

IP Masking with Generic Security Guarantees under Minimum Assumptions, and Applications

Sebastian Faust¹, Loïc Masure², Elena Micheli¹,
Hai Hoang Nguyen³, Maximilian Ortl¹, François-Xavier Standaert⁴

¹ Technical University of Darmstadt, Darmstadt, Germany

² LIRMM, Univ. Montpellier, CNRS, France

³ ETH, Zurich, Switzerland

⁴ Crypto Group, ICTEAM Institute, UCLouvain, Louvain-la-Neuve, Belgium

Abstract. Leakage-resilient secret sharing is a fundamental building block for securing implementations against side-channel attacks. In general, such schemes correspond to a tradeoff between the complexity of the resulting masked implementations, their security guarantees and the physical assumptions they require to be effective. In this work, we revisit the Inner-Product (IP) framework, where a secret y is encoded by two vectors $(\vec{\omega}, \vec{y})$, such that their inner product is equal to y . So far, the state of the art is split in two. On the one hand, the most efficient IP masking schemes (in which $\vec{\omega}$ is public but random) are provably secure with the same security notions (*i.e.*, in the abstract *probing* model) as Boolean masking, yet at the cost of a slightly more expensive implementation. Hence, their theoretical interest and practical relevance remain unclear. On the other hand, the most secure IP masking schemes (in which $\vec{\omega}$ is secret) lead to expensive implementations. We improve this state of the art by investigating the leakage resilience of IP masking with public $\vec{\omega}$ coefficients in the *bounded leakage* model, which depicts well implementation contexts where the physical noise is negligible. Furthermore, we do that without assuming independent leakage from the shares, which may be challenging to enforce in practice. In this model, we show that if m bits are leaked from the d shares $\vec{y}$ of the encoding over an n -bit field, then, with probability at least $1 - 2^{-\lambda}$ over the choice of $\vec{\omega}$, the scheme is $\mathcal{O}\left(\sqrt{2^{-(d-1) \cdot n + m + 2\lambda}}\right)$ -leakage resilient. We additionally show that in large Mersenne-prime fields, a wise choice of the public coefficients $\vec{\omega}$ can yield leakage resilience up to $\mathcal{O}(n \cdot 2^{-d \cdot n + n + d})$, in the case where one physical bit from each share is revealed to the adversary. The exponential rate of the leakage resilience we put forward significantly improves upon previous bounds in additive masking, where the past literature exhibited a constant exponential rate only. We additionally discuss the applications of our results, and the new research challenges they raise.

1 Introduction

Masking is a popular countermeasure against side-channel attacks which works by performing computations over secret-shared data. Various types of encodings can be used for this purpose (*e.g.*, additive [43], multiplicative [42], affine [41], $\dots$). As a result, research on masked implementations can, at least theoretically, be

viewed as a quest towards encodings (and addition/multiplication algorithms) that provide the best trade-off between physical security and efficiency. Nevertheless, and somewhat biased by standardization and technological constraints, the simplicity of Boolean masking has so far largely dominated the practical scene for at least two reasons. On the one hand, most current algorithms for symmetric encryption operate in binary fields and have efficient bit-level representations. This was already the case of the AES Rijndael [28], and it has been amplified with the recently selected Ascon cipher [30]. Such algorithms naturally favor simple Boolean encodings which became a de facto standard for their masked implementations [47,45,26,46,21,56]. On the other hand, whether an operation is efficient in software depends on the instructions' sets. Here as well, readily available Boolean operations favor bit-level implementations [44].

Over the years, it however also appeared that the secure implementation of Boolean masking in software and hardware is a highly non-trivial task. Theoretically, this is because the security guarantees of Boolean masking, as for example formalized in [77,33,35], strongly rely on the assumptions that (i) the shares' leakages are (sufficiently) independent and (ii) the shares' leakages are (sufficiently) noisy. Concretely, this is because none of these assumptions is easy to fulfill. The independence condition is typically contradicted by physical defaults like *glitches* that can be observed in hardware implementations [65,66] and *transitions* that can be observed in software implementations [27,7]. Dealing with such physical defaults imposes additional constraints on masked implementations [72,39], which in turn implies additional overheads.¹ As for the noise condition, it is typically not fulfilled in lightweight embedded devices [11,20], and there are limited algorithmic approaches to remedy this lack so far.²

Interestingly, the sensitivity to the independence and noise assumptions is actually not disconnected from the choice of the encodings. As far as the independence condition is concerned, it has been observed heuristically that so-called Inner Product (IP) encodings, where the secrets are written as the inner product $\langle \vec{\omega}, \vec{y} \rangle$ between shares' vectors $\vec{\omega}$ and $\vec{y}$, lead to substantially lower transition leakages [5]. As far as the noise condition is concerned, there is a growing amount of research considering so-called "noise-free" leakages, where it is only required that the leakage function is non-injective. A seed result in this direction is the Asiacrypt 2011 one of Dziembowski and Faust [37], who showed that IP masking can provide strong guarantees in a liberal setting where the leakage function is generic (i.e., can correspond to arbitrary bounded functions), adaptive (i.e., chosen by the adversary) and global (i.e., can leak on many shares jointly). In other words, this paper provides security without strong independence nor noise assumptions. Unfortunately, and despite promising from the security viewpoint, such encodings lead to large performance overheads, limiting their concrete interest [6]. This is mostly due to the fact that the IP encodings of [37] are non-linear

¹ This is for example reflected in the above list of masked hardware implementations [47,45,26,46,21,56]. Similar observations hold in the software case.

² More precisely, solutions like shuffling the operations is sometimes considered, but are in fact only effective under some noise assumptions as well [50,83,82].

with respect to $\vec{\omega}$ and $\vec{y}$, which are both secret, implying that even linear operations (like the AES' MixColumns) require expensive gadgets. Balasch et al. tried to mitigate this issue by considering IP masking with public $\vec{\omega}$ coefficients [4,5]. This allows very efficient (share-wise) masked linear operations and non-linear operations that are only marginally more expensive than with additive encodings. However, their analysis is only in the abstract probing model of Ishai et al. [54]. So despite showing promising experimental patterns, such encodings do not offer better theoretical security guarantees than Boolean masking so far.

Another seed result is the TCC 2016 one of Dziembowski et al., who showed that additive encodings in prime fields can also provide security amplification for noise-free leakages. Such encodings allow efficient multiplications, but their security guarantees for so-called local leakages (i.e., again assuming independence) are mostly asymptotic and lead to hardly usable bounds in practice.

These seed results gave rise to a rich literature on leakage-resilient secret sharing [15,16,58,63,59,60,61,55,40,62]. We next detail the three most relevant examples for our discussions. First, the Crypto 2018 results of Benhamouda et al. consider the security of linear encodings under the independent leakage assumption and for generic bounded leakages [15,16]. Their security claims do not benefit from an increase of the field size, implying a large number of shares to reach high security levels. Next, the TCC 2022 results of Maji et al. consider the security of linear encodings without independence assumption, but for a very specific class of leakage functions [59]. For this class of leakage functions, they obtain security bounds as strong as the Asiacrypt 2011 ones of Dziembowski and Faust [37]. As soon as the leakage function escapes from this class, which happens already for bit leakages, the combination of global and adaptive leakages they consider implies that security collapses (see section 4.1). Finally, the Eurocrypt 2024 results of Faust et al. consider additive prime encodings, require the independence assumption, and specialize the Crypto 2018 results of Benhamouda et al. towards (bit leakages and) Hamming weight leakages [40].³ They obtain a security bound that scales with $\log(p)^d$, where p is the field modulus and d the number of shares. While this is a natural step given the popularity of such leakages in practice [64], it remains a specific result that holds under a strong (independence) assumption, raising the question whether it can be improved, generalized and hold under weaker assumptions. This state of the art on masking / leakage-resilient secret sharing encoding is summarized in Table 1, with N the number of traces required to perform a successful side-channel attack.

In this context, our first contribution is to answer positively to the three points in the above question. The main tweak we use for this purpose is a conceptual sweet spot regarding a combination of model and assumptions that enables strong theoretical results while being practically-relevant. Namely, we analyze IP encodings under **global** and **generic** (bounded) leakage, but require that the leakage function is **non-adaptive**. Precisely, we require that the coefficients $\vec{\omega}$ used by the IP masking scheme are picked up before specifying the leakage func-

³ For bit leakages, a similar result was already put forward by Maji et al. in [60], where the authors additionally show that such leakages are worst-case.

	leakage assumptions	security guarantees	encodings	coefficients	leakage function	security bounds	field
Asiacrypt 2011 [37]	global ($\perp$)	generic	IP	random/secret	adaptive	$N \geq \sqrt{ \mathbb{F} }^d$	arbitrary
Crypto 2018 [15,16]	local ($\perp$)	generic	linear	fixed/public	adaptive	$N \geq (cst)^d$	
TCC 2022 [59]	global ($\perp$)	very specific (negligible subset)	linear	random/public	adaptive	$N \geq \sqrt{ \mathbb{F} }^d$	prime arbitrary
Eurocrypt 2024 [40]	local ($\perp$)	specific (bits) specific (HW)	additive	$\emptyset$	non-adaptive	$N \geq (\frac{\pi}{2})^d$ $N \geq \log(p)^d$	prime Mersenne prime
Contribution #1	global ($\perp$)	generic	IP	random/public	non-adaptive	$N \geq \sqrt{ \mathbb{F} }^d$	arbitrary
Contribution #2	local ($\perp$)	specific (bits)	IP	fixed/public	non-adaptive	$N \geq \mathbb{F} ^d$	Mersenne prime

Table 1: State-of-the-art masking / leakage-resilient secret sharing encodings with guarantees against noise-free leakage.

tion. This is a weak assumption, as it is widely acknowledged that the leakage function is primarily defined by the target device rather than by the adversary, and the mild adaptations that an adversary could perform with a measurement setup (*e.g.*, by moving a probe) are anyway not feasible at runtime [85]. With this tweak, we can obtain major improvements of the Eurocrypt 2024 security bounds, from $\log(p)^d$ for prime fields to $|\mathbb{F}|^{d/2}$ for any field, generalize them to any (field and) sufficiently bounded leakage function, and do not require the independence assumption. More precisely, for $|\mathbb{F}| \approx 2^n$ and for an m -bit global leakage on the secret coefficients $\vec{y}$, the number of attack traces is roughly bounded by $(2^n)^{\frac{d}{2} - \frac{m}{2n}}$. Say for example that the leakage per share is $\frac{m}{d} = \frac{n}{2}$ bits, then the bound is worth $2^{\frac{dn}{4}}$. For $\frac{m}{d} = \frac{n}{4}$ and $\frac{m}{d} = \frac{n}{8}$, it decreases to $2^{\frac{3dn}{8}}$ and $2^{\frac{7dn}{16}}$.

As a comparison, and assuming a $n = 32$ -bit software implementation leaking the Hamming weight of the shares (*i.e.*, $m \approx \log(32) = 5 \approx \frac{n}{6}$), this improves the Eurocrypt 2024 bound from 5^d to $2^{16d - \frac{5}{3}}$, reducing the number of shares to ensure security for 10 million traces from 10 to 2. Further assuming quadratic performance overheads for masking [54], this (roughly) corresponds to performance gains by a factor 25. So our results significantly amplify the interest of masking in larger fields. We refer to section 4 for a formal treatment.

In section 4.4, we additionally discuss how to choose the public coefficients $\vec{\omega}$ and the impact of keeping them constant. We also discuss the practicality of the bounded leakage assumption and how to relax it to other (more realistic) settings where the leakage is characterized based on information theoretic metrics.

Our second contribution, presented in section 5, investigates the complementary question of whether more specific assumptions on the leakage function could lead to significantly improved results. As a theoretical first step in this direction, we consider bit leakages and show that, by choosing the coefficients $\vec{\omega}$ of a (Mersenne prime) IP encoding accordingly, it is possible to improve the previous bound from $|\mathbb{F}|^{d/2}$ to p^d . Technically, this contribution is of independent interest, since it relies on different tools (*i.e.*, FFT analysis *vs.* the leftover hash lemma for the first contribution). The exponential rate of leakage resilience it puts forward also improves significantly upon previous bounds in additive masking, which have only constant exponential rate. Concretely, obtaining this result again requires the independence assumption. It leaves as an interesting open question whether fine-grain characterization for more realistic functions than bit leakages can lead to similar or better gains, in order to fully comprehend the genericity/simplicity *vs.* security tradeoff of leakage-resilient secret sharing.

2 Applications and Challenges

The careful reader probably noticed that security claims scaling as $N \geq \sqrt{|\mathbb{F}|}^d$ look too good to be true in general. The remaining caveat is that these security claims are so far for the encodings only. In other words, they do not consider what happens when computing. It turns out the vast majority of the leakage-resilient secret sharing literature is similarly restricted in this respect. This contrasts

with the situation of additive encodings for which shared computation is a thoroughly investigated topic. Algorithms for computing over linear or IP encodings are admittedly scarce: the most relevant references for our discussions are from Asiacrypt 2011 [37], Eurocrypt 2015 and Asiacrypt 2017 [4,5]. As summarized in Table 2, the first one considers (random and) secret coefficients $\vec{\omega}$: its guarantees are in the bounded leakage model but it leads to inefficient (linear and non-linear) operations. The latter ones considers (random and) public coefficients $\vec{\omega}$: their guarantees are only in the probing model and they lead to very efficient linear operations and moderately efficient non-linear operations.

	$\vec{\omega}$	linear operations		multiplications	
		leak. model	implem.	leak. model	implem.
Asiacrypt 2011 [37]	secret	bounded	inefficient	bounded	inefficient
Eurocrypt 2015 [4]	public	probing	very efficient	probing	efficient
Asiacrypt 2017 [5]					
This paper	public	bounded	very efficient	probing	efficient

Table 2: State-of-the-art algorithms for IP masked computations.

Unfortunately, in the liberal model that we consider in this paper, where we can leak jointly on the secret shares and the $\vec{\omega}$ coefficients are public, the leakage function cannot take these public coefficients as argument. That is because otherwise, the leakage function could itself compute inner products so that the strong security guarantees described above collapse to the weaker ones of additive encodings (see, again, section 4.1). Concretely, this means that our results are applicable to encodings and operations that do not manipulate the $\vec{\omega}$ coefficients, meaning linear operations only.⁴ We next discuss the practical applications of these results and the new research challenges they raise.

Nearly perfect applications. We first observe that there are contexts where our results apply quite directly. This is for example the case of fresh re-keying schemes based on key-homomorphic primitives, which aim to enable leveled implementations where only linear operations must be masked [69,68,31,38,36]. A bit more precisely, a fresh re-keying scheme uses a public random value r and multiplies it with a long-term key k to generate an ephemeral secret $k^* = r \cdot k$ to be used once in an encryption scheme. It can therefore be implemented share-wise as $k^* = r \cdot k_1 + r \cdot k_2 + \dots + r \cdot k_d$. This application is only nearly perfect, since the key shares still have to be refreshed before each execution, and such refreshing gadgets are for now only proven in the probing model. Yet, on the one hand, it is an interesting open problem to study whether refreshing gadgets could be proven in the bounded leakage model (e.g., by leveraging linear combinations of

⁴ This is trivially seen for additions and multiplications by constants, as for example needed for the NTT application below. By contrast, the square operation in [4,5] uses the $\vec{\omega}$ coefficients. It can however be implemented securely by updating these coefficients independently of the manipulation of the secret shares, as in [37].

zero encodings that could be computed without touching the $\vec{\omega}$ coefficients, or by restricting our model); and on the other hand, such an application already gives strong security guarantees against the best concrete attacks targeting fresh re-keying schemes like [13,12]. As a result, the integration of such re-keying schemes within a leakage-resistant authenticated encryption algorithm would be interesting [14,29,70]. Their application to emerging post-quantum encryption schemes that only require to mask linear operations would also be worthwhile [52].

