



Universiteit
Leiden
The Netherlands

Lattice cryptography: isomorphisms & dual attacks

Pulles, L.N.

Citation

Pulles, L. N. (2026, September 23). *Lattice cryptography: isomorphisms & dual attacks*. Retrieved from <https://hdl.handle.net/1887/4309918>

Version: Publisher's Version

License: [Licence agreement concerning inclusion of doctoral thesis in the Institutional Repository of the University of Leiden](#)

Downloaded from: <https://hdl.handle.net/1887/4309918>

Note: To cite this publication please use the final published version (if applicable).

Part II

Dual-Sieve Attack

Chapter 3

Dual-Sieve Attack Heuristics

This chapter is based on joint work with L. Ducas published at CRYPTO 2023 [DP23c].

Guo and Johansson [GJ21], and MATZOV [MAT22] have independently claimed improved attacks against various lattice-based schemes in NIST's PQC standardisation process, such as ML-KEM and ML-DSA, by using an FFT on top of the so-called Dual-Sieve attack.

However, this chapter shows that a heuristic used in above works not only theoretically contradicts with both formal theorems and well-tested heuristics in certain regimes, but also provides incorrect predictions in experiments. We conclude that this heuristic significantly overestimates the success probability of the Dual-Sieve attack.

In the process, it is shown that use of the FFT is not specific to a Dual-Sieve attack against learning with errors but is more generally useful against bounded distance decoding.

3.1 Introduction

Many post-quantum cryptosystems base their security on the hardness of the learning with errors (LWE) problem, introduced by Regev in 2005 [Reg05, Reg09], which is basically the bounded distance decoding (BDD) problem in q -ary lattices. One possible type of attack against BDD is the so-called *dual attack*, which uses the idea by Aharonov and Regev [AR04] that short dual vectors allow distinguishing points close to a lattice from points far away from the lattice. This idea can even be traced back in a pure geometric context to earlier work by Håstad [Hås88], and is not even limited to lattices, as it was already implicit in the very construction of low-density parity-check codes [Gal62]. In contrast, a lattice-based attack operating solely on the *primal* lattice is called a *primal attack*. The best dual and primal attacks are then typically used by the lattice cryptanalyst to parametrise LWE cryptosystems.

More specifically, the dual attack attempts to solve the search-BDD problem by performing a reduction to the decision-BDD problem, where one is asked to determine whether a target point in Euclidean space is either a *uniform target*, i.e., sampled uniformly modulo the lattice, or a *BDD target*, i.e., sampled from a distribution that is concentrated around lattice points. Here, a short dual vector provides a score function that assigns a high score to BDD targets and an expected score of zero to uniform targets. Only by using many dual vectors and considering the sum of individual scores [LW21], one can hope to distinguish successfully between the two target distributions with a high probability. To get exponentially many short dual vectors, Alkim, Ducas, Pöppelmann and Schwabe initially suggested [ADPS16] to use a lattice sieve [NV08, MV10, BDGL16] on the dual lattice.

We refer to this style of attack as a *Dual-Sieve* attack. The concrete cryptanalytic impact of this idea can be further improved by guessing multiple coordinates of the secret rather than just one [Alb17, EJK20], and then finding the right solution among these candidates rather than just a few.

Another development of the dual attack is also reminiscent from a cryptanalytic technique of code-based cryptography by Leveil and Fouque [LF06]. The idea is to batch the score evaluation of many algebraically related candidates via a *fast Fourier transform* (FFT). For carefully crafted parameters, the cost of computing all those scores is

barely larger than the cost of naïvely computing a single score. This led Guo and Johansson [GJ21] to claim an improved attack on various NIST post-quantum standardisation candidates, followed quickly by an independent technical report by Israel’s Center of Encryption and Information Security (MATZOV) [MAT22]. We refer to this style of attack as a *Dual-Sieve-FFT* attack. The latter has already been followed up upon, with a quantum variant [AS23] and a coding-theoretic enhanced variant [CST22]. In their analysis of the Dual-Sieve attack, all these four works use a heuristic saying that individual score functions given by a set of dual vectors, are mutually independent variables.

3.1.1 Contributions

This chapter provides three contributions.

Generalisation of the Dual-Sieve-FFT attack (Section 3.4)

The original principle of the dual attack [AR04, Hås88] is stated for the BDD problem in any lattice. However, the recent instances of the Dual-Sieve attack [ADPS16, Alb17, EJK20] and the Dual-Sieve-FFT attack [GJ21, MAT22, CST22, AS23] are stated specifically for **LWE**. One exception is the work of [LW21], which is geometrically enlightening but limited to the Dual-Sieve attack.

Our first contribution is therefore to generalise the **FFT** trick [GJ21] to the setting of **BDD**. Beyond the theoretical satisfaction of abstracting the technique to its mathematical core, this generalisation also offers further improvement over the work of [GJ21]. Namely, for the same algorithmic cost, we can further improve the shortness of the dual vectors and therefore their distinguishing power.

Contradictions to the independent score heuristic (Section 3.5)

A second observation regarding this literature is that the analysis of the Dual-Sieve attack (with or without **FFT**) relies on one specific heuristic, which has received essentially no attention so far. Namely, it is assumed that all the individual scores, given by each dual vector, are mutually independent.

We approach the analysis of this heuristic by looking at the conclusions it leads to. The geometric point of view offered by the work of Laarhoven and Walter [LW21] is pivotal in that respect: judging the reasonability of

a heuristic conclusion is very much enabled by the language of geometry. In particular, [LW21] concludes with a heuristic algorithm that solves decision-BDD, even when the noise slightly exceeds the Gaussian heuristic (GH), i.e., the expected minimal distance of a random lattice. This should raise suspicion, as a random point is on average not expected to be further than GH away from the lattice. We show this suspicion is fair: it follows from a result of Debris-Alazard, Ducas, Resch, and Tillich [DADRT23] that the above task is statistically impossible to solve, even to an unbounded attacker.

The contradiction above is, however, limited to a rather theoretical regime of the Dual-Sieve attack, which is not that of the recent concrete cryptanalytic claims [EJK20, GJ21, MAT22, CST22, AS23], where the noise is below GH. Still, here the attack needs to distinguish between a BDD sample and a uniform sample with a very high success probability. We will show that this situation also leads to a contradiction.

Namely, we argue that the heuristic probability of having a distinguishing failure is much smaller than the probability of a uniform target lying closer to the lattice than a BDD target. The claimed high success probability of distinguishing between BDD and uniform contradicts the basic principle that targets closer to the lattice get a higher score in expectation. It turns out that the parameters used in [GJ21, MAT22] specifically fall into that contradictory regime that requires a distinguisher with a very high success probability.

Experimental invalidation of prior model (Section 3.6)

To develop a more precise understanding, we zoom in on the score distribution for BDD targets and uniform targets. Extensive experiments were performed to discover that both distributions deviate from the predictions made under the independent score heuristic. First, the *body* of the distribution of scores for uniform targets is properly predicted, but not its *tail*: after a predicted rapid decrease, visually resembling a *waterfall*, this distribution hits a *floor*. This is perfectly in line with our second contradictory regime: some uniform targets will be close to the lattice, and should therefore have a high score.

The score distribution for BDD targets is also predicted incorrectly, and this is no longer just a matter of the tail. Contrary to prediction, this score is not distributed like a Gaussian. It is in fact not even symmetric around its average, and its variance appears exponentially larger than

predicted. In particular, a **BDD** target will more likely have a low score than what is heuristically predicted.

Open data

All the written code and the obtained data to generate the figures, can be found in the GitHub repository <https://github.com/ludopulles/DoesDualSieveWork>. The experiments are written in Python and use the G6K and FPyLLL libraries [ADH⁺19, FPL25], and there is a binding to some C code to accelerate the **FFT**.

3.1.2 Coding theory analogue

In coding theory, the dual attack is called “statistical decoding” [Gal62, ALJ01, Ove06]. Here, with the use of many low-weight parity-check equations $\mathcal{H}$, which is analogous to short dual lattice vectors, one can decide whether a target $\mathbf{t} \in \mathbb{F}_q^n$ is close to some codeword or not, based on a similar scoring function.

Recently, there was a work by Carrier, Debris-Alazard, Meyer-Hilfiger and Tillich [CDMT22] that claims to have an improved algorithm for statistical decoding that outperforms the state-of-the-art information set decoding algorithm [BM18] on sparse binary codes. However, this work was based on the coding-theory analogue of the **independent score heuristic** [CDMT22, Assumption 3.7], which states that individual scores $(\langle \mathbf{h}, \mathbf{t} \rangle)_{\mathbf{h} \in \mathcal{H}}$ are **i.i.d.** Bernoulli variables. Here, similarly it was soon realised this heuristic leads to a flawed analysis, and two proposed fixes [MT23] were enough to overcome the problems. Namely, first, given a $[n, k]$ -code $\mathcal{C}$, they model the number of weight w elements in a coset $\mathcal{C} + \mathbf{t}$ for $\mathbf{t} \leftarrow \mathcal{U}(\mathbb{F}_2^n)$ as a Poisson process with parameter $\lambda = \binom{n}{w}/2^{n-k}$. In addition, they use Fourier analysis on the score function, and using both fixes, they then show that the statistical decoding algorithm can be corrected while only increasing the asymptotic runtime cost by a logarithmic factor.

Remark 3.1. Our code $\mathcal{C}$ in the paragraph above is written in [MT23, Prop. 1] as a shortened $[n - s, k - s]$ -code, instead. This shortened code plays a similar rôle as the sublattice in our reduction from search-**BDD** to decision-**BDD** in Section 3.2.2.

3.2 Dual attack

3.2.1 Solving decision-BDD

The idea of using short dual vectors for solving decision-BDD can be traced back at least to [AR04] in the lattice literature, and can be viewed as the lattice analogue of an old decoding technique [Gal62, AIJ01, Ove06]. Given a BDD sample $\mathbf{t} = \mathbf{v} + \mathbf{e}$ with $\mathbf{v} \in \Lambda$, for any dual vector $\mathbf{w} \in \Lambda^\vee$ one has,

$$\langle \mathbf{t}, \mathbf{w} \rangle = \langle \mathbf{v}, \mathbf{w} \rangle + \langle \mathbf{e}, \mathbf{w} \rangle \equiv \langle \mathbf{e}, \mathbf{w} \rangle \pmod{1} .$$

In particular, if the error $\mathbf{e}$ and the dual vector $\mathbf{w}$ are of small enough ℓ_2 -norm, $\langle \mathbf{t}, \mathbf{w} \rangle$ should be close to an integer. Moreover, $\langle \mathbf{t}, \mathbf{w} \rangle \pmod{1}$ is thus well-defined for any $\mathbf{t} \in \mathbb{R}^n / \Lambda$. A natural score to consider as some indication of the target $\mathbf{t}$ being close to the lattice Λ is therefore given by,

$$f_{\mathbf{w}}(\mathbf{t}) = \frac{\chi_{\mathbf{w}}(\mathbf{t}) + \chi_{-\mathbf{w}}(\mathbf{t})}{2} = \cos(2\pi \cdot \langle \mathbf{t}, \mathbf{w} \rangle) ,$$

where $\chi_{\mathbf{w}}$ is defined in Lemma 2.73.

