerim	302
Lmath	584
lngamma	325
local	36, 104
localbitprec	18, 103
localprec	104
log	59, 61, 111, 128, 326
log1p	326
logfile	111
logfile	128
logint	215
logint0	215
LONG_MAX	161, 164, 165, 432, 443
<i>lvalue</i>	29, 32
lvalue	30

M

Map	145
mapapply	105, 106
mapdelete	106
mapget	106, 107
mapisdefined	106, 107, 312
mapput	107, 145
Mat	25, 27, 145, 275, 283
matadjoint	275, 278
matadjoint0	278
matalgtobasis	432
matbasistoalg	432
matcompanion	278
matconcat	275, 278, 280, 281
matdet	280
matdetint	280
matdetmod	280, 281
matdiagonal	281
mateigen	281, 282
matfrobenius	282
Math::Pari	57
mathess	282
mathilbert	282
mathnf	272, 282
mathnf0	284

mathnfmod	284
mathnfmodid	284
mathouseholder	285
matid	286
matimage	286
matimage0	286
matimagecompl	286
matimagemod	286, 287
matindexrank	287
matintersect	287
matinverseimage	287
matinvmod	287, 288
matisdiagonal	288
matker	288
matker0	288
matkerint	288
matkerint0	288
matkermid	288, 289
matmuldiagonal	289
matmultodiagonal	289
matpascal	289
matpermanent	289
matqpascal	289
matqr	289, 290
matrank	290
matreduce	290, 291
matrix	8, 9, 25, 50
matrix	25, 291
matrixqz	291
matrixqz0	291
matsize	292
matsnf	292
matsnf0	293
matsolve	293
matsolvemod	293, 294
matsupplement	294
mattranspose	294
max	143
member functions	47, 379, 521
mfatkin	606, 607
mfatkineigenvalues	607
mfatkininit	607, 608
mfbasis	608
mfbd	608, 609
mfbracket	609
mfcoef	609
mfcoefs	609
mfconductor	609, 610
mfcosets	610

mfcuspisregular	610	mfsturm	635
mfcusps	610	mfsymbol	635, 636
mfcuspval	610, 611	mfsymboleval	636, 637
mfCUSPwidth	611	mftaylor	637
mfDelta	605	mfTheta	606
mfderiv	611	mftobasis	637, 638
mfderivE2	611, 612	mftocoset	638, 639
mfdescribe	612	mftonew	639
mfdim	612, 613	mftraceform	639
mfdiv	613, 614	mftwist	639
mfEH	605, 606	min	143
mfeigenbasis	614, 615	minim	304
mfeigensearch	615	minim2	304
mfeisenstein	615, 616	minimal model	547
mfEk	606	minimal polynomial	294
mfembed	616	minimal vector	304
mfembed0	617	minim_raw	304
mfeval	617, 618	minim_zm	304
mffields	618	minpoly	294
mffromell	618, 619	mklist	145
mffrometaquo	619	mkoo	160
mffromlfun	619, 620	Mod	146
mffromqf	620, 621	mod	379
mfgaloisprojrep	621, 622	modpr	456
mfgaloistype	622	modprinit	444
mfhecke	622, 623	modreverse	432, 433, 467
mfheckemat	623	<i>modulus</i>	378
mfinit	623, 624, 635	Moebius	177, 215
mfisCM	624	moebius	177, 215, 216
mfisequal	624, 625	Mordell-Weil group	535, 538, 562, 564
mfisetaquo	625	mpcatalan	315
mfkohnenbasis	625	mpeuler	315
mfkohnenbijection	626, 627	mpfact	138, 199
mfkohneneigenbasis	627	mpfactr	199
mflinear	627, 628	mppi	315
mfmanin	628	mpprimerial	138
mfmul	628, 629	MPQS	177, 199
mfnumcusps	629	msatkinlehner	640
mfparams	629	mscosets	640, 641
mfperiodpol	629, 630	mscosets0	641
mfperiodpolbasis	630	mscuspidal	641
mfpetersson	630, 631	msdim	641
mfpow	631	mseisenstein	641, 642
mfsearch	631, 632	mseval	642, 643
mfshift	632	msfarey	643
mfshimura	632	msfarey0	643
mfslashexpansion	632, 634	msfromcusp	643, 644
mfspace	634	msfromell	644, 645
mfsplit	634, 635	msfromhecke	645, 646

msgetlevel	646
msgetsign	646
msgetweight	646
mshecke	647
msinit	640, 647, 648
msissymbol	648
mslattice	568, 648, 649
msnew	649, 650
msomseval	650
mspadicinit	652
mspadicL	552, 650, 652, 660
mspadicmoments	650, 652, 653, 660
mspadicseries	552, 652, 653, 654, 660
mspathgens	642, 654, 655
mspathlog	655
mspetersson	655, 656
mspolygon	656, 658
msqexpansion	658, 659
mssplit	658, 659
msstar	659, 660
mstooms	651, 660, 661
multivariate polynomial	45
my	36, 40, 104