A partial application. The most investigated context for masking is block ciphers. Since such algorithms combine many linear and non-linear operations, the application of our results is admittedly less direct in this case. Generic security guarantees under minimum assumptions hold for encodings appearing during the computations and for linear operations. Weaker (probing) security guarantees remain for the other (non-linear) operations. Despite less direct, such an application nevertheless reinforces the interest of IP masking for block ciphers, since it strengthens the security of a part of their computations. This is especially true given that the most powerful attacks against block cipher implementations usually have a hard time to exploit the leakage of secure multiplication gadgets. See [11,20] and [48], Figure 14. Besides, as already hinted with refresh gadgets for fresh re-keying schemes and probably most importantly, we show that a liberal leakage model allows strong security guarantees for linear operations. This does not preclude a mild restriction of this model to enable an extension of our results to non-linear operations, which is a great scope for further research.

A timely intermediate. Eventually, the discussion above suggests that the practical relevance of our results grows as the amount of linear operations in an algorithms increases and the security of different operations differs (i.e., the implementation is leveled). This typically suggests post-quantum cryptography algorithms, like the recently standardized CRYSTALS-Kyber [81] and CRYSTALS-Dilithium [57], as interesting applications. On the one hand, their secure implementations are indeed leveled in the sense that they rely on various (e.g., Boolean and arithmetic) encodings [71,17,19,3], which may in turn lead to different security guarantees [67]. On the other hand, these implementations require protecting the (linear) NTT against side-channel attacks, as witnessed by a large body of experimental work [76,74,1,49,2,51]. As a result, it is frequently suggested that the NTT must be masked, and sometimes shuffled on top [78]. Given the challenge of protecting post-quantum cryptography against leakage, we believe the application of our results to the NTT is an interesting direction. To make it applicable, we show in appendix A that Arithmetic to Boolean (A2B) and Boolean to Arithmetic (B2A) algorithms can be adapted to our IP encodings, which actually highly benefits from the public $\vec{\omega}$ coefficients.⁵ IP encodings could also be applied to the several (masked) matrix multiplications taking place in CRYSTALS-Kyber and -Dilithium, which are performed share-by-share.

⁵ As the refresh gadgets needed to apply IP masking to fresh re-keying schemes, these conversion algorithms are for now only proven secure in the probing model.

Overall, our conclusions are therefore twofold. On the one hand, we show that IP encodings offer strong security guarantees with minimum noise and independence requirements. Combined with the “security order amplification results” put forward in the presence of high noise [84,75,25], this gives strong incentive to consider the use of such encodings for practical applications. On the other hand, there remains a gap between the theory of leakage-resilient secret sharing, which mostly focuses on encodings so far, and the practice of side-channel countermeasures, where protecting (non-linear) computations is needed. Closing such gaps has been a rich source of research challenges and theoretical innovation and we hope our results stimulate more work in this exciting direction.

3 Background

3.1 Notation

For any set $\mathcal{X}$, denoted by calligraphic letters, we denote with $U_{\mathcal{X}}$ the uniform distribution over $\mathcal{X}$. $\mathbb{F}$ denotes a finite field, $\mathbb{F}^*$ denotes $\mathbb{F} \setminus \{0\}$, and $\mathbb{F}_p$ denotes the finite field of size p , where p is prime. We define $\gamma = \exp\left(i\frac{2\pi}{p}\right)$ as the p -th root of unity. For a vector denoted as $\vec{y} \in \mathbb{F}^d$, and for some vector $\vec{\omega} \in \mathbb{F}^d \setminus (0, \dots, 0)$, we denote by $\vec{y} + \text{Span}(\vec{\omega})^\perp$ the affine hyperplane with offset $\vec{y}$ and orthogonal vector $\vec{\omega}$. When the context does not raise any ambiguity, we may replace the vector offset $\vec{y}$ by a scalar offset y , meaning that the actual offset is $\vec{y} = (y, 0, \dots, 0)$.

3.2 Metrics

In this paper, we make use of the *total variation* distance $\text{TV}(\mathbf{p}; \mathbf{q})$ between two probability mass functions $\mathbf{p}, \mathbf{q}$ over the same support $\mathcal{X}$, defined as the quantity $\frac{1}{2} \sum_{x \in \mathcal{X}} |\mathbf{p}(x) - \mathbf{q}(x)|$. We also make use of other metrics, like the *min-entropy*.

Definition 1 (Min-Entropy). *Let Y be a random variable. The min-entropy of Y is defined as*

$$H_\infty(Y) = \min_{y \in \mathcal{Y}} \{-\log(\Pr(Y = y))\} = -\log \left(\max_{y \in \mathcal{Y}} \Pr(Y = y) \right).$$

We say that Y is a k -source if $H_\infty(Y) \geq k$.

Definition 2 (Average Min-entropy). *The average min-entropy of a random variable Y conditioned on a random variable L , denoted $\tilde{H}_\infty(Y \mid L)$, is defined as*

$$\tilde{H}_\infty(Y \mid L) = -\log_2 \left(\mathbb{E}_L \left[2^{-H_\infty(Y \mid L=\ell)} \right] \right).$$

3.3 Randomness Extractors

Definition 3. Let $\mathcal{X}, \mathcal{S}$ and $\mathcal{Y}$ be sets. Let $U_{\mathcal{S}}, U_{\mathcal{Y}}$ be the uniform random variables over $\mathcal{S}$ and $\mathcal{Y}$, respectively. We say that a function $\text{Ext} : \mathcal{X} \times \mathcal{S} \rightarrow \mathcal{Y}$ is a (k, ϵ) strong extractor if, for k -sources X over $\mathcal{X}$ such that $X, U_{\mathcal{S}}, U_{\mathcal{Y}}$ are independent, then

$$\text{TV}((\text{Ext}(X, U_{\mathcal{S}}), U_{\mathcal{S}}); (U_{\mathcal{Y}}, U_{\mathcal{S}})) \leq \epsilon.$$

Definition 4. A family $\mathcal{H}$ of hash functions from $\mathcal{X}$ to $\mathcal{Y}$ is called universal if, for every $x, x' \in \mathcal{X}$ with $x \neq x'$, then

$$\Pr_{h \leftarrow \mathcal{H}} [h(x) = h(x')] \leq \frac{1}{|\mathcal{Y}|}.$$

Theorem 1 (Leftover Hash Lemma [53]). Let X be a k -source over $\mathcal{X}$. Let $\mathcal{H} = \{h_{\omega}\}_{\omega \in \mathcal{W}}$ be a universal hash family with output space $\mathcal{Y}$. For every $x \in \mathcal{X}$ and $\omega \in \mathcal{W}$, define

$$\text{Ext}(x, \omega) = h_{\omega}(x).$$

Then Ext is a strong (k, ϵ) -extractor with $\epsilon = |\mathcal{Y}|^{\frac{1}{2}} \cdot 2^{-\frac{k}{2}-1}$.

For completeness, we include a proof of Theorem 1 in Appendix B.

3.4 Inner-Product Masking

Definition 5 (Inner-Product Masking). Let $\vec{\omega} \in (\mathbb{F}_p^*)^d$. An $\vec{\omega}$ -inner-product encoding of a secret y (or simply $\vec{\omega}$ -encoding for short) is a d -tuple $\llbracket y \rrbracket_{\vec{\omega}} = (y_1, \dots, y_d) \in \mathbb{F}_p^d$, such that

$$y = \langle \vec{\omega}, \llbracket y \rrbracket_{\vec{\omega}} \rangle = \sum_{i=1}^d \omega_i \cdot y_i,$$

i.e., such that $\llbracket y \rrbracket$ is in $y + \text{Span}(\vec{\omega})^{\perp}$, with d denoted as the masking order.

Remark 1. When the context does not raise any ambiguity, we may write $\llbracket y \rrbracket$ instead of $\llbracket y \rrbracket_{\vec{\omega}}$ for short. Moreover, as we are working over a finite field, we may assume wlog that $\omega_1 = 1$ — it suffices to divide any other entry of $\vec{\omega}$ by ω_1 .

3.5 The Leakage-Resilience Threat Model

Following the usual leakage-resilience threat model considered in the literature, an attacker is given access to some leakage function L , adversarially chosen among a given class of leakage [34]. More precisely, in a leakage-resilience context, we assume that for any leakage function of interest, there exists some (possibly randomized) function $\tau : \mathbb{F}^d \rightarrow \mathcal{L}$ such that for any $y \in \mathcal{Y}$:

$$L(y) = \tau(\text{Enc}_{\vec{\omega}}(y)) ,$$

where $\text{Enc}_{\vec{\omega}}(y)$ is a randomized function returning an $\vec{\omega}$ -encoding $\llbracket y \rrbracket_{\vec{\omega}}$ uniformly drawn over $y + \text{Span}(\vec{\omega})^\perp$. We focus on the bounded leakage model, meaning that τ is m -bounded, as defined below.

Definition 6 (Bounded Leakage). *A (possibly randomized) function $\tau : \mathbb{F}^d \rightarrow \mathcal{L}$ is m -bounded if $|\mathcal{L}| = 2^m$. We denote with Bnd_m the class of such functions.*

This threat model somewhat differs from the so-called *noisy leakage* model usually considered in the literature of provable masking [77], which characterizes the leakage with information theoretic metrics like the statistical distance. On the one hand, any m -bounded function can be seen as a particular case of noisy leakage,⁶ whereas the opposite is not always true [18, 73]. On the other hand, the noisy leakage model often requires that each share of the encoding — or each subsequent computation — leak independently from the others. Although commonly accepted in the literature, this so-called *local* leakage assumption is not always realistic, *e.g.*, due to some physical defaults [65, 72, 27, 7, 39, 23] or due to some correlated noise [24]. Therefore, relaxing this assumption is identified as a major challenge [79, Sec. 4.3.3, 4.6]. We make one step forward in this direction by removing this assumption in the particular case of bounded leakage.

In the leakage-resilient secret sharing literature, an $\vec{\omega}$ -encoding is said to be ϵ -leakage resilient against m -bounded leakage if $\sup_{\tau \in \text{Bnd}_m} \mathbf{M}_{\vec{\omega}}(\mathbf{L}) \leq \epsilon$, where

$$\mathbf{M}_{\vec{\omega}}(\mathbf{L}) = \sup_{y^{(0)}, y^{(1)} \in \mathbb{F}} \text{TV}\left(\Pr\left(\mathbf{L}(y^{(0)}) \mid \vec{\Omega} = \vec{\omega}\right); \Pr\left(\mathbf{L}(y^{(1)}) \mid \vec{\Omega} = \vec{\omega}\right)\right).$$

In a side-channel context though, most of the attacks in the state of the art use random plaintexts. This type of adversary is better reflected by a weaker metric, the so-called *statistical bias*:

$$\beta_{\vec{\omega}}(\mathbf{Y} \mid \mathbf{L}) = \mathbb{E}_{\mathbf{Y}} \left[\text{TV}\left(\Pr(\tau(\text{Enc}_{\vec{\omega}}(\mathbf{Y}))) ; \mathbb{E}_{\mathbf{Y}'} [\Pr(\tau(\text{Enc}_{\vec{\omega}}(\mathbf{Y}')))]\right) \right], \quad (1)$$

where $\mathbf{Y}'$ is an independent copy of $\mathbf{Y}$. Using Bayes' theorem, the statistical bias may also be rephrased as follows: $\beta_{\vec{\omega}}(\mathbf{Y} \mid \mathbf{L}) = \text{TV}\left(\mathbf{p}_{\mathbf{L}, \mathbf{Y}} \mid \vec{\Omega} = \vec{\omega}; \mathbf{p}_{\mathbf{L}} \mid \vec{\Omega} = \vec{\omega} \otimes \mathbf{p}_{\mathbf{Y}}\right)$.⁷ Unsurprisingly, $\mathbf{M}_{\vec{\omega}}(\mathbf{L})$ upper bounds $\beta_{\vec{\omega}}(\mathbf{Y} \mid \mathbf{L})$, where the secret is uniformly random. Choosing the secret $y \in \mathbb{F}$ can improve the bias by at most a factor $|\mathbb{F}|$. The next proposition (proven in appendix C) formalizes this intuition.

Proposition 1. *Let $\vec{\omega} \in (\mathbb{F}^*)^d$, and $\mathbf{L}$ be any random variable. Assume that $\mathbf{Y}$ is the uniform random variable over $\mathbb{F}$. Then*

$$\beta_{\vec{\omega}}(\mathbf{Y} \mid \mathbf{L}) \leq \mathbf{M}_{\vec{\omega}}(\mathbf{L}) \leq |\mathbb{F}| \cdot \beta_{\vec{\omega}}(\mathbf{Y} \mid \mathbf{L}).$$

Our goal is therefore to bound $\mathbf{M}_{\vec{\omega}}(\mathbf{L})$, either directly (in section 5), or through a reduction to the statistical bias $\beta_{\vec{\omega}}(\mathbf{Y} \mid \mathbf{L})$ (in section 4).

⁶ In a binary field, any m -bounded function is δ -noisy, for $\delta = 1 - 2^{-m}$. Note that the name “noisy leakage model” may be slightly misleading since a leakage function can be δ -noisy without physical noise if it is non-injective.

⁷ Since $\mathbf{Y}$ is uniform in the definition of the statistical bias, $\mathbf{Y}$ is independent of $\vec{\omega}$.

4 Inner-Product Masking with Randomized Coefficients

We start by describing our first result for generic bounded leakage.

4.1 The Dilemma of Adversarially-Chosen Global Leakage

In general, if an adversary can choose the leakage function depending on $\vec{\omega}$, then inner-product masking leads to the very same security as additive masking, as formalized by the following proposition and corollary, proven in appendix D.

Proposition 2. *Let $\vec{\omega}, \vec{\omega}' \in (\mathbb{F}^*)^d$. Then, for any m -bounded leakage function τ ,⁸ there exist some m -bounded leakage function τ' such that for every secret y and all observed leakage $\vec{\ell}$, we have*

$$\Pr(\tau(\text{Enc}_{\vec{\omega}}(y)) = \vec{\ell}) = \Pr(\tau'(\text{Enc}_{\vec{\omega}'}(y)) = \vec{\ell}) \quad .$$

Corollary 1. *Let $\vec{\omega} \in (\mathbb{F}^*)^d$ and $\epsilon \geq 0$. Then the $\vec{\omega}$ -encoding is ϵ -leakage resilient against m -bounded leakage if and only if the additive encoding is ϵ -leakage resilient against m -bounded leakage.*

Even worse, in our context where we do not assume the shares to leak independently from each other, we would allow the adversary to choose the function $\text{Dec}_{\vec{\omega}} : \vec{x} \mapsto \langle \vec{\omega}, \vec{x} \rangle$ as a leakage. $\text{Dec}_{\vec{\omega}}$ is m -bounded, for $m = \log |\mathbb{F}|$. Yet, by definition, $\text{Dec}_{\vec{\omega}}(\llbracket y \rrbracket_{\vec{\omega}}) = y$. In other words, leaking more than $\log |\mathbb{F}|$ bits in this model would lead to a successful attack, regardless of the number of shares d . Maji et al. circumvent this issue by restricting the adversary to a sub-class of m -bounded leakage functions, whose choice is left to the adversary, yet whose size is reflected in their security bound [59]. Hence, in order to get a non-trivial bound, the sub-class must have a negligible size compared to the class Bnd_m .