If a target $\mathbf{t}$ is indeed close to the lattice, $f_{\mathbf{w}}(\mathbf{t})$ should be close to 1, but the converse does not need to be true. To improve confidence in the fidelity of this score, one may naturally consider the total score over many dual vectors $\mathcal{W} \subset \Lambda^\vee$ given by,

$$f_{\mathcal{W}}(\mathbf{t}) = \sum_{\mathbf{w} \in \mathcal{W}} f_{\mathbf{w}}(\mathbf{t}) .$$

This function is called the simple distinguisher f_{simple} by Laarhoven and Walter [LW21], and it closely resembles the Aharonov–Regev [AR04] distinguisher, given by

$$f_{\mathcal{W}}^{\text{AR}}(\mathbf{t}) = \sum_{\mathbf{w} \in \mathcal{W}} \rho_{1/(2\pi\sigma)}(\mathbf{w}) f_{\mathbf{w}}(\mathbf{t}) .$$

By checking whether $f_{\mathcal{W}}(\mathbf{t})$ is above or below some optimally chosen threshold, one can say from which distribution $\mathbf{t}$ is drawn, effectively solving decision-BDD with a high success probability.

In carefully crafted circumstances, in particular regarding the construction of the set of short dual vectors $\mathcal{W}$, this approach can give a provable worst-case distinguisher [AR04] or certificate [Hås88].

More recent works [LW21, GJ21, MAT22] have reused this idea more heuristically, in a context where $\mathcal{W}$ simply is the set of all the dual vectors smaller than a certain radius. In such a situation, $\mathcal{W}$ is typically the output of a sieve algorithm [BDGL16]. Moreover, all the nonzero dual vectors have approximately the same weight, and [LW21] shows that the simple distinguisher $f_{\mathcal{W}}$ works asymptotically almost as good as the Aharonov–Regev one $f_{\mathcal{W}}^{\text{AR}}$, although the former can be computed faster in practice.

3.2.2 Reducing search-BDD to decision-BDD

Recent dual attacks [EJK20, GJ21, MAT22] and follow-up works reduce search-BDD on Λ to decision-BDD on a (projected) sublattice as follows. A search-BDD instance consists of a lattice Λ and a BDD sample $\mathbf{t} = \mathbf{v} + \mathbf{e}$ where $\mathbf{v} \in \Lambda$ and $\mathbf{e}$ has small norm. First, a sparsification $\Lambda' \subset \Lambda$ is chosen. Then, we try to solve decision-BDD on the sublattice Λ' with $|\Lambda/\Lambda'|$ many targets: $(\mathbf{t} - \mathbf{g})_{\mathbf{g} + \Lambda' \in \Lambda/\Lambda'}$. There is exactly one target drawn from a BDD distribution, namely the target $\mathbf{t} - \mathbf{g}$ where $\mathbf{g} \in \mathbf{v} + \Lambda'$ is the ‘correct guess.’ If decision-BDD is solved successfully, i.e., only $\mathbf{t} - \mathbf{g}$ is identified as a BDD sample for the correct guess $\mathbf{g}$, then search-BDD is solved modulo Λ' , i.e., $\mathbf{v} \equiv \mathbf{g} \pmod{\Lambda'}$. Now, Λ' and $\mathbf{t} - \mathbf{g}$ forms a new search-BDD instance that is easier because the new lattice Λ' is sparser and the distance between $\mathbf{t} - \mathbf{g}$ and Λ' is at most $\|\mathbf{e}\|$. By recursively solving this easier search-BDD instance, one may ultimately recover $\mathbf{v}$ completely. gamely, after a certain number of iterations, say ℓ , the ℓ th easier search-BDD instance $(\Lambda^{(\ell)}, \mathbf{t}^{(\ell)})$ is easy enough that Babai’s NP algorithm, upon input $\Lambda^{(\ell)}$ and $\mathbf{t}^{(\ell)}$, outputs $\mathbf{v}' \in \Lambda^{(\ell)}$ such that $\mathbf{t}^{(\ell)} = \mathbf{v}' + \mathbf{e}$ holds. In this case, $\mathbf{v} = \mathbf{v}' + \sum_{i=1}^{\ell} \mathbf{g}^{(i)}$, where $\mathbf{g}^{(i)}$ denotes the correct guess in the i th iteration, is the solution of the original search-BDD instance.

Remark 3.2. Many of the recent dual attacks, such as Guo–Johansson and MATZOV [GJ21, MAT22], however, run a sieve on a lattice $L \subset (\Lambda')^{\vee}$ of much lower rank $\beta_{\text{sieve}} < n$, from which a set of dual vectors $\mathcal{W} \subset L$ is obtained, to solve decision-BDD. Observe that the distinguisher function now satisfies

$$f_{\mathcal{W}}(\mathbf{t}) = f_{\mathcal{W}}(\pi_{\text{span}(L)}(\mathbf{t})) ,$$

where π_V denotes the projection map onto the vector space V . This shows that $f_{\mathcal{W}}$ is now merely a distinguisher for $L^{\vee}$: it assigns high scores

to all targets $\mathbf{t} \in \mathbb{R}^n$ that have their projection $\pi_{\text{span}(L)}(\mathbf{t})$ close to $L^\vee$.

More specifically, recent dual attacks first perform lattice reduction, such as **BKZ**, on Λ' to obtain some reduced basis $\mathbf{D}$ for $(\Lambda')^\vee$. Then, L is selected as the lattice generated by the basis $\mathbf{D}_{[1, \beta_{\text{sieve}}]}$. By duality, then $({}^\vee\mathbf{D})_{[n - \beta_{\text{sieve}} + 1, n]}$ is a basis for $L^\vee$, obtained by projecting Λ' away from the first $n - \beta_{\text{sieve}}$ vectors of ${}^\vee\mathbf{D}$.

Note, when $\mathbf{e}$ follows an n -dimensional Gaussian distribution with parameter σ , then $\pi_{\text{span}(L)}(\mathbf{e})$ follows a β_{sieve} -dimensional Gaussian distribution with parameter σ . On the other hand, if a guess $\mathbf{g} \notin \mathbf{v} + \Lambda'$ is incorrect, then $f_{\mathcal{W}}$ assigns a high score to the target $\mathbf{t} - \mathbf{g}$ when the point

$$\pi_{\text{span}(L)}(\mathbf{t}) = \pi_{\text{span}(L)}(\mathbf{e}) + \pi_{\text{span}(L)}(\mathbf{v} - \mathbf{g}), \quad (3.1)$$

is close to $L^\vee$. Since the lower rank lattice L was only picked after some **BKZ** reduction, and $\mathbf{v} - \mathbf{g}$ is nonzero, we presume that $\pi_{\text{span}(L)}(\mathbf{v} - \mathbf{g})$ acts as a uniform point in $\text{span}(L)/L$. In conclusion, because the distinguisher is only distinguishing for $L^\vee$, we see that the target in Equation (3.1) corresponds to a **BDD** sample if $\mathbf{g} \in \mathbf{v} + \Lambda'$, and else to a uniform sample modulo $L^\vee$.

Remark 3.3. We note that the work of [PS24] sidesteps the uniform model for incorrect targets, by using dual vectors from the full-rank lattice Λ' and assuming a priori $\|\mathbf{e}\| < \frac{1}{2}\lambda_1(\Lambda)$. In this case, incorrect targets can be proved to be far enough from the lattice in the worst case, without any statistical nor heuristic argument as it was already the case in [AR04]. This is, however, not the regime used in recent concrete attack claims [GJ21, MAT22], and is in fact separated from the contradictory regime in Figure 3.3 by a factor 2 on the error length $\|\mathbf{e}\|$.

3.3 Prior dual attack models

In this subsection, we explain how the Dual-Sieve attack is analysed in the works of [LW21, GJ21, MAT22], and explicitly state where they have used heuristics.

3.3.1 Laarhoven–Walter

First, Laarhoven and Walter analyse in [LW21, Lem. 6] the distribution of the score $f_{\mathcal{W}}(\mathbf{t})$ for **BDD** targets $\mathbf{t}$ with a distance of exactly r to

the primal lattice. In the derivation, however, they approximate this distribution by targets sampled from a continuous Gaussian of parameter $\sigma = r/\sqrt{n}$.

Lemma 3.4 ([LW21, Lem. 6]). *Let $\Lambda \subset \mathbb{R}^n$ be a full-rank lattice and $\mathbf{w} \in \Lambda^\vee$ be a dual vector.*

(a) *If $\mathbf{t} \leftarrow \mathcal{U}(\mathbb{R}^n/\Lambda)$, then $\mathbb{E}[f_{\mathbf{w}}(\mathbf{t})] = 0$, and $\mathbb{V}[f_{\mathbf{w}}(\mathbf{t})] = 1/2$,*

(b) *If $\mathbf{t} \leftarrow \mathcal{N}_n(0, \sigma^2) \pmod{\Lambda}$ with $\sigma \in \mathbb{R}_{>0}$, then*

$$\mathbb{E}[f_{\mathbf{w}}(\mathbf{t})] = e^{-2\pi^2\sigma^2\|\mathbf{w}\|^2}, \quad \text{and} \quad \mathbb{V}[f_{\mathbf{w}}(\mathbf{t})] = \frac{1}{2} - \Theta\left(e^{-4\pi^2\sigma^2\|\mathbf{w}\|^2}\right). \quad (3.2)$$

Proof. For the variance, we will use $f_{\mathbf{w}}(\mathbf{t})^2 = \frac{1}{2} + \frac{1}{2} \cos(4\pi \langle \mathbf{w}, \mathbf{t} \rangle) = \frac{1}{2} + \frac{1}{2} f_{2\mathbf{w}}(\mathbf{t})$.

(a) Integrating over a fundamental region $\mathbf{B} \cdot [0, 1]^n$ shows $\mathbb{E}[f_{\mathbf{w}}(\mathbf{t})] = 0$ since $\int_0^1 \cos(\alpha + 2\pi kx) dx = 0$ holds for all $k \in \mathbb{Z}$ and $\alpha \in \mathbb{R}$. Then by the above, it readily follows,

$$\mathbb{V}[f_{\mathbf{w}}(\mathbf{t})] = \frac{1}{2} + \frac{1}{2} \mathbb{E}[f_{2\mathbf{w}}(\mathbf{t})] = \frac{1}{2}.$$

(b) Because a Gaussian is a radial distribution,

$$\mathbb{E}[f_{\mathbf{w}}(\mathbf{t})] = \mathbb{E}_{\mathbf{t} \leftarrow \mathcal{N}(0, \sigma^2)}[\cos(2\pi x \|\mathbf{w}\|)] = \exp\left(-2\pi^2\sigma^2\|\mathbf{w}\|^2\right),$$

and

$$\mathbb{V}[f_{\mathbf{w}}(\mathbf{t})] = \frac{1}{2} + \frac{1}{2} \mathbb{E}[f_{2\mathbf{w}}(\mathbf{t})] - \mathbb{E}[f_{\mathbf{w}}(\mathbf{t})]^2 = \frac{1}{2} + \frac{1}{2} \varepsilon^4 - \varepsilon^2,$$

where $\varepsilon = \exp(-2\pi^2\sigma^2\|\mathbf{w}\|^2) \in (0, 1)$.