N

nbthreads	129
newtonpoly	433
new_galois_format	464, 465
next	53, 85
nextprime	216
<i>nf</i>	47, 375
nf	379, 524
nfadd	439
nfalgtobasis	432, 433
nfbasis	131, 434, 436, 438, 448
nfbasistoalg	432, 436
nfcertify	436, 437, 448, 466
nfcompositum	437, 438
nfdetint	438
nfdisc	131, 438
nfdiscfactors	438, 439
nfdiv	439
nfdivauc	439
nfdivmodpr	439
nfdivrem	440
nfeltadd	439
nfeltdiv	439
nfeltdiveuc	439
nfeltdivmodpr	439
nfeltdivrem	440
nfeltembed	440
nfeltispower	440
nfeltissquare	440
nfeltmod	441
nfeltmul	441
nfeltmulmodpr	441
nfeltnorm	441
nfeltpow	441
nfeltpowmodpr	441
nfeltreduce	441
nfeltreducemodpr	441
nfeltsign	442
nfelttrace	442
nfeltval	442
nffactor	196, 401, 443
nffactorback	376, 422, 443, 444
nffactormod	444
nfgaloisapply	444
nfgaloisconj	406, 445
nfgrunwaldwang	446
nfhilbert	446
nfhilbert0	446
nfhnf	287, 446
nfhnf0	446
nfhnfmod	446
nfhyperellpadicfrobenius	576
nfinit	131, 375, 406, 447, 449, 465, 467, 500
nfinit0	449
nfinitred	449
nfinitred2	449
nfinv	441
nfisideal	450
nfisincl	450, 451
nfisincl0	451
nfisisom	451, 452
nfislocalpower	452
nfispower	440
nfissquare	441
nfkermodpr	452
nflist	453, 455
nfmod	441
nfmodpr	455, 456
nfmodprinit	439, 441, 455, 456
nfmodprinit0	456
nfmodprlift	455, 456, 457
nfmul	441
nfmulmodpr	441

nfnewprec 448, 457
 nfnorm 441
 nfpolsturm 457
 nfpow 441
 nfpowmodpr 441
 nfreduce 441
 nfreducemodpr 442
 nfresolvent 457, 458
 nfroots 458
 nfrootsof1 458, 459
 nfrootsQ 458
 nfsnf 459
 nfsnf0 459
 nfsolvemodpr 459, 460
 nfsplitting 460, 461
 nfsplitting0 461
 nfsqr 441
 nfsubfield 405
 nfsubfields 461, 462
 nfsubfields0 461
 nfsubfieldscm 461, 462
 nfsubfieldsmax 462
 nftrace 442
 nfval 443
 nfvalrem 443
 nfweilheight 462
 nf_ADDZK 467
 nf_ALL 467
 nf_FORCE 384, 386
 nf_GEN 386, 431
 nf_GENMAT 386
 nf_INIT 431
 nf_NOLLL 450
 nf_ORIG 449, 467
 nf_PARTIALFACT 131, 467
 nf_RAW 467
 nf_RED 449
 no 379, 523
 norm 159
 norml2 294
 normlp 295
 not 135
 nucomp 223
 nudupl 223
 number field 21
 numbpart 174
 numdiv 216
 numer 160
 numerator 35, 160

numerical derivation 30
 numerical integration 336
 numerical summation 336
 numtoperm 174, 176
 nupow 224

O

O 245
 omega 226
 omega 216, 522, 524
 oncurve 543
 oo 160
 operator 29
 or 135
 or 152, 153
 orderell 550
 output 61, 129