As argued in introduction of the paper, an adversarially-chosen leakage function is nevertheless unlikely to happen in practice, since the leakage function is usually more a property of the target device than a property of the adversary. So to circumvent this dilemma, we may simply assume that the adversary still selects the leakage function, yet without knowing the $\vec{\omega}$ coefficients in advance. Conceptually, this means now that $\vec{\omega}$ is drawn uniformly over $(\mathbb{F}^*)^d$ once the leakage function $\tau \in \text{Bnd}_m$ is chosen by the adversary. Observe then that the metrics $\beta_{\vec{\omega}}(Y | L)$ and $M_{\vec{\omega}}(L)$ are now random variables, whose randomness is taken over the uniform choice of $\vec{\omega} \in (\mathbb{F}^*)^d$. As a result, we want to upper bound them with a high probability independent of the choice of $\tau \in \text{Bnd}_m$.

This goal motivates the introduction of the following definition.

Definition 7 (Stochastic Leakage Resilience). *For $\delta > 0$, a d -share IP encoding is $(\delta, \epsilon_{m,d})$ -leakage resilient against m -bounded leakage if for any $\delta > 0$, it is $\epsilon_{m,d}$ -leakage resilient with probability at least $1 - \delta$, or equivalently, if*

$$\sup_{\tau \in \text{Bnd}_m} \Pr(M_{\vec{\omega}}(L) \geq \epsilon_{m,d}) \leq \delta \quad ,$$

⁸ This result also applies to δ -noisy leakage functions [34].

where the randomness is taken over the uniform choice of $\vec{\omega} \in (\mathbb{F}^*)^d$.

Concretely, we will rather try to upper-bound the expectation of the statistical bias, and then apply concentration inequalities (*e.g.*, Markov) and Proposition 1. This justifies the introduction of the following metric.

Definition 8 (Bias for Random Coefficients). *Let $Y, L, \vec{\Omega}$ be random variables. The statistical bias for random coefficients induced by a leakage L over the $\vec{\Omega}$ -encoding of a secret Y is defined by*

$$\beta_{\vec{\Omega}}(Y|L) = \mathbb{E}_{\vec{\omega}} [\beta_{\vec{\omega}}(Y | L)]. \quad (2)$$

4.2 Limitations of Previous Techniques

So far, we have an objective, namely upper-bounding the – statistical or worst-case – bias for random coefficients. A naive approach to tackle this challenge would be to leverage some of the tools used in the leakage-resilient secret sharing literature, like the Cauchy-Schwarz trick, or Poisson summation formula [16]. Unfortunately, they would suffer from two drawbacks. Firstly, they rely on Fourier analysis, which implicitly requires the independence assumption. As discussed in the introduction, ensuring this assumption may be challenging and implies implementation overheads. Secondly, they would require to naively compute Equation 2 by upper-bounding each term inside the expectation. For the same reason developed in subsection 4.1, each bound would then not be able to depend on $\vec{\omega}$, unless further characterizing the leakage function.⁹ This would lead to the very same security bounds as obtained for additive masking, which would not be helpful. Therefore, an alternative approach is needed to tackle the issue of the expectation in Equation 2. This is what the following statement proposes.

Proposition 3. *The averaged statistical bias may be equivalently defined as*

$$\beta_{\vec{\Omega}}(Y|L) = \mathbb{E}_{\vec{\ell}} \left[\text{TV} \left(\mathbf{p}_{Y, \vec{\Omega}} |_{L=\vec{\ell}}, \mathbf{p}_Y \otimes \mathbf{p}_{\vec{\Omega}} \right) \right]. \quad (3)$$

Intuitively, Equation 3 respectively “swaps” the expectations over $\vec{\Omega}$ and L . We will leverage this equivalent formulation of $\beta_{\vec{\Omega}}(Y|L)$ in the next subsection.

Proof. In order to prove the proposition, we will show that the average statistical bias is equal to $\text{TV}(\mathbf{p}_{Y, L, \vec{\Omega}}; \mathbf{p}_Y \otimes \mathbf{p}_L \otimes \mathbf{p}_{\vec{\Omega}})$. Leveraging the symmetry of the roles played by L, Y and $\vec{\Omega}$ in the latter quantity, it will imply that we may permute L and $\vec{\omega}$ in the definition of the average statistical bias, hence the result.

By definition of the statistical bias given by Equation 1, we have

$$\beta_{\vec{\omega}}(Y | L) = \mathbb{E}_y \left[\text{TV} \left(\mathbf{p}_{L | Y=y}; \mathbb{E}_{y'} \left[\mathbf{p}_{L | Y=y'} \right] \right) \right] = \mathbb{E}_y \left[\text{TV} \left(\mathbf{p}_{L | Y=y}; \mathbf{p}_L \right) \right]$$

⁹ As we will see in section 5.

Using the definition of conditional probability, we rewrite the right-hand side as

$$\beta_{\vec{\omega}}(Y | L) = \text{TV}(\mathbf{p}_{L,Y}; \mathbf{p}_L \otimes \mathbf{p}_Y) \quad .$$

The probability distributions $\mathbf{p}_{L,Y}$ and $\mathbf{p}_L \otimes \mathbf{p}_Y$ implicitly depend on the value of $\vec{\omega}$, so they may be respectively written as $\mathbf{p}_{L,Y | \vec{\Omega}=\vec{\omega}}$ and $\mathbf{p}_{L,Y' | \vec{\Omega}=\vec{\omega}}$, where Y' is an independent copy of Y . So by definition of the average statistical bias

$$\begin{aligned} \beta_{\vec{\Omega}}(Y|L) &= \mathbb{E}_{\vec{\omega}} [\beta_{\vec{\omega}}(Y | L)] \\ &= \mathbb{E}_{\vec{\omega}} \left[\text{TV}(\mathbf{p}_{L,Y | \vec{\Omega}=\vec{\omega}}; \mathbf{p}_{L,Y' | \vec{\Omega}=\vec{\omega}}) \right] \\ &= \text{TV}(\mathbf{p}_{L,Y,\vec{\Omega}}; \mathbf{p}_{L,Y',\vec{\Omega}}) \quad , \end{aligned}$$

where $\vec{\Omega}'$ is an independent copy of $\vec{\Omega}$. Proposition 3 follows by observing that $\mathbf{p}_{L,Y',\vec{\Omega}'}$ is equal to $\mathbf{p}_L \otimes \mathbf{p}_Y \otimes \mathbf{p}_{\vec{\Omega}'}$, by virtue of the independence of Y' and $\vec{\Omega}'$, and by using again the definition of conditional probability. $\square$

4.3 The Leftover Hash Lemma to the Rescue

We now present the main result of this section, which estimates the amount of information given to an adversary which chooses the leakage function independently of the inner product coefficients, and gets the leakage information together with the inner product coefficients in full.

Theorem 2. *Let $\mathbb{F}$ be a field, and $d \geq 1$ be an integer. Let $Y \in \mathbb{F}$ and $\vec{\Omega} \in (\mathbb{F}^*)^d$ be two independent uniformly distributed random variables, with $d \geq 1$. Furthermore, let $L(Y) = \tau(\text{Enc}_{\vec{\Omega}}(Y))$, where $\tau \in \text{Bnd}_m$, for some $m \geq 1$. Assume the internal randomness of τ is independent of Y and $\vec{\Omega}$. Then,*

$$\beta_{\vec{\Omega}}(Y|L) \leq |\mathbb{F}|^{\frac{1-d}{2}} \cdot 2^{\frac{m-2}{2}} \cdot \left(\frac{|\mathbb{F}|}{|\mathbb{F}| - 1} \right)^d .$$

Despite Theorem 2 applies only to uniformly distributed secrets, its combination with Proposition 1 provides a bound for arbitrary distributed secrets.

Corollary 2. *For any $\delta > 0$, the d -share inner-product encoding is $(\delta, \epsilon_{m,d})$ -leakage resilient against m -bounded leakage for*

$$\epsilon_{m,d} = \delta^{-1} \cdot |\mathbb{F}|^{\frac{3-d}{2}} \cdot 2^{\frac{m-2}{2}} \cdot \left(\frac{|\mathbb{F}|}{|\mathbb{F}| - 1} \right)^d .$$

Proof. Let $\delta > 0$. By applying successively Proposition 1, Theorem 2, and Markov's inequality, we have that

$$\Pr \left(M_{\vec{\omega}}(L) \geq \delta^{-1} \cdot |\mathbb{F}|^{\frac{3-d}{2}} \cdot 2^{\frac{m-2}{2}} \cdot \left(\frac{|\mathbb{F}|}{|\mathbb{F}| - 1} \right)^d \right)$$

$$\begin{aligned}
&\leq \Pr\left(\beta_{\vec{\omega}}(Y \mid L) \geq \delta^{-1} \cdot |\mathbb{F}|^{\frac{1-d}{2}} \cdot 2^{\frac{m-2}{2}} \cdot \left(\frac{|\mathbb{F}|}{|\mathbb{F}|-1}\right)^d\right) \\
&\leq \Pr(\beta_{\vec{\omega}}(Y \mid L) \geq \delta^{-1} \cdot \beta_{\vec{\Omega}}(Y|L)) \\
&\leq \delta,
\end{aligned}$$

where the probability is taken over the choice of $\vec{\omega}$ over $(\mathbb{F}^*)^d$. $\square$

We prove Theorem 2 by connecting the theory of inner product masking to the one of extractors, following an approach similar to the work of Dziembowski and Faust [37]. First, we present the lemmata involved in the proof. In the end of this section, we prove Theorem 2 based on these lemmata.

We start by connecting the statistical distance for coefficients in $(\mathbb{F}^*)^d$ to the one for coefficients in $\mathbb{F}^d$. While the first set corresponds to the coefficients where our encoding operates, the latter guarantees universality of the inner product hash function, as we will see later in Lemma 2.

Lemma 1. *Let $Y, L, \vec{\Omega}$ and $\vec{\Omega}'$ be random variables. Assume that $\vec{\Omega}$ and $\vec{\Omega}'$ are uniformly distributed, respectively, over $(\mathbb{F}^*)^d$ and $\mathbb{F}^d$. Then,*

$$\beta_{\vec{\Omega}}(Y|L) \leq \left(\frac{|\mathbb{F}|}{|\mathbb{F}|-1}\right)^d \beta_{\vec{\Omega}'}(Y|L).$$

Proof. Using Definition 8 and the assumption that $\vec{\Omega}$ and $\vec{\Omega}'$ are uniformly distributed over $(\mathbb{F}^*)^d$ and $\mathbb{F}^d$, we can rewrite

$$\begin{aligned}
\beta_{\vec{\Omega}'}(Y|L) &= \frac{1}{|\mathbb{F}|^d} \sum_{\vec{\omega} \in \mathbb{F}^d} \beta_{\vec{\omega}}(Y|L) \\
&= \frac{1}{|\mathbb{F}|^d} \sum_{\vec{\omega} \in (\mathbb{F}^*)^d} \beta_{\vec{\omega}}(Y|L) + \frac{1}{|\mathbb{F}|^d} \sum_{\vec{\omega} \in \mathbb{F}^d \setminus (\mathbb{F}^*)^d} \beta_{\vec{\omega}}(Y|L) \\
&\geq \frac{1}{|\mathbb{F}|^d} \sum_{\vec{\omega} \in (\mathbb{F}^*)^d} \beta_{\vec{\omega}}(Y|L) \\
&= \left(\frac{(|\mathbb{F}|-1)}{|\mathbb{F}|}\right)^d \cdot \frac{1}{(|\mathbb{F}|-1)^d} \sum_{\vec{\omega} \in (\mathbb{F}^*)^d} \beta_{\vec{\omega}}(Y|L) \\
&= \left(\frac{(|\mathbb{F}|-1)}{|\mathbb{F}|}\right)^d \cdot \beta_{\vec{\Omega}}(Y|L).
\end{aligned}$$

This concludes the proof. $\square$

For every $\vec{\omega} \in \mathbb{F}^d$, let $h_{\vec{\omega}} : \mathbb{F}^d \rightarrow \mathbb{F}$ be the function that, on input any $\vec{x} \in \mathbb{F}^d$, outputs $h_{\vec{\omega}}(\vec{x}) = \langle \vec{\omega}, \vec{x} \rangle$. Define $\mathcal{H}_{\text{IP}} = \{h_{\vec{\omega}}\}_{\vec{\omega} \in \mathbb{F}^d}$. The following claim holds.

Lemma 2. *$\mathcal{H}_{\text{IP}}$ is a universal family of hash functions.*

Proof. By the definition of universal family of hash functions, we will prove that

$$\Pr_{\vec{\omega} \xleftarrow{\$} \mathbb{F}^d} [h_{\vec{\omega}}(\vec{x}) = h_{\vec{\omega}}(\vec{x}')] \leq \frac{1}{|\mathbb{F}|} \quad (4)$$

holds for any couple of distinct elements $\vec{x}, \vec{x}' \in \mathbb{F}^d$. As $h_{\vec{\omega}}(\vec{x}) = \langle \vec{\omega}, \vec{x} \rangle$, we get $[h_{\vec{\omega}}(\vec{x}) = h_{\vec{\omega}}(\vec{x}')] \Leftrightarrow [\langle \vec{\omega}, \vec{x} \rangle = \langle \vec{\omega}, \vec{x}' \rangle] \Leftrightarrow [\langle \vec{\omega}, \vec{x} - \vec{x}' \rangle = 0]$. Due to the uniform distribution of $\vec{\omega}$, we have

$$\Pr_{\vec{\omega} \xleftarrow{\$} \mathbb{F}^d} [h_{\vec{\omega}}(\vec{x}) = h_{\vec{\omega}}(\vec{x}')] = \frac{|\{\vec{\omega} \in \mathbb{F}^d \text{ with } \langle \vec{\omega}, \vec{x} - \vec{x}' \rangle = 0\}|}{|\mathbb{F}|^d}.$$

Next, we can transform the equation by replacing $|\{\vec{\omega} \in \mathbb{F}^d : \langle \vec{\omega}, \vec{x} - \vec{x}' \rangle = 0\}|$ with $|\{\vec{\omega} \in \mathbb{F}^d : \langle \vec{\omega}, \vec{x} - \vec{x}' \rangle = 0\}| = |\mathbb{F}|^{d-1}$. This holds because for every $\vec{x}, \vec{x}' \in \mathbb{F}^d$ there is at least one i such that $x_i - x'_i \neq 0$. Consequently, for every $(d-1)$ -tuple $\omega_1, \dots, \omega_{i-1}, \omega_{i+1}, \dots, \omega_d$ there is a unique ω_i such that $\langle \vec{\omega}, \vec{x} - \vec{x}' \rangle = 0$. This results in the claim – Equation 4 – since $\frac{|\mathbb{F}|^{d-1}}{|\mathbb{F}|^d} = \frac{1}{|\mathbb{F}|}$. $\square$

Since the domain of ω has been extended from $(\mathbb{F}^*)^d$ to $\mathbb{F}^d$, we now extend the definition of inner product encodings (cf Definition 5) accordingly. The extension is natural, namely an $\vec{\omega}$ -inner-product encoding $\text{Enc}_{\vec{\omega}}(y)$ of a secret y with $\vec{\omega} \in \mathbb{F}^d$ is a d -tuple $\llbracket y \rrbracket_{\vec{\omega}} = (y_1, \dots, y_d) \in \mathbb{F}_p^d$, such that

$$y = \langle \vec{\omega}, \llbracket y \rrbracket_{\vec{\omega}} \rangle = \sum_{i=1}^d \omega_i \cdot y_i.$$

Notably, every element of $\mathbb{F}^d$ is a valid 0^d -inner-product encoding, and $\vec{\omega}$ -inner-product encodings are uniformly distributed whenever y is uniform.

