



Relation collection for the Function Field Sieve

Jérémie Detrey, Pierrick Gaudry, Marion Videau

► To cite this version:

Jérémie Detrey, Pierrick Gaudry, Marion Videau. Relation collection for the Function Field Sieve. ARITH 21 - 21st IEEE International Symposium on Computer Arithmetic, Apr 2013, Austin, Texas, United States. pp.201-210, 10.1109/ARITH.2013.28 . hal-00736123v2

HAL Id: hal-00736123

<https://inria.hal.science/hal-00736123v2>

Submitted on 18 Jan 2013

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

Relation collection for the Function Field Sieve

J       Detrey, Pierrick Gaudry and Marion Videau

CARAMEL project-team, LORIA, INRIA / CNRS / Universit   de Lorraine, Vand  uvre-l  s-Nancy, France

Email: {Jeremie.Detrey, Pierrick.Gaudry, Marion.Videau}@loria.fr

Abstract—In this paper, we focus on the relation collection step of the Function Field Sieve (FFS), which is to date the best algorithm known for computing discrete logarithms in small-characteristic finite fields of cryptographic sizes. Denoting such a finite field by $\mathbb{F}_{p^n}$, where p is much smaller than n , the main idea behind this step is to find polynomials of the form $a(t) - b(t)x$ in $\mathbb{F}_p[t][x]$ which, when considered as principal ideals in carefully selected function fields, can be factored into products of low-degree prime ideals. Such polynomials are called “relations”, and current record-sized discrete-logarithm computations need billions of those.

Collecting relations is therefore a crucial and extremely expensive step in FFS, and a practical implementation thereof requires heavy use of cache-aware sieving algorithms, along with efficient polynomial arithmetic over $\mathbb{F}_p[t]$. This paper presents the algorithmic and arithmetic techniques which were put together as part of a new public implementation of FFS, aimed at medium- to record-sized computations.

Keywords—function field sieve; discrete logarithm; polynomial arithmetic; finite-field arithmetic.

I. INTRODUCTION

The computation of discrete logarithms is an old problem that has received a lot of attention after the discovery of public-key cryptography in 1976. Indeed, together with integer factorization, it is the most famous computationally hard problem on which the security of public-key algorithms relies. We are interested here in the case of discrete logarithms in finite fields of small characteristic, for which the best algorithm known is the function field sieve (FFS) first introduced by Adleman [1] in 1994. Since then, and until 2005, there have been several works on this algorithm, both on the theoretical side [2], [3], [4] and on the practical side [5], [6], [7]. This culminated with a record set by Joux and Lercier who computed discrete logarithms in $\mathbb{F}_{2^{613}}$ and surpassed the previous record set by Thom   [8] in 2001 using Coppersmith’s algorithm [9]. A long period with essentially no activity on the topic has followed; it ended recently with a new record [10], [11] set over $\mathbb{F}_{3^{6 \times 97}}$, which took advantage of the fact that the extension was composite; this particular case is of practical interest in pairing-based cryptography.

The FFS algorithm can be decomposed into several steps. Two of them are by far the most time-consuming: the relation collection step and the linear algebra step. While the latter is not specific to FFS—it is a sparse-matrix kernel computation that is essentially the same as for other discrete-logarithm algorithms based on combining relations—the relation collection, on the other hand, is specific to each family

of algorithms. The purpose of this paper is to present a new implementation of this important step for FFS, targeting in our case finite fields of small characteristic with no special Galois structure. We choose two benchmarks in the cases of characteristic 2 and 3, that are the most interesting for cryptographic applications. The extension degrees were chosen so as to match the “kilobit milestone”: we consider $\mathbb{F}_{2^{1039}}$ and $\mathbb{F}_{3^{647}}$. Note that the factorization of $2^{1039} - 1$ has been completed only recently [12]. This is relevant to our benchmark as, among others, an 80-digit (265-bit) prime factor was exhibited: had $2^{1039} - 1$ been a product of only moderately large primes, there would have been better algorithms than FFS for computing discrete logarithms in $\mathbb{F}_{2^{1039}}$.

In the literature, many implementation details about the relation collection are merely sketched. In this article, we wish to explain as many of them as possible. Some of them are “folklore” or easy adaptations of techniques used in the Number Field Sieve (NFS) for integer factorization. For instance, the notion of skewness, or the use of bucket sorting during sieving that we describe below will come as no surprise to the NFS *cognoscenti*. However, to the best of our knowledge, some techniques go beyond what was previously known. We mention two of them. The sieving step starts with a phase called *norm initialization*, where a large array is initialized with the degree of a polynomial, which is different for each cell. In several implementations, only an approximation of the size of the norm is used, in order to save time. We propose several strategies that allow us to get an exact result in this step, at a very reasonable computational cost. This includes using an analog of floating-point arithmetic for polynomials. Another original contribution lies in the heart of the relation collection, during a sieving step, which is similar in essence to the sieve of Eratosthenes: in the same large array as above, we have to visit all the positions that verify a specific arithmetic property. In FFS, the set of these positions forms a vector space for which we give the general form of a basis, along with a method to compute it quickly following some Euclidean algorithm. We also adapt the well-known Gray code walk to our case where we are only interested in monic linear combinations.

To complement the description of our implementation, we make it freely available as part of the CADO-NFS project [13]. It is used to propose runtime estimates for the relation collection step in the two kilobit-sized finite fields that we target. It reveals that these sizes are easily

reachable with moderate computing resources; it is merely a matter of months on a typical 1000-core cluster. The question of the subsequent linear algebra step is however left open for future research. Still, it is already fair to claim that the cryptosystems whose security relies on the discrete-logarithm problem in a small-characteristic finite field of kilobit size provide only marginal security. Even more so, because the relation collection and the linear algebra steps are to be done just once and for all for a given finite field.

Outline: The paper is organized as follows. A short overview of the function field sieve, along with a focus on its relation collection step, is first given in Section II. Section III then highlights several key aspects of the proposed implementation, while low-level architectural and arithmetic issues are addressed in Sections IV and V, respectively. Timing estimates for two kilobit-sized finite fields are presented in Section VI, before a few concluding remarks and perspectives in Section VII.

II. COLLECTING RELATIONS IN FFS

A. A primer on the FFS algorithm

The principle of the Function Field Sieve algorithm is very similar to that of the Number Field Sieve algorithm for computing discrete logarithms in prime fields, where the ring of integers $\mathbb{Z}$ is replaced by the ring of polynomials $K[t]$, with K a finite field of size $\#K = \kappa$.

Assume that we want to compute discrete logarithms in $\mathbb{K}$, a degree- n algebraic extension of the base field K , of cardinality $\#\mathbb{K} = \kappa^n$. We consider two polynomials f and g in $K[t][x]$ such that their resultant in the x variable contains an irreducible factor $\varphi(t)$ of degree n . The following diagram, where all maps are ring homomorphisms, is commutative:

$$\begin{array}{ccccc} & & K[t][x]/f(x,t) & \longrightarrow & K[t]/\varphi(t) \cong \mathbb{K} \\ K[t][x] & \searrow & & \nearrow & \\ & & K[t][x]/g(x,t) & \longrightarrow & \end{array}$$

The key to FFS is to study the behavior of an element of the form $a(t) - b(t)x$ in this diagram. Along the two paths, or *sides*, the element is tested for *smoothness* in the appropriate algebraic structures. Pushing the factored versions into the finite field $\mathbb{K}$, we get a linear relation between the discrete logarithms of “small” elements. Once enough relations have been collected, they are put together to form a matrix for which a kernel element gives the discrete logarithms of these “small” elements. Finally, the logarithm of a given element is obtained by a procedure called special- q descent that uses again the diagram to produce a relation between the logarithm of the given element and the logarithms of smaller elements, and recursively so until we can relate them to the known logarithms of the “small” elements.