P

p 523
 p-adic number 7, 8, 20
 padicappr 250
 padicfields 250
 padicfields0 250
 padicprec 160, 161
parametric plot 665
 parapply 119, 120
 pareval 120
 parfor 120
 parforeach 121
 parforprime 121
 parforprimestep 122
 parforstep 122
 parforstep0 122
 parforvec 122
 pariplot 663
 parisize 89, 123, 129
 parisizemax 57, 89, 123, 129, 130
 pari_err 93
 pari_malloc 84
 pari_printf 110
 pari_realloc 84
 pari_self 113
 pari_sprintf 114
 pari_strchr 86, 114
 pari_version 118
 parplot 662
 parplotexport 662, 663

parselect	122, 123	pol	379
parsum	123	polchebyshev	251
partitions	174, 175	polchebyshev1	251, 266
parvector	123	polchebyshev2	251
Pascal triangle	289	polchebyshev_eval	251
path	130	polclass	251, 252
Pauli	434	polcoef	154, 252, 253
Perl	57	polcoeff	253
Perl	36	polcompositum	462
permcycles	175	polcompositum0	463
permorder	175, 176	polcyclo	253
permsign	176	polcyclofactors	253
permtonum	174, 176	polcyclo_eval	253
Pi	315	poldegree	253
plot	663	poldisc	254
plotarc	663	poldisc0	254
plotbox	663	poldiscfactors	254
plotclip	663	poldiscreduced	254
plotcolor	127, 663, 664	polfnf	402
plotcopy	663, 664	polfromroots	254, 255
plotcursor	664	polgalois	129, 463, 465
plotdraw	664	polgraeffe	255
plotexport	664, 665	polhensellift	255
plotth	68, 665, 666	polhermite	255
plotthexport	666, 667	polhermite_eval	255
plotthraw	667	polhermite_eval0	255
plotthrawexport	667	polint	258
plotsizes	130, 667	polinterpolate	256
plotinit	667	polisclass	258
plotkill	668	poliscyclo	258
plotlines	668	poliscycloprod	258
plotlinetype	668	polisirreducible	258
plotmove	668	pollaguerre	258, 259
plotpoints	668	pollaguerre_eval	259
plotpointsize	668, 669	pollaguerre_eval0	259
plotpointtype	669	Pollard Rho	177, 199
plotrbox	669	pollead	259
plotrecth	666, 669	pollegendre	259
plotrecthraw	669	pollegendre_eval	259
plotrline	669	pollegendre_eval0	259
plotrmove	669	<i>polmod</i>	7
plotrpoint	669	polmod	8, 21
plotscale	666, 670	polmodular	259, 260, 548
plotstring	50, 670	polrecip	260
plotterm	50	polred	465
pnqn	190	polred2	465
pointell	571	polredabs	465, 466, 467
<i>pointer</i>	68	polredabs0	467
Pol	22, 146, 147	polredbest	465, 466, 467, 468

Shanks	221, 223	strsplit	114, 115
shift	143	Strtex	86
shiftmul	143	strtex	86, 115, 119
short Weierstrass equation	521	strtime	115
sigma	190, 231	strtoGEN	149
sign	143	sturm	264
sign	143, 379	sturmpart	264
signunits	389	subcyclohminus	484, 485
simplify	59, 133, 164	subcycloiwasawa	485, 488
sin	328	subcyclopclgp	488, 491
sinc	329	subfield	461
sinh	329	<i>subgroup</i>	376
sizebyte	164	subgroup	79
sizedigit	164, 165	subgrouplist	79, 392, 491
Smith normal form	79, 292, 379, 384, 388, 431, 459	subgrouplist0	491
<i>SNF generators</i>	178	subresultant algorithm	209, 254, 260
snfrank	308, 309	subst	268, 273
solve	103, 104, 359	substpol	269
solvestep	360	substvec	269
sopath	133	sum	136
sqr	329	sum	360
sqrt	329	sumalt	337, 347, 360, 361, 362, 374, 375
sqrtint	231	sumalt2	361
sqrtint0	231	sumdedekind	231
sqrtn	329	sumdigits	231, 232
sqrtnint	231	sumdigits0	232
<i>stack</i>	61, 129, 130, 134	sumdiv	231, 361
stacksize	46	sumdivk	231
Stark units	227, 400	sumdivmult	362
startup	62	sumdivmultexpr	362
Steinitz class	483	sumeulerrat	362
Stirling number	176	sumformal	269, 270
stirling	176, 177	suminf	361, 362, 363
stirling1	177	sumnum	336, 363, 366
stirling2	177	sumnumap	336, 366, 368
Str	49, 50, 149	sumnumapinit	336, 368, 369
Strchr	86	sumnuminit	336, 369
strchr	113, 150	sumnumlagrange	336, 369, 370
Strexpend	86	sumnumlagrangeinit	336, 370, 371
strexpend	86, 114	sumnummonien	337, 363, 366, 371
strftime	58, 131	sumnummonien0	372
strictargs	43, 134	sumnummonieninit	372, 373
strictmatch	134	sumnumrat	337, 373, 374
<i>string context</i>	49	sumnumsidi	337, 374
string	8, 27, 49	sumnumsidi0	374
strjoin	114	sumpos	337, 374, 375
Strprintf	86	sumpos2	375
strprintf	86, 114, 127	sunits_mod_units	389
		suppl	294

