



A Systematic Study of Sparse LWE

Aayush Jain^{1,2(✉)}, Huijia Lin^{1,2}, and Sagnik Saha^{1,2}

¹ Carnegie Mellon University, Pittsburgh, PA, USA
{aayushja,sagniks}@andrew.cmu.edu

² University of Washington, Seattle, WA, USA
rachel@cs.washington.edu

Abstract. In this work, we introduce the sparse LWE assumption, an assumption that draws inspiration from both Learning with Errors (Regev JACM 10) and Sparse Learning Parity with Noise (Alekhnovich FOCS 02). Exactly like LWE, this assumption posits indistinguishability of $(\mathbf{A}, \mathbf{s}\mathbf{A} + \mathbf{e} \bmod p)$ from $(\mathbf{A}, \mathbf{u})$ for a random $\mathbf{u}$ where the secret $\mathbf{s}$, and the error vector $\mathbf{e}$ is generated exactly as in LWE. However, the coefficient matrix $\mathbf{A}$ in sparse LPN is chosen randomly from $\mathbb{Z}_p^{n \times m}$ so that each column has Hamming weight exactly k for some small k . We study the problem in the regime where k is a constant or polylogarithmic. The primary motivation for proposing this assumption is efficiency. Compared to LWE, the samples can be computed and stored with roughly $O(n/k)$ factor improvement in efficiency. Our results can be summarized as:

- **Foundations:** We show several properties of sparse LWE samples, including: 1) The hardness of LWE/LPN with dimension k implies the hardness of sparse LWE/LPN with sparsity k and arbitrary dimension $n \geq k$. 2) When the number of samples $m = \Omega(n \log p)$, length of the shortest vector of a lattice spanned by rows of a random sparse matrix is large, close to that of a random dense matrix of the same dimension (up to a small constant factor). 3) Trapdoors with small polynomial norm exist for random sparse matrices with dimension $n \times m = O(n \log p)$. 4) Efficient algorithms for sampling such matrices together with trapdoors exist when the dimension is $n \times m = \tilde{O}(n^2)$.
- **Cryptanalysis:** We examine the suite of algorithms that have been used to break LWE and sparse LPN. While naively many of the attacks that apply to LWE do not exploit sparsity, we consider natural extensions that make use of sparsity. We propose a model to capture all these attacks. Using this model we suggest heuristics on how to identify concrete parameters. Our initial cryptanalysis suggests that concretely sparse LWE with a modest k and slightly bigger dimension than LWE will satisfy similar level of security as LWE with similar parameters.

Aayush Jain was supported by Google and Cylab of Carnegie Mellon University. Huijia Lin was supported by NSF grants CNS-1936825 (CAREER), CNS-2026774, a JP Morgan AI research Award, a Cisco research award, and a Simons Collaboration on the Theory of Algorithmic Fairness.

- **Applications:** We show that the hardness of sparse LWE implies very efficient homomorphic encryption schemes for low degree computations. We obtain the first secret key Linearly Homomorphic Encryption (LHE) schemes with *slightly super-constant*, or even *constant*, overhead, which further has applications to private information retrieval, private set intersection, etc. We also obtain secret key homomorphic encryption for arbitrary constant-degree polynomials with slightly super-constant, or constant, overhead.

We stress that our results are preliminary. However, our results make a strong case for further investigation of sparse LWE.

1 Introduction

The celebrated work by Regev [63] proposed the Learning With Errors (LWE) assumption, postulating that noisy linear samples are indistinguishable from random,

$$(\mathbf{A}, \mathbf{b} = \mathbf{s}\mathbf{A} + \mathbf{e} \bmod p) \approx_c (\mathbf{A}, \mathbf{u})$$

where the coefficient matrix $\mathbf{A} \leftarrow \mathbb{Z}_p^{n \times m}$ and the secret vector $\mathbf{s} \leftarrow \mathbb{Z}_p^n$ are sampled randomly, the entries of the noise vector $\mathbf{e} \leftarrow \mathcal{D}_{\mathbb{Z}, \sigma}^{1 \times m}$ are sampled from the discrete Gaussian distribution with width σ smaller than p , and the vector $\mathbf{u} \leftarrow \mathbb{Z}_p^{1 \times m}$ is uniformly random.

The LWE assumption is extremely versatile, and has been the basis of a plethora of constructions. This spans the development of core cryptographic tools including key agreement, public key encryption, and signature schemes that are ready to be deployed [2, 3], schemes and protocols with richer functionalities and nearly practical efficiency e.g., [18, 42, 43, 51], to advanced primitives pushing the envelop of cryptographic feasibility, such as, Fully Homomorphic Encryption [21, 26, 32, 36, 53], Attribute-Based Encryption [20, 39], Succinct Arguments for P [27] and many others [22, 40, 41, 54, 70]. The hardness of the LWE problem has been extensively studied. It was shown that the hardness of LWE is implied by the worst-case hardness of certain lattice problems such as GapSVP [25, 50, 61, 63], which is widely believed to be subexponentially hard for a range of parameters.

Many variants of the LWE problems have been explored, encompassing 1) the use of structured secrets, such as binary random secrets and entropic secrets, 2) the use of structured noise, e.g., binary random noise, uniform random noise of bounded magnitude, and rounded version known as Learning With Rounding (LWR) [17], and 3) the use of structured matrices, such as those implicit in Ring LWE (RLWE) [52] or middle-point LWE [64]. All these variants have significantly extended the landscape of functionalities, efficiency and security levels, and techniques supported by the LWE family of assumptions.

The Sparse Learning With Errors Problems. This work proposes a new variant called the k -Sparse LWE (k -sLWE) problem. We conjecture the hardness of the LWE problem when the coefficient matrices $\mathbf{A} \in \mathbb{Z}_p^{n \times m}$ are **random k -sparse**

matrices (as opposed to random dense matrices), where each column has exactly k non-zero entries for some $k \leq n$ and every non-zero entry is sampled uniformly and randomly from $\mathbb{Z}_p$. Sparsity is the new feature – the smaller the sparsity parameter k is, the more efficient it is to store and compute the samples. Each sparse LWE sample $(\mathbf{a}_i, \mathbf{s}\mathbf{a}_i + e_i)$ contains only $k + 1$ elements in $\mathbb{Z}_p$ and takes $\tilde{O}(k \log p)$ time to compute, whereas a standard LWE sample contains $n + 1$ elements and takes $\tilde{O}(n \log p)$ time to compute. At the same time, k is crucial for the hardness of the problem. A variety of settings of k is interesting; k can be constant, logarithmic, polylogarithmic, or sublinear in the dimension n . We found that k being logarithmic or polylogarithmic is the most suitable for cryptographic applications, ensuring efficiency and strong security. While the assumption with constant k and appropriately bounded number of samples might also be useful for deriving interesting theoretical results, constant k will not provide almost exponential security when the number of samples is superlinear in the dimension n .

Sparse LWE is a natural question to study. The effect of sparsity in the coefficient matrix has long been studied in the context of the Learning Parity with Noise (LPN) problem, a close cousin of LWE, which uses *sparse* random noise, instead of *small* (but dense) Gaussian noise as in LWE. Sparse LPN further uses random k -sparse coefficient matrices as described above. Variants of the sparse LPN assumption in the binary field $\mathbb{F}_2$ have been proposed and studied for at least a couple of decades in average-case complexity and are intimately connected to constraint satisfaction problems (see works such as [8, 9, 14, 29, 31, 37, 46, 59]). The work of [12] generalized the assumption to large fields $\mathbb{F}_q$. These assumptions have been used in a number of applications, for example [8, 10, 13, 28, 30, 44]. Among them, the work of [10] constructed a public-key encryption scheme from k -sparse LPN for a constant $k \geq 3$ and bounded polynomial number of samples, while the recent work of [30] constructed multi-party Homomorphic Secret Sharing and multi-party computation protocols with sublinear communication from k -sparse LPN for any polylogarithmic k and unbounded polynomial number of samples. The related assumption of local PRGs [37] has also been used in a number of works including the recent construction of a program obfuscation scheme [45].

It is somewhat surprising that the sparse variant of LWE has not been investigated, to the best of our knowledge. That is the purpose of this work. Sparse LWE is related to constraint satisfaction problems like sparse LPN and Goldreich’s one-way function and PRG, and at the same time draws connection to lattices. It is curious and interesting to ask:

How hard is the sparse LWE problem?

What properties do sparse LWE samples have?

What can the hardness of sparse LWE enable cryptographically?

We give initial answers to these questions, and ask new questions. Our preliminary results include:

- Foundations. We show several properties of sparse LWE samples, including:
 - 1) the hardness of LWE with dimension k implies the hardness of sparse LWE

with sparsity k and arbitrary dimension $n \geq k$. Our reduction is oblivious to the error model and hence the results also apply to sparse LPN.

2) When the number of samples $m = \Omega(n \log p)$, the length of the shortest vector of the lattice spanned by the rows of a random sparse matrix is large, close to that of a random dense matrix of the same dimensions (up to a small constant factor).

3) Trapdoors with small polynomial norm exist for random sparse matrices with dimension $n \times m = O(n \log p)$. An interesting consequence of the first claim from [34] is that when $m = O(n \log p)$, the distribution $(\mathbf{A}, \mathbf{A}\mathbf{r} \bmod p)$ for a vector $\mathbf{r} \in \mathbb{Z}_p^m$ is statistically close to random in $\mathbb{Z}_p^n$ provided $\mathbf{r}$ is chosen from a polynomially wide discrete Gaussian over $\mathbb{Z}^m$. This might be useful for the design of public-key primitives.

4) While we don't know how to sample random sparse matrices with trapdoors when $m = O(n \log p)$, we can design an efficient algorithm for sampling such matrices together with trapdoors when the dimension is $n \times m = O(n^2 \log n \log p)$.

- Cryptanalysis. Examining existing attacks on the LWE problem, the structure of sparse matrices does not naively speed up combinatorial [19], lattice reduction [48, 66], or algebraic attacks [15]. Moreover, many attacks that would typically apply to sparse LPN become harder to apply to sparse LWE due to the fact that the noise is dense (though small). We also explore attacks that directly leverage sparsity, and show that they are ineffective when the sparsity parameter k is polylogarithmic and number of samples m is (unbounded) polynomial. Overall, our initial cryptanalysis optimistically suggests that when all other parameters are equal, namely the modulus p , the noise width σ , and polynomial number of samples m , polylogarithmic-sparse LWE with slightly larger dimension $n' = \Theta(n)$, is comparable to LWE with dimension n . We also give some concrete parameters based on our initial cryptanalysis.
- Applications. The hardness of sparse LWE implies very efficient homomorphic encryption schemes for low degree computations. In particular, we obtain the first secret key Linearly Homomorphic Encryption (LHE) schemes with *slightly super-constant*, or even *constant*, overhead, which further has applications to private information retrieval, private set intersection, etc. We also obtain secret key homomorphic encryption for arbitrary constant-degree polynomials with slightly super-constant, or constant, overhead. In comparison, to the best of our knowledge, previous homomorphic encryption schemes have (large) polylogarithmic overhead [33]. In particular, our LHE scheme is concretely efficient.

We emphasize that our cryptanalysis is preliminary, and more thorough study must be conducted before we can rest our confidence on the hardness of sparse LWE. Nevertheless, in order to assess the potential benefits of sparse LWE and provide motivation for further study, we make tentative suggestions for concrete parameters to facilitate comparison of efficiency.

Comparison of Sparse LWE with LWE and RLWE. The main motivation behind introduction of sparse LWE is efficiency. Storing and computing sparse LWE samples are roughly $\tilde{O}(n/k)$ times more efficient compared to plain LWE. Game-changing efficiency improvements have been afforded by Ring LWE [52] due to its nice properties (such as SIMD compatibility [68]). However, there are few settings where rings may not be ideal, for instance, when we want to access data as integer elements rather than as ring elements. A good example of this is private-information retrieval applications [43]. We believe that in those settings sparse LWE might form an appealing alternative to ring LWE.

Open Problems. Our work leaves a great supply of open problems from various angles. On the complexity side, we ask if there are new attacks on sparse LWE? Can we show worst-case to average-case reductions or leakage resilience properties similar to LWE? Similarly, it is completely unclear if sparse LWE with small secrets is hard. On the concrete attacks and algorithms side, a systematic study of concrete parameters is warranted. On the applications side, it would be good to study more applications of sparse LWE beyond what we mentioned.

2 Technical Outline

We now describe our results with more context. The sLWE assumption is parameterized by a sparsity parameter k , dimension n , sample complexity m , noise parameter σ and modulus p . It posits that $(\mathbf{A}, \mathbf{s}\mathbf{A} + \mathbf{e} \bmod p)$ is indistinguishable from $(\mathbf{A}, \mathbf{u})$ where $\mathbf{A}$ is a random matrix in $\mathbb{Z}_p^{n \times m}$ such that each column is k -sparse, $\mathbf{s} \leftarrow \mathbb{Z}_p^{1 \times n}$ is randomly chosen, $\mathbf{e} \in \mathbb{Z}^{1 \times m}$ is generated from discrete Gaussian with width σ and $\mathbf{u} \in \mathbb{Z}_p^{1 \times m}$ is randomly chosen. Our results can be divided into three categories: Foundations, Cryptanalysis and Applications.

2.1 Foundations

We start by investigating some basic properties of sparse LWE. A very natural question to ask is how much is the minimum distance of the lattice $\mathcal{L} = \mathcal{L}(\mathbf{A})$ given by $\mathcal{L} = \{\mathbf{x}\mathbf{A} + p\mathbb{Z}^m : \mathbf{x} \in \mathbb{Z}^{1 \times n}\}$ as a function of k , n , m and p . This should be as large as possible, and certainly much larger than the norm of the error vector, so that given m sLWE samples the planted secret $\mathbf{s}$ is uniquely determined. In Sect. 4.1, we show that when $m = \Omega(n \log p)$ as long as $k > 1$ and p is polynomial in n , the minimum distance is $\Omega(p)$. In particular, once the number of samples exceeds $8n \log n$, the minimum distance properties of $\mathcal{L}$ behave roughly similarly as in regular LWE samples.

An immediate consequence of this from the Transference theorems of Banaszczyk [16] (and as argued in corollary 3.1) is that there exist m linearly independent vectors of bounded polynomial norm in the lattice $\mathcal{L}^\perp(\mathbf{A}) = \{\mathbf{x} \in \mathbb{Z}^m : \mathbf{A}\mathbf{x} = \mathbf{0}\}$ once $m = \Omega(n \log p)$. Thus, in principle when $m = \Omega(n \log p)$, $\mathbf{A}$ satisfies nice properties similar to a dense matrix corresponding to LWE distribution. For instance, one immediate consequence from the work of [34] is that $\mathbf{A}$

can be used to extract from Gaussian sources. In other words, if $\mathbf{r}$ is chosen from a wide enough (polynomial width) discrete Gaussian, the induced distribution on $\mathbf{A}\mathbf{r}$ is statistically close to a random vector in $\mathbb{Z}_p^n$ even if $\mathbf{A}$ is public.