The smoothness notion is as follows. The element $a(t) - b(t)x$ is mapped to an element of $K[t][x]/f(x,t)$ (substitute g for f for the other side), which is viewed as a principal

ideal in the ring of integers of the corresponding function field $K(t)[x]/f(x,t)$. Since this is a Dedekind domain, there is a unique factorization in prime ideals. And if these prime ideals have degrees less than a parameter called the *smoothness bound*, we say that the principal ideal is smooth. Just like with NFS, complications arise from the fact that the ring of integers is not necessarily principal and that there are non-trivial units. The Schirokauer maps and the notion of virtual logarithms [14] are the tools to solve these complications; we do not need to worry about this here, since it does not interfere with the relation collection step.

B. Overview of the relation collection step

In this paper, we focus on this so-called *relation collection* step: given a smoothness bound—which might be different for each side—the goal is to produce (a, b) pairs, called *relations*, that simultaneously yield smooth ideals for both sides. Solving the resulting linear system requires finding at least as many relations as there are prime ideals of degree less than the smoothness bound.

The relation collection step is independent of the other steps in the FFS algorithm as long as we allow for any kinds of polynomials f and g . This is the most time-consuming part, both in practice and in theory: optimizing the parameters for this step will optimize the overall algorithm. We remark however that the linear algebra step has a high space complexity that can become a practical issue. The usual approach to circumvent this problem is to let the relation collection run longer than what theory alone would recommend. This produces a larger set of relations, and there exist algorithms for selecting among them a subset that is better suited for linear algebra than the original set of relations obtained without this *oversieving* strategy.

Ideals and norms: The highbrow description of the smoothness notion deals with ideals. Without loss of generality, we focus here on the side of the diagram corresponding to the polynomial f , the other side being similar. For our purposes, it is enough to see a prime ideal $\mathfrak{p}$ as a pair of polynomials $(p(t), r(t))$, where $p(t)$ is a monic irreducible polynomial and $r(t)$ is a root in x of the polynomial $f(x, t)$ modulo $p(t)$. More formally, the notation $\mathfrak{p} = (p(t), r(t))$ represents the ideal $\langle p(t), x - r(t) \rangle$. Restricting to ideals of that form means that, in this article, we ignore the places at infinity and the problems that ramification can cause. Our implementation takes care of these more complicated cases, but we prefer to keep the presentation simple.

An element (a, b) that potentially leads to a relation must be tested for smoothness. It means that we want the corresponding principal ideal on the f side to be the product of “small” prime ideals. Here, we face a terminology issue, since in our setting the good notion of size of a prime ideal is the degree of its p polynomial. Therefore, we are going to talk about the degree of a prime ideal $\mathfrak{p} = (p, r)$ as the degree of p (even though, formally, the degree of $\mathfrak{p}$ is one).

The factorization of the principal ideal is dictated by the factorization of its *norm* which, in our case, takes a simple form. Denoting by $F(X, Y, t)$ the polynomial obtained by homogenization of $f(x, t)$ with respect to x , the norm corresponding to a pair (a, b) is $F(a, b)$, a polynomial in t . Then, to each irreducible factor p of $F(a, b)$ corresponds a prime ideal $\mathfrak{p} = (p, r)$ that divides the principal ideal, and such that $r \equiv a/b \pmod{p}$. (We have dropped the variable t to keep our notation simple.) As a consequence, a simple restatement of the relation search task is to find (a, b) pairs for which the two norms $F(a, b)$ and $G(a, b)$ have all their irreducible factors in $K[t]$ of degree less than a given bound.

An important remark is that only *primitive* pairs are interesting, in the sense that if (a, b) gives a valid relation, then multiplying both a and b by the same small factor yields another relation which, however, is essentially the same and is of no use to the next steps of the algorithm. Therefore, we are only interested in the case where $\gcd(a, b) = 1$ and, furthermore, where (say) b is monic to take care of multiplications by constants.

Sieving and cofactorization: In order to test the smoothness of the norms corresponding to a pair (a, b) , we chose a two-pass strategy as is now usual for NFS. The first pass consists in using a sieving process *à la* Eratosthenes in order to efficiently remove the contribution of the prime ideals of degree up to a given *factor base bound*—chosen to be smaller than the actual smoothness bound—to the norm of the principal ideals that they divide. All of those (a, b) pairs such that what then remains of their norm is small enough will have better chances to completely factor as products of prime ideals of degree less than the smoothness bound. As such, the promising pairs which have survived sieving on both sides are selected to undergo a full-blown smoothness test in the second pass of the algorithm, which is sometimes called *cofactorization* or *large prime separation*.

Note that there is no need to keep track of the actual norm of the principal ideal in the first pass. In fact, it is enough to consider only the *degree* of the norm here, to which we simply subtract $\deg p$ for every prime ideal $\mathfrak{p} = (p, r)$ which divides the corresponding principal ideal. Also note that, because of the sieving techniques used in the first pass, we cannot process each pair (a, b) one after the other: the sieving region has to be considered as a whole and allocated *a priori* as a large array, delimited by upper bounds on the degrees of a and b , and in which all the (a, b) positions will be considered simultaneously by the sieving process.

III. ORGANIZATION OF THE COMPUTATION

A. Sieving by special-q

Since the set of (a, b) pairs that are tested for smoothness is very large, we have to subdivide this search space into pieces that are handled separately. Following the current trend for NFS and FFS, we split the (a, b) space according to so-called *special-q* lattices. A special-q is just a prime ideal

$\mathfrak{q} = (q, \rho)$ corresponding to either the f or the g side, and the $\mathfrak{q}$ -lattice is exactly the set of (a, b) pairs for which the principal ideal on this side is divisible by $\mathfrak{q}$. This corresponds to the set of (a, b) pairs such that $b\rho \equiv a \pmod{q}$, which is a $K[t]$ -lattice of dimension 2 and determinant q . The relation search is then seen as a set of independent tasks, each task consisting in finding relations in a given $\mathfrak{q}$ -lattice.

Basis of the q-lattice: Let $\mathfrak{q} = (q, \rho)$ be a special-q. A basis of the $\mathfrak{q}$ -lattice is given by the two vectors $(q, 0)$ and $(\rho, 1)$. It is a very unbalanced basis, and we can improve it using Gaussian lattice reduction to obtain two vectors u_0 and u_1 whose coordinates have a degree about half of the degree of q . We then look for relations coming from (a, b) pairs of the form $iu_0 + ju_1$, with i of degree less than I and j of degree less than J . This leads to a and b of degree approximately equal to $\max(I, J) + \frac{1}{2} \deg q$.

Skewness: There are cases where the polynomials f and g are *skewed*, in the sense that the degree in t of the x^i coefficients decreases when i increases. In that case, a relevant search space for (a, b) that yields norms of more or less constant degree is a rectangle: the bound on the degree on a must be larger than the bound on the degree of b . The difference of these two bounds on the degrees is called the *skewness*. In the $\mathfrak{q}$ -lattice, to obtain the target skewness adapted to f and g , the Gaussian lattice reduction is modified so that both vectors u_0 and u_1 have a difference of degrees between their coordinates that is close to the skewness. In the implementation, this translates into a minor change in the stopping criterion of the lattice reduction.