□

However, to conclude on the behaviour of the total score $f_{\mathcal{W}}(\mathbf{t})$, one needs to resort to some heuristic. In [LW21], they resort to the following heuristic.

Heuristic 3.5 (independent score heuristic). *For any fixed set $\mathcal{W} \subset \Lambda^\vee$ and for any distribution of targets $\mathbf{t}$ considered in Lemma 3.4, the random variables $(\langle \mathbf{w}, \mathbf{t} \rangle \bmod 1)_{\mathbf{w} \in \mathcal{W}}$ are mutually independent.*

By combining the above heuristic with Heuristic 2.3—which looks fair, given the exponential size of $\mathcal{W}$ in the context of interest—one can model the total score $f_{\mathcal{W}}(\mathbf{t})$ of each type of sample as a Gaussian distribution of centre $|\mathcal{W}| \cdot E$ and variance $|\mathcal{W}| \cdot V$, where E and V are the expectation and variance given by the above Lemma 3.4.

One may then deduce the success probability of distinguishing with the score function using the following lemma.

Lemma 3.6. *Let $X \leftarrow \mathcal{N}(E_X, V_X)$ and $Y \leftarrow \mathcal{N}(E_Y, V_Y)$ be independent Gaussian random variables. Then,*

$$\Pr[X > Y] = \frac{1}{2} + \frac{1}{2} \operatorname{erf} \left(\frac{E_X - E_Y}{\sqrt{2(V_X + V_Y)}} \right). \quad (3.3)$$

Proof. Consider the variable $Z = Y - X$, which is also Gaussian, specifically $Z \sim \mathcal{N}(E_Z, V_Z)$ where $E_Z = E_Y - E_X$ and $V_Z = V_X + V_Y$. The event $X > Y$ is equivalent to $Z < 0$, so the result follows from Equation (2.2). $\square$

This lemma leads Laarhoven and Walter to the following heuristic claim.

Heuristic Claim 3.7 ([LW21, Lem. 9]). *Let $\Lambda \subset \mathbb{R}^n$ be a random unit-volume lattice, $r > 0$ and $\mathcal{W} \subset \Lambda^\vee$ a set consisting of the α^n shortest vectors of $\Lambda^\vee$, where $\alpha = \min\{\beta \mid e^2 \ln(\beta) = \beta^2 r^2\}$. Then, we have*

$$\Pr[f_{\mathcal{W}}(\mathbf{t}_{\text{BDD}}) > f_{\mathcal{W}}(\mathbf{t}_{\text{unif}})] \geq \frac{1}{2} + \frac{1}{2} \operatorname{erf} \left(\frac{1}{\sqrt{2}} \right) \approx 0.84,$$

where $\mathbf{t}_{\text{unif}} \leftarrow \mathcal{U}(\mathbb{R}^n/\Lambda)$ and $\mathbf{t}_{\text{BDD}} \leftarrow \mathcal{U}(r \operatorname{GH}(n) \mathcal{S}^{n-1})$ are sampled independently.

Heuristic Justification. First, we approximate the BDD sample $\mathbf{t}_{\text{BDD}}$ by a Gaussian distribution of parameter $\sigma = r \operatorname{GH}(n)/\sqrt{n}$. According to the Gaussian heuristic, the lengths of the vectors in $\mathcal{W}$ are concentrated around $\alpha \cdot \operatorname{GH}(n)$. By the independent score heuristic and Lemma 3.4, the score function for the BDD sample follows a Gaussian distribution $\mathcal{N}(E_X, V_X)$, where

$$E_X = \alpha^n \exp \left(-\frac{2\pi^2 \alpha^2 r^2 \cdot \operatorname{GH}(n)^4}{n} \right) = \alpha^n \exp \left(-\frac{\alpha^2 r^2 \cdot n}{2e^2} \right) = \alpha^{n/2},$$

by construction of α . The variance is $V_X \approx \frac{1}{2} \cdot \alpha^n$.

On the other hand, uniform samples give a score distribution Y following a Gaussian distribution $\mathcal{N}(E_Y, V_Y)$ where $E_Y = 0$ and $V_Y = \alpha^n/2$ by case (a) of Lemma 3.4.

Hence, by Lemma 3.6, the probability of having $X > Y$ equals

$$\frac{1}{2} + \frac{1}{2} \operatorname{erf} \left(\frac{\alpha^{n/2}}{\sqrt{2 \cdot \alpha^n}} \right) \approx 0.84. \quad \triangle$$

3.3.2 Guo–Johansson and MATZOV

The analysis proposed by Guo–Johansson [GJ21] is somewhat less explicit. Instead of analysing the score directly, they consider the statistical distance between the distribution of $\langle \mathbf{t}, \mathbf{w} \rangle \bmod 1$ for $\mathbf{t}$ uniform and Gaussian.

Using the same **independent score heuristic**, they then conclude on the statistical distance between the tuples $(\langle \mathbf{t}, \mathbf{w} \rangle \bmod 1)_{\mathbf{w} \in \mathcal{W}}$ for $\mathbf{t}$ uniform and Gaussian. While there exist optimal distinguishers (introduced in [LW21] using a lemma dating back to Neyman–Pearson [NP33]), it differs from the score function $f_{\mathcal{W}}$, but they seem to assume that the scoring function is not that far from optimal. An argument for such a statement is given by Laarhoven and Walter [LW21, Corollary 2], but it is not mentioned by Guo and Johansson [GJ21].

The analysis of **MATZOV** [MAT22] is on the contrary quite explicit on computing the distribution of scores, while taking into account the severe technical complications introduced by their *modulus switching*. Namely, they increase the number of dual vectors with a factor D_{round} in [MAT22, Section 5] to account for the effect of rounding the Fourier coefficients after performing a modulus switch. Another factor D_{arg} is also introduced, but we view it as dubious, see [DP23a, App. A.4].

In any case, we can essentially recover the key claims of [GJ21, MAT22] directly from the above analysis as well. Note that we express the result more generally in terms of **BDD** in an arbitrary lattice rather than a specific **LWE** instance. Moreover, we state this key claim here without loss of generality for a unit-volume, full-rank lattice, because of renormalisation and because the dual vectors used in distinguishing actually come from the dual of a projected sublattice. Indeed, the general setting of the Dual-Sieve attack [ADPS16, EJK20, GJ21, MAT22] first

applies **BKZ** reduction on the dual lattice $\Lambda_{\text{LWE}}^\vee$ of dimension d , and then runs a sieve on a lower-rank sublattice $\Lambda^\vee \subset \Lambda_{\text{LWE}}^\vee$, see Remark 3.2.

Heuristic Claim 3.8 (Key claim of [GJ21, MAT22], reconstructed). *Let $\Lambda \subset \mathbb{R}^n$ be a random unit-volume lattice, $\mathcal{W} \subset \Lambda^\vee$ the set consisting of the $(4/3)^{n/2}$ shortest vectors of $\Lambda^\vee$, and $\sigma > 0$. Taking $\ell = \sqrt{4/3} \cdot \text{GH}(n)$ and $\varepsilon = \exp(-2\pi^2\sigma^2\ell^2)$, we then have*

$$\Pr[f_{\mathcal{W}}(\mathbf{t}_{\text{BDD}}) > f_{\mathcal{W}}(\mathbf{t}_{\text{unif}})] \geq 1 - e^{-\frac{1}{2}|\mathcal{W}|\varepsilon^2},$$

where $\mathbf{t}_{\text{unif}} \leftarrow \mathcal{U}(\mathbb{R}^n/\Lambda)$ and $\mathbf{t}_{\text{BDD}} \leftarrow \mathcal{N}_n(0, \sigma^2)$ are sampled independently.

Heuristic Justification. Similar to the Heuristic Justification of Heuristic Claim 3.7, the score distribution for $\mathbf{t}_{\text{BDD}}$ is approximately $X \sim \mathcal{N}(\varepsilon \cdot |\mathcal{W}|, \frac{1}{2}|\mathcal{W}|)$, as lengths of vectors in $\mathcal{W}$ are concentrated around $\ell = \sqrt{4/3} \cdot \text{GH}(n)$ by the **Gaussian heuristic**. The uniform sample $\mathbf{t}_{\text{unif}}$ has a score distribution of approximately $Y \sim \mathcal{N}(0, \frac{1}{2}|\mathcal{W}|)$. Thus, by Lemma 3.6,

$$\Pr[f_{\mathcal{W}}(\mathbf{t}_{\text{BDD}}) > f_{\mathcal{W}}(\mathbf{t}_{\text{unif}})] = 1 - \frac{1}{2} \operatorname{erfc}\left(\sqrt{|\mathcal{W}|/2} \cdot \varepsilon\right).$$

If $\varepsilon\sqrt{|\mathcal{W}|} \leq 1/\sqrt{2\pi}$, the above probability is trivially at least $\frac{1}{2} \geq 1 - e^{-\frac{1}{2}|\mathcal{W}|\varepsilon^2}$. In the other case, Lemma 2.1 yields the claim:

$$\begin{aligned} \Pr[f_{\mathcal{W}}(\mathbf{t}_{\text{BDD}}) > f_{\mathcal{W}}(\mathbf{t}_{\text{unif}})] &\geq 1 - \frac{1}{\sqrt{2\pi}|\mathcal{W}| \cdot \varepsilon} e^{-\frac{1}{2}|\mathcal{W}|\varepsilon^2} \\ &> 1 - e^{-\frac{1}{2}|\mathcal{W}|\varepsilon^2}. \end{aligned} \quad \triangle$$

An important corollary of the above claim is that one can distinguish one **BDD** sample from $O(e^{\frac{1}{2}|\mathcal{W}|\varepsilon^2})$ uniform samples with constant probability, by marking the sample that evaluates to the highest score as the **BDD** sample. In the case of search-**BDD**, then with constant probability, one can recover the lattice vector $\mathbf{v}$ closest to $\mathbf{t}$ modulo the sparsified lattice Λ' , if the sparsification factor is bounded by the above quantity.

3.4 Generalisation of the Dual-Sieve-FFT attack

As established by the literature [Hås88, AR04, LW21], scoring target points to obtain information about their distance to a primal lattice Λ

using short dual vectors is very general, and not limited to **LWE** lattices. In this section, we will show that this is also the case of the extra FFT trick as proposed by Guo and Johansson [GJ21].

3.4.1 Abstracting Guo–Johansson

The general idea is as follows. Given a lattice Λ , one first crafts a sparsification Λ' of Λ , i.e., a sublattice $\Lambda' \subset \Lambda$ of finite index. This gives rise to a finite abelian group of cosets $G = \Lambda/\Lambda'$. Now, to solve **BDD** for a target $\mathbf{t} \in \mathbb{R}^n/\Lambda$ on the lattice Λ , one solves **BDD** for the target $\mathbf{t}$ on all the cosets $\Lambda' + g$, or equivalently, solve **BDD** for all the targets $\mathbf{t} - g$ with $g + \Lambda' \in G$ on the sublattice Λ' . For the correct choice of coset $g + \Lambda'$, the distance to $\mathbf{t}$ is the same as in the initial **BDD** problem, but the sublattice is sparser, making this **BDD** problem easier than the original one. However, we now have $|G|$ instances to consider.