We ask when can we sample a random sparse matrix $\mathbf{A}$ along with such a trapdoor. While we leave sampling random sparse matrices $\mathbf{A}$ along with a short trapdoor for $m = O(n \log p)$ as a great open question, in Sect. 6, we show that it is indeed possible to sample such sparse matrices along with lattice trapdoors for $m = O(n^2 \log n \log p)$. Such linearly independent short vectors in $\mathcal{L}^\perp(\mathbf{A})$ can be used to sample from a discrete Gaussian distribution on vectors in $\mathcal{L}^\perp(\mathbf{A})$ using lattice sampling methods from prior work such as [34, 58]. The current approach for sampling fully random (dense) matrices along with trapdoors works by first sampling a matrix $\mathbf{A}$ at random and then setting $\mathbf{V} = \mathbf{A}\mathbf{R} + \mathbf{G}$ for a random small norm matrix $\mathbf{R}$ and the gadget matrix $\mathbf{G}$. This gives a small norm matrix $\mathbf{X}$ so that $[\mathbf{A}|\mathbf{V}]\mathbf{X} = \mathbf{G}$. Micciancio and Peikert [58] show how to convert $\mathbf{X}$ to a trapdoor for $[\mathbf{A}|\mathbf{V}]$. This approach does not work directly for us because unlike the case with dense matrices, $\mathbf{V} = \mathbf{A}\mathbf{R} + \mathbf{G}$ is not sparse anymore. The main technical highlight of our algorithm is a novel combinatorial approach to get around this issue and sample random sparse matrices along with trapdoors.

We then ask if we can show asymptotic polynomial time hardness of sparse LWE. It would be nice to show a result of the form: “if LWE is hard, then sparse LWE is also hard”. Such results are not known even in the setting of LPN/sparse LPN which has received a lot of study. While we do not have a fully general result here, in Sect. 5, we show quite an interesting reduction.

Specifically, sparse LWE with sparsity parameter k , dimension $n > k$, sample complexity m , prime modulus p and the noise parameter σ is harder than LWE with dimension k , sample complexity m , modulus p and the noise parameter σ . Our reduction is oblivious to the error model and also applies to the case of sparse LPN/LPN. Assuming subexponential hardness for LWE, this establishes polynomial time hardness for sparse LWE provided k is polylogarithmic. We leave showing such hardness results for constant k to future work.

2.2 Cryptanalysis

We perform some preliminary cryptanalysis of this new assumption and also suggest some concrete parameters based on that. We stand on the tall shoulders of a body of great work done in analyzing security of sparse LPN [8–10, 13, 14, 29, 31, 37, 44, 46, 59] as well as LWE [1, 4, 5, 7, 62].

In sLWE sparsity creates some interesting effects. For example, when $m \geq k \binom{n}{k} + 1$ we are guaranteed to find at least $k+1$ samples supported on the same set of k variables, which might be enough to distinguish a random vector from one close to the lattice. More generally, one can potentially use an algorithm that breaks LWE for some small dimension whenever there exists a set of t samples supported on less than t distinct variables, with t being sufficiently small. However, we believe that once k is large enough, we can release polynomially many samples without compromising security. Most of our cryptanalysis deals with the case of $k = O(\log n)$ so that we do not have to worry about this effect.

We now discuss a few different types of attacks against sparse LWE. Some of our attacks are inspired by existing algorithms for sparse LPN and LWE attacks, while others leverage the sparsity of $\mathbf{A}$ more directly.

Sparse LPN Style Attacks. We first observe that many of the attacks that apply to sparse LPN do not apply in this case because our error vectors are dense as opposed to being sparse. In the main attack that still might apply how it does for sparse LPN, one aims to find sparse vectors $\mathbf{x} \in \mathbb{Z}_p^m$ so that $\mathbf{A}\mathbf{x} = \mathbf{0}$. Such a t -sparse $\mathbf{x}$ can break sparse LPN if the error probability is less than $\frac{1}{t}$. We show that if we can find a t -sparse vector $\mathbf{x}$ like this for very small t , it can lead to attacks against sparse LWE in certain cases. The idea is that one can multiply $s\mathbf{A} + \mathbf{e}$ with $\mathbf{x}$ to yield $\mathbf{e}\mathbf{x}$. Observe that $\mathbf{e}\mathbf{x}$ only touches upon t coordinates of small norm error vector $\mathbf{e}$. If t is very small and the error bound B is also very small so that $B^t \ll p$ one could potentially use brute force to iterate over all the relevant choices of error to learn with high probability the error coordinates corresponding to the t non-zero locations of $\mathbf{x}$. To rule out such attacks we propose that $B^t \gg p$. In Sect. 7.1, we show that when $k = \log n$ or above, t is linear in n as long as the sample complexity $m = \text{poly}(n)$. Hence if $p \ll B^t$, this attack will either not apply or take exponential time. We believe that this attack is easy to circumvent by choosing parameters carefully as it applies less obviously to sparse LWE than sparse LPN. For typical settings we would work with polynomial sized modulus and polynomial norm errors and so t only needs be moderately big. This can be easily ensured by setting k to be bigger than some large enough constant.

Attacks Inspired by LWE. We have a great body of work on LWE algorithms consisting of algebraic [15], combinatorial [5, 19, 49], lattice-reduction / geometric attacks [48, 65, 67] and reductions for LWE [25, 61, 63]. We observe that naively these algorithms do not exploit the sparsity structure of the coefficients directly that well. It is however plausible that some algorithms fare marginally better with sparse matrices – we leave practical investigation of this as an open question for future work to explore. We observe that the natural way for LWE based attacks to exploit sparsity would be to identify what we call a *dense minor*.

Given m samples of sparse LWE of dimension n and sparsity k , one could try to identify a special set of L variables for $L < n$ so that there are lots of LWE samples supported only over those L variables. Indeed, then one might be able to meaningfully use an LWE solver on a smaller dimension L . For instance if $L = n/2$ and say there are more than $L \log p$ samples supported over those L variables, this means that one might be able to use an LWE solver that works over much smaller dimension $n/2$ giving rise to exponential improvement over general LWE of dimension n in the case of sparse LWE.

Nevertheless, in Sect. 7.2 we show that when $k = \log n$, if the number of samples is bounded by some polynomial, with high probability there won't exist any set of up to L variables with $\geq L + 1$ samples for some $L = \Theta(n)$.

We believe that this is a very natural way of leveraging sparsity. Typical attacks that require a large number of samples generate these samples by

taking small linear combinations of the given samples. Such combinations would increase the sparsity rapidly (as the locality pattern would form an expander) and therefore it is not clear how to leverage sparsity meaningfully using this approach. Despite our serious efforts of exploiting sparsity more cleverly, we could not come up with a method that does it better than finding a dense minor. Moreover, our condition is really mild. In our *dense-minor* model, we look for minors of size L with just $L + 1$ samples (one more than the dimension). Most attacks would need $L \log p$ samples or cL samples for some constant $c > 1$ for optimal running time. At $L + 1$ samples the secret may not even be information theoretically determined.

Concrete Parameters from Dense Minor. We suggest that our dense-minor attack gives a nice way to capture all natural attacks that exploit sparsity. Then, relying on influential work on LWE parameter estimation [7], we can suggest a nice method to estimate concrete parameters. Namely, since concrete parameters are well understood thanks to highly impactful prior work [62] for LWE, based on the dense-minor model we can use it as a benchmark for finding parameters for sparse LWE. For example, for our required number of samples m , and the required level of security (say 128 bits), one could use the Lattice Estimator [6] and learn dimension n , error and the modulus parameters.

We compute a new dimension n' by setting k appropriately (say around $\log n$) and then computing the expected size of its dense minor. We want to set n' and k so that the dense minor is at least of size n given m samples. We provide a script that computes the dimension n as a function of n' , m and k based on the dense-minor model. By simply running this script for various values of n' , we can find an appropriate dimension such that dense minors of size $< n$ are highly unlikely to exist. In Sect. 7.4, we show concrete examples of this for improving SimplePIR [43]. We take their number of required samples, error probabilities and modulus and work out the dimension n' such that sparse LWE with dimension n' and m samples provides the same level of security as their parameter settings (according to our conjecture) with a sparsity k that is much smaller. Indeed, for a reasonable $k \leq 50$, we get very comparable dimensions n' close to n for all settings of sample complexity considered in [43]. Typical dimensions for 128 bit security in SimplePIR was suggested around $n = 1024$. This suggests around a factor of $n/k \approx 20$ improvement in the efficiency in certain aspects of the scheme (query generation, setup and update time). Due to this, we believe that the sparse LWE assumption has great promise and must be analyzed further.

Asymptotically, we conjecture when $k \approx \log n$, for polynomial number of samples m , sparse LWE with dimension $n' = \alpha \cdot n$ for some constant $\alpha > 1$ is concretely similarly hard as LWE with dimension n (assuming the sample complexity, error and modulus remain the same).

Exploiting Special Structure. We note that this work does not touch upon structured assumptions for sparse LWE such as circular sparse LWE or small secret

LWE. We have observed that some of these variants behave differently for sLWE as compared to LWE. We leave investigating such assumptions to future work.

2.3 Applications

The hardness of sparse LWE implies very efficient homomorphic encryption schemes for low degree computations. In particular, we obtain the first secret key Linearly Homomorphic Encryption (LHE) schemes with constant $O(1)$ or superconstant $\omega(1)$ storage and evaluation overheads, and secret key homomorphic encryption for arbitrary constant-degree polynomials with the same level of overheads. The schemes are extremely simple. For LHE, we simply consider the vanilla secret key encryption scheme where ciphertexts have form

$$\text{ct}_i^\top = (\mathbf{a}_i^\top, b_i = \mathbf{sa} + m_i + e \cdot p \bmod q),$$

for which homomorphically evaluating a linear function specified by the coefficient vector $\mathbf{l} \in \mathbb{Z}_p^\ell$ simply amounts to computing $\text{ct} = \text{CT} \cdot \mathbf{l} \bmod q$. The efficiency stems from the fact that as sparse LWE samples, all vectors $\mathbf{a}_i$ are k sparse. Therefore, each ciphertext can be stored in a succinct form using $(k \log n + (k + 1) \log q)$ bits. If the message space $\mathbb{Z}_p$ satisfies $\log p = \Omega(\log q + \log n)$, the rate of encryption is $O(k)$, which can be set to a constant if the total number of ciphertexts is a bounded polynomial and superconstant for unbounded number of ciphertexts. To analyze the evaluation overhead, we observe that the cost of one step addition $\text{ct} + \text{ct}_i \cdot l_i$ where ct_i is a fresh $(k + 1)$ -sparse ciphertext is linear in k in the Random Access Memory model: one just need to perform $(k + 1)$ random accesses into ct , costing $(k + 1) \log n$ time (the address length is $\log n$), followed by $(k + 1)$ additions mod q . Comparing with the cost of performing one multiplication and one addition mod p in the clear, the overhead is $O(k)$ when p is sufficiently large as described above. We illustrate that our efficient LHE can lead to performance improvement in the practical Private Information Retrieval scheme SimplePIR based on LWE [43], reducing both the server database update time and the client query time by a multiplicative factor of n/k .

The same phenomenon extends to evaluating constant degree polynomials. The key to small overheads is keeping the sparsity of the ciphertexts small. We show that with some optimization over the Brakerski-Gentry-Vaikuntanathan scheme [23], for computing a degree 2^D monomial, we can bound the sparsity of the ciphertexts by $k^{O(2^D)} D^{2^D}$, which is a constant when k is a constant. Using the BGV scheme directly would lead to sparsity $k^{O(2^D)} (\log q)^{2^D}$, which is a large polylog. We avoid this by using a gadget vector $\mathbf{g}$ that contains powers of a large base, namely p , instead of 2. Another issue is that the BGV modulus reduction performs rounding on the ciphertexts and requires the secret to be short, in order to bound the error introduced by rounding. When using sparse LWE, we would like to rely on large secrets, since small secret sparse LWE is *not* equivalent to sLWE. Instead, we perform flipped modulus reduction – rounding the secret instead of the ciphertexts.

We describe our LHE and constant-degree HE schemes based on sparse LWE, in Sect. 8.1 and Sect. 8.2 respectively. We also give comparisons with other lattice-based schemes proposed in the literature. In short, to the best of our knowledge, no prior HE scheme achieves constant or super-constant overhead, even for the simplest linear evaluation. Their overheads usually grows at least linearly with λ , with the exception of [33] that achieves fully homomorphic encryption with (large) polylogarithmic overheads. They are also limited to quasipolynomially large modulus, whereas our scheme can handle any message space $\mathbb{Z}_p$ that is not too small $\log p = \Omega(\log n)$.

3 Preliminaries

Let $\mathbb{N} = \{1, 2, \dots\}$ be the natural numbers, and define $[a, b] := \{a, a+1, \dots, b\}$, $[n] := [1, n]$. Our logarithms are in base 2. For a finite set S , we write $x \leftarrow S$ to denote uniformly sampling x from S . We denote the security parameter by λ ; our parameters depend on λ , e.g. $n = n(\lambda)$, and we often drop the explicit dependence. We abbreviate PPT for probabilistic polynomial-time. Our adversaries are non-uniform PPT ensembles $\mathcal{A} = \{\mathcal{A}_\lambda\}_{\lambda \in \mathbb{N}}$. We write $\text{negl}(\lambda)$ to denote negligible functions in λ . Two ensembles of distributions $\{\mathcal{D}_\lambda\}_{\lambda \in \mathbb{N}}$ and $\{\mathcal{D}'_\lambda\}_{\lambda \in \mathbb{N}}$ are computationally indistinguishable if for any non-uniform PPT adversary $\mathcal{A}$ there exists a negligible function negl such that $\mathcal{A}$ can distinguish between the two distributions with probability at most $\text{negl}(\lambda)$. For any binary outcome adversary $\mathcal{A}$ we define its distinguishing advantage in distinguishing $\{\mathcal{D}_\lambda\}_{\lambda \in \mathbb{N}}$ and $\{\mathcal{D}'_\lambda\}_{\lambda \in \mathbb{N}}$ as:

$$\left| \Pr_{x \leftarrow \mathcal{D}_\lambda} [\mathcal{A}_\lambda(1^\lambda, x) = 1] - \Pr_{x \leftarrow \mathcal{D}'_\lambda} [\mathcal{A}_\lambda(1^\lambda, x) = 1] \right|$$

We will use matrices and vectors throughout. Both of them will be represented in boldface letters. Matrices are represented using capital bold letters (such as $\mathbf{A}, \mathbf{B}$) and vectors using lowercase bold letters (such as $\mathbf{a}, \mathbf{b}$). To refer to i^{th} index of any vector $\mathbf{v}$, we denote it by $\mathbf{v}[i]$. For matrices, we denote that by $\mathbf{A}[i, j]$.

3.1 Lattice Preliminaries

A *lattice* $\mathcal{L}$ is a discrete subgroup of $\mathbb{R}^m$, or equivalently the set

$$\mathcal{L}(\mathbf{b}_1, \dots, \mathbf{b}_n) = \left\{ \sum_{i=1}^n x_i \mathbf{b}_i : x_i \in \mathbb{Z} \right\}$$

of all integer combinations of n linearly independent vectors $\mathbf{b}_1, \dots, \mathbf{b}_n \in \mathbb{R}^m$. Such $\mathbf{b}_i$'s form a *basis* of $\mathcal{L}$.