Lemma 3. *Let Y be uniformly random, and $\vec{\omega}$ be any value in $\mathbb{F}^d$. Then $\text{Enc}_{\vec{\omega}}(Y)$ is uniformly random over $\mathbb{F}^d$.*

Proof. If $\vec{\omega} = 0^d$, then $\text{Enc}_{\vec{\omega}}(Y)$ can be any value in $\mathbb{F}^d$, and therefore we just sample at random from $\mathbb{F}^d$. Otherwise, there exists $\omega_i \neq 0$. Then sampling a random encoding is equivalent to sampling a vector $\vec{X}$ as follows

- $X_j \xleftarrow{\$} \mathbb{F}$ for every $j \neq i$,
- $X_i = \omega_i^{-1} \left(Y - \sum_{j \neq i} \omega_j X_j \right)$, which is also uniformly random under the assumption that Y is uniformly random.

This makes $\vec{X}$ uniformly random. $\square$

Lemma 4. *Let $\vec{\Omega}'$ and Y be two independent uniform random variables over $\mathbb{F}^d$ and $\mathbb{F}$, respectively. Furthermore, let $L(Y) = \tau(\text{Enc}_{\vec{\Omega}'}(Y))$, where $\tau \in \text{Bnd}_m$. Assume that the internal randomness of τ is independent of Y and $\vec{\Omega}'$. Let $\vec{\ell} \in \mathcal{L}$ such that $\Pr[L = \vec{\ell}] \neq 0$, and let $\vec{X}_{\vec{\ell}} = \left(\text{Enc}_{\vec{\Omega}'}(Y) \mid L = \vec{\ell} \right)$. Then, it holds that $\vec{X}_{\vec{\ell}}$ and $\vec{\Omega}'$ are independent.*

Proof. For any $\vec{x} \in \mathbb{F}^d$, and any $\vec{\omega} \in \mathbb{F}^d$ we need to prove that the probability

$$p_{\vec{\ell}, \vec{\omega}} = \Pr(\text{Enc}_{\vec{\Omega}'}(Y) = \vec{x} \mid \vec{\Omega}' = \vec{\omega}, L = \vec{\ell})$$

does not depend on $\vec{\omega}$. Using Bayes' theorem, we have

$$p_{\vec{\ell}, \vec{\omega}} = \frac{\Pr(L = \vec{\ell} \mid \text{Enc}_{\vec{\Omega}'}(Y) = \vec{x}, \vec{\Omega}' = \vec{\omega}) \cdot \Pr(\text{Enc}_{\vec{\Omega}'}(Y) = \vec{x} \mid \vec{\Omega}' = \vec{\omega})}{\Pr(L = \vec{\ell} \mid \vec{\Omega}' = \vec{\omega})}.$$

Let us inspect every factor in the right-hand side.

- The first factor, *i.e.* $\Pr(L = \vec{\ell} \mid \text{Enc}_{\vec{\Omega}'}(Y) = \vec{x}, \vec{\Omega}' = \vec{\omega})$, can be rewritten as $\Pr(\tau(\vec{x}) = \vec{\ell})$. Since we assumed the internal randomness of τ to be independent of Y and $\vec{\Omega}'$, this probability is also independent of $\vec{\omega}$.
- We rewrite the second factor as

$$\Pr(\text{Enc}_{\vec{\Omega}'}(Y) \mid \vec{\Omega}' = \vec{\omega}) = \Pr(\text{Enc}_{\vec{\omega}}(Y) = \vec{x}).$$

Given that Y is assumed uniform, Lemma 3 guarantees that so is $\text{Enc}_{\vec{\omega}}(Y)$, *i.e.*, it does not depend on $\vec{\omega}$.

- Using the total probability formula, the third factor $\Pr(L = \vec{\ell} \mid \vec{\Omega}' = \vec{\omega})$ can be expressed as the sum over $\vec{x} \in \mathbb{F}^d$ of the product of the first two ones. Since they do not depend on $\vec{\omega}$, their sum does not either.

Thus, none of the factors depend on $\vec{\omega}$, hence $p_{\vec{\ell}, \vec{\omega}}$ is independent of $\vec{\omega}$. $\square$

Similarly to what is done by Dziembowski and Faust [37, Lemma 20], we need to estimate the entropy of the encoding after leakage. For sake of tightness we use a better approximation that follows from the works of Dodis [32].

Lemma 5. *Let Y be uniformly distributed over $\mathbb{F}$, let $\tau \in \text{Bnd}_m$, and define $L = \tau(\text{Enc}_{\vec{\Omega}'}(Y))$. Then*

$$\tilde{H}_{\infty}(\text{Enc}_{\vec{\Omega}'}(Y) \mid L) \geq d \log_2 |\mathbb{F}| - m.$$

Proof. [32, Lemma 2.2, item b] guarantees that for every couple of random variables A, B such that B takes values in a space $\mathcal{B}$ of 2^m elements, then $\tilde{H}_{\infty}(A \mid B) \geq H_{\infty}(A) - m$. Setting $A = \text{Enc}_{\vec{\Omega}'}(Y)$ and $B = L$, the lemma follows with $H_{\infty}(\text{Enc}_{\vec{\Omega}'}(Y)) = d \log_2 |\mathbb{F}|$ by Lemma 3. $\square$

Putting all these lemmas together, Theorem 2 follows.

Proof of Theorem 2.

$$\beta_{\vec{\Omega}}(Y|L) \leq \left(\frac{|\mathbb{F}|}{|\mathbb{F}| - 1} \right)^d \beta_{\vec{\Omega}'}(Y|L) \quad (\text{Lemma 1})$$

$$\begin{aligned}
&= \left(\frac{|\mathbb{F}|}{|\mathbb{F}| - 1} \right)^d \sum_{\vec{\ell} \in \mathcal{L}} \Pr(L = \vec{\ell}) \cdot \text{TV}(\mathbf{p}_{Y, \vec{\Omega}'} \mid_{\{L=\vec{\ell}\}}; \mathbf{p}_Y \otimes \mathbf{p}_{\vec{\Omega}'}) \\
&\leq \left(\frac{|\mathbb{F}|}{|\mathbb{F}| - 1} \right)^d \sum_{\vec{\ell} \in \mathcal{L}} \Pr(L = \vec{\ell}) \cdot |\mathbb{F}|^{\frac{1}{2}} \cdot 2^{-1 - H_\infty(\text{Enc}_{\vec{\Omega}'}(Y) \mid L=\vec{\ell})/2} \quad (\text{LHL}) \\
&= \frac{1}{2} \left(\frac{|\mathbb{F}|}{|\mathbb{F}| - 1} \right)^d \cdot |\mathbb{F}|^{\frac{1}{2}} \cdot \sum_{\vec{\ell} \in \mathcal{L}} \Pr(L = \vec{\ell}) \cdot 2^{-H_\infty(\text{Enc}_{\vec{\Omega}'}(Y) \mid L=\vec{\ell})/2} \\
&\leq \frac{1}{2} \left(\frac{|\mathbb{F}|}{|\mathbb{F}| - 1} \right)^d \cdot |\mathbb{F}|^{\frac{1}{2}} \cdot \left(\sum_{\vec{\ell} \in \mathcal{L}} \Pr(L = \vec{\ell}) \right)^{1/2} \\
&\quad \cdot \left(\sum_{\vec{\ell} \in \mathcal{L}} \Pr(L = \vec{\ell}) \cdot 2^{-H_\infty(\text{Enc}_{\vec{\Omega}'}(Y) \mid L=\vec{\ell})/2} \right)^{1/2} \quad (\text{Cauchy-Schwarz}) \\
&\leq \frac{1}{2} \left(\frac{|\mathbb{F}|}{|\mathbb{F}| - 1} \right)^d \cdot |\mathbb{F}|^{\frac{1}{2}} \cdot 2^{-\tilde{H}_\infty(\text{Enc}_{\vec{\Omega}'}(Y) \mid L)/2} \\
&\leq \frac{1}{2} \left(\frac{|\mathbb{F}|}{|\mathbb{F}| - 1} \right)^d \cdot |\mathbb{F}|^{\frac{1}{2}} \cdot 2^{(-d \log_2 |\mathbb{F}| + m)/2} \quad (\text{Lemma 5}) \\
&= |\mathbb{F}|^{\frac{1-d}{2}} \cdot 2^{\frac{m-2}{2}} \cdot \left(\frac{|\mathbb{F}|}{|\mathbb{F}| - 1} \right)^d,
\end{aligned}$$

which completes the proof. It is important to note that the approximation step involving the leftover hash lemma (LHL) requires independence of the inputs to the inner product extractor, namely $\vec{X}_{\vec{\ell}} = (\text{Enc}_{\vec{\Omega}'}(Y) \mid L = \vec{\ell})$ and $\vec{\Omega}'$. The independence of these inputs is shown in Lemma 4. $\square$

4.4 Discussion

We conclude this section by discussing two important aspects of our theoretical results, and we put them in perspective with practical considerations.

Drawing $\vec{\omega}$ once for all. Our threat model assumes that the public coefficients $\vec{\omega}$ are drawn independently of the physical leakage behaviour of the device, *e.g.*, after it comes out of the foundry, but before manipulation by the leaking device. This means that if we now consider a threat model in which the device operates several executions of an algorithm, the public coefficients $\vec{\omega}$ must be refreshed before every execution, in order to stick to the theoretical assumptions of Theorem 2 and Corollary 2. Despite this may be achievable with limited overheads, it questions what would happen in the even more convenient situation where $\vec{\omega}$ is not refreshed, *i.e.*, $\vec{\omega}$ is drawn once for all before the very first execution. This would naturally represent a gap with our security model but, as we argue next, may not significantly impact the concrete security of an IP encoding.

First, the probability that such a coefficient leads to weak security by chance is negligible, as proven above. So the only thing that could happen is that the adversary now tries to exploit the possibility to tweak the leakage function to fall in one of the statistically unlikely worst-cases. Concretely, this would mean adapting the setup in order to shape the leakage function as required.

While this is theoretically feasible, we believe that the extent of adversarial control on side-channel measurement setups is in general insufficient to lead to significant security degradations. The practical investigation of such attacks is an interesting scope for further research. Note that in case these adaptive attacks are possible, it is likely that adapting the leakage function to $\vec{\omega}$ is not instantaneous and various tradeoffs could be considered, between updating these coefficients for every new encryption and keeping them constant all the time.

From Bounded to Noisy Leakages. So far, we have established our results in a threat model where the leakage function has a bounded range. This is aligned with similar assumptions in the leakage-resilient secret sharing literature [16]. Yet, the difficulty of setting a reasonable value m for the function's range is a frequent criticism of the bounded leakage model, because it is essentially determined by the attacker's measurement ability — *e.g.*, the resolution on the oscilloscope — which can improve with reasonable effort and budget.

Interestingly, we can relax the bounded leakage assumption to the noisy setting — and for free. The only part of the proof of Theorem 2 and Corollary 2 relying on the bounded leakage assumption is Lemma 5, which lower bounds the average min-entropy left after leakage. In the noisy setting, the output of the leakage function can be arbitrarily long but must retain some noise-hiding sensitive information. Therefore, this becomes more a property of the leakage, rather than a property of the adversary's measurement ability. If we quantify this information as ℓ in the U-noisy model (see Definition 6 in [18]), it means that we can replace m with ℓ in all the results of this section.¹⁰

On semi adaptive leakage. Our threat model assumes an attacker that chooses the leakage function τ entirely independently from the coefficients $\vec{\omega}$, and later gets $\vec{\omega}$ in full together with the leakage $\tau(\text{Enc}_{\vec{\omega}}(y))$. As discussed in subsection 4.1, some limitations on the leakage function are necessary, as a fully adaptive adversary could directly decode within the leakage function. However, one might be interested in a semi-adaptive scenario, where the adversary can issue adaptive leakage queries to either $\vec{\omega}$ or $\text{Enc}_{\vec{\omega}}(y)$, as long as the overall leakage is λ bits from each. This model is considered in Section 2.3 of Dziembowski and Faust's work [37], where the coefficients $\vec{\omega}$ are even generalized to a matrix, but there is no security guarantee when $\vec{\omega}$ is revealed in the end. A semi-adaptive variant of our model can be derived by restricting their setting to the case where $\vec{\omega}$ remains a vector (*i.e.* “ $m = 1$ ” in their encoding scheme), but is revealed afterward. To argue about security in this model, we quickly adapt the proof of

¹⁰ Generic reductions proposed in [18,73] could also be considered.

their Lemma 1. This results in a weaker, but still meaningful, security bound of $N \gtrsim |\mathbb{F}|^{d/4}$, in contrast to our bound of $N \gtrsim |\mathbb{F}|^{d/2}$. We make this observation formal in the following corollary, and provide a proof sketch in Appendix E.

Corollary 3. *Consider an adversary that adaptively leaks λ bits from the coefficients $\vec{\omega}$ and the encoding $\text{Enc}_{\vec{\omega}}(y)$, i.e. each bit of $\mathbf{L}$ is a function of the previous bits and either $\vec{\omega}$ or $\text{Enc}_{\vec{\omega}}(y)$. Then for any $\delta > 0$, the d -share inner-product encoding is $(\delta, \epsilon_{m,d})$ -leakage resilient against m -bounded leakage for*

$$\epsilon_{m,d} = \delta^{-1} \cdot |\mathbb{F}|^{1-\frac{d}{4}} \cdot 2^{\frac{\lambda+1}{2}}.$$

5 Beyond the Square-Root Frontier

The main result of section 4, namely Theorem 2, resolves the dilemma mentioned in subsection 4.1. It allows to circumvent the case where the attacker is able to adversarially choose the leakage function depending on $\vec{\omega}$ and the leakage function does not have these coefficients as arguments (See Section 2).

Interestingly, the literature on inner-product masking schemes has thoroughly investigated the opposite threat scenario: with enough characterization, the leakage function might be considered as fixed and known, so that designers may choose the coefficients constructively, i.e., to maximize the side-channel security. This is for example at the core of the line of research on *security order amplification* for IP masking, provided that the scheme operates over a binary field and that the leakage function is linear, in a high-noise regime [84,75,25].

In this section, we draw our interest on a similar threat scenario, where the leakage function is fixed, and where the designer may choose the coefficients accordingly. Our intuition follows from the probabilistic argument in section 4: if we get good bounds on average for random coefficients, there must be some vector $\vec{\omega}_{\text{bad}}$ for which the security bound is worse, as well as some vector $\vec{\omega}_{\text{good}}$ for which the security guarantee is better. It is not that hard to find instances of bad coefficients. Take the Least Significant Bit (LSB) as an example. It has been shown in the literature that $\vec{\omega}_{\text{bad}} = (1, 1, \dots, 1)$ results in *much* worse security bounds than what could expect from Theorem 2, whether it is in a binary field [67], or in a prime field [60]. Does that mean that we could also get *much* better coefficients? We answer positively to this question, by investigating the use-case of physical bit leakage applied to each share of an $\vec{\omega}$ -encoding in a prime field. To this end, we first need to introduce some Fourier analysis tools that will serve in order to derive a more refined security bound than the generic one from Theorem 2, for some well-chosen coefficients.

5.1 Background on Fourier Analysis

We first recall the definition of the discrete Fourier transform of a function. Then, we recall the Poisson summation formula that has been used in [16].