From now on, we consider that the sieving space for a given special-q is a κ^I by $\left(\frac{\kappa^J - 1}{\kappa - 1} + 1\right)$ rectangle in the (i, j) plane, covering all the polynomial pairs of degrees less than I and J , respectively, and such that the j coordinate is either 0 or monic. Indeed, if both i and j are multiplied by the same constant, so are a and b and we get duplicate relations. The correspondence with (a, b) is of the form $a = ia_0 + ja_1$ and $b = ib_0 + jb_1$, where $u_0 = (a_0, b_0)$ and $u_1 = (a_1, b_1)$ is the skew-reduced basis of the $\mathfrak{q}$ -lattice.

In our current implementation, I and J are fixed and equal. It could make sense to let them vary, for instance giving them a smaller value for larger special-q's so as to keep more or less always the same bound on the (a, b) pairs.

B. Initializing the norms

Initializing the norm at a given position is a simple task; however it can become costly if done naively. One could use an estimate of the degree (e.g., an upper bound) but it would impact the number of positions to test for smoothness by either losing or keeping too many positions. To avoid these drawbacks, we decided to work on being precise and even exact on the degree computation at a reasonable cost.

We present our method with the polynomial f ; it works similarly for the polynomial g . Recall that F denotes the homogenization of the polynomial f . For each special-q, we

start by computing a linear transformation on f to get f_q such that $F_q(i, j) = F(a, b)$. We then have to compute the degree (in t) of $F_q(i, j)$ and put it into the array to initialize the corresponding position. We assume from now on that the loop is organized by rows, *i.e.*, that j is fixed and i varies.

Saving many degree computations: Recall that, for a pair (i, j) , the norm that has to be tested for smoothness is given by $F_q(i, j) = \sum_{0 \leq k \leq d} f_{q,k} i^k j^{d-k}$, where each term is a polynomial in t of degree $\deg f_{q,k} + k \deg i + (d-k) \deg j$. When a position (i_0, j_0) has been initialized, it is possible to deduce the initialization of many following positions $(i_0 + i', j_0)$ since for small values of i' the degree of the norm does not change. To precisely characterize these i' , we define the notion of *gap*.

Definition 1. *With the previous notation, the gap of F_q at (i_0, j_0) , denoted by γ , is defined by $\gamma = -1$ if only one term $f_{q,k} i_0^k j_0^{d-k}$ has maximal degree and, otherwise, by*

$$\gamma = \max_{0 \leq k \leq d} \deg(f_{q,k} i_0^k j_0^{d-k}) - \deg(F_q(i_0, j_0)).$$

If there is only one term $f_{q,k} i_0^k j_0^{d-k}$ of maximal degree, this property will not change as long as $(i_0 + i')$ has the same degree as i_0 . Thus, if we enumerate the polynomials in an order that respects the degree, we get the degree of the norm for free until the degree of $(i_0 + i')$ increases.

If there are several terms of maximal degree, a change in the degree of the norm only depends on the $\gamma + 1$ most significant coefficients of $(i_0 + i')$. Therefore, for all i' such that $\deg i' < \deg i_0 - \gamma$, we also get the degree of the norm for free. (Note that the case $\gamma = 0$ can occur.)

Since computing the gap comes at almost no additional cost when computing the degree of the norm, this method saves many computations.

Computing the degree of the norm: If there is only one term, let say $f_{q,k_0} i_0^{k_0} j_0^{d-k_0}$, of maximal degree, then $\deg(F_q(i_0, j_0)) = \deg f_{q,k_0} + k_0 \deg i_0 + (d - k_0) \deg j_0$, $\gamma = -1$ and we do not have to actually compute the norm to obtain its degree.

If there are several terms of maximal degree, we need to take care of the possible cancellations. As a cheaper alternative to computing the actual norm, we resort to an adaptive-precision approximation, using an analog of the floating-point number system for polynomials. We define the precision N representation of a polynomial $p(t) = \sum p_k t^k$ by the pair $(\deg p, \text{mant}_N(p))$, where the mantissa $\text{mant}_N(p) = p_{\deg p} t^{N-1} + \dots + p_{\deg p - N + 1}$ is the polynomial of degree $N-1$ corresponding to the N most significant coefficients of p .

Starting from an initial precision $N = N_0$, we compute the “floating-point” approximations of the terms $f_{q,k} i_0^k j_0^{d-k}$, using truncated high products for the multiplications. Before summing those terms, we first align them to the maximal degree, as is the case for regular floating-point additions. This can be seen as a conversion to a fixed-point representation.

Finally, from the degree of the sum of the terms, we deduce the degree of the norm and the gap, unless the result is zero, meaning that the precision N was not high enough to conclude. Therefore, the computation takes the form of a loop where the precision N increases iteratively until the computation of the norm with N significant coefficients is enough to infer its degree.

A few more improvements can be added to this strategy. When the precision N is small, it might happen that several terms have a degree that is too small to contribute to the final sum, so that their computation can be skipped from the beginning. Also, when summing the terms, it can make sense to do so starting with the terms of highest degrees. If there are not too many cancellations, it can be the case that we can then guess the final degree before having to take the remaining terms into account.

We finally remark that the first step where we consider only the degrees of the terms can be seen as the particular case $N = 0$ of the subsequent loop.

C. Bases for the $\mathfrak{p}$ -lattices

Let $\mathfrak{p} = (p, r)$ be a prime ideal. The goal of this section is to describe the set of (i, j) positions for which the corresponding ideal is divisible by $\mathfrak{p}$ (hence whose norm is divisible by p). We recall that in terms of (a, b) position, it translates into the condition $a - br \equiv 0 \pmod{p}$, and therefore, the set of such (a, b) positions is a $K[t]$ -lattice for which $(p, 0)$ and $(r, 1)$ is a basis. Using the definition of a and b in terms of i and j , the condition becomes $i(a_0 - rb_0) \equiv -j(a_1 - rb_1) \pmod{p}$. This is a K -linear condition on the polynomials i and j , and therefore the set of solutions is a K -vector space which is of finite dimension since the degree of i must be less than I and the degree of j less than J . More precisely, the linear system corresponding to the condition has $I + J$ unknowns and $\deg p$ equations, so that the dimension of the space of solutions is at least $I + J - \deg p$. We will present a basis for this vector space that takes a different shape, depending on whether $\deg p$ is smaller or larger than I .

As a warm-up, we deal with the case where p divides $a_0 - rb_0$. The main condition on i and j simplifies to $j \equiv 0 \pmod{p}$. A basis is then given by the set of vectors $\mu_k = (t^k, 0)$ for $0 \leq k < I$, and $\nu_\ell = (0, t^\ell p)$ for $0 \leq \ell < J - \deg p$. The dimension is then $I + J - \deg p$.

From now on, we will assume that $a_0 - rb_0$ is invertible modulo p , and we define

$$\lambda_p = -(a_1 - rb_1)/(a_0 - rb_0) \pmod{p},$$

so that the main condition becomes $i \equiv \lambda_p j \pmod{p}$.

1) Case of small primes:

Lemma 2. *Let $\mathfrak{p} = (p, r)$ be a prime ideal of degree $L < I$. The vector space of (i, j) positions for which the corresponding ideal is divisible by $\mathfrak{p}$ is of dimension $I + J - L$*

and admits the basis given by the vectors $\mu_k = (t^k p, 0)$ for $0 \leq k < I - L$, and $\nu_\ell = (t^\ell \lambda_p \bmod p, t^\ell)$ for $0 \leq \ell < J$.