The lattice $\mathcal{L}$ is said to be *full-rank* if $n = m$. We denote by $\lambda_1(\mathcal{L})$ the first minimum of $\mathcal{L}$, defined as the length of a shortest non-zero vector of $\mathcal{L}$.

Definition 3.1 (Successive Minima). *Given a lattice $\mathcal{L}$ and an integer $1 \leq k \leq n$, define the k -th successive minimum $\lambda_k(\mathcal{L})$ as the smallest possible length of a set of k linearly independent vectors in $\mathcal{L}$.*

Equivalently, $\lambda_k(\mathcal{L})$ is the smallest value r such that a ball of radius r (centered on the origin) contains k linearly independent lattice points.

For a rank n lattice $\mathcal{L} \subset \mathbb{R}^n$, the *dual lattice*, denoted $\mathcal{L}^*$, is defined as the set of all points in $\text{span}(\mathcal{L})$ that have integer inner products with all lattice points,

$$\mathcal{L}^* = \{\mathbf{w} \in \text{span}(\mathcal{L}) : \forall \mathbf{y} \in \mathcal{L}, \langle \mathbf{w}, \mathbf{y} \rangle \in \mathbb{Z}\}.$$

Similarly, for a lattice basis $\mathbf{B} = (\mathbf{b}_1, \dots, \mathbf{b}_n)$, we define the dual basis $\mathbf{B}^* = (\mathbf{b}_1^*, \dots, \mathbf{b}_n^*)$ to be the unique set of vectors in $\text{span}(\mathcal{L})$ satisfying $\langle \mathbf{b}_i^*, \mathbf{b}_j \rangle = 1$ if $i = j$, and 0, otherwise. It is easy to show that $\mathcal{L}^*$ is itself a rank n lattice and $\mathbf{B}^*$ is a basis of $\mathcal{L}^*$. Given a lattice $\mathbf{B} = (\mathbf{b}_1, \dots, \mathbf{b}_n)$, we denote $\|\mathbf{B}\|_2 = \max_i \|\mathbf{b}_i\|$.

An important property of the dual lattice is the following transference theorem, shown in Theorem 2.1 of [16].

Lemma 3.1 (Transference Theorem). *For any n -rank lattice $\mathcal{L}$, $1 \leq \lambda_1(\mathcal{L}) \cdot \lambda_n(\mathcal{L}^*) \leq n$.*

Discrete Gaussian and Related Distributions. For any $s > 0$, define

$$\rho_s(\mathbf{x}) = \exp(-\pi \|\mathbf{x}\|^2 / s^2)$$

for all $\mathbf{x} \in \mathbb{R}^n$. We write ρ for ρ_1 . For a discrete set S , we extend ρ to sets by $\rho_s(S) = \sum_{\mathbf{x} \in S} \rho_s(\mathbf{x})$. Given a lattice $\mathcal{L}$, the *discrete Gaussian* $\mathcal{D}_{\mathcal{L},s}$ is the distribution over $\mathcal{L}$ such that the probability of a vector $\mathbf{y} \in \mathcal{L}$ is proportional to $\rho_s(\mathbf{y})$:

$$\Pr_{\mathbf{x} \leftarrow \mathcal{D}_{\mathcal{L},s}}[X = \mathbf{y}] = \frac{\rho_s(\mathbf{y})}{\rho_s(\mathcal{L})}.$$

Using standard sub-Gaussian tail-bounds, one can show we can show the following claim.

Claim. Let $m \in \mathbb{N}$, $\sigma > 0$, then it holds that:

$$\Pr_{\mathbf{e} \leftarrow \mathcal{D}_{\mathbb{Z}^m, \sigma}}[\|\mathbf{e}\| > m\sigma] < \exp(-\tilde{\Omega}(m)).$$

We define the truncated discrete Gaussian as a distribution statistically close to discrete Gaussian where one samples the vectors from discrete Gaussian conditioned on their norm being upper bounded by $m\sigma$.

Trapdoor Sampling. We will need the following definition of a lattice trapdoor [35, 58, 69]. For $\mathbf{A} \in \mathbb{Z}_p^{n \times m}$, we define the rank m lattice

$$\mathcal{L}^\perp(\mathbf{A}) = \{\mathbf{z} \in \mathbb{Z}^m : \mathbf{A}\mathbf{z} = \mathbf{0} \pmod{q}\}.$$

A lattice trapdoor for $\mathbf{A}$ is a set of short linearly independent vectors in $\mathcal{L}^\perp(\mathbf{A})$.

Definition 3.2. *A matrix $\mathbf{T} \in \mathbb{Z}^{m \times m}$ is a β -good lattice trapdoor for a matrix $\mathbf{A} \in \mathbb{Z}_p^{n \times m}$ if*

1. $\mathbf{AT} = \mathbf{0} \pmod{p}$.
2. For each column vector $\mathbf{t}_i$ of $\mathbf{T}$, $\|\mathbf{t}_i\|_\infty \leq \beta$.
3. $\mathbf{T}$ has rank m over $\mathbb{R}$.

Theorem 3.1. [35, 58, 69] *There is an efficient algorithm that, on input $1^n, p, m \geq 10n \log p$, outputs a matrix $\mathbf{A}$ distributed statistically close to uniformly on $\mathbb{Z}_p^{n \times m}$, and a $O(m)$ -good lattice trapdoor $\mathbf{T}$ for $\mathbf{A}$.*

We also observe that due to the transference theorems for any matrix $\mathbf{A} \in \mathbb{Z}_p^{n \times m}$ a lower bound on $\lambda_1(\mathcal{L})$ implies an upper bound on $\lambda_m(\mathcal{L}^\perp)$.

Corollary 3.1. (of Transference Theorem 3.1) *Given any $\mathbf{A} \in \mathbb{Z}_p^{n \times m}$, if $\lambda_1(\mathcal{L}) \geq \beta$, then $\lambda_m(\mathcal{L}^\perp) \leq m \cdot \frac{p}{\beta}$. Here $\mathcal{L} = \mathcal{L}(\mathbf{A})$ and $\mathcal{L}^\perp = \mathcal{L}^\perp(\mathbf{A})$.*

Proof. Recall that the transference theorem states:

$$1 \leq \lambda_1(\mathcal{L}) \cdot \lambda_m(\mathcal{L}^*) \leq m,$$

where $\mathcal{L}^*$ is the dual of $\mathcal{L}$. Observe that $\mathcal{L}^\perp = p\mathcal{L}^*$ and so $\lambda_m(\mathcal{L}^*) = \frac{\lambda_m(\mathcal{L}^\perp)}{p}$. The claim follows from using the right inequality using the lower bound on $\lambda_1(\mathcal{L}) \geq \beta$.

4 Sparse Learning with Error

We now formally define our sparse LWE assumption. Our assumption is exactly identical to LWE with a crucial difference that each coefficient vector is exactly k -sparse for some k that for the most part should be thought of as a small function of the dimension n (ideally a constant or polylogarithmic in n).

Definition 4.1 (Coefficient Distribution $\mathcal{D}_{\text{coeff},n,k,p}$). *For a parameter $n \in \mathbb{N}$, and a $k = k(n)$, a modulus p , we define the distribution $\mathcal{D}_{\text{coeff},n,k,p}$ that first samples a uniformly random set or a “locality pattern” $S \subset [n]$ so that S has exactly size k (we abuse notation to denote this by saying $S \leftarrow \binom{[n]}{k}$). Then, the distribution samples a random vector as follows. Sample $\mathbf{a} \in \mathbb{Z}_p^n$ so that each coordinate $\mathbf{a}[i]$ is set to 0 if $i \notin S$. Otherwise $\mathbf{a}[i]$ is chosen at random from $\mathbb{Z}_p$.*

Definition 4.2 (Locality Set). *For a modulus p , and any vector $\mathbf{v} \in \mathbb{Z}_p^n$ we denote by the locality set of $\mathbf{v}$ (denoted as $\text{LocSet}(\mathbf{v})$) as a set of indices $S \subset [n]$ so that $\mathbf{v}[i]$ is non-zero for every $i \in S$.*

Definition 4.3 (Sparse LWE Distinguishing Problem). *For an integer dimension n , sparsity parameter k , a sample complexity m , modulus p , and noise parameter $\sigma = \sigma(n)$ we define the $\text{sLWE}_{n,k,m,p,\sigma}$ problem as a distinguishing problem. In the first distribution, we generate m sparse LWE samples as follows:*

- Sample m coefficient vectors $\mathbf{a}_1, \dots, \mathbf{a}_m$ where each $\mathbf{a}_i \leftarrow \mathcal{D}_{\text{coeff},n,k,p}$.
- Sample a secret vector $\mathbf{s} \leftarrow \mathbb{Z}_p^n$.
- Sample m noise values from $\{e_i \leftarrow \mathcal{D}_{\mathbb{Z},\sigma}\}_{i \in [m]}$.
- The first distribution consists of $\{(\mathbf{a}_i, b_i = \langle \mathbf{s}, \mathbf{a}_i \rangle + e_i \pmod{p})\}$

The second distribution is exactly the same except $\{b_i\}_{i \in [m]}$ are randomly sampled from $\mathbb{Z}_p$.

We also define the LWE distinguishing problem to be identical to a sparse LWE distinguishing problem as above except that the vectors $\mathbf{a}_i$ are chosen randomly from $\mathbb{Z}_p^n$ as opposed to being sampled as sparse vectors.

Definition 4.4 (Sparse LWE Assumption). Let $k = k(n) \in [0, n]$, $p = p(n) : \mathbb{N} \rightarrow \mathbb{N}$, $\sigma = \sigma(n) : \mathbb{N} \rightarrow \mathbb{N}$, $m = m(n) : \mathbb{N} \rightarrow \mathbb{N}$ be efficiently computable functions (computable in time polynomial in n).

We say that the sparse LWE assumption with respect to sparsity $k(n)$, modulus $p(n)$, noise parameter $\sigma = \sigma(n)$ and sample complexity $m(n)$ holds, denoted as the $\text{sLWE}_{n,k,m,p,\sigma}$ assumption, if for every non-uniform p.p.t adversary $\mathcal{A}$, the distinguishing advantage in winning the $\text{sLWE}_{n,k,m,p,\sigma}$ distinguishing problem (Definition 4.3) is bounded by some negligible function $\text{negl}(n)$.

We say that the sparse LWE assumption with respect to sparsity $k(n)$, modulus $p(n)$, noise parameter $\sigma = \sigma(n)$ and unbounded polynomial number of samples holds, denoted as the $\text{sLWE}_{n,k,p,\sigma}$ assumption, if the above condition holds for every non-uniform p.p.t adversary $\mathcal{A}$, and every polynomial $m(n)$.

We say that the subexponential $\text{sLWE}_{n,k,m,p,\sigma}$ (respectively $\text{sLWE}_{n,k,p,\sigma}$) assumption holds if there exists a constant $\epsilon \in (0, 1)$, such that, every non-uniform probabilistic $\text{poly}(2^{n^\epsilon})$ -time adversary $\mathcal{A}$ (and respectively every polynomial $m(n)$), the distinguishing advantage in winning the $\text{sLWE}_{n,k,m,p,\sigma}$ distinguishing problem is bounded by $1/\text{poly}(2^{n^\epsilon})$.

Remark 4.1. For notational brevity, we will also denote sparse LWE samples as $(\mathbf{A}, \mathbf{s}\mathbf{A} + \mathbf{e} \bmod p)$ where the coefficient vectors $\{\mathbf{a}_i\}$ form the columns of the matrix $\mathbf{A} \in \mathbb{Z}_p^{n \times m}$. $\mathbf{s}$ is interpreted to live in $\mathbb{Z}_p^{1 \times m}$ and the error vector $\mathbf{e}$ is in $\mathbb{Z}^{1 \times m}$.

4.1 Minimum Distance of the Sparse LWE Lattice

We now analyze bounds on $\lambda_1(\mathcal{L})$ for the lattice:

$$\mathcal{L}(\mathbf{A}) = \{\mathbf{x}\mathbf{A} + p\mathbb{Z}^m : \mathbf{x} \in \mathbb{Z}_p^{1 \times n}\}.$$

We compute these bounds as a function of dimension n , sparsity parameter $k = k(n)$, sample complexity $m = m(n)$ and the modulus $p = p(n)$. These bounds will be useful for identifying when short lattice trapdoors for $\mathbf{A}$ might exist. We will also use this calculation for understanding the running time of various lattice-reduction based attacks on sparse LWE.

We start by recalling the minimum distance in the case when matrix $\mathbf{A}$ is randomly chosen from $\mathbb{Z}_p^n$ (in other words, it corresponds to plain LWE). In this case, it can be shown that as n grows, the minimum distance $\lambda_1(\mathcal{L})$ is roughly around $p^{1-n/m}$ for $m \gg n$ (up to polynomial factors, which we ignore). When $m = \Omega(n \log p)$, this is roughly p .

For sparse LWE, since the matrix has a lot of zeros, one would expect the minimum distance to be smaller. However, we show that for interesting parameter regimes (i.e. with $k > \log n$, and p being polynomial in the size n , and for polynomial m) the minimum distance with high probability for the case of sparse LWE is also around p . We give a lower-bound on the minimum distance.

Our main result is the following:

Theorem 4.1. *Let $\mathbf{A}$ be a $n \times m$ sparse matrix each of whose columns are sampled i.i.d. from $\mathcal{D}_{\text{coeff},n,k,p}$. With probability $1 - o(1)$, the length of the shortest nonzero vector in $\mathcal{L}(\mathbf{A})$ is $\Theta(p)$ assuming $m \geq 8n \log n$ and $p = \text{poly}(n)$.*

Please refer to the full version of this paper for a proof.

5 Reduction from LWE to Sparse LWE

We now show that given an oracle that breaks sparse LWE of dimension n and sparsity k with m samples, we can break LWE of dimension k with m samples as long as the modulus p is a prime.

We get as input $\mathbf{A}, \mathbf{b} = \mathbf{sA} + \mathbf{e}$, and we will effectively convert it to a tuple of the form $\mathbf{A}', \mathbf{b}' = \mathbf{s'A}' + \mathbf{e}$ with the same error (if our input $\mathbf{b}$ is sampled randomly instead of from the LWE distribution, our reduction will keep it randomly distributed). The distribution for $(\mathbf{A}', \mathbf{b}')$ will be identical to that of average-case sparse LWE. The nice thing about our reduction is that it does not manipulate the error $\mathbf{e}$ and therefore also applies to the setting of learning parity with noise.

We now describe the conversion procedure. Our rough idea is that given a $k \times m$ dense matrix $\mathbf{A}$ we transform each column into a sample in a related matrix (that we shortly show how to build) $\mathbf{A}'$ where $\mathbf{A}'$ is a $n \times m$ sparse matrix with k nonzero entries per column. So far we keep the $\mathbf{b}$ vector unchanged; it can be written as $\mathbf{b} = \mathbf{zA}' + \mathbf{e}$. At this point, the secret $\mathbf{z}$ is not uniformly random. But we can randomize it by adding $\mathbf{vA}'$ to $\mathbf{b}$ for a randomly chosen $\mathbf{v}$. The new secret $\mathbf{v} + \mathbf{z}$ is random and independent of $\mathbf{A}'$.