Definition 9 (Discrete Fourier Transform). Let $f : \mathbb{F}_p \rightarrow \mathbb{C}$ be a function over $\mathbb{F}_p$ seen as a cyclic group. The discrete Fourier transform of f is the mapping

$$\alpha \in \mathbb{F}_p \mapsto \widehat{f}(\alpha) = \frac{1}{p} \sum_{x \in \mathbb{F}_p} f(x) \gamma^{\alpha \cdot x} ,$$

and for any $x \in \mathbb{F}_p$, the inverse discrete Fourier transform is the mapping

$$x \in \mathbb{F}_p \mapsto f(x) = \sum_{\alpha \in \mathbb{F}_p} \widehat{f}(\alpha) \cdot \gamma^{-\alpha \cdot x} .$$

Lemma 6 (Poisson Summation Formula). Let $f_1, \dots, f_d$ be functions $\mathbb{F}_p \rightarrow \mathbb{C}$, and let $\vec{y} + C \subseteq \mathbb{F}_p^d$ be an affine subspace of $\mathbb{F}_p^d$, with offset $\vec{y} \in \mathbb{F}_p^d$ and kernel C . Then the following equality holds:

$$\mathbb{E}_{\vec{x} \stackrel{\$}{\leftarrow} \vec{y} + C} \left[\prod_{i=1}^d f_i(x_i) \right] = \sum_{\vec{\alpha} \in C^\perp} \left(\prod_{i=1}^d \widehat{f}_i(\alpha_i) \right) \cdot \gamma^{-\langle \vec{\alpha}, \vec{y} \rangle} .$$

Proof. Expressing every f_i through its inverse Fourier transform, we get

$$\begin{aligned} \mathbb{E}_{\vec{x} \stackrel{\$}{\leftarrow} \vec{y} + C} \left[\prod_{i=1}^d f_i(x_i) \right] &= \mathbb{E}_{\vec{x} \stackrel{\$}{\leftarrow} \vec{y} + C} \left[\prod_{i=1}^d \sum_{\alpha_i \in \mathbb{F}_p} \widehat{f}_i(\alpha_i) \cdot \gamma^{-\alpha_i \cdot x_i} \right] \\ &= \mathbb{E}_{\vec{x} \stackrel{\$}{\leftarrow} \vec{y} + C} \left[\sum_{\vec{\alpha} \in \mathbb{F}_p^d} \left(\prod_{i=1}^d \widehat{f}_i(\alpha_i) \cdot \gamma^{-\alpha_i \cdot x_i} \right) \right] \\ &= \sum_{\vec{\alpha} \in \mathbb{F}_p^d} \left(\prod_{i=1}^d \widehat{f}_i(\alpha_i) \right) \cdot \mathbb{E}_{\vec{x} \stackrel{\$}{\leftarrow} \vec{y} + C} \left[\gamma^{-\langle \vec{\alpha}, \vec{x} \rangle} \right] , \end{aligned}$$

where the second equality comes from distributing the product of sums, and the third equality holds thanks to the linearity of the expectation. Let us now focus on the term $\mathbb{E}_{\vec{x} \stackrel{\$}{\leftarrow} \vec{y} + C} [\gamma^{-\langle \vec{\alpha}, \vec{x} \rangle}]$, for a fixed $\vec{\alpha} \in \mathbb{F}_p^d$, we can express any $\vec{x} \in \vec{y} + C$ as the sum $\vec{y} + (\vec{x} - \vec{y})$, where the term inside the parenthesis belongs to C . Hence,

$$\begin{aligned} \mathbb{E}_{\vec{x} \stackrel{\$}{\leftarrow} \vec{y} + C} \left[\gamma^{-\langle \vec{\alpha}, \vec{x} \rangle} \right] &= \mathbb{E}_{\vec{x} \stackrel{\$}{\leftarrow} \vec{y} + C} \left[\gamma^{-\langle \vec{\alpha}, \vec{y} + (\vec{x} - \vec{y}) \rangle} \right] \\ &= \gamma^{-\langle \vec{\alpha}, \vec{y} \rangle} \cdot \mathbb{E}_{\vec{x} \stackrel{\$}{\leftarrow} C} \left[\gamma^{-\langle \vec{\alpha}, \vec{x} \rangle} \right] , \end{aligned}$$

where the second equality is obtained by a change of variable $\vec{x} - \vec{y} \mapsto \vec{x}$. If $\vec{x} \perp \vec{\alpha}$, then $\gamma^{-\langle \vec{\alpha}, \vec{x} \rangle} = 1$, otherwise, the expectation equals zero. As a consequence, we may restrict the sum over the orthogonal subspace of C . $\square$

5.2 Inner-Product Masking in the Independence Assumption

For the remaining of the paper, we focus on a class of leakage functions occurring on a single encoding, verifying the so-called *independence assumption*.

Definition 10 (Independence Assumption). *A (possibly randomized) function $\tau : \mathbb{F}^d \rightarrow \mathcal{L}$ verifies the independence assumption if there exist some (possibly randomized) functions $\varphi_1, \dots, \varphi_d$ such that for any input $\vec{x} \in \mathbb{F}^d$, (1) the random variables $\varphi_i(x_i)$ are mutually independent; and (2) $\tau(\vec{x})$ and $(\varphi_1(x_1), \dots, \varphi_d(x_d))$, are identically distributed.*

The independence assumption is a condition required in the *noisy leakage* model [77,34,35]. When applied to an encoding generated from inner-product masking, one may rephrase the security metric, as stated hereafter.¹¹

Proposition 4. *Let $\vec{\omega} \in (\mathbb{F}^*)^d$. Let $\tau : \mathbb{F}^d \rightarrow \mathcal{L}^d$ verifying the independence assumption (Definition 10), and denote by $\varphi_1, \dots, \varphi_d$ the functions $\mathbb{F} \rightarrow \mathcal{L}$ such that for any $\vec{x} \in \mathbb{F}^d$, $\tau(\vec{x}) = (\varphi_1(x_1), \dots, \varphi_d(x_d))$. Let $\mathbf{p}_{\varphi_i, \ell_i}$ be the mapping $x \mapsto \Pr(\varphi_i(x) = \ell_i)$.¹² For any $y^{(0)}, y^{(1)} \in \mathbb{F}_p$, let $L(y^{(i)}) = \tau(\text{Enc}_{\vec{\omega}}(y^{(i)}))$, for $i \in \{0, 1\}$, and denote its probability distribution by $\mathbf{p}_i$. Then,*

$$\text{TV}(\mathbf{p}_0; \mathbf{p}_1) = \frac{1}{2} \sum_{\vec{\ell} \in \mathcal{L}^d} \left| \sum_{\alpha \in \mathbb{F}_p^*} \left(\prod_{i=1}^d \widehat{\mathbf{p}_{\varphi_i, \ell_i}}(\omega_i \cdot \alpha) \right) \cdot \left(\gamma^{-\alpha \cdot y^{(0)}} - \gamma^{-\alpha \cdot y^{(1)}} \right) \right|.$$

Proof. By definition of the total variation,

$$\text{TV}(\mathbf{p}_0; \mathbf{p}_1) = \frac{1}{2} \sum_{\vec{\ell} \in \mathcal{L}^d} \left| \Pr(L(y^{(0)}) = \vec{\ell}) - \Pr(L(y^{(1)}) = \vec{\ell}) \right|.$$

Denote the latter quantity by Δ for short. The independence assumption (see Definition 10) implies that for any $\vec{\ell} \in \mathcal{L}$, and any $\vec{x} \in \mathbb{F}_p^d$,

$$\Pr(\tau(\vec{x}) = \vec{\ell}) = \prod_{i=1}^d \Pr(\varphi_i(x_i) = \ell_i).$$

Injecting this into Δ gives

$$\Delta = \frac{1}{2} \sum_{\vec{\ell} \in \mathcal{L}^d} \left| \mathbb{E}_{\llbracket y \rrbracket_{\vec{\omega}} \stackrel{\$}{\leftarrow} y^{(0)} + \text{Span}(\vec{\omega})^\perp} \left[\prod_{i=1}^d \Pr(\varphi_i(y_i) = \ell_i) \right] - \mathbb{E}_{\llbracket y \rrbracket_{\vec{\omega}} \stackrel{\$}{\leftarrow} y^{(1)} + \text{Span}(\vec{\omega})^\perp} \left[\prod_{i=1}^d \Pr(\varphi_i(y_i) = \ell_i) \right] \right|.$$

¹¹ Proposition 4 adapts the statement of Faust *et al.* [40, Prop. 3] to IP masking.

¹² This is not a probability mass function as it is a function of y instead of ℓ .

Applying Poisson summation formula of Lemma 6 to both terms inside the absolute value, respectively with $\vec{y} = (y^{(0)}, 0, \dots, 0)$, $C = \text{Span}(\vec{\omega})$ for the first term, and $\vec{y} = (y^{(1)}, 0, \dots, 0)$, $C = \text{Span}(\vec{\omega})$ for the second term. It follows that

$$\Delta = \frac{1}{2} \sum_{\vec{\ell} \in \mathcal{L}^d} \left| \sum_{\vec{\alpha} \in \text{Span}(\vec{\omega})^\perp} \left(\prod_{i=1}^d \widehat{\mathbf{p}_{\varphi_i, \ell_i}}(\alpha_i) \right) \cdot \left(\gamma^{-\langle \vec{\alpha}, \llbracket y^{(0)} \rrbracket \rangle} - \gamma^{-\langle \vec{\alpha}, \llbracket y^{(1)} \rrbracket \rangle} \right) \right|.$$

Noticing that the sum over $\text{Span}(\vec{\omega})^\perp$ may be replaced by a simple sum over $\mathbb{F}_p$,

$$\Delta = \frac{1}{2} \sum_{\vec{\ell} \in \mathcal{L}^d} \left| \sum_{\alpha \in \mathbb{F}_p} \left(\prod_{i=1}^d \widehat{\mathbf{p}_{\varphi_i, \ell_i}}(\omega_i \cdot \alpha) \right) \cdot \left(\gamma^{-\alpha \cdot \omega_1 \cdot y^{(0)}} - \gamma^{-\alpha \cdot \omega_1 \cdot y^{(1)}} \right) \right|.$$

The proof ends by simplifying $\omega_1 = 1$, and observing that the summand with $\alpha = 0$ equals 0. $\square$

Proposition 5 (Upper-Bound Lemma). *Let $\vec{\omega} \in (\mathbb{F}_p^\star)^d$, and let L be defined as in Proposition 4. Then the following inequality holds:*

$$M_{\vec{\omega}}(L) \leq \sum_{\vec{\ell} \in \mathcal{L}^d} \sum_{\alpha \in \mathbb{F}_p^\star} \prod_{i=1}^d |\widehat{\mathbf{p}_{\varphi_i, \ell_i}}(\omega_i \cdot \alpha)|.$$

Proof. Apply the triangle inequality in the right-hand side of Proposition 4 and observe that $|\gamma^{-\alpha \cdot y^{(0)}} - \gamma^{-\alpha \cdot y^{(1)}}| \leq 2$ for every $y^{(0)}, y^{(1)} \in \mathbb{F}_p$. $\square$

5.3 Application: the Physical Bit Leakage Model

The statements introduced in the previous subsection hold for any leakage model verifying the independence assumption. We now focus on a specific sub-class of such leakages, namely the *physical-bit leakage* model. In this model, each share is encoded and stored as a binary string in a register, whose one bit is revealed to the adversary. This leakage model is considered an object of interest in the leakage-resilient secret sharing literature [67,60]. Our goal is to find good coefficients $\vec{\omega}$ such that the worst-case bias can be much better bounded than when the coefficients are chosen randomly. To this end, we leverage Proposition 5 and recall two facts on the Fourier spectrum of the physical-bit leakage.

Proposition 6 (Fourier spectrum of the LSB model [60, Claim 15]).

The Fourier spectrum of the lsb leakage function is given by:

$$\widehat{\mathbf{1}_{\text{lsb},0}}(\alpha) = -\widehat{\mathbf{1}_{\text{lsb},1}}(\alpha) = \frac{1}{2p} \cdot \left(\cos\left(\frac{\alpha \cdot \pi}{p}\right) \right)^{-1} \cdot e^{\frac{i\pi\alpha}{p}}.$$

Proposition 7 (Reduction from ksb to lsb [40, p.15]). *Let ksb be the function that maps a value $x \in \mathbb{F}_p$ to its $(k+1)$ -th significant bit — i.e., the bit of weight 2^k — where p is a Mersenne number. Then, the Fourier spectrum of the ksb function verifies*

$$\widehat{\mathbf{1}_{\text{ksb},0}}(\alpha) = \widehat{\mathbf{1}_{\text{lsb},0}}(2^{-k} \cdot \alpha).$$

By injecting the Fourier transform of the ksb into Proposition 5, a product of cosine appears. It turns out that for some well-chosen coefficient $\vec{\omega}$, this product can be simplified, as stated by the following lemma.

Lemma 7. *Let p be an odd prime, let $\alpha \in \mathbb{Z}_p^*$, and let d be an integer. Then the following equality holds:*

$$\prod_{i=1}^d \cos\left(\frac{2^{i-1} \cdot \alpha \cdot \pi}{p}\right) = \frac{1}{2^d} \cdot \frac{\sin\left(\frac{2^d \cdot \alpha \cdot \pi}{p}\right)}{\sin\left(\frac{\alpha \cdot \pi}{p}\right)}.$$

Proof. Let

$$\mathcal{C} = \prod_{i=1}^d \cos\left(\frac{2^{i-1} \cdot \alpha \cdot \pi}{p}\right), \text{ and } \mathcal{S} = \prod_{i=1}^d \sin\left(\frac{2^{i-1} \cdot \alpha \cdot \pi}{p}\right).$$

By rearranging the factors in the product $\mathcal{C} \cdot \mathcal{S}$, and leveraging the identity $\sin(2x) = 2 \cos(x) \sin(x)$, we get that

$$\mathcal{C} \cdot \mathcal{S} = \frac{1}{2^d} \cdot \prod_{i=1}^d \sin\left(\frac{2^i \cdot \alpha \cdot \pi}{p}\right) = \frac{1}{2^d} \cdot \prod_{i=2}^{d+1} \sin\left(\frac{2^{i-1} \cdot \alpha \cdot \pi}{p}\right).$$

Introducing an extra factor for $i = 1$, and extracting the factor of index $i = d+1$ in the product, we get that

$$\mathcal{C} \cdot \mathcal{S} = \frac{1}{2^d} \cdot \sin\left(\frac{2^d \alpha \pi}{p}\right) \cdot \sin\left(\frac{\alpha \pi}{p}\right)^{-1} \cdot \prod_{i=1}^d \sin\left(\frac{2^{i-1} \cdot \alpha \cdot \pi}{p}\right).$$

We conclude by identifying the product in the latter right-hand side as $\mathcal{S}$. Since p is prime, there is no value of $\alpha \in \mathbb{F}_p^*$ such that $\mathcal{S} = 0$. Hence, we may divide each side by $\mathcal{S}$. $\square$

Proposition 8 (Good even coefficients). *Let $\text{ksb}_1, \dots, \text{ksb}_d$ be respectively the $(k_1 + 1)$ -th, $\dots$, $(k_d + 1)$ -th significant-bit leakage functions, for $k_1, \dots, k_d \in \llbracket 0, n - 1 \rrbracket$. Let $\vec{\omega} = (2^{k_1}, \dots, 2^{k_1 + i - 1}, \dots, 2^{k_d + d - 1})$, and let $L : y \mapsto \tau(\text{Enc}_{\vec{\omega}}(y))$, where $\tau(\vec{x}) = (\text{ksb}_1(x_1), \dots, \text{ksb}_d(x_d))$. Then,*

$$M_{\vec{\omega}}(L) \leq \left(\frac{2}{p}\right)^d \cdot p \cdot \left(\log \left\lfloor \frac{p-1}{2} \right\rfloor + 1\right).$$