Proof: The given vectors are indeed in the target set of positions, and they are linearly independent, due to their echelonized degrees. To see that it is a basis, let us consider (i, j) with $\deg i < I$, $\deg j < J$, and $i \equiv j \lambda_p \pmod{p}$. Let then i_0 be the polynomial of degree less than p such that $i_0 \equiv j \lambda_p \pmod{p}$. Then the vector (i_0, j) can be obtained by combining the ν_ℓ 's with the coefficients of j . Furthermore i and i_0 differ by a multiple of p of degree less than I , so that (i, j) can be written as a linear combination of the claimed basis which, consequently, is indeed a basis. ■

2) *Case of large primes:* We assume now that $\mathfrak{p}$ is of degree L greater than or equal to I . In fact, we can assume furthermore that its degree is also greater than or equal to J . Indeed, the symmetric role of I and J makes it possible to write a similar basis as before if the degree of J is larger than the degree of $\mathfrak{p}$. In practice, this would raise memory locality issues when traversing a vector space given by a basis of this shape, but for the moment, we always take $I = J$ in our code.

The $K[t]$ -lattice of the valid (i, j) positions corresponding to $\mathfrak{p}$ is generated by $(p, 0)$ and $(\lambda_p, 1)$. To study this lattice, we consider the sequence computed by the Euclidean algorithm starting with these two vectors: let $v_\ell = (\sigma_\ell, \tau_\ell)$, defined by $v_0 = (p, 0)$, $v_1 = (\lambda_p, 1)$ and, for all $\ell \geq 1$, $\sigma_{\ell+1} = \sigma_{\ell-1} - \sigma_\ell q_\ell$ and $\tau_{\ell+1} = \tau_{\ell-1} - \tau_\ell q_\ell$, where q_ℓ is the quotient in the Euclidean division of $\sigma_{\ell-1}$ by σ_ℓ . Since p is prime, λ_p is coprime to p and the final vector of the sequence is $(0, p)$. The finite family of vectors (v_ℓ) is a K -free family of vectors with the degrees of both coordinates less than or equal to L . In the following lemma it is shown how to complete it to form a basis of such vectors.

Lemma 3. *Let $\mathfrak{p} = (p, r)$ be a prime ideal of degree L and $(v_\ell = (\sigma_\ell, \tau_\ell))$ the corresponding Euclidean sequence. The vector space of (i, j) positions for which the corresponding ideal is divisible by $\mathfrak{p}$ with degrees at most L in each coordinate has dimension $L + 2$ and admits the basis $\{v_0\} \cup (\cup_{\ell \geq 1} \{t^s v_\ell \mid 0 \leq s < \deg \sigma_{\ell-1} - \deg \sigma_\ell\})$.*

Proof: With the convention that the degree of 0 is -1 , the claimed basis is such that the degree in the first coordinate takes all values from -1 to L exactly once, and the same for the second coordinate. Indeed, since, by induction, the degrees of the σ_ℓ 's are strictly decreasing, and those of the τ_ℓ 's are strictly increasing, we have that $\deg \tau_{\ell+1} - \deg \tau_\ell = \deg q_\ell$, which is itself, by construction, equal to $\deg \sigma_{\ell-1} - \deg \sigma_\ell$. Therefore, when iterating from $v_{\ell-1}$ to v_ℓ , a drop δ in the degree of σ_ℓ is directly followed by an identical increase in the degree of τ_ℓ in the next iteration, thus showing that the vectors $t^s v_\ell$, for $0 \leq s < \delta$, bridging the gaps between $v_{\ell-1}$, v_ℓ and $v_{\ell+1}$, will all have distinct degrees in both of their coordinates, as illustrated in

Figure 1. The vectors obtained during the Euclidean algorithm (black dots), completed with some of their multiples (gray dots) to form the full K -basis of vectors of the $\mathfrak{p}$ -lattice with coordinates of degree at most $\deg p$.

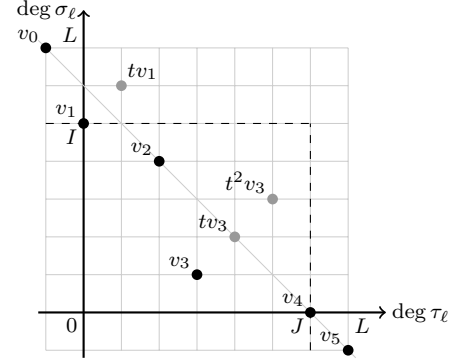


Figure 1. Consequently, the proposed family of vectors is K -free and contains $L + 2$ elements.

In order to get the exact dimension of the vector space, we just count its elements. Consider the vector space $\{(i, j) \mid i \equiv \lambda_p j \pmod{p}, \text{ with } \deg i \leq L, \deg j \leq L\}$, and split it into κ^{L+1} disjoint subsets S_j , one for each possible value of j with $\deg j \leq L$. For any j , S_j has the same cardinality as the set of polynomials i of degree at most L such that $i \equiv j \lambda_p \pmod{p}$. Since L is precisely the degree of p , this set has the same cardinality κ as the base field K . Adding up the contributions, we obtain a vector space of cardinality κ^{L+2} , thus of dimension $L + 2$ as claimed. ■

For our purposes, we are only interested in those elements for which the degrees of their coordinates are (strictly) bounded by I and J respectively. Due to the echelonized form of the basis, it is enough to keep only those basis elements whose degrees satisfy these constraints in order to generate all the relevant $\mathfrak{p}$ -lattice elements.

D. Enumerating the elements of a $\mathfrak{p}$ -lattice

Once the K -basis of a $\mathfrak{p}$ -lattice is known, the next step is to enumerate all the (i, j) pairs in this lattice and to subtract $\deg p$ to the degree of the norm of the corresponding ideals. Obviously, this task can be achieved by considering all the possible K -linear combinations of the basis vectors. However, attention must be paid to a few details.

In order to address these issues for all possible types of $\mathfrak{p}$ -lattices at once, whether $\mathfrak{p}$ be a small or a large prime ideal, or even when $a_0 - r b_0 \equiv 0 \pmod{p}$, we introduce the following notations, designed to cover all cases without loss of generality: given a prime ideal $\mathfrak{p}$, we denote the basis of the corresponding lattice seen as a K -vector space, as constructed in Section III-C, by the vectors $(\mu_0, \dots, \mu_{d_0-1}, \nu_0, \dots, \nu_{d_1-1})$, where the μ_k 's and the ν_ℓ 's are of the form $\mu_k = (\gamma_k, 0)$ and $\nu_\ell = (\alpha_\ell, \beta_\ell)$, for $0 \leq k < d_0$ and $0 \leq \ell < d_1$, respectively, and where the degrees of the γ_k 's and of the β_ℓ 's are echelonized: $0 \leq \deg \gamma_0 < \deg \gamma_1 < \dots < \deg \gamma_{d_0-1} < I$ and $0 \leq \deg \beta_0 < \deg \beta_1 < \dots < \deg \beta_{d_1-1} < J$.

As we impose j to be monic (see Section III-A), we first normalize the ν_ℓ vectors of the basis so that their second coordinate β_ℓ is monic. Using the fact that the β_ℓ 's are echelonized and sorted by increasing degree, enumerating only those lattice vectors whose j -coordinates are monic boils down to considering only *monic* linear combinations of the ν_ℓ 's, that is to say, K -linear combinations whose last nonzero coefficient, if any, is 1. On the other hand, since the j -coordinates of the μ_k 's are all zero, all K -linear combinations of these vectors should be considered. We first explain this easy case and come back to the ν_ℓ 's later.