We now describe the transformation from $\mathbf{A}$ to $\mathbf{A}'$. We want to ensure that the equation $\mathbf{CA}' = \mathbf{A}$ holds for some $k \times n$ matrix $\mathbf{C}$. This would allow us to write $\mathbf{b} = \mathbf{sA} + \mathbf{e} = (\mathbf{sC})\mathbf{A}' + \mathbf{e}$. Now if we set $\mathbf{A}'$ randomly and try to find $\mathbf{C}$, this might be hard. Therefore, we do the reverse. We first fix $\mathbf{C}$ and then sample only the locality pattern $S_1, \dots, S_m$ at random from $\binom{n}{k}$. Then, we choose the non-zero coordinates given by S_i so that $\mathbf{CA}'[i]$ is equal to $\mathbf{A}[i]$. This can be done since we can view the above matrix equation as a system of linear equations over $\mathbb{Z}_p$ with k variables corresponding to the nonzero entries of $\mathbf{A}'[i]$ and k equations corresponding to the k rows of $\mathbf{A}[i]$. We will choose $\mathbf{C}$ to be a $k \times n$ Vandermonde matrix which has the property that any k columns are linearly independent. This allows us to guarantee the system of equations described above correspond to an invertible matrix, and is hence solvable.

We have to show that the above process yields a random $\mathbf{A}'$. Observe that each column $\mathbf{A}'[i]$ can be described by its locality pattern, and the sequence of

nonzero entries in the indices corresponding to the locality pattern. $\mathbf{A}_i$ is random means the nonzero entries of $\mathbf{A}'_i$ must be random. We chose the locality pattern independently at random to begin with. Therefore, the whole column $\mathbf{A}'[i]$ is distributed uniformly in $\mathcal{D}_{\text{coeff},n,k,p}$ and is independent of the other columns.

Formally, we show the following theorem:

Theorem 5.1. *Assume there exists a PPT algorithm which solves average-case $\text{sLWE}_{n,k,m,p,\sigma}$ instances in time T with success probability γ , where the modulus p is a prime and the number of samples $m < p$. Then we can solve average-case $\text{LWE}_{k,m,p,\sigma}$ instances in time $T + O(n^3 m \log p)$ with success probability γ .*

Please refer to the full version of this paper for a proof.

Remark 5.1. Our reduction does not preserve dimension. If we choose k to be small (constant or logarithmic in n), we end up with an algorithm for solving very low dimensional LWE which is anyway easy. On the other hand, if k is very large (say $\Theta(n)$), we conclude that sLWE is very hard if the columns are not too sparse (only a constant fraction of each column vector is required to be zero). The intermediate range, where k is polylogarithmic or polynomial in n , yields more interesting conclusions.

6 Trapdoor Sampling for Sparse LWE

We are going to present an algorithm for sampling pairs of matrices $\mathbf{A}, \mathbf{T}$ where $\mathbf{A}$ is a $n \times m$ sparse matrix and $\mathbf{T}$ is a trapdoor for $\mathbf{A}$. The distribution of each column of $\mathbf{A}$ should be statistically close to independent samples from $\mathcal{D}_{\text{coeff},n,k,p}$.

Our approach stands on the shoulders of amazing prior work of [57]. Recall that they showed how to sample a matrix $\mathbf{V} \in \mathbb{Z}_p^{n \times m}$ for $m = O(n \log p)$ that is statistically close to random along with a full-rank square matrix $\mathbf{T}_{\mathbf{V}}$ of norm $O(m)$ in the nullspace of $\mathbf{V}$ ($\mathbf{T}_{\mathbf{V}}$ can be converted to a basis generically as shown by lemma 7.1 of [56] to a basis for $\mathcal{L}^\perp(\mathbf{V})$). The idea is that first we sample a matrix $\mathbf{A} \leftarrow \mathbb{Z}_p^{n \times m'}$ for $m' > 2n \log p$ at random. Then we sample a random binary matrix $\mathbf{R} \leftarrow \{0, 1\}^{m' \times n \lceil \log p \rceil}$, and consider $\mathbf{V} = [\mathbf{A} | \mathbf{A}\mathbf{R} + \mathbf{G}]$ where $\mathbf{G}$ is the gadget matrix defined as $\mathbf{I}_n \otimes \mathbf{g}$ where $\mathbf{g} = [1 \ 2 \ 4 \ \dots \ 2^{\lceil \log p \rceil - 1}]$.

Notice that there is a small-norm square matrix $\mathbf{D}$ which satisfies $\mathbf{V}\mathbf{D} = [\mathbf{X} | \mathbf{G}]$. In particular, we can pick

$$\mathbf{D} := \left[\begin{array}{c|c} \mathbf{I} & -\mathbf{R} \\ \hline \mathbf{0} & \mathbf{I} \end{array} \right]$$

This is enough to obtain a short trapdoor; consider the matrix

$$\mathbf{E} := \left[\begin{array}{c|c} \mathbf{0} & \mathbf{I} \\ \hline \mathbf{T}_{\mathbf{G}} & -\mathbf{G}^{-1}(\mathbf{X}) \end{array} \right] \text{ where } \mathbf{T}_{\mathbf{G}} \text{ is a trapdoor for } \mathbf{G} \quad (1)$$

It is easy to check that the product $\mathbf{D}\mathbf{E}$ is short (since both $\mathbf{D}$ and $\mathbf{E}$ have low norm), has full rank (since both $\mathbf{D}$ and $\mathbf{E}$ are nonsingular), and lies in the nullspace of $\mathbf{V}$.

Note that this approach fails completely for sparse matrices. Even if we sample $\mathbf{A}$ to be sparse, $\mathbf{A}\mathbf{R}$ will not be sparse anymore. Our idea is to start exactly as in the step above to construct $\mathbf{V}$. However, this matrix is not sparse at all. We come up with a procedure to “sparsify” $\mathbf{V}$ to obtain a wider k -sparse matrix $\mathbf{C}$. We take each column of $\mathbf{V}$ and replace it by a $n \times n \log n$ sparse matrix whose columns add up to the column we are replacing.

We construct a matrix $\mathbf{C}$ of the form $[\mathbf{C}_1 | \dots | \mathbf{C}_L]$ as follows (L is the width of $\mathbf{V}$).

- Each column of $\mathbf{C}_i$ has a random locality pattern chosen from $\binom{n}{k}$.
- The columns of $\mathbf{C}_i$ sum up to the i^{th} column $\mathbf{v}_i$ of $\mathbf{V}$.
- Subject to the first two constraints $\mathbf{C}_i$ is random.

Such a $\mathbf{C}$ can be easily found provided the width of each $\mathbf{C}_i$ is large enough ($n \log n$ suffices). At this point, $\mathbf{C}$ is from the right distribution. Its locality pattern is randomly chosen, and because the distribution of $\mathbf{V}$ is statistically close to random, so is the distribution of $\mathbf{C}$.

This $\mathbf{C}$ is indeed our matrix. We will now show how to construct the trapdoor. We try to find a small-norm full rank square matrix $\mathbf{D}$ such that $\mathbf{C}\mathbf{D}$ has the form $[\mathbf{X} | \mathbf{G}]$. Notice that if we can find this $\mathbf{D}$, we are already done (similar to [57]). We again use the same $\mathbf{E}$ from Eq. (1). The trapdoor we are looking for is the product $\mathbf{D}\mathbf{E}$, since it satisfies the following 3 properties:

- $\mathbf{C}\mathbf{D}\mathbf{E} = \mathbf{0}$
- $\mathbf{D}\mathbf{E}$ has small norm.
- The columns of $\mathbf{D}\mathbf{E}$ are linearly independent.

Let us now describe how we can find $\mathbf{D}$. Recall that we started with a matrix $\mathbf{V}$ such that

$$\mathbf{V} \cdot \begin{bmatrix} -\mathbf{R} \\ \mathbf{I} \end{bmatrix} = \mathbf{G}$$

By construction, we also have $\mathbf{C}_i [1 \ 1 \ \dots \ 1]^T = \mathbf{v}_i$ for each column $\mathbf{v}_i$ of $\mathbf{V}$. Therefore, for any matrix $\mathbf{U}$, we can write $\mathbf{V}\mathbf{U} = \mathbf{C}(\mathbf{U} \otimes [1 \ 1 \ \dots \ 1]^T)$. So now our last equation gives us

$$\mathbf{C} \cdot \left(\begin{bmatrix} -\mathbf{R} \\ \mathbf{I} \end{bmatrix} \otimes \begin{bmatrix} 1 \\ 1 \\ \vdots \\ 1 \end{bmatrix} \right) = \mathbf{G}$$

So now we have found a matrix $\mathbf{F}$ which satisfies $\mathbf{C}\mathbf{F} = \mathbf{G}$, which was the difficult part. To extend it so that we can obtain $\mathbf{C}\mathbf{D} = [\mathbf{X} | \mathbf{G}]$, we just choose an arbitrary full rank matrix not in the column span of $\mathbf{F}$ to concatenate to the left of $\mathbf{F}$ until the resulting matrix becomes square. We can do this by picking columns from the identity matrix until there is a linear dependence.

Formally, we show the following analog of theorem 3.1 for sparse LWE:

Theorem 6.1. *There is a PPT algorithm that on input $1^n, p, m \geq 3n^2 \log p \log n$, $k > 1$ outputs a matrix $\mathbf{A}$ distributed statistically close to $\mathcal{D}_{\text{coeff}, n, k, p}^m$, and a $O(m)$ -good lattice trapdoor $\mathbf{T}$ for $\mathbf{A}$ with success probability $1 - o(1)$.*

Please refer to the full version of this paper for a proof.

Remark 6.1. Our algorithm requires a much higher number of samples as compared to the regular LWE trapdoor sampling process, $O(n^2 \log n)$ instead of $O(n)$.

7 Cryptanalysis

7.1 Sparse Vectors in the Kernel

One potential approach to attack sparse LWE could be (in the same vein as sparse Learning Parity with Noise) finding sparse vectors $\mathbf{x} \in \mathbb{Z}_p^{m \times 1}$ (not necessarily small) so that $\mathbf{A}\mathbf{x} = \mathbf{0}$. If this is true, given a t sparse $\mathbf{x}$ and a sLWE tuple $(\mathbf{A}, \mathbf{b} = \mathbf{sA} + \mathbf{e} \bmod p)$, one can observe that:

$$\mathbf{b} \cdot \mathbf{x} = \mathbf{e} \cdot \mathbf{x}.$$

If the error coordinates for $\mathbf{e}$ lie in $[-B, B]$ for some B and t is sufficiently small so that $(2B+1)^t \ll p$, then this will give an attack that runs in time roughly $B^{O(t)}$. The idea is that since $\mathbf{x}$ is t sparse, the product $\mathbf{e} \cdot \mathbf{x}$ touches upon only t coordinates of $\mathbf{e}$. There are $2B + 1$ possible choices for each coordinate. One could brute-force and check if there is a setting that yields the inner product $\mathbf{b} \cdot \mathbf{x}$. One can show that for a typical $\mathbf{x}$ only one choice of $\mathbf{e}$ will agree with p when $(2B + 1)^t$ is sufficiently smaller than p . To rule out such attacks we ask for what sparsity t does there exist solutions to $\mathbf{A}\mathbf{x} = \mathbf{0}$. One can show a naive bound by lifting off analysis for the sparse LPN assumption. Such analysis has been done for sparse LPN over $\mathbb{Z}_2$ that can be helpful here.

The above attack can be significant if there exist very sparse vectors in the kernel. For instance, if t is a constant and B is polynomial, this attack could be implemented in polynomial time. Notice that when the number of samples $m \geq k \binom{n}{k} + 1$, by pigeonhole principle there must exist $k + 1$ coefficient vectors supported on the same set of k input variables. Since there are at least $k + 1$ vectors supported on the same set of k variables they must be linearly dependent. This allows us to find a t -sparse vector in the nullspace of $\mathbf{A}$ for $t \leq k + 1$.

Inspired by the literature on random k -XOR in the field of average case complexity, we analyze a certain graph expansion condition corresponding to the locality set of the column vectors of $\mathbf{A}$. If this condition is met, we would be guaranteed with high probability that t is not too small. As we will describe later, this condition turns out to be conservative.

Locality Graph Expansion. Our main theorem is the following:

Theorem 7.1. *Let $\mathbf{A}$ be a $\text{sLWE}_{n,k,m,p,\sigma}$ matrix. With probability $1 - o(1)$, there does not exist any vector $\mathbf{x}$ in the kernel of $\mathbf{A}$ which has L or fewer nonzero entries in either of the following cases:*

- $k \geq 3$ is a constant, $m = n^{\delta \frac{k}{2}}$ for some $\delta < 1$, and $L = \Theta(n^{1-\delta})$.
- $k = \text{poly}(\log n)$, $m = \text{poly}(n)$, and $L = \Theta(n)$

We observe that if $\mathbf{A}\mathbf{x} = \mathbf{0}$ for a t -sparse $\mathbf{x}$, the following holds. Let us denote by $S_1, \dots, S_m \subset [n]$ the locality sets (see Definition 4.2) corresponding to the column vectors of $\mathbf{A} \in \mathbb{Z}_p^{n \times m}$.

Observe that if there exist a t sparse $\mathbf{x}$ so that $\mathbf{A}\mathbf{x} = \mathbf{0}$ this necessarily means that the following expansion property is satisfied:

$$\left| \bigcup_{i: \mathbf{x}[i] \neq 0} S_i \right| < \frac{k \cdot t}{2}. \quad (2)$$

This is because each column of $\mathbf{A}$ is exactly k sparse. A t sparse combination that yields the all 0 must cancel any variable that appears in the combination. Thus, each variable must appear at least twice. We also state that this condition is a bit conservative. It is plausible that each variables appear multiple times yet the column vectors are linearly independent. This might be more prominent for large p 's. We discuss this issue shortly.

We now compute an upper-bound on the probability of the above condition (Eq. 2 being violated) as a function of n, k, m (note that the condition is independent of p). From this we can get the desired tradeoffs for n, k, m . This is a routine calculation.

Lemma 7.1. *For n, m, k and $t > k$ if the sets $S_1, \dots, S_m$ are chosen at random from $\binom{[n]}{k}$ then the probability that the condition in Eq. 2 is violated is bounded by $f_{\text{SPARSE}}(k, n, m)^t$ where $f_{\text{SPARSE}}(k, n, m) = e^{1+3k/2} \cdot (k/2)^{k/2} \cdot \frac{m \cdot t^{k/2-1}}{n^{k/2}}$.*

Please refer to the full version of this paper for a proof.

Our intention should be to set k and m as a function of n so that f_{SPARSE} is small for a large t .

