

Lattice-Based Cryptography: From Practice to Theory to Practice

Vadim Lyubashevsky, ENS Paris

Lattice-based cryptography is currently seen as one of the most promising alternatives to cryptography based on number theory. The major advantages of lattice-based protocols is that they are faster than ones based on number theory and they also seem to be resistant against quantum attacks. The origins of lattice-based cryptography, in the mid 90's, can be attributed to the “practical” NTRU cryptosystem and the “theoretical” constructions of Ajtai. Initially, these areas developed independently, but work done in the past few years showed surprising connections between the two and there has been a lot of recent success in bridging the gap between the practical and theoretical aspects of lattice-based cryptography. On the one hand, we have used the theoretical ideas to construct new practical, provably-secure schemes, and on the other hand, we have also been able to prove that (with a few modifications) the early practical schemes are actually provably secure. In this talk, I will explain the current hardness assumptions used for constructing efficient lattice-based schemes as well as present some sample constructions.

Analyse de BKZ

Xavier Pujol – travail conjoint avec Guillaume Hanrot et Damien Stehlé
ÉNS de Lyon - LIP/CNRS/INRIA/U. Lyon/UCBL

11 février 2011

La réduction forte de réseaux est la clé de la plupart des attaques contre les cryptosystèmes basés sur les réseaux. Entre la réduction HKZ, forte mais trop coûteuse, et la réduction LLL, plus faible mais rapide, existent plusieurs tentatives d'obtenir des compromis efficaces. Celle qui semble atteindre le meilleur rapport temps/qualité est la réduction BKZ, introduite par Schnorr et Euchner en 1991. Cependant, aucune borne de complexité raisonnable sur cet algorithme n'était connue jusqu'à maintenant. Nous prouvons qu'après $\tilde{O}(\frac{n^3}{k^2})$ appels à une routine de HKZ-réduction, BKZ_k renvoie une base telle que la norme du premier vecteur est bornée par $\approx \gamma_k^{\frac{n}{2(k-1)}} (\det L)^{\frac{1}{n}}$. Le principal outil de la preuve est l'analyse d'un système dynamique linéaire associé à cet algorithme.

Introducing L1 a quasi-linear LLL

Andy Novocin, LIP, ENS Lyon

The famous LLL lattice reduction algorithm of 1982 has applications in many computational fields which include but aren't limited to cryptography. The most common use for the algorithm is as a quick-and-dirty way to approximate solutions to the shortest vector problem (which is NP-complete in general). Certain lattices can be devised with low dimension and large bit size for which short vectors can lead to attacks on the RSA cryptography scheme. Classical LLL has a complexity bound which is polynomial and cubic in the bit-size of the input lattice. We present the first lattice reduction algorithm with a polynomial complexity bound which is quasi-linear in the bit-size of the input lattice.

Our approach mimics the Schönhage-Strassen approach for quasi-linear GCD computations. To prove the validity of our attack we developed a class of independently interesting lattice reduction tools. We present this new framework which allows better analysis for reductions of lattices which are shifts of reduced lattices. We then present a general strategy by which lattice reduction can be performed with a sequence of these shifted-reductions.

Modélisations de l'algorithme LLL

Résumé pour C2, Avril 2011

Mariya Georgieva

LLL est un algorithme de réduction des réseaux inventé en 1982 par L. Lenstra, A. Lenstra et L. Lovász. Cet algorithme, initialement présenté dans le cadre de la factorisation des polynômes, a trouvé de nombreuses applications en théorie des nombres (conjecture de Mersenne, approximations diophantiennes, ...), en programmation linéaire entière (résolution d'inéquations linéaires), en algèbre (factorisation de polynômes ou d'entiers), etc. Depuis une trentaine d'années, les réseaux et l'algorithme LLL ont connu un véritable essor en cryptologie: cryptanalyse de protocoles de type RSA ou sac-à-dos, conception de cryptosystèmes (Ajtai-Dwork, NTRU, GGH...), preuves de sécurité, conception de protocoles homomorphes, ...