With the help of the **DFT**, the score function $f_{\mathcal{W}}$ can be computed for all targets $\mathbf{t} - g$ in a batch. That is, applying the DFT_G in Equation (2.10) on a sequence $(\chi_{\mathbf{w}'}(g - \mathbf{t}))_{g \in G}$ for some $\mathbf{w}' \in (\Lambda')^\vee$ gives another sequence that at index $\chi_{\mathbf{w}} \in \widehat{G}$ has a value of,

$$\sum_{g \in G} \chi_{\mathbf{w}'}(g - \mathbf{t}) \overline{\chi_{\mathbf{w}}(g)} = \chi_{\mathbf{w}'}(-\mathbf{t}) \cdot \sum_{g \in G} \frac{\chi_{\mathbf{w}'}(g)}{\chi_{\mathbf{w}}(g)}, \quad (3.4)$$

where $\mathbf{w} + \Lambda^\vee \in (\Lambda')^\vee/\Lambda^\vee$ as $\widehat{G}$ is isomorphic to $(\Lambda')^\vee/\Lambda^\vee$ by Lemmata 2.73 and 2.74. Note that $\chi_{\mathbf{w}'}(g)$ is well-defined for $g \in \Lambda/\Lambda'$ as $\chi_{\mathbf{w}'}$ is Λ' -periodic. By the orthogonality of characters, note that

$$\sum_{g \in G} \frac{\chi_{\mathbf{w}'}(g)}{\chi_{\mathbf{w}}(g)} = \begin{cases} |G|, & \text{if } \mathbf{w}' \in \mathbf{w} + \Lambda^\vee, \\ 0, & \text{otherwise.} \end{cases}$$

Hence, Equation (3.4) is zero everywhere except at index $\chi_{\mathbf{w}'}$, where it is equal to $|G| \cdot \chi_{\mathbf{w}'}(-\mathbf{t})$.

By $\mathbb{C}$ -linearity of the **DFT**, one can obtain an expression for the **DFT** of $f_{\mathcal{W}}(g - \mathbf{t})$ for any finite set of dual vectors $\mathcal{W} \subset (\Lambda')^\vee$. More specifically, if for all $\mathbf{w} \in \mathcal{W}$ we have $-\mathbf{w} \in \mathcal{W}$, i.e., it is *symmetric*, we have

$$\text{DFT}_G\left((f_{\mathcal{W}}(\mathbf{t} - g))_{g \in G}\right) = |G| \cdot \left(\sum_{\mathbf{w}' \in \mathcal{W} \cap (\mathbf{w} + \Lambda^\vee)} f_{\mathbf{w}'}(\mathbf{t}) \right)_{\mathbf{w} + \Lambda^\vee \in \widehat{G}}.$$

Neglecting this scalar $|G|$, we therefore construct a batch of score functions, by performing an inverse FFT on the sequence $\sum_{\mathbf{w}' \in \mathbf{w} + \Lambda^\vee} f_{\mathbf{w}'}(-\mathbf{t})$ for all (dual) cosets $\mathbf{w} + \Lambda^\vee \in (\Lambda')^\vee / \Lambda^\vee$. Then, the entry with $g + \Lambda' \in G$ that has the highest score is most likely the coset $g + \Lambda'$ containing the lattice point in Λ that is closest to $\mathbf{t}$.

3.4.2 Implementation

In this section, we will give a concrete implementation of an algorithm that performs the general Dual-Sieve-FFT attack on a lattice Λ .

Concretely, the lattice Λ is specified by a basis $\mathbf{B} = [\mathbf{b}_1, \dots, \mathbf{b}_n]$, and one can take a simple sparsification such as Λ' being specified by the basis $[d_1 \mathbf{b}_1, \dots, d_n \mathbf{b}_n]$ for suitable $d_1, \dots, d_n \in \mathbb{N}$.

In fact, any sparsification is, after a basis change, of this shape. When the sublattice $\Lambda' = \mathbf{B}' \cdot \mathbb{Z}^n \subset \mathbf{B} \cdot \mathbb{Z}^n = \Lambda$ is described by a matrix $\mathbf{B}'$, we can express the basis $\mathbf{B}'$ in terms of $\mathbf{B}$, i.e., find the matrix $\mathbf{A} \in \mathbb{Z}^{n \times n}$ such that $\mathbf{B}' = \mathbf{B} \cdot \mathbf{A}$.

Then, put $\mathbf{A}$ in **SNF**, i.e., find matrices $\mathbf{S}, \mathbf{T} \in \text{GL}_n(\mathbb{Z})$ and a diagonal matrix $\mathbf{D}$ such that $\mathbf{A} = \mathbf{SDT}$ and thus, we have $\mathbf{B}'\mathbf{T}^{-1} = (\mathbf{BS}) \cdot \mathbf{D}$. As $\mathbf{A}$ was full-rank, $\mathbf{D}$ is full-rank and invertible over $\mathbb{Q}$, so here Λ' is described by the basis $[d_1 \mathbf{b}'_1, \dots, d_n \mathbf{b}'_n]$, where $\text{diag}(d_1, \dots, d_n) = \mathbf{D}$ and $\mathbf{BS} = [\mathbf{b}'_1, \dots, \mathbf{b}'_n]$ is a basis for Λ . Note that the **SNF** is efficiently computable [Sto96].

Hence, without loss of generality, we have a sparsification $\Lambda' \subset \Lambda$, where $\mathbf{B}$ is a basis for Λ and $\mathbf{B}' = [d_1 \mathbf{b}_1, \dots, d_n \mathbf{b}_n]$ is a basis for Λ' . Then Algorithm 3.1 will find the coset of Λ' that is most likely to contain a lattice vector closest to target $\mathbf{t}$.

Structure of the quotient group

From a geometric perspective, concerning the length of the vectors in $\mathcal{W}$, the structure of the group $G = \Lambda/\Lambda'$ does not appear to matter at all, only its size does. On the other hand, while asymptotically all group structures allow computing DFT_G in time $O(|G| \log |G|)$, the structure of the group matters quite a lot in practice and the case $G = (\mathbb{Z}/2\mathbb{Z})^k$, i.e., the Walsh–Hadamard Transform, should definitely be the best choice. That is, one should construct the sublattice Λ' as generated by $\mathbf{B}' = [2\mathbf{b}_1, \dots, 2\mathbf{b}_k, \mathbf{b}_{k+1}, \dots, \mathbf{b}_n]$, which has index 2^k in Λ .

Algorithm 3.1. DUALFFT($\mathbf{B}, \mathbf{B}', \mathcal{W}, \mathbf{t}$): Dual-Sieve-FFT attack

Input: two bases $\mathbf{B}', \mathbf{B}$ for full-rank lattices $\Lambda' \subset \Lambda \subset \mathbb{R}^n$, a set of short dual vectors $\mathcal{W} \subset (\Lambda')^\vee$, and a target $\mathbf{t} \in \mathbb{R}^n/\Lambda'$ such that $\mathbf{B}' = \mathbf{B} \cdot \text{diag}(d_1, \dots, d_n)$ holds for certain $d_1, \dots, d_n \in \mathbb{Z}_{\geq 1}$

Output: the lattice coset $\mathbf{g} + \Lambda' \in \Lambda/\Lambda'$ closest to $\mathbf{t}$

- 1: initialise a table $\mathbf{T}$ of dimension $d_1 \times d_2 \times \dots \times d_n$ with zeros
 - 2: **for** $\mathbf{w} \in \mathcal{W}$ **do**
 - 3: compute $j_i \in \{0, 1, \dots, d_i - 1\}$ s.t. $\mathbf{w} \in \frac{j_1}{d_1} \mathbf{b}_1^\vee + \dots + \frac{j_n}{d_n} \mathbf{b}_n^\vee + \Lambda^\vee$
 - 4: $\mathbf{T}[j_1, j_2, \dots, j_n] \leftarrow \mathbf{T}[j_1, j_2, \dots, j_n] + \cos(2\pi \langle \mathbf{w}, \mathbf{t} \rangle)$
 - 5: $\mathbf{S} = \text{DFT}_{\mathbb{Z}/d_1\mathbb{Z} \times \dots \times \mathbb{Z}/d_n\mathbb{Z}}^{-1}(\mathbf{T})$
 - 6: $(j_1, j_2, \dots, j_n) \leftarrow \arg \max_{\substack{k_1, k_2, \dots, k_n \\ 0 \leq k_i < d_i}} \{\mathbf{S}[k_1, k_2, \dots, k_n]\}$
 - 7: **return** $j_1 \mathbf{b}_1 + \dots + j_n \mathbf{b}_n$
-

Randomised sparsification

Note that the analysis of the length of the vectors in $\mathcal{W}$ requires applying **Gaussian heuristic** to the densification of the dual, induced by the sparsification of the primal.

This might require care. Indeed, if the basis is well reduced before we apply the dual densification $\text{diag}(\frac{1}{d_1}, \dots, \frac{1}{d_n})$, this might create a dual lattice which is not random-looking; in particular it might contain a few vectors shorter than predicted by **Gaussian heuristic**, with an unclear impact on the rest of $\mathcal{W}$.

We do not expect this to be an issue if the basis $\mathbf{B}$ is adequately randomised before constructing the sparsification.

3.4.3 Advantages

Not only is it theoretically more satisfying to apply the FFT trick to the general decoding problem rather than to the specific **LWE** problem, it also makes recursion more straightforward.

Shorter dual vectors

The algorithm of Guo and Johansson [GJ21] seems to perfectly fit this formalisation, where the basis $\mathbf{B}$ is the standard basis associated with the q -ary representation of the lattice, and $\mathbf{B}' = \mathbf{B} \cdot \text{diag}(\gamma, \dots, \gamma, 1, \dots, 1)$