One can examine that for a constant k , t is some polynomial provided $m = O(n^{k/2(1-\epsilon)})$ for some constant $\epsilon > 0$. When $m = \Omega(n^{k/2})$, one can find a constant number of subsets that violate Eq. 2.

We now do some (asymptotic) case analysis for a constant k and setting $m = n^{k/2(1-\epsilon)}$. To set f_{SPARSE} to be $o(1)$, we consider another function g_{SPARSE} that is only bigger than f_{SPARSE} but more amenable to analysis.

$$g_{\text{SPARSE}}(k, n, m) = \frac{m \cdot (e^3 \cdot k \cdot t)^{k/2}}{n^{k/2}}. \quad (3)$$

To ensure g_{SPARSE} is small, we set:

$$m^{2/k} e^3 \cdot k \cdot t \ll n$$

Thus, this holds for:

$$t < \frac{n}{m^{2/k} \cdot e^3 \cdot k} \quad (4)$$

Proof (Proof of theorem 7.1). The theorem statement follows directly from Eq. (4) by plugging in the appropriate values of n, m, k and setting L to the highest value of t that satisfies the inequality.

For concrete parameter estimation one can compute exact values of Eq. (4) and find out how this attack performs in practice. We anticipate if $B^t \gg p$, then, the attack does not apply.

Large p Effect It is clear that for the analysis for this attack we were overly conservative. Not every collection of t sets with less than $\frac{kt}{2}$ neighbours is likely to give a sparse vector in the kernel. Moreover, if p increases the chances of this happening should be really small. Nevertheless as described before while the calculation predicts security against polynomial time adversaries with $m \ll n^{\frac{k}{2}}$ when k is a constant, system over any p can be broken when $m > (k+1) \binom{n}{k}$. Thus, the prediction in terms of sample complexity is potentially only off by a square-root factor. We leave the analysis that utilizes the largeness of p for future work as a great open question. We show how to loosely incorporate this above so that it is better for concrete parameter estimation. Finding precise bounds is a great open question.

Claim. If $m \leq n^{k-3}$, the following inequality holds with probability $1 - \text{negl}(n)$ for all $t \leq L$ where $L = \text{poly}(n)$ and k is a constant.

$$\left| \bigcup_{i: \mathbf{x}[i] \neq 0} S_i \right| > 2 \cdot t \quad \text{where } \mathbf{x} \text{ is } t\text{-sparse} \quad (5)$$

Please refer to the full version of this paper for a proof.

If the locality pattern of $\mathbf{A}$ satisfies the above combinatorial property, we can show that there does not exist a L -sparse $\mathbf{x}$ in the kernel of $\mathbf{A}$ with high probability. Since L above is polynomial in n , this rules out the sparse vector in kernel attack. Recall that sparse LPN over $\mathbb{F}_2$ is broken for $m = n^{k/2}$ samples due to this attack (see [9, 46] etc.). However, our result holds as long as the sample complexity m is less than n^{k-3} , which is a substantial improvement over theorem 7.1. This is one of the ways increasing the modulus from 2 to some $p = \omega(m)$ increases security.

We can now prove the following result:

Theorem 7.2. *Let $\mathbf{A}$ be a sLWE matrix where the sample complexity m , dimension n and sparsity parameter $k = O(1)$ satisfy the relation $m < n^{k-3}$. If the modulus p is $\omega(m)$, then there does not exist any vector $\mathbf{x}$ in the kernel of $\mathbf{A}$ which has L or fewer nonzero entries for some $L = \text{poly}(n)$ with probability $1 - o(1)$.*

Please refer to the full version of this paper for a proof.

7.2 Dense Minor Lower-Bound Model

We now introduce a lower-bound model to provide an estimate of lattice attacks that exploit the sparsity structure of the equations. We call our attack the Dense-Minor attack. This captures a variety of different attacks as we will describe later and enables us to reason about security.

The rough idea is that all attacks that break LWE over dimension n requires at least $n+1$ samples or else secrets are not information theoretically defined. The problem, with the low-locality structure of our assumption is that each sample is constant sparse (say k). Thus, it is plausible that exploiting the low locality, one can find $L+1$ samples supported over L coordinates for some small L , and then use an algorithm that breaks LWE with dimension L and $L+1$ samples. For example, if the number of samples is at least $k\binom{n}{k} + 1$, the pigeonhole principle guarantees the existence of some $k+1$ columns which are supported on the same set of k variables. If k is a constant, we can solve the resulting k -dimensional LWE problem in constant time; this points to an attack against sparse LWE with constant sparsity and polynomially many samples.

We show how to set parameters to avoid this. In particular, setting the number of samples a reasonable polynomial such as $m = n^2$ or n^4 for target applications, we will aim to find k such that L turns out to be at least linear in n e.g. $0.9n$ for the above attack strategy. This roughly states that any meaningful attack exploiting the sparsity must involve $\Omega(n)$ variables. Thus, one can appeal to concrete cryptanalysis to make a reasonable conjecture on the security of sparse LWE.

Conjecture 7.1. $\text{sLWE}_{n,k,m,p,\sigma}$ is at least as hard as $\text{LWE}_{L,f(L),p,\sigma}$ for $L = \Theta(n)$, $k = \Omega(\log n)$, and f is some small polynomial.

We will use the above heuristic for parameter estimation purposes. We show some examples in Sect. 7.4.

Definition 7.1 (Dense-Minor of order r). *For any matrix $\mathbf{A} \in \mathbb{Z}_p^{n \times m}$, we say that it has a dense minor of size $L \in [0, n]$, if there exists a subset $T \subseteq [n]$ of size L such that there are $r \cdot L$ column vectors supported entirely over indices given by T . Moreover, L is the minimum such number.*

We will typically be interested in r as $1 + 1/L$ or some constant (like 2), or even $r = \log p$.

Problem with a Small Dense Minor. We observe that if for $\mathbf{A} \in \mathbb{Z}_p^{n \times m}$ there is a dense minor of size L and order r , then, this means that we have $r \cdot L$ columns all supported at some L indices. This means that if one is given a sparse LWE sample $\mathbf{A}_i \mathbf{s} + \mathbf{e}_i \pmod p$ one can pick out these $r \cdot L$ equations of the form $\{\mathbf{a}_i, \langle \mathbf{a}_i, \mathbf{s} \rangle + e_i \pmod p\}_{i \in S}$ for some set $S \subseteq [m]$ of size at least $r \cdot L$. Here we assume $\mathbf{a}_i$'s are column vectors of $\mathbf{A}$. This new system is an LWE type system with $r \cdot L$ samples of dimension L . This can potentially be broken with access to an LWE oracle that solves LWE on dimension L , sample complexity $r \cdot L$ and the same modulus p .

Theorem 7.3. *Let n be the dimension, $k = \log n$ be the sparsity parameter and $m = n^\beta$ be the sample complexity for some constant $\beta > 1$. Let p be the modulus. We have that, with probability $1 - o(1/n)$ over choice of $\mathbf{A} \leftarrow \mathcal{D}_{\text{coeff}, n, k, p}^m$ (sampling each column independently from the distribution $\mathcal{D}_{\text{coeff}, n, k, p}$ specified in definition 4.1), $\mathbf{A}$ the minimum L for which $\mathbf{A}$ has a dense minor of order $1 + 1/L$ is $L = \Omega(n)$.*

Remark 7.1. We remark that most attacks care about $L \log p$ samples or cL samples for some $c > 1$. In the above theorem, we rule out minors with just $L + 1$ samples. This only makes our model stronger.

Please refer to the full version of this paper for a proof.

The expression above is quite loose as we make several approximations regarding Binomial coefficients. If one computes the expression

$$f_{\text{DENSEMINOR}} = \sum_{t \in [k, L]} \left(\frac{\binom{t}{k}}{\binom{n}{k}} \right)^{t+1} \cdot \binom{n}{t} \cdot \binom{m}{t+1}$$

it is possible to find tradeoffs exactly using a computer program. We use it to estimate the exact size of the dense minor. For more discussion on this, see Sect. 7.4.

Now we describe how our dense-minor model captures attacks such as Arora-Ge, BKW and Lattice Reductions.

Capturing Arora-Ge Style Attacks. Recall that for the Arora-Ge attack [15] on LWE, we assume that the infinity norm of the error vector is at most B for some $B < p/2$. Each LWE sample $(\mathbf{a}, b)$ gives rise to the following equation on the entries of the secret vector:

$$\prod_{j=-B}^B \left(\sum_{i=1}^n a_i s_i - b - j \right) = 0$$

The degree of each equation is $2B + 1$. Since there are n variables, the total number of terms is $\binom{n+2B+1}{n}$. We treat each unique term as a separate variable, and we can solve the resulting linear system directly as long as the number of equations, m , exceeds the total number of terms. If we have fewer equations, we

first multiply each equation by every possible term of degree $\leq d$. This increases the number of equations to $m \binom{n+d}{d}$ and the total number of terms to $\binom{n+2B+1+d}{n}$, and we can solve this system as soon as the former exceeds the latter.

Let us now consider what happens if we try to apply the same strategy to sparse LWE. The only difference is the additional constraint that each of the original equations are supported on k variables instead of n . Since our dense minor attack already provides an algorithm for solving sLWE when we are given more than $k \binom{n}{k}$ samples, we will assume that $m < k \binom{n}{k} < \binom{n+1}{k+1}$ below.

- If $2B + 1 < k$, there is virtually no difference between LWE and sparse LWE for this approach. This is because in either case, the original equations do not have terms containing more than $2B + 1$ distinct variables. We will therefore focus on the case where k is smaller than $2B + 1$.
- The total number of terms is bounded above by the number of all degree $2B + 1$ monomials supported on at most k out of n variables. This equals

$$\binom{n+1}{k+1} \binom{2B+1+k}{k} > \binom{n+1}{k+1} > m$$

Clearly, the number of terms is greater than the number of equations, and trying to linearize the system naively does not work.

- We can consider increasing the number of equations by following the exact recipe used by Arora-Ge. For some d , multiply each equation by every possible monomial of degree at most d . The problem is that this breaks the locality pattern that distinguishes sLWE from LWE in the Arora-Ge approach, since each equation can be now supported on many more variables.
- We can instead multiply each equation by all the monomials of degree at most d which are supported on the same k variables as the equation. This preserves the locality, but does not suffice since the number of terms continue to exceed the number of equations

$$\binom{n+1}{k+1} \binom{2B+d+1+k}{k} > m \binom{d+k}{k}$$

Another natural way to utilize the sparsity is to try to find a large number of samples supported on the same set of t variables for some $t < n$, and then run the Arora-Ge algorithm on those samples with an effective dimension of t . The dense minor lower bound rules out this approach for any $t = o(n)$. Therefore, the effective dimension in this approach will be linear in n , and we would still require many more samples than what we have available.

Capturing BKW Style Attacks. The BKW family of attacks ([5, 19, 49]) look for a short vector $\mathbf{x}$ such that the product $\mathbf{A}\mathbf{x}$ is either 0 or very sparse. They first partition the columns of $\mathbf{A}$ into blocks of size B . New samples are then created by finding collisions in these blocks and reducing those blocks to zero by subtraction. Iteratively applying this process for each block, we can find vectors of magnitude $2^{n/B}$ in the nullspace of $\mathbf{A}$.

To execute this attack naively, we need access to a very large number of samples. For **LWE**, we can generate new samples by taking random linear combinations with small norm. However, new samples generated by combining columns of $\mathbf{A}$ destroy sparsity really fast. Therefore, the search step of **BKW** attack does not advantage sparse **LWE** in any way over standard **LWE**.

We can try to take advantage of the fact that the columns of sparse **LWE** are mostly filled with zeros, and hence it will be easier to find collisions in blocks. While this is indeed true for the first few blocks, observe that we subtract columns from each other at every iteration of **BKW**. Each subtraction operation would increase the number of nonzero entries in the remaining blocks exponentially, and hence the required sample complexity for the attack will at least be the same as that of regular **LWE** with a constant factor smaller dimension.

So a natural way to leverage sparsity is therefore to work with a dense minor of order $\log p$ of size L . Our theorem suggests that $L = \Omega(n)$, and we would again require more samples than we have available.

Capturing Lattice-Style Attacks. Typically, lattice attacks are applied by applying lattice reduction to solve either:

- Finding a short $\mathbf{x}$ so that $\mathbf{Ax} = \mathbf{0}$. We already discussed why this approach doesn't have an advantage for **sLWE** over **LWE** in the **BKW** section.
- Solving **BDD** with respect to lattice $\mathcal{L} = \{\mathbf{xA}\}$ and target point $\mathbf{b} = \mathbf{sA} + \mathbf{e}$. The hardness of this depends on the ratio γ of $\lambda_1(\mathcal{L})$ and $\|\mathbf{e}\|$.

In Sect. 4.1, we argued that the size of the shortest nonzero vector in $\mathcal{L}$ is similar to what it would have been if $\mathbf{A}$ was chosen randomly. In fact, the minimum distance is smaller than a random lattice. We therefore, expect γ to only be smaller. So if anything, so called primal attacks against **LWE** work worse than expected when they are given sparse **LWE** instances.

On apparent inspection of lattice reduction algorithms, naive application of those don't seem to exploit sparsity at all. That said, most of our column vectors are mutually perpendicular since the dot product of two column vectors is nonzero only if their locality patterns intersect (these columns are however large in norm roughly $\sqrt{kp}$, and so they do not necessarily represent a good basis). So, although unlikely, it is plausible that lattice reduction based methods turn out to do better. We leave the study of this phenomenon to future work.

The most natural way to exploit the small sparsity of the system is to find a dense minor and apply lattice algorithms to an instance with a smaller dimension. As long as the noise ratio is inverse polynomial in the dimension n , the best lattice based attacks take $2^{O(n)}$ time (see the fantastic thesis of Rachel Player [62] or Albrecht et.al. [6, 7] for an excellent survey on concrete running times). Thus if we are able to find dense minor of size (say) $n/2$ of order $\log p$, that is an exponential improvement in the running time over a dense matrix. Our dense-minor model captures such improvements. In fact, we rule out attacks that makes use of only $L + 1$ samples in L variables. Typical lattice attacks require at least cL (where $c > 1$ is some constant) or even $L \log p$ samples.

7.3 Open Questions

We leave some exciting open problems for future work:

1. Are there other ways of leveraging sparsity besides the dense-minor method?
2. Analyze the performance of current lattice reduction algorithms with sparse LWE matrices.
3. Build new lattice reduction algorithms that naturally exploit sparsity

Exploiting Special Structure. One aspect that we did not touch at all in this work is whether structured assumptions for sparse LWE are secure and/or useful for cryptographic constructions. Some variants include circular sparse LWE and small secret LWE. While we did not observe vulnerabilities in circular sparse LWE, we note that one has to be careful working with small secret sparse LWE. Unlike LWE where these two variants are equivalent [11, 38, 55], it is not the case here.

For instance, if each secret coordinate is boolean, a single sample of k -sparse LWE would leak a lot of information about the secret vector. There are 2^k possible choices of secret coordinates occurring in each equation. If k is constant or logarithmic, this could effectively lead to a search algorithm for sparse LWE. We note that however, some applications might need a small secret and it might be plausible that once $(\|\mathbf{s}\|_\infty)^k$ is large enough, this variant becomes hard. We leave understanding small secret LWE as a great open problem.