Assuming that K is a prime field, the enumeration of all K -linear combinations can be carried out by means of a κ -ary Gray code, as already suggested by Gordon and McCurley in the binary case [15]. Let us consider the infinite sequence $\Delta = (\delta_i)_{i \geq 1}$ of the t -adic valuations of the nonzero polynomials of $K[t]$ sorted by increasing lexicographic order. For a given integer $d \geq 0$, the prefix sequence $\Delta_d = (\delta_1, \dots, \delta_{\kappa^d-1})$ can be computed via a recursive definition as $\Delta_0 = ()$ (the empty sequence), and $\Delta_{i+1} = (\Delta_i, i, \Delta_i, i, \dots, i, \Delta_i)$, where Δ_i and i are repeated κ and $\kappa - 1$ times, respectively. For instance, over the base field $K = \mathbb{F}_3$, one would have $\Delta = (0, 0, 1, 0, 0, 1, 0, 0, 2, 0, 0, 1, \dots)$.

In fact, Δ_d corresponds to the transition sequence of a d -digit κ -ary Gray code: starting from the polynomial $\pi_1 = 0$, and repeating the construction $\pi_{i+1} = \pi_i + t^{\delta_i}$, we end up after $\kappa^d - 1$ iterations with the set $\{\pi_1, \dots, \pi_{\kappa^d}\}$ covering precisely all the polynomials of $K[t]$ of degree less than d . Additionally, given a K -free family of vectors $(\mu_k)_{0 \leq k < d}$, we can use the iteration $\pi_{i+1} = \pi_i + \mu_{\delta_i}$ instead, this way enumerating the whole set of K -linear combinations of the μ_k 's, at the cost of only one vector addition per combination. This technique can therefore be used to efficiently go through all the κ^{d_0} linear combinations of the μ_k basis vectors.

However, as far as the ν_ℓ 's are concerned, this method cannot be applied directly, as only the *monic* linear combinations of these basis vectors need be considered. To that intent, we define a so-called *monic* variant of κ -ary Gray codes by considering the infinite sequence $\Delta' = (\delta'_i)_{i \geq 1}$ defined as the t -adic valuations of the nonzero *monic* polynomials of $K[t]$ sorted by increasing lexicographic order. As for Δ , the prefix sequence $\Delta'_d = (\delta'_1, \dots, \delta'_{(\kappa^d-1)/(\kappa-1)})$ can be computed, for any nonnegative integer d , as $\Delta'_0 = ()$ and $\Delta'_{i+1} = (\Delta'_i, i, \Delta'_i)$. For instance, over $K = \mathbb{F}_3$, one would have $\Delta' = (0, 1, 0, 0, 2, 0, 0, 1, 0, 0, 1, 0, 0, 3, \dots)$.

As in the non-monic case above, Δ'_d can also be interpreted as the transition sequence of a so-called d -digit κ -ary *monic* Gray code. Starting from 0, the construction $\pi_{i+1} = \pi_i + t^{\delta'_i}$ would indeed end up, after $\frac{\kappa^d-1}{\kappa-1}$ iterations, enumerating all the monic polynomials in $K[t]$ of degree less than d . Alternatively, the iteration $\pi_{i+1} = \pi_i + \nu_{\delta'_i}$ would instead enumerate all the monic K -linear combinations of

any K -free family of monic vectors $(\nu_\ell)_{0 \leq \ell < d}$.

E. Cofactorization

We did not put a lot of efforts in the implementation of this step. The classical approach, coming back to Coppersmith and Davenport [9] has been used. To test whether a polynomial $P(t)$ is smooth with respect to a smoothness bound B , we compute $P'(t) \prod_{k=\lceil (B+1)/2 \rceil}^B (t^{\kappa^k} - t) \bmod P(t)$. If $P(t)$ is B -smooth, this quantity is zero, and the converse is true, unless we are in the unfortunate case where there is a large irreducible factor that occurs with a multiplicity that is a multiple of κ . To speed-up the reductions modulo $P(t)$, its preinverse is precomputed once and for all. Then, each product modulo $P(t)$ can be done at a cost of one polynomial product in degree $\deg P$, one low short product and one high short product (see for instance [16]). For the time being, the short products are implemented as regular products, which leaves room for improvement.

Finally, at least in characteristic 2, the overall cost of cofactorization is small compared to the sieving step. Therefore, we did not yet fully implement a resieving step. For the moment, the input given to the cofactorization step is the full norm, in which we do not remove the small irreducible factors that have been detected during the sieving step. Again, this leaves room for improvement, especially in characteristic 3 where the cofactorization step takes a much larger share of the total runtime (see Section VI below).

F. Sub-lattices

Since we want to consider only primitive (i, j) pairs, a classical improvement [7] is to consider the congruence class of i and j modulo one or a few small irreducible polynomials. Let h be an irreducible polynomial of degree d ; the probability that it divides both i and j is κ^{-2d} . This probability gives the proportion of the sieving space that is useless. To save some sieving time, we can therefore decompose the sieving space into the κ^{2d} translates of the sub-lattice corresponding to h , and sieve all but one of them.

The efficient enumeration of the elements of the $\mathfrak{p}$ -lattices can easily be modified to consider only its intersection with (i, j) pairs that are congruent to some prescribed (i_0, j_0) modulo h . The key observation is that once a first point has been found, the others are obtained by adding the elements of the vector spaces that have been described in Lemmata 2 and 3. The modifications to make are therefore mostly concerned with the initial position for the enumeration that is no longer the point $(0, 0)$ but a point that is deduced thanks to a few computations modulo h . These computations of the starting point are cheap, but they are to be done for each $\mathfrak{p}$ and for each of the $\kappa^{2d} - 1$ translates of the sub-lattice that we consider. As a consequence, only polynomials h of small degree must be considered. However, it is possible to combine several h of small degrees.

The case where this strategy gives the best practical improvement is when the base field K is $\mathbb{F}_2$. Then the two irreducible polynomials of degree 1 are t and $t + 1$. Using them, we need only sieve 9/16 of the original sieving space at the price of a reasonable overhead in the arithmetic cost of the initialization. Going further would require to consider the only irreducible polynomial of degree 2, namely $t^2 + t + 1$. In theory, this would allow us to further decrease the sieving space by a 1/16 factor. On the other hand, this would multiply by 15 the overhead and, with our current implementation, we estimate that it is not worth the effort.

The next interesting case is when $K = \mathbb{F}_3$. The irreducible polynomials of degree 1 are t , $t + 1$ and $t + 2$. Each of them can be used to avoid sieving 1/9 of the (i, j) pairs. Using the three of them in conjunction would thus allow us to sieve only $(8/9)^3 = 512/729$ of the original sieving space. Again, we consider that with our current implementation the potential gain would be cancelled by the overhead. This could however change in the future with further optimization of the small arithmetic building blocks.

For larger base fields, the overall probability for a polynomial to be divisible by an irreducible polynomial h of degree 1 decreases. It makes this strategy less and less attractive. For instance, in [11] where the authors took $K = \mathbb{F}_{27}$, it is not considered.

IV. MEMORY LOCALITY ISSUES

When enumerating the elements of a given $\mathfrak{p}$ -lattice, the visited positions are scattered across the (i, j) plane. It results in random accesses to the array containing the degree of the norm. For instance, for a kilobit-sized field K , the number of visited positions lies between a few hundreds of millions and a few billions, meaning that data locality rapidly becomes an important issue on typical modern-day computers with strongly hierarchized memory. We thus detail in the following paragraphs two techniques designed to improve the spatial locality of memory accesses, one for small primes and the other for large primes.