Tous ces travaux sont liés à la notion de réduction d'un réseau. Un réseau est l'ensemble des combinaisons linéaires entières de vecteurs linéairement indépendants et appelés base du réseau. Un réseau admet une infinité de bases. La réduction d'un réseau consiste à trouver une "bonne base", avec des vecteurs assez courts et assez orthogonaux à partir d'une base quelconque. En cryptanalyse après linéarisation et construction d'un réseau adéquat, casser un système est équivalent à trouver un vecteur le plus court dans le réseau. Ce problème, connu sous le nom de SVP (Shortest Vector Problem), est prouvé NP-Difficile sous certaines conditions ce qui démontre la difficulté de réduire une base.

L'algorithme LLL construit en temps polynomial une base dite LLL-réduite où la norme des vecteurs et leur orthogonalité sont maîtrisées. La dynamique de l'algorithme est encore bien mal connue. Les comportements probabilistes de ses principaux paramètres (temps d'exécution, configuration de sortie) sont difficiles à analyser en grande dimension. Le temps d'exécution est souvent meilleur que le pire des cas et il dépend beaucoup des entrées et de la stratégie. L'analyse en moyenne des algorithmes de réduction en dimension 1 et 2 (algorithmes du PGCD et algorithme de Gauss) ont été réalisées par le groupe d'Analyse Dynamique de Caen. Il semble difficile d'étendre la méthodologie à l'algorithme LLL tant sa dynamique est compliquée. L'idée est donc de passer par des modèles simplifiés qui conservent le même comportement global de LLL.

Dans un premier temps, je décrirai des modèles simplifiés à base de tas de sable (systèmes dynamiques discrets) visant à mieux comprendre la dynamique de l'algorithme LLL. A partir de ces modèles, je présenterai des conjectures qui seront ensuite confirmées ou infirmées par des expérimentations.

Le comportement probabiliste de LLL diffère aussi selon le type des entrées. Les réseaux qui interviennent en cryptologie ont une forme très particulière, et ne peuvent être considérées comme génériques. Ces réseaux cryptographiques sont très différents selon leur provenance. Les réseaux de Coppersmith, issus de la méthode de Coppersmith (qui permet de trouver des petites racines de polynômes univariés ou multivariés modulo un entier) conduisent à un comportement de LLL très différent des réseaux de Ajtai (qui correspondent à des instances difficiles), des réseaux de type sac-à-dos (instances "faciles") ou bien des réseaux liés au système NTRU.

La deuxième partie de l'exposé sera consacrée à l'étude de ces entrées particulières, de leurs propriétés géométriques et des conséquences algorithmiques.

Un nouveau protocole d'authentification zero-knowledge basé sur la théorie des codes correcteurs optimisant le coût des communications.

Carlos Aguilar, Philippe Gaborit, Julien Schrek
Université de Limoges, XLIM-DMI,
123, Av. Albert Thomas 87060 Limoges Cedex France
{carlos.aguilar, philippe.gaborit, julien.schrek}@xlim.fr

Nous présentons ici un nouveau protocole d'authentification zero-knowledge 5 passes dont la sécurité est basée sur le problème du *syndrome decoding* binaire ainsi que l'une de ses variantes utilisant les codes doublement circulants. Le paramètre le plus important pour les algorithmes zero-knowledge sur les codes correcteurs est le coût total des communications (le nombre de bits échangé lors du protocole) qui est de l'ordre de la centaine de millier de bits. Ce nouvel algorithme permet de réduire ces coûts de communication de 30% par rapport aux meilleurs protocoles basés sur le même problème pour une sécurité équivalente. Pour y parvenir, les deux idées présentées consistent d'une part à baisser la probabilité de tricher à $\frac{1}{2}$, d'autre part à adopter une gestion plus économique des engagements.

La première idée permet de diminuer la probabilité de triche à une valeur proche de $\frac{1}{2}$ et consiste à stocker un nombre plus important d'instance du *syndrome decoding problem* dans les clés en vue d'augmenter considérablement le nombre de challenges. Ce grand nombre de challenge permet d'avoir une probabilité de tricher proche de $\frac{1}{2}$ tout en restant en binaire alors que les protocoles de Stern [3] et Véron [4] ont une probabilité de tricher proche de $\frac{2}{3}$ en binaire et que le protocole de Cayrel-Véron [1] à une probabilité de tricher proche de $\frac{1}{2}$ sur des corps plus gros. Ce protocole s'associe très bien avec le fait de choisir une matrice doublement circulante. En effet on peut dériver plusieurs instances du *syndrome decoding problem* à partir d'une seule à l'aide de cette matrice comme on peut le voir ci-dessous.

$$H(e_1|e_2)^t = s \Rightarrow H(\sigma(e_1)|\sigma(e_2))^t = \sigma(s)$$

avec σ une permutation doublement circulante et H une matrice doublement circulante. De plus, les clés du protocole se représentent de façon beaucoup plus compacte comme il a déjà été développé dans [2].