7.4 Concrete Parameter Estimation

In this section, we focus on concrete security. Based on the cryptanalysis so far, we identify two major avenues of attack against sLWE specifically:

- Find a t -sparse vector in the kernel. To rule out this attack, we merely need to set $p < B^t$.
- Find a dense minor of order at least $1 + 1/L$. We write a script that estimates the size of the dense minor as a function of the dimension n , sparsity k and sample complexity m . To see the code, please refer to the full version of this paper.

To get a desired level of security using sLWE, we first look at the parameters for regular LWE to obtain the necessary dimension n' for m samples, modulus p and noise σ . We then compute a small k and a dimension n for which the size of dense minor in a $\text{sLWE}_{n,k,m,p,\sigma}$ matrix exceeds n' . We set concrete parameters based on the following assumption:

Conjecture 7.2. Let the expected size of the dense minor in $\text{sLWE}_{n,k,m,p,\sigma}$ matrix be $L > n'$. Then the $\text{sLWE}_{n,k,m,p,\sigma}$ problem is at least as hard as $\text{LWE}_{n',m,p,\sigma}$

We show a few recommended parameter settings in the table below. We use LWE dimension $n' = 2^{10}$ and database size N between 2^{26} and 2^{42} as per [43]. Note that the number of samples $m = \sqrt{N}$. For all of these settings we suggest

Table 1. The recommended dimension for sLWE whose security guarantees are equivalent to LWE with dimension 1024

Number of samples m Sparsity parameter k	2^{13}	2^{17}	2^{21}
20	1218	1425	1656
30	1143	1265	1395
40	1110	1195	1285
50	1090	1158	1234

setting the modulus to either 2^{16} , 2^{32} or 2^{64} since there is generally hardware support for arithmetic modulo these values. We also set the noise parameter $\sigma = 6.4$. Neither σ nor p depends on any of n, m or k (Table 1).

We need not worry about the sparse vector in kernel attack since the size of the dense minor in each of the above settings is 1023, which is much larger than the logarithm of the modulus. Furthermore, it is evident from the above table that for typical parameters used in practice, we only need to increase the dimension by a small constant factor < 1.5 in most cases to achieve equivalent security guarantees.

8 Applications of Sparse LWE

Our main applications lie in (secret-key) encryption schemes that support homomorphism (please refer to the full version of this paper for formal definition).

8.1 Linearly Homomorphic Encryption

We recall the Brakerski-Vaikuntanathan secret key homomorphic encryption scheme from LWE [24, 26]. We then focus on analyzing the efficiency improvement when the scheme is instantiated with sparse LWE and used for computing linear functions, in comparison with instantiation using LWE or Ring LWE.

The Basic Secret Key Homomorphic Encryption Scheme

Let $p = p(\lambda)$ be a modulus and the message space be $\mathbb{Z}_p$.

$\text{pp} \leftarrow \text{sklhe.Setup}(1^\lambda, 1^D)$: Based on the security parameter λ , choose a modulus q , where $\gcd(p, q) = 1$, a dimension n , a B -bounded distribution of noise elements χ . Output $\text{pp} := (q, n, \chi)$.

$(\text{sk}, \text{evk}) \leftarrow \text{sklhe.KeyGen}()$: Sample a secret key $\mathbf{s}' \leftarrow \mathbb{Z}_q^{1 \times n}$ and set $\mathbf{s} = (\mathbf{s}', -1)$. Output $\text{sk} = \mathbf{s}$.

$\text{ct} \leftarrow \text{sklhe.Enc}(\text{sk}, m)$: Parse the message m as an element of $\mathbb{Z}_p$. Sample a random k -sparse coefficient vector $\mathbf{a} \in \mathbb{Z}_q^n$, and a noise $e \leftarrow \chi$. Output $\text{ct}^\top := (\mathbf{a}^\top, b = \mathbf{s}^\top \mathbf{a} + m + e \cdot p \bmod q)$.

$m \leftarrow \text{sklhe.Dec}(\mathbf{s}, \text{ct})$: Parse $\text{ct} = (\mathbf{a}, b) \in \mathbb{Z}_q$. Output $m := (b - \mathbf{s}^\top \mathbf{a}) \bmod q, p$.

$\text{ct} \leftarrow \text{skhe.Add}(\text{ct}_1, \text{ct}_2, l)$: This algorithm takes two ciphertexts and scalar $l \in \mathbb{Z}_p$ and homomorphically computes $m_1 + l \cdot m_2 \bmod p$. Output $\text{ct} := \text{ct}_1 + l \cdot \text{ct}_2 \bmod q$.

Using the above operation, one can evaluate any linear function L specified by a coefficient vector $\mathbf{l} \in \mathbb{Z}_p^\ell$, on fresh ciphertexts $\text{ct}_1, \dots, \text{ct}_\ell$ with noises bounded by B , and obtain and output ciphertext ct with noise bounded by $\bar{B} = \ell \cdot p \cdot B$. Decryption correctness is guaranteed if $p \cdot \bar{B} \leq q$, that is, $q > \ell \cdot p^2 B$.

Sparse Representation and Storage. In the above scheme, while the sk is a dense vector in $\mathbb{Z}_q^{1 \times n}$, the ciphertext ct is only $(k+1)$ -sparse. We can store ct using the following succinct representation. For every integer vector $\mathbf{a}$ in $\mathbb{Z}_Q^n$, define

$$\begin{aligned} \text{indicator set:} \quad I_{\mathbf{a}} &= \{i_j \mid i_j \in [n] \text{ s.t. } a_{i_j} \neq 0\}, \\ \text{sparsity of } \mathbf{a}: \quad k_{\mathbf{a}} &= |I_{\mathbf{a}}| \\ \text{succinct representation:} \quad \tilde{\mathbf{a}} &:= \text{SparseRep}(\mathbf{a}) = \{(i_j, a_{i_j}) \mid i_j \in S_{\mathbf{a}}\} \end{aligned}$$

The size of the sparse representation is only $k_{\mathbf{a}}(\log n + \log Q)$, as opposed to $n \log Q$. As a result, encryption scheme has a good rate as shown below. In particular, when $\log p = \Omega(\log n + \log q)$, the rate is $O(k)$, which could be as small as a super-constant, or even constant if the number of ciphertexts encrypted under $\mathbf{s}$ is a bounded polynomial.

$$\text{rate} = \frac{|\text{ct}|}{|m|} = \frac{O(k(\log n + \log q))}{\log p} = O(k) \text{ if } \log p = \Omega(\log n + \log q)$$

Table 2. Efficiency comparison of different linearly homomorphic schemes.

Scheme & Assumptions	# of ciphertexts	rate	encryption overhead	evaluation overhead
Elgamal & DDH ¹	unbounded	$O(\lambda)$	$\tilde{O}(\lambda)$	$O(\lambda)$
Elgamal & DCR [60] ²	unbounded	$O(1)$	$\tilde{O}(\lambda)$	$O(\lambda)$
BV [26] & LWE ³	unbounded	$O(n)$	$O(n \log \log q)$	$O(n)$
BV & RLWE ³	unbounded	$O(n)$	$\tilde{O}(n \log \log q)$	$O(n)$
[33] & RLWE	unbounded	$O(1)$	$O(1)$	$\text{polylog}(n)$ ⁴
BV & $O(1)$ -sparse LWE (Ours) ³	bounded poly	$O(1)$	$O(\log \log q)$	$O(1)$
BV & $\omega(1)$ -sparse LWE (Ours) ³	unbounded	$\omega(1)$	$\omega(\log \log q)$	$\omega(1)$

¹ Use a group G with $|G| = 2^{\Omega(\lambda)}$ and generator g . For $\mathbf{m} \in \{0, 1\}^\ell$, $\text{ct} = g^r, g^{sr+\mathbf{m}}$. Linear function $\mathbf{l} \in \mathbb{Z}_{|G|}^\ell$ is evaluated over $\mathbb{Z}_{|G|}$ and output must be small in order to decrypt.

² Use group $\mathbb{Z}_{N^2}^*$ for an RSA integer $N = 2^{\Omega(\lambda)}$, with hard group generator g . For $\mathbf{m} \in \mathbb{Z}_{N^2}^\ell$, $\text{ct} = g^r, g^{sr(1+N)^{\mathbf{m}}}$. Linear function $\mathbf{l} \in \mathbb{Z}_N^\ell$ is evaluated over $\mathbb{Z}_N$.

³ For all the BV schemes, we consider the case $\log p = \Omega(\log n + \log q)$.

⁴ The scheme [33] has large polylogarithmic overhead.

Efficiency Analysis. We now analyze the efficiency overhead, in particular, we measure the encryption overhead as ratio between the encryption time and the number of bits encrypted, and the evaluation overhead as the ratio between the evaluation time and the time for evaluating the linear function in the clear.

The sparsity of the ciphertext can be leveraged for efficiency in the RAM model, where one can access an element in a vector or matrix at index $i \in [n]$ in logarithmic time $O(\log n)$ (note the index already has length $O(\log n)$). In contrast, in the circuit model this takes linear time. Consider encryption, where the dense secret key $\mathbf{s}$ is multiplied with a k -sparse coefficient vector $\mathbf{a}$, in the RAM model, this can be implemented in time $O(k(\log n + \log q \log \log q))$, involving fetching all the elements s_j for $j \in I_{\mathbf{a}}$, and performing a k -wise inner product between $\mathbf{s}_{I_{\mathbf{a}}}$ and $\mathbf{a}_{I_{\mathbf{a}}}$. The evaluation of linear function $\mathbf{l} \in \mathbb{Z}_p^\ell$ can be done using ℓ invocation of `sklhe.Add(ct, cti, li)`, where ct_{*i*} is the ciphertext of the i^{th} input element. Since ct_{*i*} is $(k + 1)$ -sparse, this can be implemented in time $O(k(\log n + \log q \log \log q))$.

$$\begin{aligned} \text{encryption overhead} &= \frac{T_{\text{Enc}}}{|m|} = \frac{O(k(\log n + \log q \log \log q))}{\log p} = O(k \log \log q) \\ \text{evaluation overhead} &= \frac{T_{\text{Eval}}}{T_L} = \frac{O(\ell k(\log n + \log q \log \log q))}{O(\ell \log p \log \log p)} = O(k) \end{aligned}$$

The right hand side equality holds when $\log p = \Omega(\log n + \log q)$. By setting k to be super-constant or even constant, we get nearly constant overheads.

Comparison with other Secret-Key LHE. To the best of our knowledge, the BV scheme instantiated with sparse LWE is the only scheme achieving nearly constant overheads in all aspects, storage, encryption time, and evaluation time. A comparison with previous schemes is provided in Table 2.

The most competitive scheme is by [33] instantiated by RLWE, which presented a polylogarithmic overhead FHE scheme. To achieve this low overhead, they use the BV/BGV encryption scheme to encrypt a vector of messages $\mathbb{Z}_p^n$ in the same ciphertext, resulting in a “packed” ciphertexts. These packed ciphertexts natively support SIMD computation i.e., coordinate-wise addition and multiplication. To support general, non SIMD, computation, [33] introduced novel methods for permuting elements encrypted in the same ciphertexts, and combining elements in different ciphertexts into the same ciphertexts. Combined with the use of permutation networks, they can implement arbitrary permutation over elements encrypted in multiple ciphertexts with polylogarithmic amortized overhead. Since general circuit computation can be implemented using SIMD operation and permutation, they achieve homomorphic evaluation with polylogarithmic amortized overhead. The only drawback is that the polylogarithmic overhead is very high and unlikely going to be practical.

Applications to Private Information Retrieval. A nice application that could help test concrete efficiency gains afforded by sparse LWE could be private-

information retrieval schemes. In particular, consider the practical scheme SimplePIR [43] that builds on top of [47] and instantiates the linearly-homomorphic encryption scheme using LHE from LWE. The idea in that scheme is that the data $\tilde{\mathbf{D}} \in \{0, 1\}^N$ is processed into a square shape yielding a matrix $\mathbf{D} \in \{0, 1\}^{\sqrt{N} \times \sqrt{N}}$. For $i \in [\sqrt{N}]$, we let $\mathbf{D}_i$ as the i^{th} column of $\mathbf{A}$.

The client samples a matrix $\mathbf{A} \in \mathbb{Z}_p^{n \times \sqrt{N}}$ at random from the LWE distribution. To fetch the data index (i, j) the client computes a query $\mathbf{b} = \mathbf{s}\mathbf{A} + 2 \cdot \mathbf{e} + \mathbf{v}_i$ where $\mathbf{v}_i \in \{0, 1\}^{1 \times \sqrt{N}}$ is an indicator vector with 1 at the i^{th} coordinate. The server computes $(\mathbf{b}\mathbf{D}_1, \dots, \mathbf{b}\mathbf{D}_{\sqrt{N}})$ and $(\mathbf{A}\mathbf{D}_1, \dots, \mathbf{A}\mathbf{D}_{\sqrt{N}})$ and sends these to the client. Note that $\mathbf{b}\mathbf{D}_j \approx \mathbf{s}\mathbf{A}\mathbf{D}_j + \mathbf{D}[i, j]$ upto small and even error. The client can therefore derive $\mathbf{D}[i, j]$ using the secret $\mathbf{s}$. The scheme can be proven to be secure due to the security of LWE. Note that the total communication from client to the server is $\tilde{O}(n\sqrt{N})$ where n is the dimension of LWE as the client needs to send $\mathbf{A}$ along with $\mathbf{b}$. The query compute time is also $\tilde{O}(n\sqrt{N})$ which is also the time to compute $\mathbf{b}$. We note that the client communication can be improved to $\tilde{O}(\sqrt{N})$ if $\mathbf{A}\mathbf{D}$ can be precomputed. We will assume this in the following discussion. For server, given the preprocessing, the computation consists of evaluating $\mathbf{b}\mathbf{D}$ which takes $\tilde{O}(N)$ times but the communication is $\tilde{O}(\sqrt{N})$. The downside to this scheme is that if a data entry needs to be changed, it takes $\tilde{O}(n)$ time to change the preprocessing $\mathbf{A}\mathbf{D}$.

If one used sparse LWE instead, it will reduce the client query generation time to $\tilde{O}(k\sqrt{N})$ as computing LWE samples will now be faster. This might be useful if the client is low-complexity. Further, the update time reduces to $\tilde{O}(k)$ as opposed to $\tilde{O}(n)$. This might be useful in dynamic settings where the update needs to be faster. To summarize, we have the following comparison.

	LWE	k -sparse LWE
Total Communication	$\tilde{O}(\sqrt{N})$	$\tilde{O}(\sqrt{N})$
Update Time	$\tilde{O}(n)$	$\tilde{O}(k)$
Client Query Time	$\tilde{O}(n \cdot \sqrt{N})$	$\tilde{O}(k \cdot \sqrt{N})$

We leave concretely implementing the PIR scheme for a future update.