Proof. Applying Proposition 5 to the ksb leakage model, we have

$$M_{\vec{\omega}}(L) \leq \sum_{\vec{\ell} \in \mathcal{L}^d} \sum_{\alpha \in \mathbb{F}_p^*} \prod_{i=1}^d |\widehat{\text{psb}_{\text{ksb}_i, \ell_i}}(\omega_i \cdot \alpha)|.$$

Proposition 7 and Proposition 6 imply that for any $\vec{\ell}$, we have

$$\sum_{\alpha \in \mathbb{F}_p^*} \prod_{i=1}^d |\widehat{\text{psb}_{\text{ksb}_i, \ell_i}}(\omega_i \cdot \alpha)| = \sum_{\alpha \in \mathbb{F}_p^*} \prod_{i=1}^d \frac{1}{2p} \cdot \left| \cos\left(\frac{2^{i-1} \cdot \alpha \cdot \pi}{p}\right) \right|^{-1}, \quad (5)$$

hence,

$$\begin{aligned} M_{\vec{\omega}}(L) &\leq 2^d \cdot \sum_{\alpha \in p^*} \prod_{i=1}^d \frac{1}{2p} \cdot \left| \cos\left(\frac{2^{i-1} \cdot \alpha \cdot \pi}{p}\right) \right|^{-1} \\ &= \frac{1}{p^d} \cdot \sum_{\alpha \in p^*} \prod_{i=1}^d \left| \cos\left(\frac{2^{i-1} \cdot \alpha \cdot \pi}{p}\right) \right|^{-1}. \end{aligned}$$

Invoking Lemma 7, we get that:

$$\mathrm{TV}\left(L(y^{(0)}); L(y^{(1)})\right) \leq \left(\frac{2}{p}\right)^d \cdot \sum_{\alpha \in \mathbb{F}_p^*} \left| \frac{\sin\left(\frac{\alpha \cdot \pi}{p}\right)}{\sin\left(\frac{2^d \cdot \alpha \cdot \pi}{p}\right)} \right|.$$

We conclude the proof by using Lemma 8. □

Lemma 8. *Let p be a prime number. Then, for any integer d ,*

$$\sum_{\alpha \in \mathbb{F}_p^*} \left| \frac{\sin\left(\frac{\alpha \cdot \pi}{p}\right)}{\sin\left(\frac{2^d \cdot \alpha \cdot \pi}{p}\right)} \right| \leq p \cdot \left(\log \left\lfloor \frac{p-1}{2} \right\rfloor + 1 \right).$$

In order to prove the lemma, we need the following identity.

Claim. Let p be a prime number, and let $x, y \in \mathbb{F}_p$. Then,

$$\left| \sin\left(x \cdot y \cdot \frac{\pi}{p}\right) \right| = \sin\left((x \otimes y) \cdot \frac{\pi}{p}\right),$$

where $\otimes$ denotes the field multiplication in $\mathbb{F}_p$.

Proof. Observe that the function $t \mapsto |\sin(t)|$ is π -periodic, so for any $x, y \in \mathbb{Z}_p$,

$$\begin{aligned} \left| \sin\left(\frac{x \cdot y}{p} \cdot \pi\right) \right| &= \sin\left(\left(\frac{x \cdot y}{p} \cdot \pi\right) [\pi]\right) \\ &= \sin\left(\frac{x \cdot y}{p} [1] \cdot \pi\right) \\ &= \sin\left((x \cdot y) [p] \cdot \frac{\pi}{p}\right) \\ &= \sin\left((x \otimes y) \cdot \frac{\pi}{p}\right). \end{aligned}$$

□

Proof of Lemma 8. Let S be the sum to bound. By upper bounding the numerator of each summand by one, we get that

$$S \leq \sum_{\alpha \in \mathbb{F}_p^*} \left| \sin\left(\alpha \cdot 2^d \cdot \frac{\pi}{p}\right) \right|^{-1}.$$

Now we use the previous claim to make the change of variable $\alpha \leftarrow \alpha \otimes 2^d$ in the latter sum. As a result,

$$S \leq \sum_{\alpha \in \mathbb{F}_p^*} \sin\left(\alpha \otimes 2^d \cdot \frac{\pi}{p}\right)^{-1} = \sum_{\alpha \in \mathbb{F}_p^*} \sin\left(\frac{\alpha \cdot \pi}{p}\right)^{-1}.$$

We now use the symmetry of the sine with respect to $\frac{\pi}{2}$:

$$S \leq 2 \cdot \sum_{\alpha=1}^{\lfloor \frac{p-1}{2} \rfloor} \sin\left(\frac{\alpha \cdot \pi}{p}\right)^{-1}.$$

Observe now that for any $x \in [0, \frac{\pi}{2}]$, $\sin(x) \geq \frac{2x}{\pi}$, i.e., $\sin(x)^{-1} \leq \frac{\pi}{2x}$. Hence,

$$S \leq 2 \cdot \sum_{\alpha=1}^{\lfloor \frac{p-1}{2} \rfloor} \frac{\pi}{2 \cdot \left(\frac{\alpha\pi}{p}\right)} = p \cdot \sum_{\alpha=1}^{\lfloor \frac{p-1}{2} \rfloor} \frac{1}{\alpha} \leq p \cdot \left(\log \left\lfloor \frac{p-1}{2} \right\rfloor + 1\right).$$

□

We have emphasized one good combination of even inner-product coefficients, using powers of two. The following proposition finally allows us to derive a combination of odd coefficients leading to the same result.

Proposition 9. *If $\vec{\omega}$ is a good choice of even coefficients, then $-\vec{\omega}$ is an as good choice of odd coefficients.*

Proof. Observe that since p is prime, $-\omega$ is odd if and only if ω is even. Then, leverage the parity of the cosine function in Equation 5. □

Corollary 4 (Good odd coefficients). *Let $\text{ksb}_1, \dots, \text{ksb}_d$ be respectively the $(k_1 + 1)$ -th, $\dots$, $(k_d + 1)$ -th significant-bit leakage functions, for $k_1, \dots, k_d \in \llbracket 0, n-1 \rrbracket$.*

Let $\vec{\omega} = -(2^{k_1}, \dots, 2^{k_1+1}, \dots, 2^{k_d+d-1})$, and let $L : y \mapsto \tau(\text{Enc}_{\vec{\omega}}(y))$, where $\tau(\vec{x}) = (\text{ksb}_1(x_1), \dots, \text{ksb}_d(x_d))$. Then,

$$M_{\vec{\omega}}(L) \leq \left(\frac{2}{p}\right)^d \cdot p \cdot \left(\log \left\lfloor \frac{p-1}{2} \right\rfloor + 1\right).$$

In other words, upon a good choice of $\vec{\omega}$, we may improve the security bound from $\mathcal{O}\left(2^n \cdot \sqrt{2^{-nd+n+d}}\right)$ with Theorem 2 to $\mathcal{O}(n \cdot 2^{-dn+n+d})$, by leveraging a finer-grain analysis. In particular, these asymptotic bounds suggest that the designer has an almost similar incentive in increasing the bit-size as increasing the masking order. This may be helpful if one is cheaper than the other.

Interestingly, such an improvement of the security bound by wisely choosing the coefficients may happen because the bound obtained with additive masking was much worse. It has been emphasized that the amplitude of the Fourier spectrum of the lsb leakage function is extremely concentrated over two points [60, 40],

which was the root cause of a tight security bound independent of the field size. By carefully choosing the coefficients in Proposition 8 and Corollary 4, we somehow leverage Proposition 4 to break this concentration. This suggests that a similar analysis may be applied to other leakage functions suffering from bad security bounds for some given inner-product coefficients. It is likely that they would suffer from the concentration of the Fourier spectrum as well. If so, a careful Fourier analysis may give some hints in order to derive good inner-product coefficients. Overall, this case study depicts how a designer could take profit from leakage characterization to derive better security bounds (holding with probability one) than that of Theorem 2, for some well-chosen coefficients.

Acknowledgments. François-Xavier Standaert is a research director of the Belgian Fund for Scientific Research (FNRS-F.R.S.). Hai H. Nguyen received funding in part from the Zurich Information Security and Privacy Center (ZISC), ETH Zurich. This work was supported by the German Research Foundation (DFG) via the DFG CRC 1119 CROSSING (project S7), by the German Federal Ministry of Education and Research and the Hessen State Ministry for Higher Education, Research and the Arts within their joint support of the National Research Center for Applied Cybersecurity ATHENE via the Projects LEAK and CoVeriNoisy, by the France 2030 program, managed by the French National Research Agency under grant agreement No. ANR-22-PETQ-0008 PQ-TLS, and by the European Research Council (ERC) Consolidator Grant 101044770 (CRYPTOLAYER) and Advanced Grant 101096871 (BRIDGE) under the European Union’s Horizon 2020 and Horizon Europe research and innovation programs. Views and opinions expressed are those of the authors and do not necessarily reflect those of the European Union or the ERC. Neither the European Union nor the granting authority can be held responsible for them.

References

1. Dorian Amiet, Andreas Curiger, Lukas Leuenberger, and Paul Zbinden. Defeating NewHope with a Single Trace. In *PQCrypto*, volume 12100 of *Lecture Notes in Computer Science*, pages 189–205. Springer, 2020.
2. Guilh  m Assael, Philippe Elbaz-Vincent, and Guillaume Reymond. Improving Single-Trace Attacks on the Number-Theoretic Transform for Cortex-M4. In *HOST*, pages 111–121. IEEE, 2023.
3. Melissa Azouaoui, Olivier Bronchain, Ga  tan Cassiers, Cl  ment Hoffmann, Yulia Kuzovkova, Joost Renes, Tobias Schneider, Markus Sch  nauer, Fran  ois-Xavier Standaert, and Christine van Vredendaal. Protecting dilithium against leakage revisited sensitivity analysis and improved implementations. *IACR Trans. Cryptogr. Hardw. Embed. Syst.*, 2023(4):58–79, 2023.
4. Josep Balasch, Sebastian Faust, and Benedikt Gierlichs. Inner product masking revisited. In *EUROCRYPT (1)*, volume 9056 of *Lecture Notes in Computer Science*, pages 486–510. Springer, 2015.
5. Josep Balasch, Sebastian Faust, Benedikt Gierlichs, Clara Paglialonga, and Fran  ois-Xavier Standaert. Consolidating inner product masking. In *ASIACRYPT*

- (1), volume 10624 of *Lecture Notes in Computer Science*, pages 724–754. Springer, 2017.
6. Josep Balasch, Sebastian Faust, Benedikt Gierlichs, and Ingrid Verbauwhede. Theory and practice of a leakage resilient masking scheme. In *ASIACRYPT*, volume 7658 of *Lecture Notes in Computer Science*, pages 758–775. Springer, 2012.
7. Josep Balasch, Benedikt Gierlichs, Vincent Grosso, Oscar Reparaz, and François-Xavier Standaert. On the cost of lazy engineering for masked software implementations. In *CARDIS*, volume 8968 of *Lecture Notes in Computer Science*, pages 64–81. Springer, 2014.
8. Gilles Barthe, Sonia Belaïd, François Dupressoir, Pierre-Alain Fouque, Benjamin Grégoire, and Pierre-Yves Strub. Verified proofs of higher-order masking. In Elisabeth Oswald and Marc Fischlin, editors, *Advances in Cryptology - EUROCRYPT 2015 - 34th Annual International Conference on the Theory and Applications of Cryptographic Techniques, Sofia, Bulgaria, April 26-30, 2015, Proceedings, Part I*, volume 9056 of *Lecture Notes in Computer Science*, pages 457–485. Springer, 2015.
9. Gilles Barthe, Sonia Belaïd, François Dupressoir, Pierre-Alain Fouque, Benjamin Grégoire, Pierre-Yves Strub, and Rébecca Zucchini. Strong non-interference and type-directed higher-order masking. In Edgar R. Weippl, Stefan Katzenbeisser, Christopher Kruegel, Andrew C. Myers, and Shai Halevi, editors, *Proceedings of the 2016 ACM SIGSAC Conference on Computer and Communications Security, Vienna, Austria, October 24-28, 2016*, pages 116–129. ACM, 2016.
10. Gilles Barthe, Sonia Belaïd, Thomas Espitau, Pierre-Alain Fouque, Benjamin Grégoire, Mélissa Rossi, and Mehdi Tibouchi. Masking the GLP lattice-based signature scheme at any order. *J. Cryptol.*, 37(1):5, 2024.
11. Alberto Battistello, Jean-Sébastien Coron, Emmanuel Prouff, and Rina Zeitoun. Horizontal side-channel attacks and countermeasures on the ISW masking scheme. In *CHES*, volume 9813 of *Lecture Notes in Computer Science*, pages 23–39. Springer, 2016.
12. Sonia Belaïd, Jean-Sébastien Coron, Pierre-Alain Fouque, Benoît Gérard, Jean-Gabriel Kammerer, and Emmanuel Prouff. Improved side-channel analysis of finite-field multiplication. In *CHES*, volume 9293 of *Lecture Notes in Computer Science*, pages 395–415. Springer, 2015.
13. Sonia Belaïd, Pierre-Alain Fouque, and Benoît Gérard. Side-channel analysis of multiplications in $\text{GF}(2^{128})$ - application to AES-GCM. In *ASIACRYPT (2)*, volume 8874 of *Lecture Notes in Computer Science*, pages 306–325. Springer, 2014.
14. Davide Bellizia, Olivier Bronchain, Gaëtan Cassiers, Vincent Grosso, Chun Guo, Charles Momin, Olivier Pereira, Thomas Peters, and François-Xavier Standaert. Mode-level vs. implementation-level physical security in symmetric cryptography - A practical guide through the leakage-resistance jungle. In *CRYPTO (1)*, volume 12170 of *Lecture Notes in Computer Science*, pages 369–400. Springer, 2020.
15. Fabrice Benhamouda, Akshay Degwekar, Yuval Ishai, and Tal Rabin. On the local leakage resilience of linear secret sharing schemes. In *CRYPTO (1)*, volume 10991 of *Lecture Notes in Computer Science*, pages 531–561. Springer, 2018.
16. Fabrice Benhamouda, Akshay Degwekar, Yuval Ishai, and Tal Rabin. On the local leakage resilience of linear secret sharing schemes. *J. Cryptol.*, 34(2):10, 2021.
17. Joppe W. Bos, Marc Gourjon, Joost Renes, Tobias Schneider, and Christine van Vredendaal. Masking kyber: First- and higher-order implementations. *IACR Trans. Cryptogr. Hardw. Embed. Syst.*, 2021(4):173–214, 2021.
18. Gianluca Brian, Antonio Faonio, Maciej Obremski, João Ribeiro, Mark Simkin, Maciej Skórski, and Daniele Venturi. The mother of all leakages: How to simulate