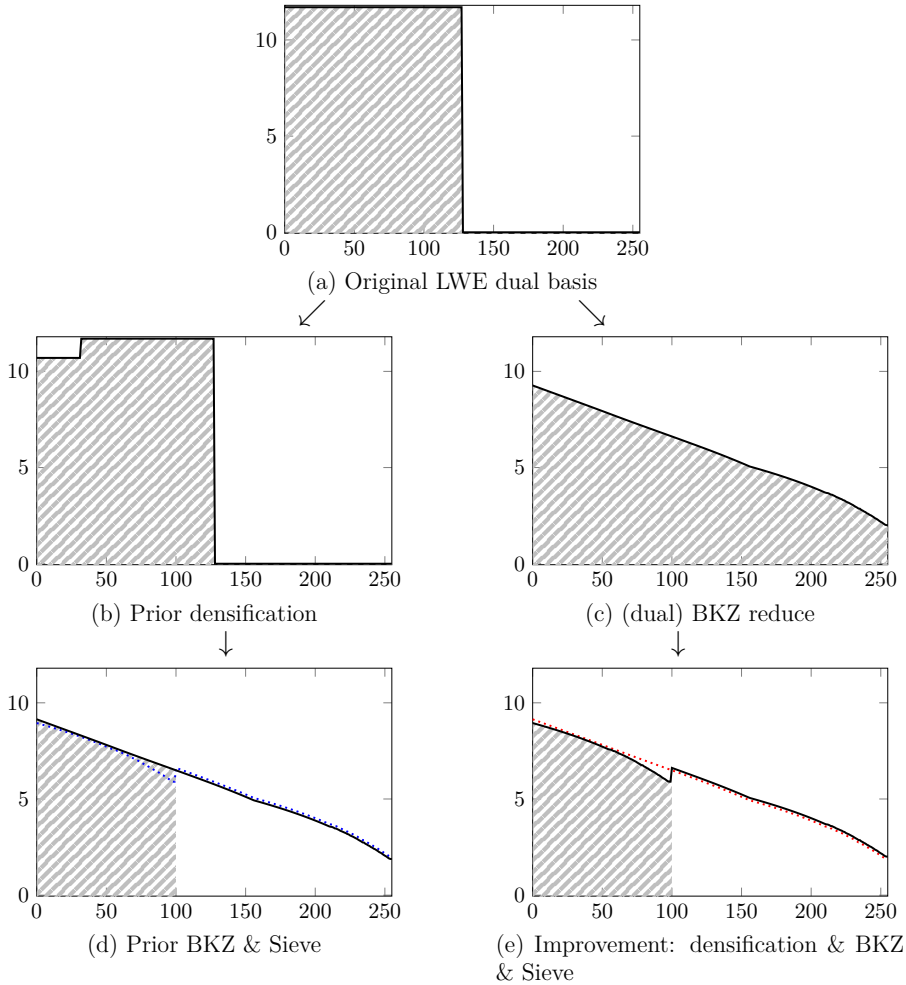


Figure 3.1: Simulation of dual basis profile, i.e., $(\log_2 \|\mathbf{d}_i^*\|)_{i=0}^{n-1}$ as a function of the basis index i using the **BKZ** 2.0 simulator [CN11]. Used parameters are $n = 128, \gamma = 2, k = 32, \beta_{\text{BKZ}} = \beta_{\text{sieve}} = 100, q = 3329$. Prior work follows the path (a) $\rightarrow$ (b) $\rightarrow$ (d), while the suggested improvement follows the path (a) $\rightarrow$ (c) $\rightarrow$ (e). The filled regions in (d) and (e) represent the log volume of the sublattice used for sieving, and the dotted line shows (e) and (d) respectively for comparison. Note that the densification in (e) shortens $(\mathbf{B})_{[\beta-k, \beta-1]}$ by a factor γ . See paragraph § Randomised sparsification (p. 99) regarding the lengths of dual vectors in $\mathcal{W}$.

with k many γ 's. The set of short dual vectors is obtained by first running **BKZ** reduction with block size β_{BKZ} on the dual of $\mathbf{B}'$, and then sieving in the sublattice generated by the first β_{sieve} vectors of this reduced dual basis. The impact of the sparsification $\mathbf{B}' = \mathbf{B} \cdot \text{diag}(\gamma, \dots, \gamma, 1, \dots, 1)$ on the length of the vectors in $\mathcal{W}$ is that these are shortened by a factor $\gamma^{k/n}$. That is, the sparsification has been *diluted* over n dimensions.

Instead, consider first applying dual-**BKZ** reduction with block size β_{BKZ} on $\mathbf{B}$, and then taking $\mathbf{B}' = \mathbf{B} \cdot \text{diag}(1, \dots, 1, \gamma, \dots, \gamma)$. In this way, the densification of the reversed dual basis is given by ${}^\vee(\mathbf{B}') = {}^\vee\mathbf{B} \cdot \text{diag}(\frac{1}{\gamma}, \dots, \frac{1}{\gamma}, 1, \dots, 1)$ remains concentrated on the k first dual vectors!¹ Therefore, the sparsification now makes the vectors in $\mathcal{W}$ shorter by a factor $\gamma^{k/\beta_{\text{sieve}}}$ (assuming $k \leq \beta_{\text{sieve}}$).

We leave the concrete exploitation of this improvement to the Dual-Sieve-FFT attack as future work, because the next section will invalidate the analysis of [GJ21, MAT22], so we first need a fixed analysis before thinking about this possible improvement. A visualisation of this improvement can be found in Figure 3.1.

Remark 3.9. The algorithm of **MATZOV** [MAT22] differs quite a bit from that of Guo–Johansson [GJ21] by resorting to a modulus switching technique, and it is claimed that this technique allows to decrease dimension at the cost of some extra error. We remain rather circumspect regarding a perceived superiority of the variant of **MATZOV** [MAT22] over that of Guo–Johansson [GJ21]. For more details, see [DP23a, App. A.6].

3.5 Contradictions to the independent score heuristic

In this section, we consider two regimes, in which we show that the analyses of [LW21, GJ21, MAT22] give rise to absurd conclusions when using the **independent score heuristic**. Both regimes are concerned with distinguishing between a **BDD** target and a uniform target.

In the first regime, we consider **BDD** samples where the expected distance to the lattice exceeds the **Gaussian heuristic**, a task that was recently proved statistically impossible in a random lattice [DADRT23].

The second regime is concerned with the case of finding a planted solution among many candidates. For certain parameters, the analysis

¹The warning on randomised sparsification from Section 3.4.2 still applies here; the basis randomisation should be applied to the block $\mathbf{B}_{[n-\beta_{\text{sieve}}+1, n]}$.

of [GJ21, MAT22] predicts a successful guess of the desired target, despite the existence of many other candidates that are even closer to the lattice than the planted solution. We would like to stress that this contradiction, regardless whether the FFT trick is used or not, merely arises from the large number of candidates. Phrased differently, the works of [GJ21, MAT22] predict a much lower failure probability on distinguishing between a BDD and uniform target than what one can reasonably assume.

Recall that in this section, as in [LW21], we implicitly renormalise the lattice Λ to have volume 1.

3.5.1 Distinguishing the indistinguishable

Set up

Recall that the main result of Laarhoven and Walter [LW21, Lem. 9], reformulated as Heuristic Claim 3.7, provides an algorithm to distinguish between a BDD instance at distance $r \cdot \text{GH}(n)$, and a uniform target modulo the lattice. After a precomputation depending solely on the lattice Λ , this algorithm has exponential complexity α^n , where α satisfies $e^2 \ln(\alpha) = \alpha^2 r^2$, and the heuristic analysis claims that it is successful with constant probability. That is, the algorithm outputs ‘BDD’ with constant probability if its input is a sampled BDD instance $\mathbf{t}$, and it outputs ‘uniform’ with constant probability if its input $\mathbf{t}$ is sampled uniformly modulo the lattice.²

Lemma 3.10. *The equation $e^2 \ln(x) = x^2 r^2$ admits a real solution in x if and only if $r^2 \leq e/2$.*

Proof. The statement and its proof are illustrated in Figure 3.2. First, note that $x \mapsto e^2 \ln(x)$ is concave, while $x \mapsto x^2 r^2$ is convex for any $r \in \mathbb{R}$. We discuss three cases.

Case 1: $r^2 = e/2$.

The parabola $y = r^2 x^2$ intersects tangentially the curve $y = e^2 \ln x$

²Note that this algorithm gets only the sample $\mathbf{t}$ as input, whereas Heuristic Claim 3.7 considers an algorithm that gets as input $\mathbf{t}_{\text{BDD}}$ and $\mathbf{t}_{\text{unif}}$ without knowing which is which. By saying ‘BDD’ if the score for $\mathbf{t}$ is at least the threshold value $\frac{1}{2} \mathbb{E}[f_{\mathcal{W}}(\mathbf{t}_{\text{BDD}})]$ and ‘uniform’ otherwise, one derives a success probability of $\frac{1}{2} + \frac{1}{2} \text{erf}(\frac{1}{2}) \approx 0.76$.

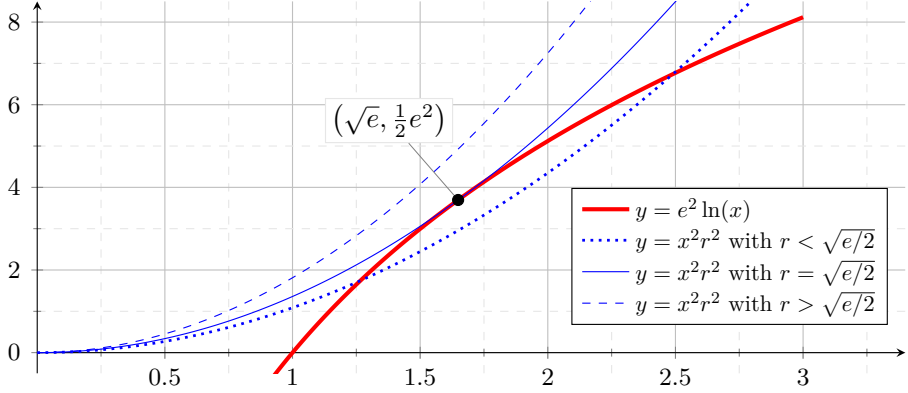


Figure 3.2: Equation $e^2 \ln(x) = x^2 r^2$ for various r

at $(x, y) = (\sqrt{e}, \frac{1}{2}e^2)$, with slope $\left. \frac{dy}{dx} \right|_{x=\sqrt{e}} = e^{3/2}$. By convexity, we have $e^2 \ln(x) < x^2 r^2$ for any $x \neq \sqrt{e}$.

Case 2: $r^2 > e/2$.

Note that $r^2 x^2$ is strictly increasing in r^2 for any x . Reusing the first case, we see $e^2 \ln(x) < x^2 r^2$ holds for all x , so there are no solutions in this case.

Case 3: $r^2 < e/2$.

We have $e^2 \ln(x) > x^2 r^2$ at $x = \sqrt{e}$. We also have $e^2 \ln(x) < x^2 r^2$ when $x = 1$. The Intermediate Value Theorem then tells there is a solution with $x \in (1, \sqrt{e})$. $\square$

This lemma suggests that the algorithm works successful supposedly for all $r < \sqrt{e/2} \approx 1.1658$. It is suspicious that the algorithm works beyond $r > 1$ because a uniform target often has a lattice point within a distance of $r \text{ GH}(n)$. Namely, every lattice Λ satisfies by Proposition 2.11,

$$\mathbb{E}_{\mathbf{t} \leftarrow \mathbf{U}(\mathbb{R}^n/\Lambda)} \left[\left| \left\{ \mathbf{v} \in \Lambda \mid \|\mathbf{v} - \mathbf{t}\| \leq r \text{ GH}(n) \right\} \right| \right] = r^n .$$

Thus, a nearby lattice point is guaranteed for a **BDD** target and likely for a uniform target, making distinguishing hard.

Still, one could imagine a scenario where with small probability a target has few close vectors, but most likely it will not, making distinguishing statistically possible.

The contradiction

It has been shown recently by Debris-Alazard et al. [DADRT23] that the above scenario does not occur with random lattices. More specifically, for a formally defined notion of random lattices and a fixed $r > 1$, it is proved that for errors from a uniform distribution on the ball of radius $r \text{GH}(n)$, the statistical distance between the error modulo the lattice and $U(\mathbb{R}^n/\Lambda)$ is exponentially small as a function of the dimension [DADRT23, Prop. 4.3].

That is, for all $r > 1$, no algorithm, whatever its complexity, can even succeed with probability greater than $\frac{1}{2} + O(1) \cdot r^{-n/2}$. Yet, [LW21, Lem. 9] (reformulated as Heuristic Claim 3.7) claims a constant success probability. ζ

Discussion

One could counter-argue that the claim [LW21, Lem. 9] is given for a uniform distribution on a sphere, a case not contradicted by Debris-Alazard et al. [DADRT23]. However, the actual analysis in [LW21] is done for a Gaussian distribution, a case which *is* also covered by Debris-Alazard et al. [DADRT23, Thm. 4.6].