8.2 Constant Degree Homomorphic Encryption

We now describe how to perform homomorphic multiplication. In order to achieve the best efficiency for evaluating constant-degree polynomials, we will make two modifications to the Brakerski-Gentry-Vaikuntanathan [24] homomorphic multiplication. First, instead of using the gadget matrix consisting of powers of 2, we will use the gadget matrix consisting of powers a larger base, in particular, p the modulus of the message space. Second, instead of performing modulus reduction on the ciphertexts, we will perform the modulus reduction

on the secret instead. The second modification is necessary since we want to rely on large secret sparse LWE instead of small secret sparse LWE. In contrast, standard modulus reduction (performed on ciphertexts) relies on having small secrets in order to ensure that errors introduced by rounding is bounded. While using small secrets is without loss of generality for LWE, it is not so for sparse LWE.

Let $\mathbf{g}_{Q,p} = (1, p, p^2, \dots, p^{\lfloor \log_p Q \rfloor})$ be the gadget matrix in $\mathbb{Z}_Q$ with base p . Correspondingly, define operation $\text{Decomp}(\mathbf{x}, p)$ that decomposes the integer vector $\mathbf{x}$ in base p . We will also make use of the subroutine Scale defined in [24, 26].

Definition 8.1. For integer vector $\mathbf{x} \in \mathbb{Z}_q$ and integers $q \geq p \geq r$, we define $\mathbf{x}' \leftarrow \text{Scale}(\mathbf{x}, q, p, r)$ to be the $\mathbb{Z}_p$ -vector closest to $(p/q) \cdot \mathbf{x}$ that satisfies $\mathbf{x}' = \mathbf{x} \bmod r$.

They proved the following key lemma about the scale operation, showing that if the inner product between two vectors $\text{ip} \mathbf{u}, \mathbf{v}$ over $\mathbb{Z}_q$ is not too large, then inner product with the scaled vector $\text{ip} \mathbf{u}' = \text{Scale}(q, p, r), \mathbf{v}$ over $\mathbb{Z}_p$ is congruent to $\text{ip} \mathbf{u}, \mathbf{v} \bmod r$, and is roughly q/p times smaller than $\text{ip} \mathbf{u}, \mathbf{v}$ with an rounding error bounded by $r \|\mathbf{v}\|_1 / 2$.

Lemma 8.1. Let $q > p > r$ be positive integers satisfying $q = p = 1 \bmod r$. Let $\mathbf{u} \in \mathbb{Z}_q^n$ and $\mathbf{u}' \leftarrow \text{Scale}(\mathbf{u}, q, p, r) \in \mathbb{Z}_p^n$. Then, for any $\mathbf{v} \in \mathbb{Z}_q^n$, the following holds

$$\text{If,} \quad [\langle \mathbf{u}, \mathbf{v} \rangle]_q < q/2 - (q/p) \cdot (r/2) \cdot \|\mathbf{v}\|_1,$$

$$\text{Then,} \quad [\langle \mathbf{u}', \mathbf{v} \rangle]_{p,r} = [\langle \mathbf{u}, \mathbf{v} \rangle]_{q,r} \quad \text{and} \quad [\langle \mathbf{u}', \mathbf{v} \rangle]_p < \frac{q}{p} [\langle \mathbf{u}, \mathbf{v} \rangle]_q + \frac{1}{2} \cdot r \cdot \|\mathbf{v}\|_1$$

In the BGV scheme, the inner product between the secret key and a ciphertext $\langle \mathbf{s}, \text{ct} \rangle$ over $\mathbb{Z}_q$ is the approximate message $m + p \cdot E$. The modulus reduction step scales the ciphertext to a smaller modulus q' with respect to p . By the above lemma, they have that $\langle \mathbf{s}, \text{Scale}(\text{ct}, q, q', p) \rangle \bmod q' = m + p \cdot e$, where the noise is bounded by $|e| \leq \frac{q}{q'} |E| + p \cdot \|\mathbf{s}\|_1$, which is small if $\mathbf{s}$ is short.

In our setting, we use large secret $\mathbf{s}$ and instead scale the secret. To start with, observe that $\langle \mathbf{s} \otimes \mathbf{g}_{q,p}, \text{Decomp}(\text{ct}, p) \rangle \bmod q = m + p \cdot E$. Therefore, we have $\langle \text{Scale}(\mathbf{s} \otimes \mathbf{g}_{q,p}, q', p), \text{Decomp}(\text{ct}, p) \rangle \bmod q' = m + p \cdot e$, where the noise is bounded by $|e| \leq \frac{q}{q'} |E| + p \cdot \|\text{Decomp}(\text{ct}, p)\|_1$. In the actual scheme described below, we apply this scaling w.r.t. the tensor product of two ciphertexts $\text{ct} = \text{ct}_1 \otimes \text{ct}_2$, and the secret key is $\mathbf{s} \otimes \mathbf{s}$. The operation is summarized in the following corollary.

Corollary 8.1. Let $q > q' > p$ be positive integers satisfying $q = q' = 1 \bmod p$. For every $\mathbf{s} \in \mathbb{Z}_q^{1 \times (n+1)}$, and $\text{ct}_1, \text{ct}_2 \in \mathbb{Z}_q^{(n+1)}$.

$$\begin{aligned}
& \text{If,} \quad m + p \cdot E = [(s \otimes s) \cdot (\text{ct}_1 \otimes \text{ct}_2)]_q = [(s \otimes s \otimes g_{q,p}) \cdot (\text{Decomp}(\text{ct}_1 \otimes \text{ct}_2, p))]_q \\
& < q/2 - (q/q') \cdot (p/2) \cdot (n^2 \cdot p \cdot \log q / \log p), \\
& \text{Let} \quad \bar{t} = \text{Scale}(s \otimes s \otimes g_{q,p}, q', p) \\
& \quad m' + p \cdot e = [\bar{t} \cdot (\text{Decomp}(\text{ct}_1 \otimes \text{ct}_2, p))]_{q'} \\
& \text{Then,} \quad m = m', \quad \text{and} \quad |m' + p \cdot e| \leq \frac{q}{q'} |m + p \cdot E| + (p/2) \cdot (n^2 \cdot p \cdot \log q / \log p)
\end{aligned}$$

Constant Degree Homomorphic Encryption Scheme

pp $\leftarrow$ **skhe.Setup**($1^\lambda, 1^D$): Based on the security parameter λ and maximal depth D , choose a tower of decreasing modulus $q = q^{(D)} \geq q^{(D-1)} > q^{(D-2)} > \dots > q^{(0)} > p$, where $\gcd(p, q^{(j)}) = 1$, $q^{(d)} = 1 \bmod p$, and $q^d / q^{d-1} > \Delta$. Also choose a dimension n , a B -bounded distribution of noise elements χ . Output **pp** := (q, n, χ, p) , where q is the vector of modulus.

sk, evk $\leftarrow$ **skhe.KeyGen**(): For every $d \in [D]$, sample a random k -sparse coefficient matrix $\mathbf{B}^{(d)}$ over $\mathbb{Z}_{q^{(d)}}$, a noise vector $e^{(d)} \leftarrow \chi^{1 \times m}$, and a secret key $s'^{(d)} \leftarrow \bar{\mathbb{Z}}_p^{1 \times n}$. Let $s^{(d)} := (s'^{(d)}, -1)$. Compute the level d evaluation key as follows:

$$\begin{aligned}
t^{(d+1)} &:= s^{(d+1)} \otimes s^{(d+1)} \otimes g_{q^{(d+1)}, p} \bmod q^{(d+1)} && \in \mathbb{Z}_{q^{(d+1)}}^{1 \times m^{(d+1)}} \\
\bar{t}^{(d+1)} &:= \text{Scale}(t, q^{(d+1)}, q^{(d)}, p) && \in \mathbb{Z}_{q^{(d)}}^{1 \times m^{(d+1)}} \\
\text{evk}^{(d)} &:= \left(s^{(d)} \mathbf{B}^{(d)} + \bar{t}^{(d+1)} + p \cdot e^{(d)} \bmod q^{(d)} \right) && \in \mathbb{Z}_{q^{(d)}}^{(n+1) \times m^{(d+1)}}
\end{aligned}$$

Above dimension $m^{(d)} = \Theta(n \log q^{(d+1)} / \log p) = O(n \log q^{(D)} / \log p)$.
Output (**sk** = $\{s^d\}$, **evk** = $\{\text{evk}^d\}$).

ct^(D) $\leftarrow$ **skhe.Enc**(**sk**, m): Parse the message m as an element of $\mathbb{Z}_p$. Sample a random k -sparse coefficient vector $a \in \mathbb{Z}_{q^{(D)}}$, and a noise $e \leftarrow \chi$.

Output **ct**^(D)^T := $(a^T, b = s'^{(D)} a + m + e \cdot p \bmod q^{(D)})$.

$m \leftarrow$ **skhe.Dec**(**sk**, **ct**^(d), d): Parse **ct**^(d) = $(a, b) \in \mathbb{Z}_{q^{(d)}}$. Output $m := (b - s^{(d)} a) \bmod q^{(d)}, p$.

ct^(d+1) $\leftarrow$ **skhe.Mult**(**ct**₁^(d+1), **ct**₂^(d+1)): This algorithm takes two ciphertexts at level $d+1$ in $\mathbb{Z}_{q^{(d+1)}}^{n+1}$ and outputs a ciphertext of dimension $(n+1)^2$.

Output **ct**^(d+1) := **ct**₁^(d+1) $\otimes$ **ct**₂^(d+1) $\bmod q^{(d+1)}$.

ct^(d) $\leftarrow$ **skhe.Reduce**(**evk**, **ct**^(d+1)): This algorithm takes the evaluation key and a ciphertext at level $d+1$ of dimension $(n+1)^2$ and returns a ciphertext at level d of dimension $n+1$.

$$\text{ct}^{(d)} := \text{evk}^{(d)} \cdot \text{Decomp}(\text{ct}^{(d+1)}, p) \bmod q^{(d)}.$$

Evaluating Constant-Degree Polynomials and Noise Bounds. Using the above operations, we can homomorphically evaluate a constant-degree polynomial, expressed as a sum of degree 2^D monomials, by homomorphic evaluation of all monomial followed by homomorphic additions. To homomorphically compute a degree 2^{D+1-d} monomial $\prod_{j \in [2^{D+1-d}]} m_{i_j}$ (initially $d = D$, and at the last step $d = 1$), suppose we have already obtained ciphertext ct_1 encrypting $\prod_{j \in [2^{D-d}]} m_{i_j}$ and ct_2 encrypting $\prod_{j \in [2^{D-d+1}, 2^{D+1-d}]} m_{i_j}$, both at level d . We can obtain a ciphertext ct encrypting $\prod_{j \in [2^{D+1-d}]} m_{i_j}$ at level $d - 1$ by first invoking skhe.Mult on ct_1, ct_2 obtaining an intermediate ciphertext of dimension $(n + 1)^2$, followed by skhe.Reduce to reduce the dimension to n while also reducing the modulus to $q^{(D-d)}$, i.e.,

$$\text{ct} = \text{evk}^{(d-1)} \cdot \text{Decomp}(\text{ct}_1 \otimes \text{ct}_2) \bmod q^{(d-1)}.$$

For a formal proof of correctness as well as runtime analysis, please refer to the full version of this paper.

A comparison of the efficiency of our scheme with previous schemes is provided in Table 3.

Table 3. Efficiency comparison of different homomorphic schemes supporting constant-degree polynomials.

Scheme & Assumptions	# of ciphertexts	rate	encryption overhead	evaluation overhead
BV [26] & LWE ¹	unbounded	$O(n)$	$O(n \log \log q)$	$\Omega(n^3)$
BV & RLWE ¹	unbounded	$O(n)$	$\tilde{O}(n \log \log q)$	$\Omega(n \log n)$
[33] & RLWE	unbounded	$O(1)$	$O(1)$	$\text{polylog}(n)$ ²
BV & $O(1)$ -sparse LWE (Ours) ¹	bounded poly	$O(1)$	$O(\log \log q^{(D)})$	$O(1)$
BV & $\omega(1)$ -sparse LWE (Ours) ¹	unbounded	$\omega(1)$	$\omega(\log \log q^{(D)})$	$\omega(1)$

¹ For all the BV schemes, we consider the case $\log p = \Omega(\log n + \log q)$.

² The scheme [33] has large polylogarithmic overhead.

References

1. Aggarwal, D., Dadush, D., Regev, O., Stephens-Davidowitz, N.: Solving the shortest vector problem in $2n$ time using discrete gaussian sampling. In: Proceedings of the Forty-Seventh Annual ACM Symposium on Theory of Computing, pp. 733–742 (2015)
2. Alagic, G., et al.: Status Report on the Third Round of The NIST Post-Quantum Cryptography Standardization Process. US Department of Commerce, NIST (2022)

3. Albrecht, M., et al.: Homomorphic encryption security standard. Tech. rep., HomomorphicEncryption.org, Toronto, Canada (2018)
4. Albrecht, M., Cid, C., Faugère, J.C., Fitzpatrick, R., Perret, L.: Algebraic algorithms for LWE problems (2014)
5. Albrecht, M.R., Cid, C., Faugère, J.-C., Fitzpatrick, R., Perret, L.: On the complexity of the BKW algorithm on LWE. *Des. Codes Crypt.* **74**(2), 325–354 (2013). <https://doi.org/10.1007/s10623-013-9864-x>
6. Albrecht, M.R., et al.: Estimate all the LWE, NTRU schemes! In: Catalano, D., De Prisco, R. (eds.) SCN 2018. LNCS, vol. 11035, pp. 351–367. Springer, Cham (2018). https://doi.org/10.1007/978-3-319-98113-0_19
7. Albrecht, M.R., Player, R., Scott, S.: On the concrete hardness of learning with errors. *J. Math. Cryptol.* **9**(3), 169–203 (2015)
8. Alekhnovich, M.: More on average case vs approximation complexity. In: 44th FOCS, pp. 298–307. IEEE Computer Society Press (2003)
9. Allen, S.R., O'Donnell, R., Witmer, D.: How to refute a random CSP. In: Guruswami, V. (ed.) 56th FOCS, pp. 689–708. IEEE Computer Society Press (2015)
10. Applebaum, B., Barak, B., Wigderson, A.: Public-key cryptography from different assumptions. In: Schulman, L.J. (ed.) 42nd ACM STOC, pp. 171–180. ACM Press (2010)
11. Applebaum, B., Cash, D., Peikert, C., Sahai, A.: Fast cryptographic primitives and circular-secure encryption based on hard learning problems. In: Halevi, S. (ed.) CRYPTO 2009. LNCS, vol. 5677, pp. 595–618. Springer, Heidelberg (2009). https://doi.org/10.1007/978-3-642-03356-8_35
12. Applebaum, B., Damgård, I., Ishai, Y., Nielsen, M., Zichron, L.: Secure arithmetic computation with constant computational overhead. In: Katz, J., Shacham, H. (eds.) CRYPTO 2017. LNCS, vol. 10401, pp. 223–254. Springer, Cham (2017). https://doi.org/10.1007/978-3-319-63688-7_8
13. Applebaum, B., Ishai, Y., Kushilevitz, E.: On pseudorandom generators with linear stretch in NC^0 . In: Díaz, J., Jansen, K., Rolim, J.D.P., Zwick, U. (eds.) APPROX/RANDOM -2006. LNCS, vol. 4110, pp. 260–271. Springer, Heidelberg (2006). https://doi.org/10.1007/11830924_25
14. Applebaum, B., Lovett, S.: Algebraic attacks against random local functions and their countermeasures. In: Wichs, D., Mansour, Y. (eds.) 48th ACM STOC, pp. 1087–1100. ACM Press (2016)
15. Arora, S., Ge, R.: New algorithms for learning in presence of errors. In: Aceto, L., Henzinger, M., Sgall, J. (eds.) ICALP 2011. LNCS, vol. 6755, pp. 403–415. Springer, Heidelberg (2011). https://doi.org/10.1007/978-3-642-22006-7_34
16. Banaszczyk, W.: New bounds in some transference theorems in the geometry of numbers. *Math. Ann.* **296**(1), 625–635 (1993)
17. Banerjee, A., Peikert, C., Rosen, A.: Pseudorandom functions and lattices. In: Pointcheval, D., Johansson, T. (eds.) EUROCRYPT 2012. LNCS, vol. 7237, pp. 719–737. Springer, Heidelberg (2012). https://doi.org/10.1007/978-3-642-29011-4_42
18. Beullens, W., Lyubashevsky, V., Nguyen, N.K., Seiler, G.: Lattice-based blind signatures: Short, efficient, and round-optimal. *Cryptology ePrint Archive*, Report 2023/077 (2023). <https://eprint.iacr.org/2023/077>
19. Blum, A., Kalai, A., Wasserman, H.: Noise-tolerant learning, the parity problem, and the statistical query model. *J. ACM (JACM)* **50**(4), 506–519 (2003)