- noisy leakages via bounded leakage (almost) for free. In *EUROCRYPT (2)*, volume 12697 of *Lecture Notes in Computer Science*, pages 408–437. Springer, 2021.
19. Olivier Bronchain and Gaëtan Cassiers. Bitslicing arithmetic/boolean masking conversions for fun and profit with application to lattice-based kems. *IACR Trans. Cryptogr. Hardw. Embed. Syst.*, 2022(4):553–588, 2022.
 20. Olivier Bronchain and François-Xavier Standaert. Breaking masked implementations with many shares on 32-bit software platforms or when the security order does not matter. *IACR Trans. Cryptogr. Hardw. Embed. Syst.*, 2021(3):202–234, 2021.
 21. Gaëtan Cassiers, Benjamin Grégoire, Itamar Levi, and François-Xavier Standaert. Hardware private circuits: From trivial composition to full verification. *IEEE Trans. Computers*, 70(10):1677–1690, 2021.
 22. Gaëtan Cassiers and François-Xavier Standaert. Trivially and efficiently composing masked gadgets with probe isolating non-interference. *IEEE Trans. Inf. Forensics Secur.*, 15:2542–2555, 2020.
 23. Gaëtan Cassiers and François-Xavier Standaert. Provably secure hardware masking in the transition- and glitch-robust probing model: Better safe than sorry. *IACR Trans. Cryptogr. Hardw. Embed. Syst.*, 2021(2):136–158, 2021.
 24. Suresh Chari, Josyula R. Rao, and Pankaj Rohatgi. Template attacks. In Burton S. Kaliski Jr., Çetin Kaya Koç, and Christof Paar, editors, *Cryptographic Hardware and Embedded Systems - CHES 2002, 4th International Workshop, Redwood Shores, CA, USA, August 13-15, 2002, Revised Papers*, volume 2523 of *Lecture Notes in Computer Science*, pages 13–28. Springer, 2002.
 25. Wei Cheng, Sylvain Guilley, Claude Carlet, Sihem Mesnager, and Jean-Luc Danger. Optimizing inner product masking scheme by a coding theory approach. *IEEE Trans. Inf. Forensics Secur.*, 16:220–235, 2021.
 26. Thomas De Cnudde, Oscar Reparaz, Begül Bilgin, Svetla Nikova, Ventzislav Nikov, and Vincent Rijmen. Masking AES with $d+1$ shares in hardware. In *CHES*, volume 9813 of *Lecture Notes in Computer Science*, pages 194–212. Springer, 2016.
 27. Jean-Sébastien Coron, Christophe Giraud, Emmanuel Prouff, Soline Renner, Matthieu Rivain, and Praveen Kumar Vadnala. Conversion of security proofs from one leakage model to another: A new issue. In *COSADE*, volume 7275 of *Lecture Notes in Computer Science*, pages 69–81. Springer, 2012.
 28. Joan Daemen and Vincent Rijmen. *The Design of Rijndael - The Advanced Encryption Standard (AES), Second Edition*. Information Security and Cryptography. Springer, 2020.
 29. Christoph Dobraunig, Maria Eichlseder, Stefan Mangard, Florian Mendel, Bart Mennink, Robert Primas, and Thomas Unterluggauer. Isap v2.0. *IACR Trans. Symmetric Cryptol.*, 2020(S1):390–416, 2020.
 30. Christoph Dobraunig, Maria Eichlseder, Florian Mendel, and Martin Schläffer. Ascon v1.2: Lightweight authenticated encryption and hashing. *J. Cryptol.*, 34(3):33, 2021.
 31. Christoph Dobraunig, François Koeune, Stefan Mangard, Florian Mendel, and François-Xavier Standaert. Towards fresh and hybrid re-keying schemes with beyond birthday security. In *CARDIS*, volume 9514 of *Lecture Notes in Computer Science*, pages 225–241. Springer, 2015.
 32. Yevgeniy Dodis, Rafail Ostrovsky, Leonid Reyzin, and Adam Smith. Fuzzy extractors: How to generate strong keys from biometrics and other noisy data. *SIAM Journal on Computing*, 38(1):97–139, January 2008.

33. Alexandre Duc, Stefan Dziembowski, and Sebastian Faust. Unifying leakage models: From probing attacks to noisy leakage. In *EUROCRYPT*, volume 8441 of *Lecture Notes in Computer Science*, pages 423–440. Springer, 2014.
34. Alexandre Duc, Stefan Dziembowski, and Sebastian Faust. Unifying leakage models: From probing attacks to noisy leakage. *J. Cryptol.*, 32(1):151–177, 2019.
35. Alexandre Duc, Sebastian Faust, and François-Xavier Standaert. Making masking security proofs concrete - or how to evaluate the security of any leaking device. In *EUROCRYPT (1)*, volume 9056 of *Lecture Notes in Computer Science*, pages 401–429. Springer, 2015.
36. Sébastien Duval, Pierrick Méaux, Charles Momin, and François-Xavier Standaert. Exploring crypto-physical dark matter and learning with physical rounding towards secure and efficient fresh re-keying. *IACR Trans. Cryptogr. Hardw. Embed. Syst.*, 2021(1):373–401, 2021.
37. Stefan Dziembowski and Sebastian Faust. Leakage-resilient cryptography from the inner-product extractor. In *ASIACRYPT*, volume 7073 of *Lecture Notes in Computer Science*, pages 702–721. Springer, 2011.
38. Stefan Dziembowski, Sebastian Faust, Gottfried Herold, Anthony Journault, Daniel Masny, and François-Xavier Standaert. Towards sound fresh re-keying with hard (physical) learning problems. In *CRYPTO (2)*, volume 9815 of *Lecture Notes in Computer Science*, pages 272–301. Springer, 2016.
39. Sebastian Faust, Vincent Grosso, Santos Merino Del Pozo, Clara Paglialonga, and François-Xavier Standaert. Composable masking schemes in the presence of physical defaults & the robust probing model. *IACR Trans. Cryptogr. Hardw. Embed. Syst.*, 2018(3):89–120, 2018.
40. Sebastian Faust, Loïc Masure, Elena Micheli, Maximilian Ortl, and François-Xavier Standaert. Connecting leakage-resilient secret sharing to practice: Scaling trends and physical dependencies of prime field masking. In *EUROCRYPT (4)*, volume 14654 of *Lecture Notes in Computer Science*, pages 316–344. Springer, 2024.
41. Guillaume Fumaroli, Ange Martinelli, Emmanuel Prouff, and Matthieu Rivain. Affine masking against higher-order side channel analysis. In *Selected Areas in Cryptography*, volume 6544 of *Lecture Notes in Computer Science*, pages 262–280. Springer, 2010.
42. Jovan Dj. Golic and Christophe Tymen. Multiplicative masking and power analysis of AES. In *CHES*, volume 2523 of *Lecture Notes in Computer Science*, pages 198–212. Springer, 2002.
43. Louis Goubin and Jacques Patarin. DES and differential power analysis (the "duplication" method). In *CHES*, volume 1717 of *Lecture Notes in Computer Science*, pages 158–172. Springer, 1999.
44. Dahmun Goudarzi and Matthieu Rivain. How fast can higher-order masking be in software? In *EUROCRYPT (1)*, volume 10210 of *Lecture Notes in Computer Science*, pages 567–597, 2017.
45. Hannes Groß and Stefan Mangard. Reconciling $d+1$ masking in hardware and software. In *CHES*, volume 10529 of *Lecture Notes in Computer Science*, pages 115–136. Springer, 2017.
46. Hannes Groß and Stefan Mangard. A unified masking approach. *J. Cryptogr. Eng.*, 8(2):109–124, 2018.
47. Hannes Groß, Stefan Mangard, and Thomas Korak. An efficient side-channel protected AES implementation with arbitrary protection order. In *CT-RSA*, volume 10159 of *Lecture Notes in Computer Science*, pages 95–112. Springer, 2017.

48. Qian Guo, Vincent Grosso, François-Xavier Standaert, and Olivier Bronchain. Modeling soft analytical side-channel attacks from a coding theory viewpoint. *IACR Trans. Cryptogr. Hardw. Embed. Syst.*, 2020(4):209–238, 2020.
49. Mike Hamburg, Julius Hermelink, Robert Primas, Simona Samardjiska, Thomas Schamberger, Silvan Streit, Emanuele Strieder, and Christine van Vredendaal. Chosen Ciphertext k-Trace Attacks on Masked CCA2 Secure Kyber. *IACR Trans. Cryptogr. Hardw. Embed. Syst.*, 2021(4):88–113, 2021.
50. Christoph Herbst, Elisabeth Oswald, and Stefan Mangard. An AES smart card implementation resistant to power analysis attacks. In *ACNS*, volume 3989 of *Lecture Notes in Computer Science*, pages 239–252, 2006.
51. Julius Hermelink, Silvan Streit, Emanuele Strieder, and Katharina Thieme. Adapting Belief Propagation to Counter Shuffling of NTTs. *IACR Trans. Cryptogr. Hardw. Embed. Syst.*, 2023(1):60–88, 2023.
52. Clément Hoffmann, Benoît Libert, Charles Momin, Thomas Peters, and François-Xavier Standaert. POLKA: towards leakage-resistant post-quantum cca-secure public key encryption. In *Public Key Cryptography (1)*, volume 13940 of *Lecture Notes in Computer Science*, pages 114–144. Springer, 2023.
53. Russell Impagliazzo, Leonid A. Levin, and Michael Luby. Pseudo-random generation from one-way functions (extended abstracts). In *STOC*, pages 12–24. ACM, 1989.
54. Yuval Ishai, Amit Sahai, and David A. Wagner. Private circuits: Securing hardware against probing attacks. In *CRYPTO*, volume 2729 of *Lecture Notes in Computer Science*, pages 463–481. Springer, 2003.
55. Ohad Klein and Ilan Komargodski. New bounds on the local leakage resilience of shamir’s secret sharing scheme. In *CRYPTO (1)*, volume 14081 of *Lecture Notes in Computer Science*, pages 139–170. Springer, 2023.
56. David Knichel, Amir Moradi, Nicolai Müller, and Pascal Sasdrich. Automated generation of masked hardware. *IACR Trans. Cryptogr. Hardw. Embed. Syst.*, 2022(1):589–629, 2022.
57. Vadim Lyubashevsky, Léo Ducas, Eike Kiltz, Tancrede Lepoint, Peter Schwabe, Gregor Seiler, Damien Stehlé, and Shi Bai. CRYSTALS-Dilithium Algorithm Specifications and Supporting Documentation. *NIST Post-Quantum Cryptography Standard*, 2022.
58. Hemanta K. Maji, Hai H. Nguyen, Anat Paskin-Cherniavsky, Tom Suad, and Mingyuan Wang. Leakage-resilience of the shamir secret-sharing scheme against physical-bit leakages. In *EUROCRYPT (2)*, volume 12697 of *Lecture Notes in Computer Science*, pages 344–374. Springer, 2021.
59. Hemanta K. Maji, Hai H. Nguyen, Anat Paskin-Cherniavsky, Tom Suad, Mingyuan Wang, Xiuyu Ye, and Albert Yu. Leakage-resilient linear secret-sharing against arbitrary bounded-size leakage family. In *TCC (1)*, volume 13747 of *Lecture Notes in Computer Science*, pages 355–383. Springer, 2022.
60. Hemanta K. Maji, Hai H. Nguyen, Anat Paskin-Cherniavsky, Tom Suad, Mingyuan Wang, Xiuyu Ye, and Albert Yu. Tight estimate of the local leakage resilience of the additive secret-sharing scheme & its consequences. In *ITC*, volume 230 of *LIPICs*, pages 16:1–16:19. Schloss Dagstuhl - Leibniz-Zentrum für Informatik, 2022.
61. Hemanta K. Maji, Hai H. Nguyen, Anat Paskin-Cherniavsky, and Mingyuan Wang. Improved bound on the local leakage-resilience of shamir’s secret sharing. In *ISIT*, pages 2678–2683. IEEE, 2022.
62. Hemanta K. Maji, Hai H. Nguyen, Anat Paskin-Cherniavsky, and Xiuyu Ye. Constructing leakage-resilient shamir’s secret sharing: Over composite order fields. In

- EUROCRYPT (4)*, volume 14654 of *Lecture Notes in Computer Science*, pages 286–315. Springer, 2024.
63. Hemanta K. Maji, Anat Paskin-Cherniavsky, Tom Suad, and Mingyuan Wang. Constructing locally leakage-resilient linear secret-sharing schemes. In *CRYPTO (3)*, volume 12827 of *Lecture Notes in Computer Science*, pages 779–808. Springer, 2021.
 64. Stefan Mangard, Elisabeth Oswald, and Thomas Popp. *Power analysis attacks - revealing the secrets of smart cards*. Springer, 2007.
 65. Stefan Mangard, Thomas Popp, and Berndt M. Gammel. Side-channel leakage of masked CMOS gates. In *CT-RSA*, volume 3376 of *Lecture Notes in Computer Science*, pages 351–365. Springer, 2005.
 66. Stefan Mangard, Norbert Pramstaller, and Elisabeth Oswald. Successfully attacking masked AES hardware implementations. In *CHES*, volume 3659 of *Lecture Notes in Computer Science*, pages 157–171. Springer, 2005.
 67. Loïc Masure, Pierrick Méaux, Thorben Moos, and François-Xavier Standaert. Effective and efficient masking with low noise using small-mersenne-prime ciphers. In *EUROCRYPT (4)*, volume 14007 of *Lecture Notes in Computer Science*, pages 596–627. Springer, 2023.
 68. Marcel Medwed, Christophe Petit, Francesco Regazzoni, Mathieu Renaud, and François-Xavier Standaert. Fresh re-keying II: securing multiple parties against side-channel and fault attacks. In *CARDIS*, volume 7079 of *Lecture Notes in Computer Science*, pages 115–132. Springer, 2011.
 69. Marcel Medwed, François-Xavier Standaert, Johann Großschädl, and Francesco Regazzoni. Fresh re-keying: Security against side-channel and fault attacks for low-cost devices. In *AFRICACRYPT*, volume 6055 of *Lecture Notes in Computer Science*, pages 279–296. Springer, 2010.
 70. Bart Mennink. Beyond birthday bound secure fresh rekeying: Application to authenticated encryption. In *ASIACRYPT (1)*, volume 12491 of *Lecture Notes in Computer Science*, pages 630–661. Springer, 2020.
 71. Vincent Migliore, Benoît Gérard, Mehdi Tibouchi, and Pierre-Alain Fouque. Masking dilithium - efficient implementation and side-channel evaluation. In *ACNS*, volume 11464 of *Lecture Notes in Computer Science*, pages 344–362. Springer, 2019.
 72. Svetla Nikova, Vincent Rijmen, and Martin Schl  ffer. Secure hardware implementation of nonlinear functions in the presence of glitches. *J. Cryptol.*, 24(2):292–321, 2011.
 73. Maciej Obremski, Jo  o Ribeiro, Lawrence Roy, Fran  ois-Xavier Standaert, and Daniele Venturi. Improved reductions from noisy to bounded and probing leakages via hockey-stick divergences. In *CRYPTO (6)*, volume 14925 of *Lecture Notes in Computer Science*, pages 461–491. Springer, 2024.
 74. Peter Pessl and Robert Primas. More Practical Single-Trace Attacks on the Number Theoretic Transform. In *LATINCRYPT*, volume 11774 of *Lecture Notes in Computer Science*, pages 130–149. Springer, 2019.
 75. Romain Poussier, Qian Guo, Fran  ois-Xavier Standaert, Claude Carlet, and Sylvain Guilley. Connecting and improving direct sum masking and inner product masking. In *CARDIS*, volume 10728 of *Lecture Notes in Computer Science*, pages 123–141. Springer, 2017.
 76. Robert Primas, Peter Pessl, and Stefan Mangard. Single-Trace Side-Channel Attacks on Masked Lattice-Based Encryption. In *CHES*, volume 10529 of *Lecture Notes in Computer Science*, pages 513–533. Springer, 2017.