The suspect heuristic

We note that this counter-argument applies only to the heuristic [LW21, Lem. 9], that is given a single sample, and not to [LW21, Lem. 8]. Indeed, in the context of the heuristic [LW21, Lem. 8] where exponentially many samples (either all uniform or all BDD) are given, the exponentially small statistical distance can be compensated for using numerous samples, as discussed between both Lemmata.

We note in particular that [LW21, Lem. 8] does not require the **independent score heuristic**, as it uses only one dual vector. In fact, after close inspection of the reasoning behind [LW21, Lem. 8], we could not identify any step that should be too hard to make formally provable, up to minor conditions and small losses in the concrete efficiency of the distinguisher. Indeed, with enough effort, it appears that all the other heuristics and approximations could be dealt with formally.

This sets the **independent score heuristic** as the prime suspect leading to the erroneous [LW21, Lem. 9].

3.5.2 Candidates closer than the solution (asymptotic)

Set up

Recall that the key claim of [GJ21] and [MAT22], reconstructed as Heuristic Claim 3.8 considers the case where the set of dual vectors comes from a lattice sieve [NV08, MV10, BDGL16], i.e., it consists of $N = (4/3)^{n/2}$ dual vectors of length $\ell = \sqrt{4/3} \cdot \text{GH}(n)$. Given one uniform sample and one BDD sample with an error sampled from a Gaussian distribution of parameter σ , it was claimed that one can distinguish with a failure probability that is exponentially small in $N\varepsilon^2$, whenever $N \geq 1/\varepsilon^2$, where $\varepsilon = \exp(-2\pi^2\sigma^2\ell^2)$. Let us consider the constraint the other way around by asking the following question.

Question 3.11. How large can one take σ to have a failure probability of at most p_{fail} ?

It can be seen that the failure probability in Heuristic Claim 3.8 is at most p_{fail} when $\varepsilon \geq \sqrt{\frac{\ln(1/p_{\text{fail}}^2)}{N}}$. This constrains σ to satisfy

$$\sigma = \sqrt{\frac{\ln(1/\varepsilon)}{2\pi^2\ell^2}} \leq \frac{\sqrt{\ln N - \ln \ln(1/p_{\text{fail}}^2)}}{2\pi\ell} = \frac{\sqrt{\frac{n}{2} \ln \frac{4}{3} - \ln \ln(1/p_{\text{fail}}^2)}}{2\pi\sqrt{\frac{4}{3}} \cdot \text{GH}(n)}. \quad (3.5)$$

With the approximation $\text{GH}(n) \approx \sqrt{\frac{n}{2\pi e}}$, one then arrives at $\sigma \leq \sqrt{C - \frac{C' \ln \ln(1/p_{\text{fail}}^2)}{n}}$ for some constants $C = \frac{3e \ln(4/3)}{16\pi} \approx 0.047$ and $C' = \frac{3e}{8\pi} \approx 0.32$. This means that Heuristic Claim 3.8 supposedly still distinguishes a BDD sample with an expected distance of $\sqrt{C \cdot n} \approx 0.89 \text{GH}(n)$ from a uniform sample with a failure probability that is *double-exponentially small*: $p_{\text{fail}} = \exp(-\frac{1}{2} \exp(n^{.99}))$.

The contradiction

We will show that the above heuristic claim leads to a contradiction, already for a single exponentially small failure probability of $p_{\text{fail}} = (0.48)^n$.

Lemma 3.12. *Let Λ be a unit-volume lattice, and $r > 0$ such that $r < \frac{\lambda_1(\Lambda)}{2 \text{GH}(n)}$. Then, for a target $\mathbf{t}$ uniform in $\mathbb{R}^n/\Lambda$, it holds with probability r^n that $\mathbf{t}$ is at distance at most $r \text{GH}(n)$ from the lattice.*

Proof. Note that the volume of the ball of radius $r \cdot \text{GH}(n)$ is exactly r^n by definition of $\text{GH}(n)$. Furthermore, because $r \cdot \text{GH}(n) < \lambda_1(\Lambda)/2$, all translations of this ball by points in Λ are disjoint. Said otherwise, this ball does not intersect itself modulo the lattice. More formally, its projection onto the torus $\mathbb{R}^n/\Lambda$ is injective. Hence, the ball modulo the lattice also has volume r^n in $\mathbb{R}^n/\Lambda$. The probability of a uniform target $\mathbf{t}$ falling into that ball is therefore $r^n / \text{Vol}_n(\mathbb{R}^n/\Lambda) = r^n / \det(\Lambda) = r^n$. $\square$

Let us use this lemma in the case of a random unit-volume lattice, or more specifically, one where we expect $\lambda_1(\Lambda) \approx \text{GH}(n)$. Using Lemma 3.12 with $r = 0.49$, the probability that $\mathbf{t}_{\text{unif}}$ lies in the ball of radius $r \cdot \text{GH}(n)$ is r^n . Assume this event happens. In this case, the uniform target lies at a distance $r \cdot \text{GH}(n)$ from the lattice, which is much shorter than the distance of $0.89 \cdot \text{GH}(n)$ the **BDD** sample was away from the lattice. However, the score function $f_{\mathcal{W}}$ is precisely meant to associate a larger score to closer targets, so in this case we expect to have $f_{\mathcal{W}}(\mathbf{t}_{\text{unif}}) > f_{\mathcal{W}}(\mathbf{t}_{\text{BDD}})$, with at least a constant probability. Thus, in fact the failure probability of distinguishing is close to 0.49^n , which is much more likely than the claimed 0.48^n . $\nmid$

Discussion

One might counter-argue that $f_{\mathcal{W}}$ only probabilistically classifies vectors by their distance to the lattice and might somehow still give the particularly close uniform sample a lower score than the **BDD** sample. However, the probability that the uniform target $\mathbf{t}$ lies at distance $\text{GH}(n)/n^2 = \Theta(n^{-3/2})$ away from the lattice, is n^{-2n} , which is (only) super-exponentially small. In this case we have $\langle \mathbf{t}, \mathbf{w} \rangle \leq O(1/n)$ for any $\mathbf{w} \in \mathcal{W}$ output by a sieve, and approximating the cosine we know the score $f_{\mathcal{W}}(\mathbf{t})$ for this target $\mathbf{t}$ will be at least $N(1 - O(1/n^2))$, which is essentially maximal.

The suspect heuristic

The discussion above points to the same suspect as Section 3.5.1, namely, the **independent score heuristic**. Indeed, under independence the probability of one uniform target reaching a constant fraction of the maximal score N should decrease as fast as $\exp(-\Theta(N))$, but we have shown that this probability is in fact at least $\exp(-\Theta(n \log n))$. Independence cannot

hold when N grows asymptotically at least as fast as $\omega(n \log n)$, and this should be visible in the tail of the score distribution for uniform targets.

3.5.3 Candidates closer than the solution (concrete)

Set up

In the contradiction above, we choose p_{fail} as small as 0.48^n to be able to invoke Lemma 3.12 to have the uniform sample at distance $r \text{ GH}(n) < \frac{1}{2} \lambda_1(\Lambda)$, quantifying the probability that a random target in $\mathbb{R}^n/\Lambda$ falls close to the lattice. This leads to an *extreme contradiction*: we analysed the probability that the uniform target $\mathbf{t}_{\text{unif}}$ is at a distance $0.49 \text{ GH}(n)$ from the lattice, much closer than $\mathbf{t}_{\text{BDD}}$ at distance $0.89 \text{ GH}(n)$. However, even when there is a uniform sample at distance e.g. $0.88 \text{ GH}(n)$ from the lattice, $\mathbf{t}_{\text{unif}}$ can get a higher score than the $\mathbf{t}_{\text{BDD}}$.

To extend Lemma 3.12 up to radii $r < 1$, we will resort to a heuristic instead. In this regime, translations of the ball by lattice points may start to intersect. In practice, however, the volume of this intersection remains rather small and should not affect the volume so much.

Heuristic Claim 3.13. *Let Λ be a random unit-volume lattice, and $r \in (0, 1)$. For a target $\mathbf{t}$ sampled uniformly in the torus $\mathbb{R}^n/\Lambda$, we have a probability of $r^n \cdot (1 - n^{O(1)} r^n)$ that $\mathbf{t}$ is at distance at most $r \text{ GH}(n)$ from the lattice.*

Heuristic Justification. Contrary to the proof of Lemma 3.12, we now have to subtract from r^n the probability that a target $\mathbf{t}$ has at least two lattice points at distance less than $r \text{ GH}(n)$ away, which by Heuristic Claim 2.48 happens with probability $O(n\sqrt{n})r^{2n}$. $\triangle$

The contradiction

The idea is to arrive at a contradiction whenever we instantiate the claim from above with the largest possible $r > 0$ up to which it is more likely to have a uniform target within $r \text{ GH}(n)$ away from the lattice than the heuristically claimed failure probability p_{fail} .

For a given dimension n and relative distance $r \in (0, 1)$, Heuristic Claim 3.13 lower bounds the probability that we fail to distinguish a BDD sample from a uniform sample. Indeed, when the uniform sample happens to be closer to the lattice than the BDD sample, with constant

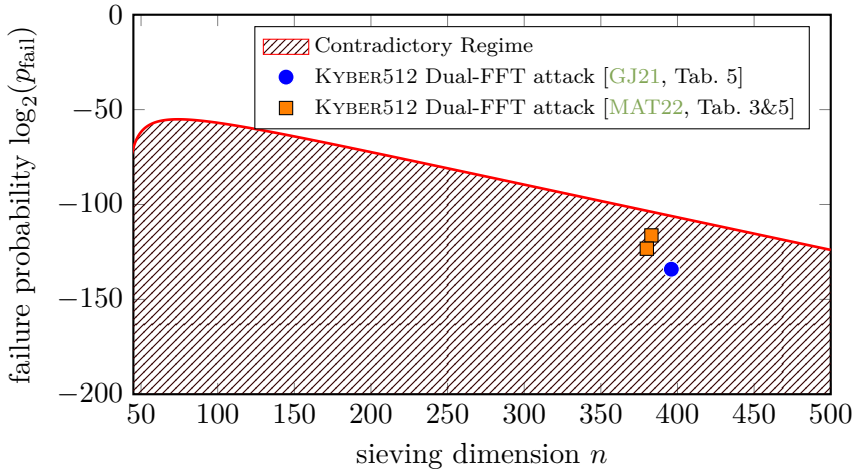


Figure 3.3: The concrete contradictory regime, in which Heuristic Claim 3.8 *overpredicts* the success probability of distinguishing in dimension n , compared to a *simple* lower bound derived from Heuristic Claim 3.13. By determining the largest $\sigma \in [0, \text{GH}(n)/\sqrt{n}]$ for which the contradiction arises, we know the upper boundary of the contradictory regime by computing the probability for that σ . For any pair (n, p_{fail}) falling in the concrete contradictory regime, a distinguisher cannot distinguish, with failure probability at most p_{fail} , the uniform distribution from a **BDD** distribution having σ satisfy Equation (3.5). Data was obtained with the script `volumetric_contr.py`.

probability a higher score will be assigned to it. Now by identifying all the $r \in (0, 1)$ for which the lower bound is *bigger* than the upper bound from Heuristic Claim 3.8, we have found a regime in which the two heuristics are in clear contradiction with each other. We will call this regime the ‘*concrete contradictory regime*’, and it is depicted in Figure 3.3. ζ

Contradictory regime in the context of concrete attacks against LWE

Above, we have determined the contradictory regime when obtaining the set $\mathcal{W}$ by a sieve over the full lattice, and assuming its volume was 1. It is not hard to see that scaling the lattice up or down is not going to affect the conclusion. Indeed, the **Gaussian heuristic** of the primal, σ and r will scale with the lattice, while the length ℓ of the dual vectors

will scale inversely; this leaves ε and p unaffected.