La deuxième idée présentée est une façon générale de faire baisser la taille des communications d'une authentification ou la taille d'une signature par une meilleur gestion des engagements. Il s'agit en fait de cumuler tous les engagements dans un seul haché et d'envoyer seulement les engagements utiles avant la vérification. Cette méthode permet d'envoyer seulement un haché à chaque tour.

Une comparaison entre les différents protocoles basés sur la théorie des codes correcteurs sera faite pour montrer les avantages de notre protocole, à savoir des communications 30% plus faibles (19305 bits pour une sécurité de 2^{-16}) pour des tailles de clés similaires aux autres protocoles (clé privée de 698 et une clé publique de 1057).

D'autre part un autre avantage de ce protocole est qu'il n'utilise que très peu de ressources : 5360 de RAM pour un temps d'exécution très rapide.

Références

- [1] Pierre-Louis Cayrel and Pascal Véron. Improved code-based identification scheme. *CoRR*, abs/1001.3017, 2010.
- [2] Philippe Gaborit and Marc Girault. Lightweight code-based identification and signature. In *Information Theory, 2007. ISIT 2007. IEEE International Symposium on*, pages 191 –195, 06 2007.
- [3] Jacques Stern. A new identification scheme based on syndrome decoding. In Douglas R. Stinson, editor, *CRYPTO*, volume 773 of *Lecture Notes in Computer Science*, pages 13–21. Springer, 1993.
- [4] Pascal Véron. Improved identification schemes based on error-correcting codes. *Appl. Algebra Eng. Commun. Comput.*, 8(1) :57–69, 1996.

Un schéma de fingerprinting asymétrique utilisant les codes de Tardos

Ana Charpentier, Caroline Fontaine, Teddy Furon et Ingemar Cox

Nous parlons ici de *fingerprinting*, désigné aussi par *transactional watermarking*. Le problème est le suivant : un distributeur de données numériques (images, audio, vidéo,...) distribue des copies d'un contenu à n utilisateurs. On insère dans chaque copie un identifiant afin de repérer d'éventuels fraudeurs. Quelques uns de ces utilisateurs appelés *colluders* sont malhonnêtes, et mélangent leurs contenus pour créer une copie pirate qui est redistribuée illégalement. On souhaite retrouver leur identité en analysant la copie pirate. Les meilleurs codes pour générer des identifiants qui permettent de contrer de telles coalitions sont les codes de Tardos. On s'intéresse ici à la conception d'un schéma de fingerprinting dit *asymétrique*, dans lequel les utilisateurs participent à la génération de leur identifiant, et le vendeur n'a pas connaissance de l'identifiant associé à chaque utilisateur. Cela empêche le vendeur de pouvoir accuser sciemment un utilisateur innocent. De tels schéma ont déjà été étudiés [3], mais sans s'attacher à des codes particuliers. Or, le choix des codes est déterminant si l'on souhaite obtenir un schéma efficace dans la pratique. La particularité de notre travail est de proposer un protocole asymétrique utilisant les codes anti-collusion de Tardos.

Code de Tardos. En 2003, G. Tardos a présenté une famille de codes de fingerprinting probabilistes très efficaces [1]. Leur intérêt réside surtout dans leur longueur réduite et dans la facilité de génération des mots de code, qui servent d'identifiants. B. Skoric *et al* ont proposé une version améliorée de ces codes, sur laquelle nous nous appuyons [2]. Soient n le nombre d'utilisateurs, m la longueur du code et c le nombre maximum de *colluders*. Le distributeur de contenu tire une fois pour toutes m valeurs $p_i \in [0, 1]$ i.i.d. suivant la pdf $f(p) = \frac{1}{\pi\sqrt{p(1-p)}}$. Ce vecteur $\mathbf{p}$ ne doit pas être connu des utilisateurs. On note $\mathbf{X}$ la matrice contenant l'ensemble des mots du code, $\mathbf{X}_j = (X_{j1}, X_{j2}, \dots, X_{jm})$ désignant le mot de code de l'utilisateur j . Les X_{ji} sont tirés de façon aléatoire, en suivant $\mathbb{P}(X_{ji} = 1) = p_i$. Chaque utilisateur reçoit une copie du contenu (image, vidéo) contenant une marque différente. Lors de l'accusation, soit $\mathbf{Y}$ la marque extraite de la copie pirate. Avec cette marque et le mot de code $\mathbf{X}_j$ on calcule un score S_j d'accusation pour l'utilisateur j . Le score est calculé de telle sorte que les plus grands scores sont en espérance ceux des *colluders* : $S_j = \sum_{i=1}^m g_{Y_i, X_{ji}}(p_i)$. Les fonctions $g_{Y_i, X_{ji}}$ ainsi que les valeurs de m et du seuil au-delà duquel un score accuse un utilisateur sont spécifiées dans [2].