A. Sieving by rows

Let us first consider the case of small prime ideals $\mathfrak{p}$, of degree $L < I$. From Lemma 2, the first $I - L$ vectors $(\mu_k)_{0 \leq k < I-L}$ of the basis of the corresponding $\mathfrak{p}$ -lattice all have their second coordinates equal to 0. Consequently, all the κ^{I-L} linear combinations of these vectors lie in the same horizontal line of the (i, j) plane. Therefore, if this (i, j) plane is stored as a row-major array, and assuming that at least a single row—i.e., κ^I elements—can fit into the L1 data cache, then memory locality can be drastically improved by sieving row by row. This technique was first described by Pollard in the case of the Number Field Sieve [17].

Furthermore, since the remaining J vectors $(\nu_\ell)_{0 \leq \ell < J}$ of the basis have their j -coordinates equal to t^ℓ , sieving rows by increasing lexicographic order on their second coordinates

is simply a matter of enumerating in the same order monic polynomials over K of degree less than J , then using their coefficients for computing linear combinations of the ν_ℓ 's.

However, directly using the monic Gray code technique presented in Section III-D to speed up the enumeration process would break the lexicographic ordering of the j -coordinates. In order to maintain this ordering, we propose a simple change of basis which would allow us to still benefit from the efficiency of Gray code enumeration at the expense of a slight overhead.

Thus, let $(\nu'_\ell)_{0 \leq \ell < J}$ be the family of vectors defined by $\nu'_\ell = \sum_{i=0}^{\ell} \nu_i$. For any $0 \leq \ell < J$, the second coordinate of ν'_ℓ is $\sum_{i=0}^{\ell} t^i$, from which it follows that the ν'_ℓ 's form a K -free family, and that $(\mu_0, \dots, \mu_{I-L-1}, \nu'_0, \dots, \nu'_{J-1})$ is also a K -basis of the $\mathfrak{p}$ -lattice.

One can already note that a non-monic Gray code enumeration of the ν'_ℓ 's would go through all linear combinations thereof, sorted by increasing lexicographic order in their j -coordinates. Indeed, intuitively, the second coordinate $\sum_{i=0}^{\ell} t^i$ of ν'_ℓ emulates the propagation of the “carry” through the low-weight coefficients of j , just as a lexicographic enumeration would do right before reaching a polynomial j of t -adic valuation ℓ . For instance, over $K = \mathbb{F}_3$, the sequence $(0, \nu'_0, 2\nu'_0, \nu'_1 + 2\nu'_0, \nu'_1, \nu'_1 + \nu'_0, 2\nu'_1 + \nu'_0, 2\nu'_1 + 2\nu'_0, 2\nu'_1, \nu'_2 + 2\nu'_1, \dots)$, obtained thanks to a ternary Gray code enumeration, corresponds in fact to a lexicographic enumeration in the j -coordinates.

However, this change of basis is not sufficient when only monic j -coordinates should be considered. One can easily verify that, when applying the previous technique with monic Gray codes, the first enumerated vector to have its second coordinate of degree equal to ℓ will be $\nu_\ell + 2\nu_{\ell-1}$, and not ν_ℓ as required by the lexicographic order. Since this happens only $J - 1$ times throughout the whole enumeration of the $\mathfrak{p}$ -lattice, a simple fix to tackle this issue is to subtract $2\nu_{\ell-1}$ from the current vector whenever the degree of its j -coordinate has just increased from $\ell - 1$ to ℓ .

Note that sieving by rows is also compatible with the case $a_0 - rb_0 \equiv 0 \pmod{p}$, in which the $\mathfrak{p}$ -lattice is the rows of the (i, j) plane whose j -coordinates are multiples of p .

B. Bucket sieving

When the degree L of the ideal $\mathfrak{p}$ increases, the corresponding lattice positions are much more sparsely scattered over the (i, j) plane, since at most one hit per row can be expected when $L \geq I$. A different strategy is therefore necessary in the case of large prime ideals.

The main idea is to coalesce the norm updates in the (i, j) array according to their j -coordinates, using an algorithm inspired from bucket sorting. Partitioning the sieving area into small groups of consecutive rows, or *bucket regions*, it is possible to go through all the large prime ideals and fill *buckets* with all the hits falling into each corresponding bucket region. The actual updating of the norms in the (i, j)

array is deferred to a second phase, where the buckets are processed individually and sequentially. This ensures that all the norm updates corresponding to a same bucket region are applied before moving to the next region, therefore enforcing spatial locality of the memory accesses. This method was described in details by Aoki and Ueda in [18].

There is however a trade-off to consider when dimensioning the buckets: if smaller bucket regions will fit better in the data cache during the norm update phase, this also implies more buckets to fill when enumerating the hit positions, which might end up exhausting the TLB (Translation Lookaside Buffer) of the CPU. Careful tuning is therefore necessary to balance cache misses against TLB misses.

Finally, since this technique forces the sieving area to be processed one bucket region at a time, it is possible to split the other steps of the relation collection algorithm (norm initialization, sieving by rows and cofactorization) so that they also operate on only one bucket region at a time. This way, there is no need to store the whole (i, j) array in memory, but only one bucket region.

V. A LIBRARY FOR POLYNOMIAL ARITHMETIC

Were it be for computing the bases of the q -lattices, for initializing the degrees of the norms of the principal ideals in the (i, j) array, for enumerating the elements of the p -lattices, or for testing promising ideals for smoothness, polynomial arithmetic over $K[t]$ plays a central role throughout the whole relation collection phase in FFS. Efficient implementation of this arithmetic is therefore crucial, and should not be overlooked. However, compared to existing, general-purpose libraries such as NTL [19], ZEN [20], or `gf2x` [21], we have the following specific constraints:

Memory footprint: We need types for polynomials of various fixed sizes, and we cannot accept to loose too much in term of memory. *i.e.*, no padding to 64 bits, no additional integer storing the degree, *etc.*.

Efficiency: The library should have a reasonably fast fallback for all operations, while also allowing one to write specific code for critical operations.

Compile-time optimization for a given base field K : Even if the library should be able to support different base fields, the chosen field is known at compile time. This knowledge must be exploited so as to properly inline and optimize all the field-specific low-level primitives.

To meet all these requirements, we have developed our own library for polynomial arithmetic. Written in C, it defines types and related functions for 16-, 32-, and 64-coefficient polynomials, along with multiprecision polynomials. Most of the provided functions are independent of the size of K , and only a few field-specific core functions, such as addition or multiplication by an element of K , should be defined in order to add support for another base field to the library. For the time being, only $\mathbb{F}_2$ and $\mathbb{F}_3$ are supported.

A. Representation

The current version of the library only supports bitsliced representation of the polynomials. Therefore, if a single element of K can be represented on k bits (typically, $k = \lceil \log_2 \kappa \rceil$), then an ℓ -coefficient polynomial will be represented as an array of k ℓ -bit words.

For $K = \mathbb{F}_2$, this is equivalent to evaluating the corresponding integer polynomial at 2. Over $\mathbb{F}_3$, an ℓ -coefficient polynomial $p(t) = \sum_{i=0}^{\ell-1} p_i t^i$ will be represented by an array `p` of two ℓ -bit words `p[0]` and `p[1]` such that $p_i = 2p[1]_i + p[0]_i$, where $p[j]_i$ denotes the i -th bit of word `p[j]`. Note that $p[0]_i = p[1]_i = 1$ is not allowed.

B. Basic operations

On top of being quite a compact format (as opposed to packed representations, where extra zeros have to be inserted between coefficients), bitsliced representation allows us to use bitwise CPU instructions to perform coefficient-wise operations on the polynomials.