20. Boneh, D., et al.: Fully key-homomorphic encryption, arithmetic circuit ABE and compact garbled circuits. In: Nguyen, P.Q., Oswald, E. (eds.) EUROCRYPT 2014. LNCS, vol. 8441, pp. 533–556. Springer, Heidelberg (2014). https://doi.org/10.1007/978-3-642-55220-5_30
21. Brakerski, Z.: Quantum FHE (almost) as secure as classical. In: Shacham, H., Boldyreva, A. (eds.) CRYPTO 2018. LNCS, vol. 10993, pp. 67–95. Springer, Cham (2018). https://doi.org/10.1007/978-3-319-96878-0_3
22. Brakerski, Z., Christiano, P., Mahadev, U., Vazirani, U.V., Vidick, T.: A cryptographic test of quantumness and certifiable randomness from a single quantum device. In: Thorup, M. (ed.) 59th FOCS, pp. 320–331. IEEE Computer Society Press (2018)
23. Brakerski, Z., Gentry, C., Vaikuntanathan, V.: (leveled) fully homomorphic encryption without bootstrapping. In: Innovations in Theoretical Computer Science 2012, Cambridge, MA, USA, 8–10 January 2012, pp. 309–325 (2012)
24. Brakerski, Z., Gentry, C., Vaikuntanathan, V.: (Leveled) fully homomorphic encryption without bootstrapping. In: Goldwasser, S. (ed.) ITCS 2012, pp. 309–325. ACM (2012)
25. Brakerski, Z., Langlois, A., Peikert, C., Regev, O., Stehlé, D.: Classical hardness of learning with errors. In: Boneh, D., Roughgarden, T., Feigenbaum, J. (eds.) 45th ACM STOC, pp. 575–584. ACM Press (2013)
26. Brakerski, Z., Vaikuntanathan, V.: Efficient fully homomorphic encryption from (standard) LWE. In: Ostrovsky, R. (ed.) 52nd FOCS, pp. 97–106. IEEE Computer Society Press (2011)
27. Choudhuri, A.R., Jain, A., Jin, Z.: SNARGs for $\mathcal{P}$ from LWE. In: 62nd FOCS, pp. 68–79. IEEE Computer Society Press (2022)
28. Couteau, G., Rindal, P., Raghuraman, S.: Silver: silent VOLE and oblivious transfer from hardness of decoding structured LDPC codes. In: Malkin, T., Peikert, C. (eds.) CRYPTO 2021. LNCS, vol. 12827, pp. 502–534. Springer, Cham (2021). https://doi.org/10.1007/978-3-030-84252-9_17
29. Cryan, M., Miltersen, P.B.: On pseudorandom generators in NC^0 . In: Sgall, J., Pultr, A., Kolman, P. (eds.) MFCS 2001. LNCS, vol. 2136, pp. 272–284. Springer, Heidelberg (2001). https://doi.org/10.1007/3-540-44683-4_24
30. Dao, Q., Ishai, Y., Jain, A., Lin, H.: Multi-party homomorphic secret sharing and sublinear MPC from sparse LPN. In: Handschuh, H., Lysyanskaya, A. (eds.) CRYPTO 2023, Part II. LNCS, vol. 14082, pp. 315–348. Springer, Heidelberg (Aug (2023)). https://doi.org/10.1007/978-3-031-38545-2_11
31. Feige, U.: Relations between average case complexity and approximation complexity. In: 34th ACM STOC, pp. 534–543. ACM Press (2002)
32. Gentry, C.: Fully homomorphic encryption using ideal lattices. In: Mitzenmacher, M. (ed.) 41st ACM STOC, pp. 169–178. ACM Press (2009)
33. Gentry, C., Halevi, S., Smart, N.P.: Fully homomorphic encryption with polylog overhead. In: Pointcheval, D., Johansson, T. (eds.) EUROCRYPT 2012. LNCS, vol. 7237, pp. 465–482. Springer, Heidelberg (2012). https://doi.org/10.1007/978-3-642-29011-4_28
34. , Gentry, C., Peikert, C., Vaikuntanathan, V.: Trapdoors for hard lattices and new cryptographic constructions. In: Ladner, R.E., Dwork, C. (eds.) 40th ACM STOC, pp. 197–206. ACM Press (2008)
35. Gentry, C., Peikert, C., Vaikuntanathan, V.: Trapdoors for hard lattices and new cryptographic constructions. In: STOC (2008). <https://eprint.iacr.org/2007/432>

36. Gentry, C., Sahai, A., Waters, B.: Homomorphic encryption from learning with errors: conceptually-simpler, asymptotically-faster, attribute-based. In: Canetti, R., Garay, J.A. (eds.) CRYPTO 2013. LNCS, vol. 8042, pp. 75–92. Springer, Heidelberg (2013). https://doi.org/10.1007/978-3-642-40041-4_5
37. Goldreich, O.: Candidate one-way functions based on expander graphs. In: Goldreich, O. (ed.) Studies in Complexity and Cryptography. Miscellanea on the Interplay between Randomness and Computation. LNCS, vol. 6650, pp. 76–87. Springer, Heidelberg (2011). https://doi.org/10.1007/978-3-642-22670-0_10
38. Goldwasser, S., Kalai, Y.T., Peikert, C., Vaikuntanathan, V.: Robustness of the learning with errors assumption. In: Yao, A.C. (ed.) Innovations in Computer Science - ICS 2010, Tsinghua University, Beijing, China, 5-7 January 2010. Proceedings, pp. 230–240. Tsinghua University Press (2010)
39. Gorbunov, S., Vaikuntanathan, V., Wee, H.: Attribute-based encryption for circuits. In: Boneh, D., Roughgarden, T., Feigenbaum, J. (eds.) 45th ACM STOC, pp. 545–554. ACM Press (2013)
40. Goyal, R., Koppula, V., Waters, B.: Lockable obfuscation. In: Umans, C. (ed.) 58th FOCS, pp. 612–621. IEEE Computer Society Press (2017)
41. Goyal, R., Koppula, V., Waters, B.: Collusion resistant traitor tracing from learning with errors. In: Diakonikolas, I., Kempe, D., Henzinger, M. (eds.) 50th ACM STOC, pp. 660–670. ACM Press (2018)
42. Henzinger, A., Dauterman, E., Corrigan-Gibbs, H., Zeldovich, N.: Private web search with tiptoe. In: Flinn, J., Seltzer, M.I., Druschel, P., Kaufmann, A., Mace, J. (eds.) Proceedings of the 29th Symposium on Operating Systems Principles, SOSP 2023, Koblenz, Germany, October 23–26, 2023, pp. 396–416. ACM (2023), <https://doi.org/10.1145/3600006.3613134>
43. Henzinger, A., Hong, M.M., Corrigan-Gibbs, H., Meiklejohn, S., Vaikuntanathan, V.: One server for the price of two: Simple and fast single-server private information retrieval. In: Calandrino, J.A., Troncoso, C. (eds.) 32nd USENIX Security Symposium, USENIX Security 2023, Anaheim, CA, USA, August 9–11, 2023, pp. 3889–3905. USENIX Association (2023). <https://www.usenix.org/conference/usenixsecurity23/presentation/henzinger>
44. Ishai, Y., Kushilevitz, E., Ostrovsky, R., Sahai, A.: Cryptography with constant computational overhead. In: Ladner, R.E., Dwork, C. (eds.) 40th ACM STOC, pp. 433–442. ACM Press (2008)
45. Jain, A., Lin, H., Sahai, A.: Indistinguishability obfuscation from well-founded assumptions. In: Khuller, S., Williams, V.V. (eds.) 53rd ACM STOC, pp. 60–73. ACM Press (2021)
46. , Kothari, P.K., Mori, R., O’Donnell, R., Witmer, D.: Sum of squares lower bounds for refuting any CSP. In: Hatami, H., McKenzie, P., King, V. (eds.) 49th ACM STOC, pp. 132–145. ACM Press (2017)
47. Kushilevitz, E., Ostrovsky, R.: Replication is NOT needed: SINGLE database, computationally-private information retrieval. In: 38th FOCS, pp. 364–373. IEEE Computer Society Press (1997)
48. Lenstra, A.K., Lenstra, H.W., Lovász, L.: Factoring polynomials with rational coefficients. *Math. Ann.* **261**(ARTICLE), 515–534 (1982)
49. Lyubashevsky, V.: The parity problem in the presence of noise, decoding random linear codes, and the subset sum problem. In: Chekuri, C., Jansen, K., Rolim, J.D.P., Trevisan, L. (eds.) APPROX/RANDOM -2005. LNCS, vol. 3624, pp. 378–389. Springer, Heidelberg (2005). https://doi.org/10.1007/11538462_32

50. Lyubashevsky, V., Micciancio, D.: On bounded distance decoding, unique shortest vectors, and the minimum distance problem. In: Halevi, S. (ed.) CRYPTO 2009. LNCS, vol. 5677, pp. 577–594. Springer, Heidelberg (2009). https://doi.org/10.1007/978-3-642-03356-8_34
51. Lyubashevsky, V., Nguyen, N.K., Plançon, M.: Lattice-based zero-knowledge proofs and applications: shorter, simpler, and more general. In: Dodis, Y., Shrimpton, T. (eds.) CRYPTO 2022, Part II. LNCS, vol. 13508, pp. 71–101. Springer, Heidelberg (Aug (2022))
52. Lyubashevsky, V., Peikert, C., Regev, O.: On ideal lattices and learning with errors over rings. In: Gilbert, H. (ed.) EUROCRYPT 2010. LNCS, vol. 6110, pp. 1–23. Springer, Heidelberg (2010). https://doi.org/10.1007/978-3-642-13190-5_1
53. Mahadev, U.: Classical homomorphic encryption for quantum circuits. In: Thorup, M. (ed.) 59th FOCS, pp. 332–338. IEEE Computer Society Press (2018)
54. Mahadev, U.: Classical verification of quantum computations. In: Thorup, M. (ed.) 59th FOCS, pp. 259–267. IEEE Computer Society Press (2018)
55. Micciancio, D.: On the hardness of learning with errors with binary secrets. *Theory Comput.* **14**(13), 1–17 (2018). <https://theoryofcomputing.org/articles/v014a013>
56. Micciancio, D., Goldwasser, S.: Complexity of lattice problems: a cryptographic perspective, vol. 671. Springer, New York (2002). <https://doi.org/10.1007/978-1-4615-0897-7>
57. Micciancio, D., Peikert, C.: Trapdoors for lattices: simpler, tighter, faster, smaller. In: Pointcheval, D., Johansson, T. (eds.) EUROCRYPT 2012. LNCS, vol. 7237, pp. 700–718. Springer, Heidelberg (2012). https://doi.org/10.1007/978-3-642-29011-4_41
58. Micciancio, D., Peikert, C.: Hardness of SIS and LWE with small parameters. In: Canetti, R., Garay, J.A. (eds.) CRYPTO 2013. LNCS, vol. 8042, pp. 21–39. Springer, Heidelberg (2013). https://doi.org/10.1007/978-3-642-40041-4_2
59. Mossel, E., Shpilka, A., Trevisan, L.: On e-biased generators in NC0. In: FOCS, pp. 136–145 (2003)
60. Paillier, P.: Public-key cryptosystems based on composite degree residuosity classes. In: Stern, J. (ed.) EUROCRYPT 1999. LNCS, vol. 1592, pp. 223–238. Springer, Heidelberg (1999). https://doi.org/10.1007/3-540-48910-X_16
61. Peikert, C.: Public-key cryptosystems from the worst-case shortest vector problem: extended abstract. In: Mitzenmacher, M. (ed.) Proceedings of the 41st Annual ACM Symposium on Theory of Computing, STOC 2009, Bethesda, MD, USA, May 31 - June 2, 2009, pp. 333–342. ACM (2009)
62. Player, R.: Parameter selection in lattice-based cryptography. Ph.D. thesis, Royal Holloway, University of London (2018)
63. Regev, O.: On lattices, learning with errors, random linear codes, and cryptography. In: STOC, pp. 84–93 (2005)
64. Roşca, M., Sakzad, A., Stehlé, D., Steinfeld, R.: Middle-product learning with errors. In: Katz, J., Shacham, H. (eds.) CRYPTO 2017. LNCS, vol. 10403, pp. 283–297. Springer, Cham (2017). https://doi.org/10.1007/978-3-319-63697-9_10
65. Schnorr, C.P.: A hierarchy of polynomial time lattice basis reduction algorithms. *Theoret. Comput. Sci.* **53**(2–3), 201–224 (1987)
66. Schnorr, C.: Block reduced lattice bases and successive minima. *Comb. Probab. Comput.* **3**, 507–522 (1994)
67. Schnorr, C.P., Euchner, M.: Lattice basis reduction: improved practical algorithms and solving subset sum problems. *Math. Program.* **66**, 181–199 (1994)
68. Smart, N.P., Vercauteren, F.: Fully homomorphic SIMD operations. *Des. Codes Crypt.* **71**(1), 57–81 (2012). <https://doi.org/10.1007/s10623-012-9720-4>

69. Vaikunthanathan, V.: CS 294 lecture 4: worst-case to average-case reductions for LWE. Class at UC Berkeley (2020). <http://people.csail.mit.edu/vinodv/CS294/lecture6.pdf>
70. Wichs, D., Zirdelis, G.: Obfuscating compute-and-compare programs under LWE. In: Umans, C. (ed.) 58th FOCS, pp. 600–611. IEEE Computer Society Press (2017)