Attaque d'un vendeur malhonnête lors du processus d'accusation. Nous avons montré qu'il était possible au vendeur de tricher lors du processus d'accusation en modifiant légèrement les valeurs du vecteur $\mathbf{p}$. Cela a pour effet d'augmenter la valeur de tous les scores. Ainsi, le score d'un innocent peut très bien se trouver finalement dépasser le seuil et donner lieu à une accusation à tort. Cette étude montre qu'il faut prendre beaucoup de précautions lors de l'utilisation du vecteur $\mathbf{p}$, et mettre en œuvre des moyens pour garantir son intégrité.

Un protocole asymétrique. La principale difficulté pour utiliser des codes de Tardos dans un protocole asymétrique vient du fait que l'identifiant doit être généré par l'utilisateur, conformément au vecteur $\mathbf{p}$ qu'il ne doit pas connaître. Ainsi l'utilisateur doit générer un bit qui suit une certaine distribution ($\mathbb{P}(X_{ji} = 1) = p_i$) sans connaître la valeur de p_i . Nous résolvons ce problème à l'aide d'un protocole d'Oblivious Transfer (OT). L'Oblivious Transfer est une primitive cryptographique qui permet à Bob de choisir k éléments au hasard dans une liste de N éléments possédée par Alice, de telle sorte que : (1) Bob obtient des éléments qui sont réellement dans la liste ; (2) Bob n'obtient pas d'information sur les éléments qu'il n'a pas choisis ; (3) Alice ne sait pas quels sont les éléments choisis par Bob. Nous quantifions ici chaque valeur de p_i en $p_i = L_i/N$ pour donner une liste de bits telle que dans une liste de longueur N , il y a L_i bits de valeur '1'. Nous discutons l'adéquation des protocoles OT existants, en les adaptant à nos contraintes, parfois même en les améliorant. Certains sont issus de la communauté crypto (et garantissent une sécurité plus formelle), tandis que d'autres relèvent d'autres approches, comme le Commutative Encryption Scheme [4], dont la sécurité est moins formalisée, mais qui s'adapte mieux à nos contraintes.

Mais le vendeur doit quand même posséder de l'information sur l'utilisateur afin de lancer son protocole d'accusation en cas de découverte d'une copie falsifiée. C'est pourquoi nous avons une deuxième phase dans laquelle l'utilisateur révèle une partie (qu'il ne choisit pas lui-même) de son identifiant au vendeur, en utilisant à nouveau un OT. Lors de cet échange, nous nous assurons que les protagonistes obtiennent bien ni plus ni moins que ce qu'ils doivent apprendre.

Une troisième phase concerne l'insertion dans le medium (image, vidéo) de l'identifiant, par un procédé de tatouage qui permet au vendeur d'insérer dans le document un identifiant sans le connaître. Nous utilisons à cette étape des schémas existants reposant sur du chiffrement homomorphe.

Références

- [1] G.Tardos. *Optimal probabilistic fingerprint codes*. Proc. of the 35th annual ACM symposium on theory of computing, 2003
- [2] B.Skoric, S.Katzenbeisser et M.Celik. *Symmetric Tardos fingerprinting codes for arbitrary alphabet sizes*. DCC, 2008.
- [3] B. Pfitzmann, M.Schunter : *Asymmetric fingerprinting*, EUROCRYPT 96.
- [4] , Bao, F. and Deng, R.H. and Feng, P., *An efficient and practical Scheme for Privacy Protection in the E-Commerce of Digital Goods*, ICISC 2000, 2001, 2015, LNCS, 162-170, Springer-Verlag,