The default implementation of the polynomial multiplication follows a simple serial/parallel (also known as shift-and-add) scheme. It is however possible to plug in optimized functions to accelerate this critical operation. For instance, over $K = \mathbb{F}_2$, one can build the library against `gf2x` and use the faster `gf2x_mul1` function as a drop-in replacement.

Other supported functions include degree computation (using fast leading-zero counting primitives such as GCC's `__builtin_clz`), conversions between types of different sizes, division, modular reduction, GCD, and so on.

C. Enumeration and κ -ary Gray codes

As already mentioned in the previous sections, enumerating (monic) polynomials by increasing lexicographic order is required at various stages of the relation collection process. If, over $\mathbb{F}_2$, this operation boils down to incrementing the ℓ -bit word representing the polynomial, it rapidly becomes trickier over larger base fields. For instance, given $p(t) \in \mathbb{F}_3[t]$, computing the next polynomial $r(t)$ requires 7 operations:

```
t0 = p[1] + 1;   t1 = p[1] ^ t0;   t2 = (t1 >> 1) ^ t1;
r[0] = p[0] ^ t2; r[1] = (r[0] & t2) ^ t0;
```

Enumerating (monic) polynomials in such a way also computes the transition sequence Δ (Δ' , respectively) of the corresponding (monic) κ -ary Gray code, as it is the sequence of the t -adic valuations of the successive (monic) polynomials taken in that order. In the previous example, the sequence Δ can be retrieved simply by computing the number of trailing zeros of the successive values of `t2`.

D. Addressing the (i, j) array

When sieving by rows or when applying norm updates from a bucket, the accesses to the (i, j) array are mostly random, albeit restricted to a few rows. Converting a given position in the (i, j) plane to an integer offset in the memory

representation of this plane is thus critical to the sieving process.

Over the binary field $K = \mathbb{F}_2$, this conversion is trivial, since the bitsliced representation of a polynomial $p(t) \in K[t]$ already corresponds to the integer $\tilde{p}(2)$, where $\tilde{p}(t) \in \mathbb{Z}[t]$ denotes the corresponding integer polynomial. One can then simply address the (i, j) array by offsets of the form $\tilde{p}_{(i,j)}(2)$, with $p_{(i,j)}(t) = i(t) + t^I j(t)$.

The situation is however not as simple over $K = \mathbb{F}_3$. Given the representation p of a polynomial $p(t) \in \mathbb{F}_3[t]$ as two ℓ -bit words, interleaving the bits of $p[0]$ and $p[1]$ would correspond to evaluating $\tilde{p}(t)$ at 4, which would then result in only $(3/4)^\ell$ of the 2ℓ -bit integers being valid representations of polynomials. On the other hand, if evaluating $\tilde{p}(t)$ at 3 would provide us with a compact integer representation, it would be far more expensive to compute.

In the current version of our library, we have settled for an intermediate solution between those two extremes: noting that $3^5 = 243 \approx 256 = 2^8$, we can split $p(t)$ into 5-coefficient chunks (or *pentatrits*) and, thanks to a simple look-up table, convert the 10 bits representing each chunk into its evaluation at 3, which then fits on a single octet (8 bits). A ternary polynomial of degree less than d can then be represented using at most $8\lceil d/5 \rceil$ bits, which is much more compact than using 2 bits for each trit. It is also possible to easily adapt this method to account for monic polynomials by using a specific look-up table for the most-significant pentatrit. Addressing the (i, j) array is then just a matter of converting $p_{(i,j)}(t) = i(t) + t^I j(t)$ in the same manner.

VI. BENCHMARKS AND TIME ESTIMATES

We propose two finite fields $\mathbb{K}$ as benchmarks for our implementation: $\mathbb{F}_{2^{1039}}$ and $\mathbb{F}_{3^{647}}$. The extension degrees are prime, so that the improvements of [22], [11] based on the Galois action are not available. The sizes of both problems are similar and correspond to the “kilobit” milestone.

We performed a basic polynomial selection, based on an exhaustive search of a polynomial with many small roots. This step would benefit from much further research but this is not the topic of the present paper. We thus give without further details the polynomials used for this benchmark, where hexadecimal numbers encode polynomials in $\mathbb{F}_2[t]$ (resp. $\mathbb{F}_3[t]$) represented by their evaluation at 2 (resp. 4):

$\mathbb{F}_{2^{1039}}$	$f = x^6 + 0x7x^5 + 0x6x + 0x152a$ $g = x + t^{174} + 0x1ef9a3$
$\mathbb{F}_{3^{647}}$	$f = x^6 + 0x2x^2 + 0x11211$ $g = x + t^{109} + 0x1681446166521980$

It can be checked that the resultants of f and g are polynomials in t that have irreducible factors of degree 1039 and 647, respectively, and are therefore suitable for discrete-logarithm computations in the target fields.

Parameters: For our choices of polynomials, when taking special- q ’s on the g side (the so-called *rational side*), norms

have more or less similar degrees on both f and g sides. We thus take the parameters given in the following table:

	$\mathbb{F}_{2^{1039}}$	$\mathbb{F}_{3^{647}}$
Sieving range $I = J$	15	9
Skewness	3	1
Max. deg. of sieved primes (factor base bound)	25	16
Max. deg. of large primes (smoothness bound)	33	21
Threshold degree for starting cofactorization	99	63

Timings: We give the running time of our code with these parameters on one core of an Intel Core i5-2500. The code was linked against the `gfm2x` library version 1.1 so that we can take advantage of the available `PCLMULQDQ` instruction for polynomial multiplication over $\mathbb{F}_2$. For various degrees of special- q ’s, we give the average time to compute one relation and the average number of relations per special- q . Since the number of special- q ’s of a given degree is close to the number of monic irreducible polynomials of this degree, this gives enough information to deduce how many relations can be computed in how much time. The yield and the number of relations vary a lot from one special- q to the other, even at the same degree. The average values given in the following table have been obtained by letting the program run with special- q ’s of the same degree until we reach a point where the measured yield is statistically within a $\pm 3\%$ interval with a confidence level of 95.4 %.

$\mathbb{F}_{2^{1039}}$	deg q	26	27	28	29	30	31	32
	Yield (s/rel)	0.59	0.71	0.88	1.07	1.31	1.70	2.05
	Rels per special- q	32.6	26.2	20.6	16.5	13.3	10.2	8.3

$\mathbb{F}_{3^{647}}$	deg q	17	18	19	20
	Yield (s/rel)	2.24	2.88	3.52	4.71
	Rels per special- q	23.5	15.7	11.8	8.0

For $\mathbb{F}_{2^{1039}}$, since the smoothness bound is set to degree 33, we can give a rough estimate of $2 \times 2^{34}/33 \approx 1.04 \cdot 10^9$ for the number of ideals in both factor bases. Collecting relations up to special- q ’s of degree 30 will provide about $1.19 \cdot 10^9$ relations in about 28 600 days. Taking into account the fact that not all ideals are involved but that there are duplicate relations, this should be just enough to get a full set of relations. Collecting relations for all special- q ’s up to degree 31 will provide about $1.90 \cdot 10^9$ relations in about 51 000 days, which will give a lot of excess and plenty of room to decrease the pressure on the linear algebra. The tuning of this amount of oversieving cannot be done without a linear algebra implementation and is, therefore, out of the scope of this paper. Nevertheless, we expect that in practice the relation collection will be done for a large proportion of special- q ’s of degree 31. For $\mathbb{F}_{3^{647}}$, similar estimates based on the large prime bound give a theoretical number of ideals of $1.53 \cdot 10^9$. Collecting relations up to special- q ’s of degree 19 will provide about $1.24 \cdot 10^9$ relations in about 45 300 days while going to degree 20 will yield about $2.63 \cdot 10^9$ relations in about 121 000 days.