sylvestermatrix 266
symmetric powers 266
system 50, 100, 115, 133

T

t2 380
Tamagawa number 535, 545
tan 330
tanh 330
Taniyama-Shimura-Weil conjecture 526
tate 523
tayl 270
Taylor series 137
taylor 270
teich 331
teichmuller 330, 331
teichmullerinit 331
tex2mail 129, 130
TeXstyle 115, 123, 128
theta 331
thetanullk 331
threadsize 134
threadsizemax 134
thue 270, 271, 272
thueinit 270, 271, 272
time expansion 58
timer 135
trace 309
Trager 401
trap 50, 116
trap0 117
trueeta 320
trunc0 165
truncate 158, 159, 165, 248, 263
tschirnhaus 468
tu 380
tutorial 60
type 117
type0 117
t_CLOSURE 8, 28
t_COL 8, 23
t_COMPLEX 7, 20
t_ERROR 8, 28
t_FFELT 7, 19
t_FRAC 7, 19
t_INFINITY 8, 29
t_INT 7, 17
t_INTMOD 7, 18

t_LIST 8, 27
t_MAT 8, 25
t_PADIC 7, 20
t_POL 7, 22
t_POLMOD 7, 21
t_QFB 8, 23
t_QUAD 7, 20
t_REAL 7, 17
t_RFRAC 7, 23
t_SER 7, 22
t_STR 8, 27
t_VEC 8, 23
t_VECSMALL 8, 27

U

ulimit 46
unexport 117
unexportall 117
uninline 117
until 86
user defined functions 38

V

valuation 165
van Hoeij 401, 443
varhigher 34, 165, 166, 167, 168
variable (priority) 21, 34
variable scope 36
variable 21, 32
variable 34, 166
variables 167
variables_vec 167
variables_vecsmall 168
varlower 165, 168, 169
Vec 23, 24, 27, 149
veceint1 319
vecextract 287, 309
vecmax 143
vecmax0 143
vecmin 143, 144
vecmin0 144
vecprod 310
Vecrev 24, 150
vecsearch 310, 311, 312
vecsmall 8
Vecsmall 27, 150
vecsort 310, 311
vecsort0 312

vecsum	312
vecethanullk	331
vecethanullk_tau	331
vector	9
vector	24, 312, 313
vectorsmall	313
vectorv	24, 313
version number	61
version	117
Vi	64

W

warning	118
warning0	118
weber	332
weber0	332
weberf	332
weberf1	332
weberf2	332
Weierstrass $\wp$ -function	571
Weierstrass equation	521
Weil curve	566
whatnow	50, 118
while	86
write	50, 58, 62, 118
write0	118
write1	118
writebin	118
writetex	119

X

x[,n]	154
x[m,n]	154
x[m,]	154
x[n]	154

Y

ybessel	318
---------	-----

Z

Zassenhaus	248
zbrent	360
zell	560
zero	9
zeropadic	245
zeroser	245
zeta function	43

zeta	332
zetahurwitz	332, 333
zetamult	333, 334
zetamultall	334, 335
zetamultconvert	335
zetamultdual	335
zetamult_interpolate	334
Zideallog	426
zk	380
zkst	380
ZM_det	280
ZM_ker	288
znchar	232
zncharconductor	232, 233
znchardecompose	233
znchargauss	233, 234
zncharinduce	234, 235
zncharisodd	235
znchartokronecker	235
znchartoprimitive	235, 236
znconreychar	236, 237, 586
znconreyconductor	237, 238
znconreyexp	238, 239
znconreylog	236, 238, 240, 586
zncoppersmith	240, 241
znlog	206, 241, 242, 243, 426, 545
znlog0	242
znorder	242
znprimroot	242
znstar	241, 242
znstar0	243
znsubgroupgenerators	243, 245
Zp_appr	250
ZX_hyperellred	577