In the context of the cryptanalytic literature [EJK20, GJ21, MAT22], the set $\mathcal{W}$ does not come from the full dual lattice $\Lambda_{\text{LWE}}^\vee \subset \mathbb{R}^n$, but comes from a dual lattice $L \subset \Lambda_{\text{LWE}}^\vee$ of rank β_{sieve} , by Remark 3.2. Now, targets are projected onto the β_{sieve} -dimensional space spanned by all $\mathcal{W}$. Hence, we effectively run the distinguisher here over a projected sublattice of dimension β_{sieve} .

In this scenario, the contradictory regime is solely determined by β_{sieve} and the needed failure probability p_{fail} , and not by other quantities such as the LWE parameters and β_{BKZ} . Indeed, the LWE parameters and β_{BKZ} are going to influence the volume of the lattice on which we run the final sieve to obtain $\mathcal{W}$.

This might not perfectly be representative of the exact analysis of MATZOV [MAT22] in that we do not make a special analysis of the modulus switching effect on the score distribution. Instead, this treats modulus switching as adding an implicit error, increasing σ . This remains a strong signal on the credibility of the heuristic analysis in that regime.

Another point raising discussion is the fact that our contradiction is established in the case where the uniform targets $\mathbf{t}$ are independent. This is not formally the case when those targets come from a partial enumeration, though such a heuristic has been used in the past, for example underlying the analysis of the hybrid attack [How07]. More critically, we see no mention of such dependence and how they would affect the algorithm in the existing analysis [EJK20, GJ21, MAT22]. While we do not claim that it is impossible, we view the notion that such dependencies could fix the algorithm as quite doubtful, and requiring specific substantiation with analysis and experiments.

In other words, while our contradiction does not formally disprove the recent claims on the Dual-Sieve attack [EJK20, GJ21, MAT22], it does invalidate the reasoning leading to these claims. Because we do not see an obvious reason why this or that detail would solve the issues raised here, it seems reasonable to presume these claims are incorrect.

The parameters of Guo–Johansson and MATZOV

We now turn to two concrete instantiations of the dual attack [GJ21, MAT22] against the KYBER512 [SAB⁺22] parameter set with a focus towards the ‘asymptotic model’ [MAT22, Assumption 7.2]. This ‘asymptotic model’ is an optimistic estimate of the number of dimensions for

free [Duc18], whereas the other used ‘G6K model’ [GJ21, MAT22] is debated in [DP23a, App. A.2].

In [GJ21, Table 5], we find a sieving dimension $\beta_{\text{sieve}} = 396$ (where all the dual vectors come from), a guessing dimension $t_1 = 20$, and an FFT dimension $t = 78$. The guessing part considers all 7 possible values $\{-3, \dots, 3\}$ of each coordinate, while the FFT is done with $\gamma = 2$, giving rise to $T = 7^{20} \cdot 2^{78} \approx 2^{134.1}$ targets of which all are uniform samples except for one.

In [MAT22, Table 3], we find a sieving dimension $\beta_{\text{sieve}} = 380$ (where all the dual vectors come from), a guessing dimension $k_{\text{enum}} = 19$, and an FFT dimension $k_{\text{fft}} = 34$. The guessing part enumerates over $\{-3, \dots, 3\}^{k_{\text{enum}}}$ in order of decreasing probability from the used binomial distribution, while the FFT is done with $p = 5$, giving rise to, according to [MAT22], $T = 2^{19 \cdot H(\chi_s)} \cdot 5^{34} \approx 2^{123.3}$ many targets. Using an improved cost metric, they also give [MAT22, Table 5] another set of parameter where $\beta_{\text{sieve}} = 383$ and $T = 2^{17 \cdot H(\chi_s)} \cdot 5^{33} \approx 2^{116.3}$.³

To get a constant success probability in the dual attack, the distinguisher needs to have a failure probability p_{fail} of the order $1/T$: taking $p_{\text{fail}} = 1/T$ yields a total success probability of $(1 - p_{\text{fail}})^T \geq e^{-p_{\text{fail}}T} = 1/e$ of distinguishing one BDD sample from T (independent) uniform samples.

For both instantiations [GJ21, MAT22] the dual attack is used rather deep in its contradictory regime, as depicted in Figure 3.3. ⚡

3.6 Experiments

In this section, we provide further substantiation of the concrete contradiction (Sec. 3.5.3) with experimental evidence. We hope that the experiments will provide insight on what exactly is the issue with the prior analysis, and will show how the issue with the analysis can be resolved. In our analysis, we focus on the case where $\mathcal{W}$ is the output of a full sieve with a saturation radius of $r_{\text{sat}} = \sqrt{4/3}$ and a saturation ratio of $f_{\text{sat}} = 0.90$.

We look at two distributions: the score for uniform targets, and the score for BDD targets with a Gaussian error. There are two plausible

³We are not quite sure how this quantity was derived, and it seems incorrect. See [DP23a, App. A.5].

diagnoses for how the contradictory regime appears: the **BDD** scores are smaller than predicted, or the uniform scores are higher than predicted.

Because the concrete contradictory regime of Figure 3.3 starts when p_{fail} is very small, like 2^{-55} in dimension 70, these unpredicted high scores might be very rare. However, these events are important to determine the tail of the uniform score distribution, so the experiments require numerous samples. Naïvely, it would take a long time to run such large scale experiments, but the same FFT trick from Section 3.4 makes it feasible to run experiments on this scale!

3.6.1 Implementation details

We used the **G6K** software [ADH⁺19] for running the experiments, using Python on a high-level with a binding to some C code for computing the **WHT**. For the uniform targets we wrote the script `unif_scores.py`, which computes scores for many points sampled uniformly from $(\mathbb{Z}/q\mathbb{Z})^n$, where $\mathcal{W}$ is the output of a full sieve on the dual of

$$\mathbf{B} \cdot \text{diag}(2, \dots, 2, 1, \dots, 1),$$

with the number of twos equal to the FFT dimension. This setup allowed us to get roughly 2^{25} samples per second per **CPU** core. The scores were stored in buckets of width 1 while the exceptionally high scores were kept in a list. Here, we sieve using the built-in dual mode of **G6K**, which works only implicitly with the dual basis, see <https://github.com/fplll/fplll/wiki/FPLLL-Days-5-Summary>.

In addition, **BDD** scores are obtained using the script `bdd_scores.py`. The script randomly samples a q -ary lattice for the dual lattice $\Lambda^\vee$, so the script did not use the dual mode in **G6K**, making the implementation a little bit easier. Then, it computed the score function for the **BDD** distribution being a Gaussian distribution of parameter $\sigma = f_{\text{GH}} \cdot \text{GH}(n) / \sqrt{nq}$ for some parameter $f_{\text{GH}} \in (0, 1]$, where $1/\sqrt{q}$ is a normalisation factor needed because the dual lattice has determinant $q^{n/2}$.

For the uniform scores, we extracted the $2N = f_{\text{sat}} r_{\text{sat}}^n$ shortest dual vectors from the database. Because **G6K** returns at most one out of $\{\mathbf{w}, -\mathbf{w}\}$, **G6K** provides only N dual vectors of length below $r_{\text{sat}} \cdot \text{GH}(n)$. The other N dual vectors that were extracted, are therefore slightly larger than $r_{\text{sat}} \cdot \text{GH}(n)$, and the whole database is rather concentrated around this length. This does *not* affect our conclusions, as the prediction under

the heuristic analysis plotted in Figure 3.4 does not depend on the length of dual vectors.

Remark 3.14. This halving of the number of short vectors, which is also in Heuristic 2.114, was missed in the prior works of [GJ21, MAT22] leading to some irrelevant complication in the analysis of [MAT22], as discussed in [DP23a, App. A.5].

3.6.2 Uniform targets

We measured the score distribution for uniform targets over lattices of various dimension, and plotted our result in Figure 3.4. On each of these curves, we see a clear deviation from prediction for rare events: large scores are more likely to occur than predicted. After following a *waterfall* shape, i.e., a quadratic decay in logarithmic scale, the score probability seems to reach a *floor*, where it decays much slower than a normal distribution predicts. This is perfectly in accordance with the contradiction discussed in Sections 3.5.2 and 3.5.3: we start encountering vectors that are quite close to the lattice, which should have a rather high score.

Conclusion

Comparing Figure 3.3 with the floors appearing in Figure 3.4, we conclude that the floor seems to appear earlier than expected by the contradictory regime, vindicating the notion that the analysis might fail even in an earlier regime than our predicted contradiction. As we will see in the next section, the floor behaviour is not the only thing that the *independent score heuristic* predicts incorrectly.

3.6.3 BDD targets

We measured the score distribution for *BDD* targets sampled from a Gaussian distribution of parameter $\sigma = 0.7 \cdot \text{GH}(n)/\sqrt{n}$, over lattices of various dimensions n , and plotted our result in Figure 3.5. The predicted score distribution is based on the *independent score heuristic*, but also takes into account the exact lengths of each dual vector in Equation (3.2), instead of approximating all the lengths to be equal to $\sqrt{4/3} \cdot \text{GH}(n)$, to make the prediction more accurate.