To convert these large numbers into more practical notions, let us assume that we have a cluster of 1000 cores similar to the one we used. Then a full set of relations with

enough redundancy can be computed in a bit more than a month for $\mathbb{F}_{2^{1039}}$, and in about 4 months for $\mathbb{F}_{3^{647}}$.

Runtime breakdown: For $\mathbb{F}_{2^{1039}}$ (resp. $\mathbb{F}_{3^{647}}$), we give details on where the time is spent for special- q 's of degree 30 (resp. degree 19), in percentage and in average number of cycles per (a, b) candidate.

Step	$\mathbb{F}_{2^{1039}}, \deg q = 30$		$\mathbb{F}_{3^{647}}, \deg q = 19$	
	Cycles/pos	Percentage	Cycles/pos	Percentage
Initialize norms	1.10	2.04 %	43.65	6.17 %
Sieve by rows	9.73	18.15 %	180.11	25.45 %
Fill buckets	31.73	59.21 %	211.23	29.84 %
Apply buckets	2.74	5.12 %	4.53	0.64 %
Cofactorization	7.43	13.87 %	266.53	37.66 %
Total	53.59	100.00 %	707.77	100.00 %

From this table, it is clear that we have not yet spent as much time on optimizing the case $K = \mathbb{F}_3$ as we have for $\mathbb{F}_2$. We expect that there is a decent room for improvement in characteristic 3. On the other hand, the arithmetic cost is intrinsically higher over $\mathbb{F}_3$ than over $\mathbb{F}_2$: the addition of two polynomials requires at least 6 instructions [23] instead of just one, and the multiplication in characteristic 2 is implemented in hardware whereas in characteristic 3 this is a software implementation based on the addition. This difference is due to our choice to implement on general purpose CPUs, but would decrease or even disappear if we allowed ourselves to resort to specific hardware.

VII. CONCLUSION

We have presented in this paper a new implementation of the relation collection step for the function field sieve algorithm, which combines state-of-the-art algorithmic, arithmetic and architectural techniques from previous works on both FFS and NFS (where applicable), along with original contributions on several key steps of the algorithm. The source code of our implementation is released under a public license. To the best of our knowledge, this is the first FFS public code for handling record sizes. We hope this will serve as a reference point for future developments on this subject as well as for dimensioning key-sizes for DLP-based cryptosystems.

Of course, this project is still under heavy development, especially for base fields K larger than $\mathbb{F}_2$. Among other perspectives, we plan to investigate the relevance of delegating some parts of the computation to dedicated hardware accelerators such as GPUs or FPGAs. Finally, we plan to integrate our code for the relation collection step into a full implementation of FFS, in order to complete the ongoing computations of discrete logarithms over $\mathbb{F}_{2^{1039}}$ and $\mathbb{F}_{3^{647}}$, whence setting kilobit-sized records as new landmarks.

Acknowledgements. The authors wish to thank the other members of the CAMEL group for numerous interesting discussions regarding preliminary versions of this work, in particular Razvan Barbulescu, Cyril Bouvier, Emmanuel Thomé and Paul Zimmermann.

REFERENCES

- [1] L. M. Adleman, “The function field sieve,” in *ANTS I*, LNCS 877, pp. 108–121, 1994.
- [2] L. M. Adleman and M.-D. A. Huang, “Function field sieve method for discrete logarithms over finite fields,” *Inf. Comput.*, 151(1–2):5–16, 1999.
- [3] R. Matsumoto, “Using C_{ab} curves in the function field sieve,” *IEICE Trans. Fund.*, E82-A(3):551–552, 1999.
- [4] R. Granger, “Estimates for discrete logarithm computations in finite fields of small characteristic,” in *Cryptography and Coding*, LNCS 2898, pp. 190–206, 2003.
- [5] R. Granger, A. J. Holt, D. Page, N. P. Smart, and F. Vercauteren, “Function field sieve in characteristic three,” in *ANTS VI*, LNCS 3076, pp. 223–234, 2004.
- [6] A. Joux and R. Lercier, “Discrete logarithms in $GF(2^{607})$ and $GF(2^{613})$,” Posting to the Number Theory List, 2005.
- [7] —, “The function field sieve is quite special,” in *ANTS V*, LNCS 2369, pp. 343–356, 2002.
- [8] E. Thomé, “Computation of discrete logarithms in $\mathbb{F}_{2^{607}}$,” in *ASIACRYPT 2001*, LNCS 2248, pp. 107–124.
- [9] D. Coppersmith, “Fast evaluation of logarithms in fields of characteristic two,” *IEEE Trans. Inf. Theory*, 30(4):587–594, 1984.
- [10] T. Hayashi, N. Shinohara, L. Wang, S. Matsuo, M. Shirase, and T. Takagi, “Solving a 676-bit discrete logarithm problem in $GF(3^{6n})$,” in *PKC 2010*, LNCS 6056, pp. 351–367.
- [11] T. Hayashi, T. Shimoyama, N. Shinohara, and T. Takagi, “Breaking pairing-based cryptosystems using η_T pairing over $GF(3^{97})$,” in *ASIACRYPT 2012*, LNCS 7658, pp. 43–60.
- [12] K. Aoki, J. Franke, T. Kleinjung, A. Lenstra, and D. A. Osik, “A kilobit special number field sieve factorization,” in *ASIACRYPT 2007*, LNCS 4833, pp. 1–12.
- [13] S. Bai, A. Filbois, P. Gaudry, A. Kruppa, F. Morain, E. Thomé, and P. Zimmermann, “CADO-NFS: Crible algébrique: Distribution, optimisation – Number field sieve,” <http://cado-nfs.gforge.inria.fr/>.
- [14] O. Schirokauer, “Discrete logarithms and local units,” *Phil. Trans. R. Soc. A*, 345(1676):409–423, 1993.
- [15] D. M. Gordon and K. S. McCurley, “Massively parallel computation of discrete logarithms,” in *CRYPTO ’92*, LNCS 740, pp. 312–323.
- [16] L. Hars, “Applications of fast truncated multiplication in cryptography,” *EURASIP J. Embed. Syst.*, 2007:061721, 2007.
- [17] J. M. Pollard, “The lattice sieve,” in A. K. Lenstra and H. W. Lenstra, eds., *The development of the number field sieve*, LNM 1554, pp. 43–49, 1993.
- [18] K. Aoki and H. Ueda, “Sieving using bucket sort,” in *ASIACRYPT 2004*, LNCS 3329, pp. 92–102.
- [19] V. Shoup, “NTL: A library for doing number theory,” <http://www.shoup.net/ntl/>.
- [20] F. Chabaud and R. Lercier, “ZEN: A toolbox for fast computation in finite extension over finite rings,” <http://zenfact.sourceforge.net/>.
- [21] R. P. Brent, P. Gaudry, E. Thomé, and P. Zimmermann, “gf2x: A library for multiplying polynomials over the binary field,” <http://gf2x.gforge.inria.fr/>.
- [22] A. Joux and R. Lercier, “The function field sieve in the medium prime case,” in *EUROCRYPT 2006*, LNCS 4004, pp. 254–270.
- [23] Y. Kawahara, K. Aoki, and T. Takagi, “Faster implementation of η_T pairing over $GF(3^m)$ using minimum number of logical instructions for $GF(3)$ -addition,” in *Pairing 2008*, LNCS 5209, pp. 282–296.