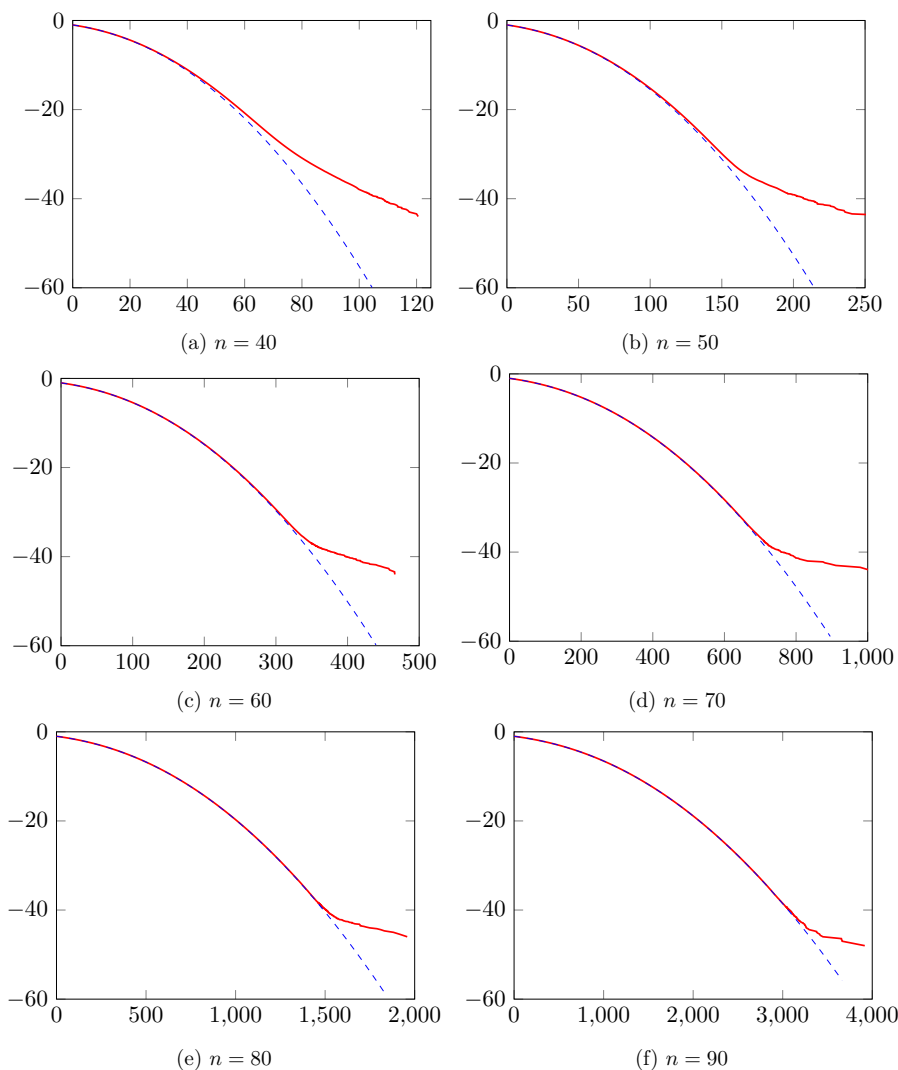


Figure 3.4: Binary logarithm of survival function (SF) of the score distribution for uniform targets in various dimensions n , according to prediction from the Heuristic Analysis (dashed blue line), and experiments (red line). Each experimental curve has 2^{45} samples, which were obtained with `code/unif_scores.py` and are listed in `data/unif_scores_nX.csv` of the auxiliary files.

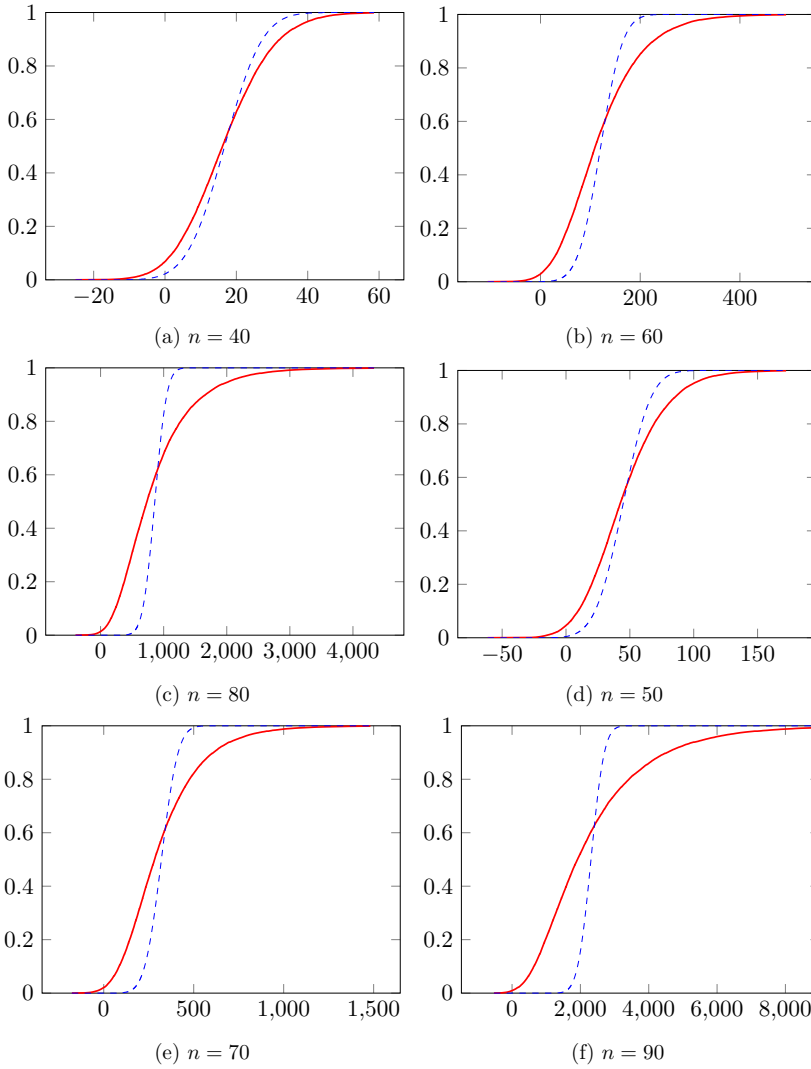


Figure 3.5: CDF of the score for Gaussian BDD targets in various dimensions n , with parameter $\sigma = 0.7 \cdot \text{GH}(n)/\sqrt{n}$, according to prediction from the Heuristic Analysis (dashed blue line), and experiments (red line). Each experimental curve has 2^{15} samples, which were obtained with `code/bdd_scores.py` and are listed in `data/bdd_scores_nX.csv` of the auxiliary files.

Table 3.1: Experimental BDD score distribution versus prediction

n	40	50	60	70	80	90
pred. std. dev. σ	8.31	17.19	35.41	72.80	149.52	307.01
meas. std. dev. σ	11.86	30.24	80.19	221.03	614.16	1732.72
ratio pred./meas.	1.43	1.76	2.26	3.04	4.11	5.64
pred. mean μ	16.71	44.68	120.26	320.53	859.99	2310.28
meas. mean μ	16.75	45.32	120.60	322.51	863.15	2325.99
ratio pred./meas.	1.00	1.01	1.00	1.01	1.00	1.01
meas. mean μ	16.75	45.32	120.60	322.51	863.15	2325.99
meas. median	16.13	42.27	108.74	282.54	735.47	1914.36
ratio med./mean	0.96	0.93	0.90	0.88	0.85	0.82

The first thing one notices is that the distribution is significantly more spread out than predicted. The variance is significantly higher. In fact, as shown in Table 3.1, the ratio between the actual and predicted standard deviation, i.e., square root of variance, appears to grow drastically with the dimension.

One might also want to consider the average. However, since the average is linear, its prediction does not require the *independent score heuristic*, and the prediction is indeed close to the measured average.

A more interesting statistic is the median. According to the heuristic analyses of [LW21, MAT22] reconstructed in Section 3.3, the median is predicted to be equal to the average. In practice, however, the median is noticeably lower. This partially implies that the distribution is quite asymmetric around its average, contrary to the analysis' prediction.

Conclusion

All in all, it is fair to say that the heuristic analyses of [LW21, MAT22], reconstructed in Section 3.3 of the dual attack is completely off when it comes to predicting the score for BDD targets, with or without modulus switching. The distribution is definitely *not* Gaussian, nor even symmetric around its average, and its variance is hugely underestimated.

3.7 Conclusion

Our theoretical contradictions in Section 3.5 and the experiments in Section 3.6 both demonstrate that the **independent score heuristic**, which has been used abundantly recently to analyse the Dual-Sieve attack, is invalid. Two independent phenomena uncovered by the experiments indicate that the success probability of the Dual-Sieve attack (with or without **FFT**) is presumably overestimated by this **independent score heuristic**, at least in certain regimes of interest.

In particular, the concrete cryptanalytic claims of numerous works, including [ADPS16, EJK20, GJ21, MAT22, CST22, AS23], should be considered at least unsubstantiated, as these are currently based on this flawed heuristic. Still, some of those claims might not be that far from reality, but those of [GJ21, MAT22, CST22, AS23] are so deep in the contradictory regime that these are presumably significantly far away from reality.

After mentioning some relevant pitfalls to avoid when fixing the analysis of the dual attack, we conclude this chapter with further improvements to the dual attack, which may help with overcoming the issues of independent score heuristic.

3.7.1 Possible pitfalls

Note that pulling the parameters of an attack outside the contradictory regime in Figure 3.3 is not a convincing way to prevent mispredictions while using the **independent score heuristic**. Indeed, the contradictory regime shows a fundamental flaw with this heuristic, and experiments indicate that there may well be an impact outside the regime. That is, a sound model needs to be used with theoretically justified predictions matching experimental behaviour measured in Figures 3.4 and 3.5 and beyond.

Moreover, we warn against a flawed argument for a fix, that would consist of constructing the set of dual vectors $\mathcal{W} = \mathcal{W}_1 \cup \mathcal{W}_2$ from two **BKZ** reductions and sieves instead of just one, yielding the score function $f_{\mathcal{W}}(\mathbf{t}) = f_{\mathcal{W}_1}(\mathbf{t}) + f_{\mathcal{W}_2}(\mathbf{t})$. One might argue that dual distinguishing now happens in a lattice of dimension 2β instead of β , which would push the attack possibly far outside the contradictory regime of Figure 3.3. By considering the experiments in Figure 3.4, however, one can see this argument is invalid. Each score distribution $f_{\mathcal{W}_i}$ is going to hit its floor,

and thus $f_{\mathcal{W}}$ will also have a floor starting from essentially the same score. Indeed, it suffices to hit the floor of *at least one* of the functions to hit the floor of the aggregate.

Instead, a more credible approach is the following. Run two (or a few) **BKZ** reductions and sieves, obtain two sets of short dual vectors $\mathcal{W}_1, \mathcal{W}_2$ and then consider the aggregate score as the *minimum* of both scores $f' = \min(f_{\mathcal{W}_1}, f_{\mathcal{W}_2})$ rather than its sum. Now to hit the floor of this new aggregate function f' , a uniform target would need to hit *both* floors simultaneously. If $f_{\mathcal{W}_1}, f_{\mathcal{W}_2}$ are sufficiently independent (an assumption that would need substantiation), such event is much rarer than hitting either floor. However, note that taking such a minimum aggregate of scores might also amplify the issues with low scores for **BDD** targets observed in Section 3.6.3. A robust model for all the distributions at hand is therefore still required.

Note this list of pitfalls is not comprehensive, and thus any potential fix of an attack, needs to be motivated at least by theoretical arguments or experimental validation.

3.7.2 Possible improvements

The experiments on the **independent score heuristic** show that current parametrisations of the attacks in [GJ21, MAT22, CST22, AS23] will result in many false positives. However, if this number of false positives is still reasonable, one might mitigate the issue at hand, although the attack cost will be somewhat higher. Indeed, one may consider the Dual-Sieve-**FFT** technique as a first ‘filtering stage’ in an attack with multiple stages: the remaining problem is still the problem of finding one **BDD** target among many candidates, but in an easier lattice, i.e., of lower dimension and/or sparser. If this remaining problem is sufficiently easier to accommodate all these targets, the Dual-Sieve-**FFT** attack might be salvaged, although the cost of such an attack will be higher than what was listed in [GJ21, MAT22, CST22, AS23]. Also note that first filtering and then filtering the ‘survivors’ again, is conceptually not that far off from the idea of taking the smallest of both scores.

Alternatively, one could potentially improve the dual attack further by applying different weights to the individual scores, to achieve better separation between the **BDD** and uniform distributions [AR04, LW21].