77. Emmanuel Prouff and Matthieu Rivain. Masking against side-channel attacks: A formal security proof. In *EUROCRYPT*, volume 7881 of *Lecture Notes in Computer Science*, pages 142–159. Springer, 2013.
78. Prasanna Ravi, Anupam Chattopadhyay, Jan-Pieter D’Anvers, and Anubhab Baksi. Side-channel and fault-injection attacks over lattice-based post-quantum schemes (kyber, dilithium): Survey and new results. *ACM Trans. Embed. Comput. Syst.*, 23(2):35:1–35:54, 2024.
79. Matthieu Rivain. On the provable security of cryptographic implementations : Habilitation thesis. Personal website, 2022. <https://www.matthieurivain.com/hdr.html>.
80. Tobias Schneider, Clara Paglialonga, Tobias Oder, and Tim Güneysu. Efficiently masking binomial sampling at arbitrary orders for lattice-based crypto. In *Public Key Cryptography (2)*, volume 11443 of *Lecture Notes in Computer Science*, pages 534–564. Springer, 2019.
81. Peter Schwabe, Roberto Avanzi, Joppe Bos, Léo Ducas, Eike Kiltz, Tancrede Lepoint, Vadim Lyubashevsky, John M Schanck, Gregor Seiler, Damien Stehlé, and Jintai Ding. CRYSTALS-Kyber Algorithm Specifications and Supporting Documentation. *NIST Post-Quantum Cryptography Standard*, 2022.
82. Balazs Udvarhelyi, Olivier Bronchain, and François-Xavier Standaert. Security analysis of deterministic re-keying with masking and shuffling: Application to ISAP. In *COSADE*, volume 12910 of *Lecture Notes in Computer Science*, pages 168–183. Springer, 2021.
83. Nicolas Veyrat-Charvillon, Marcel Medwed, Stéphanie Kerckhof, and François-Xavier Standaert. Shuffling against side-channel attacks: A comprehensive study with cautionary note. In *ASIACRYPT*, volume 7658 of *Lecture Notes in Computer Science*, pages 740–757. Springer, 2012.
84. Weijia Wang, François-Xavier Standaert, Yu Yu, Sihang Pu, Junrong Liu, Zheng Guo, and Dawu Gu. Inner product masking for bitslice ciphers and security order amplification for linear leakages. In *CARDIS*, volume 10146 of *Lecture Notes in Computer Science*, pages 174–191. Springer, 2016.
85. Yu Yu, François-Xavier Standaert, Olivier Pereira, and Moti Yung. Practical leakage-resilient pseudorandom generators. In *CCS*, pages 141–151. ACM, 2010.

A Conversion Algorithms

In this section, we discuss conversion algorithms that enable transformations between various IP encodings. This includes conversions not only within the same underlying field but also between distinct fields, such as algebraic and Boolean extension fields.

Algorithm 1 $\text{IPconversion}_{\vec{\omega}, \vec{\omega}'}^d$ (NI)

Input: $\llbracket y \rrbracket_{\vec{\omega}} = (y_1, \dots, y_d) \in \mathbb{F}_p^d$

Output: $\llbracket y \rrbracket_{\vec{\omega}'} = (y'_1, \dots, y'_d) \in \mathbb{F}_p^d$

- 1: **for** $i = 1, \dots, d - 1$ **do**
 - 2: $y'_i \leftarrow (\omega_i \cdot \omega_i'^{-1}) \cdot y_i$ ▷ All $\omega_i \cdot \omega_i'^{-1}$ are public and fix values.
 - 3: **end for**
-

Lemma 9. *The gadget $\text{IPconversion}_{\vec{\omega}, \vec{\omega}'}^d$ is correct.*

Proof. The gadget is correct as it holds that

$$\langle \vec{\omega}', (y'_1, \dots, y'_d) \rangle = \sum_{i=1}^d \omega'_i \cdot y'_i = \sum_{i=1}^d \omega'_i \cdot (\omega_i \cdot \omega_i'^{-1}) \cdot y_i = \sum_{i=1}^d \omega_i \cdot y_i = \langle \vec{\omega}, (y_1, \dots, y_d) \rangle.$$

□

Further, the gadget operates in a share-wise manner, i.e, each output share depends only on the corresponding input share with the same index. This preserves the isolation of the shares in the circuit when propagating probes within the gadget, as specified in [22]. Therefore, the gadget fulfills the requirements of both NI [8,9] and PINI [22].

Note that this gadget also enables conversion between simple additive and IP encodings. In particular, we can set $\vec{\omega} = (1, \dots, 1)$ to convert an additive encoding into an IP encoding, yielding $\llbracket y \rrbracket_{\vec{\omega}}$. Conversely, to transform an IP encoding $\llbracket y \rrbracket_{\vec{\omega}}$ into an additive encoding, we simply set $\vec{\omega}' = (1, \dots, 1)$.

The gadget $\text{IPconversion}_{\vec{\omega}, \vec{\omega}'}^d$ is quite generic and works over any field, as long as the input and output IP encodings are defined over the same field. However, we can also use this gadget to build more advanced constructions that convert between encodings defined over different underlying fields. For instance, we can compose this gadget with existing additive-to-Boolean (A2B), e.g., [80,19] and Boolean-to-additive (B2A) [10,19] conversion gadgets, denoted $\mathbf{G}_{\text{A2B}}$ and $\mathbf{G}_{\text{B2A}}$, to support field conversion for IP masking.

For instance, we can construct the following composed gadget:

$$\mathbf{G}_{\text{A2B}}^{\omega'', \omega'}(\cdot) = \text{IPconversion}_{\vec{\omega}, \vec{\omega}'}^d(\mathbf{G}_{\text{A2B}}(\text{IPconversion}_{\vec{\omega}', \vec{\omega}'''}^d(\cdot)))$$

with $\vec{\omega} = \vec{\omega}''' = (1, \dots, 1)$, to build a gadget that transforms an ω'' -IP encoding over an arithmetic field into an ω' -IP encoding over a Boolean field. A similar construction applies for:

$$\mathbf{G}_{\text{B2A}}^{\omega'', \omega'}(\cdot) = \text{IPconversion}_{\vec{\omega}, \vec{\omega}'}^d(\mathbf{G}_{\text{B2A}}(\text{IPconversion}_{\vec{\omega}', \vec{\omega}'''}^d(\cdot)))$$

We keep the security discussion somewhat high-level, since the focus of this paper is on the IP encoding itself. Nevertheless, we note that our IP conversion gadget is a share-wise operating gadget, with only one input and output encoding. Consequently, the standard security properties—such as PINI and (S)NI—that apply to the well-studied $\mathbf{G}_{\text{A2B}}$ and $\mathbf{G}_{\text{B2A}}$ gadgets also hold for $\mathbf{G}_{\text{A2B}}^{\omega'', \omega'}(\cdot)$ and $\mathbf{G}_{\text{B2A}}^{\omega'', \omega'}(\cdot)$. Or in other words the only additional values that pop up in the new gadgets in $\mathbf{G}_{\text{A2B}}^{\omega'', \omega'}(\cdot)$ and $\mathbf{G}_{\text{B2A}}^{\omega'', \omega'}(\cdot)$ are the output and input probes of $\mathbf{G}_{\text{A2B}}^{\omega'', \omega'}(\cdot)$ and $\mathbf{G}_{\text{B2A}}^{\omega'', \omega'}(\cdot)$ multiplied by a public constant.

B Leftover Hash Lemma

Definition 11. *The collision probability $\text{CP}(X)$ of a random variable X is defined as the probability that two independent samples of X are equal, i.e.,*

$$\text{CP}(X) = \Pr(X = X')$$

where X' is an independent copy of X .

Proposition 10 (Simple Properties of the Collision Probability). *Let X be a random variable over $\mathcal{X}$.*

1. *Let p_X be the probability mass function of X . Then $\text{CP}(X) = \|p_X\|_2^2$.*
2. *$\text{CP}(X) \leq 2^{-H_\infty(X)}$.*

Proof. item 1 follows from the chain of equalities below.

$$\text{CP}(X) = \mathbb{P}_{X \sim X'}[X = X'] = \sum_{x \in \mathcal{X}} \mathbb{P}[X = x]^2 = \sum_{x \in \mathcal{X}} p_X(x)^2 = \|p_X\|_2^2.$$

We prove item 2 as follows

$$\text{CP}(X) = \sum_{x \in \mathcal{X}} \mathbb{P}[X = x]^2 \leq \max_{x \in \mathcal{X}} \mathbb{P}[X = x] \cdot \sum_{x \in \mathcal{X}} \mathbb{P}[X = x] = 2^{-H_\infty(X)}.$$

□

Proof of Theorem 1. The proof¹³ is divided into three main parts. First, we upper bound the collision probability of the extractor. Next, we use this to bound the euclidean distance between the extractor's output and true randomness. Finally, we convert the euclidean distance into total variation.

Part 1. The collision probability for the joint random variable $(h_\Omega(X), \Omega)$ is

$$\begin{aligned} \text{CP}(h_\Omega(X), \Omega) &= \Pr(h_\Omega(X) = h_{\Omega'}(X'), \Omega = \Omega') \\ &= \Pr(\Omega = \Omega') \cdot \Pr(h_\Omega(X) = h_{\Omega'}(X') \mid \Omega = \Omega') \end{aligned}$$

Let us elaborate on the second factor of the right-hand side, that we rephrase as $\Pr(h_\Omega(X) = h_\Omega(X'))$. The total probability formula tells us that

$$\begin{aligned} \Pr(h_\Omega(X) = h_\Omega(X')) &= \Pr(h_\Omega(X) = h_\Omega(X') \mid X \neq X') \cdot \Pr(X \neq X') \\ &\quad + \Pr(h_\Omega(X) = h_\Omega(X') \mid X = X') \cdot \Pr(X = X') . \end{aligned}$$

Since $\mathcal{H}$ is universal, we may bound the first term in the previous sum as follows:

$$\Pr(h_\Omega(X) = h_\Omega(X') \mid X \neq X') \cdot \Pr(X' \neq X) \leq \frac{1}{|\mathcal{Y}|} .$$

¹³ The proof comes from Reyzin's lecture notes, we have just adapted the notation.

As per the second term of the sum, it can be trivially upper-bounded by the collision probability $\text{CP}(X)$, which in turn can be upper bounded using its min entropy (Proposition 10, item 2). As a result,

$$\text{CP}(h_\Omega(X), \Omega) \leq \frac{1}{|\mathcal{W}|} \left(\frac{1}{|\mathcal{Y}|} + \frac{1}{2^k} \right) = \frac{1 + 4\epsilon^2}{|\mathcal{W}| \cdot |\mathcal{Y}|} \quad .$$

Here, the equality comes from defining $\epsilon = |\mathcal{Y}|^{\frac{1}{2}} \cdot 2^{-\frac{k}{2}-1}$.

Part 2. Let ν be the vector corresponding to the difference between the probability distributions of $(h_\Omega(X), \Omega)$ and $(U_{\mathcal{Y}}, U_{\mathcal{W}})$. This means that

$$\begin{aligned} \|\nu\|_2^2 &= \|\Pr(h_\Omega(X), \Omega) - \Pr(U_{\mathcal{Y}}, U_{\mathcal{W}})\|_2^2 \\ &= \sum_{y, \omega} \left(\Pr(h_\Omega(X) = y, \Omega = \vec{\omega}) - \frac{1}{|\mathcal{Y}|} \cdot \frac{1}{|\mathcal{W}|} \right)^2 \\ &= \sum_{y, \omega} \Pr(h_\Omega(X) = y, \Omega = \vec{\omega})^2 \\ &\quad - 2 \cdot \sum_{y, \omega} \Pr(h_\Omega(X) = y, \Omega = \vec{\omega}) \cdot \frac{1}{|\mathcal{Y}|} \cdot \frac{1}{|\mathcal{W}|} \\ &\quad + \sum_{y, \omega} \left(\frac{1}{|\mathcal{Y}|} \cdot \frac{1}{|\mathcal{W}|} \right)^2 \\ &= \text{CP}(h_\Omega(X), \Omega) - \frac{2}{|\mathcal{Y}| \cdot |\mathcal{W}|} + \frac{1}{|\mathcal{Y}| \cdot |\mathcal{W}|} \\ &\leq \frac{4 \cdot \epsilon^2}{|\mathcal{Y}| \cdot |\mathcal{W}|} \quad , \end{aligned}$$

where the first two equalities hold by definition of norm 2, the third equality is obtained by distributing the square for each term of the sum, the fourth equality holds by Proposition 10, item 1, and the last inequality follows from part 1.

Part 3. Let u_i be the sign of the i -th entry of the vector ν , *i.e.*, such that ν can be expressed as $\nu = (\dots, u_i \cdot |\nu_i|, \dots)$. Then Cauchy-Schwarz inequality states that

$$\|\nu\|_1 = \langle \nu, u \rangle \leq \|\nu\|_2 \cdot \|u\|_2 \quad .$$

Given that $\|\nu\|_2 \leq \frac{2 \cdot \epsilon}{\sqrt{|\mathcal{Y}| \cdot |\mathcal{W}|}}$, and $\|u\|_2 = \sqrt{|\mathcal{Y}| \cdot |\mathcal{W}|}$, we get that

$$\text{TV}(\Pr(h_\Omega(X), \Omega); \Pr(U_{\mathcal{Y}}, U_{\mathcal{W}})) = \frac{1}{2} \|\nu\|_1 \leq \epsilon \quad ,$$

which completes the proof. $\square$

C Reduction from Non-Uniform to Uniform Secrets

Proof of Proposition 1. The first inequality is proven by using Jensen's inequality, thanks to the convexity of the total variation, and by bounding the expectation by the supremum. The second inequality is proven as follows: let $\vec{\ell} \in \mathcal{L}$, and denote $\Delta_{\vec{\ell}} = \left| \Pr(\mathbf{L}(y^{(0)}) = \vec{\ell} \mid \vec{\Omega} = \vec{\omega}) - \Pr(\mathbf{L}(y^{(1)}) = \vec{\ell} \mid \vec{\Omega} = \vec{\omega}) \right|$ for short. Then we have

$$\begin{aligned} \Delta_{\vec{\ell}} &= \left| \Pr(\mathbf{L} = \vec{\ell} \mid \mathbf{Y} = y^{(0)}, \vec{\Omega} = \vec{\omega}) - \Pr(\mathbf{L} = \vec{\ell} \mid \mathbf{Y} = y^{(1)}, \vec{\Omega} = \vec{\omega}) \right| \\ &= |\mathbb{F}| \cdot \left| \Pr(\mathbf{L} = \vec{\ell}, \mathbf{Y} = y^{(0)} \mid \vec{\Omega} = \vec{\omega}) - \Pr(\mathbf{L} = \vec{\ell}, \mathbf{Y} = y^{(1)} \mid \vec{\Omega} = \vec{\omega}) \right| \\ &\leq |\mathbb{F}| \cdot \left| \Pr(\mathbf{L} = \vec{\ell}, \mathbf{Y} = y^{(0)} \mid \vec{\Omega} = \vec{\omega}) - \Pr(\mathbf{L} = \vec{\ell} \mid \vec{\Omega} = \vec{\omega}) \cdot \Pr(\mathbf{Y} = y^{(0)}) \right| \\ &\quad + |\mathbb{F}| \cdot \left| \Pr(\mathbf{L} = \vec{\ell} \mid \vec{\Omega} = \vec{\omega}) \cdot \Pr(\mathbf{Y} = y^{(1)}) - \Pr(\mathbf{L} = \vec{\ell}, \mathbf{Y} = y^{(1)} \mid \vec{\Omega} = \vec{\omega}) \right| \\ &\leq |\mathbb{F}| \cdot \sum_{y \in \mathbb{F}} \left| \Pr(\mathbf{L} = \vec{\ell}, \mathbf{Y} = y \mid \vec{\Omega} = \vec{\omega}) - \Pr(\mathbf{L} = \vec{\ell} \mid \vec{\Omega} = \vec{\omega}) \cdot \Pr(\mathbf{Y} = y) \right| \end{aligned}$$

Here, the first equality holds by making the knowledge of $y^{(0)}, y^{(1)}$ explicit. The second equality holds by definition of the conditional probability, and using the fact that $\mathbf{Y}$ is considered uniform here. The first inequality comes from the triangle inequality, while the second inequality holds since we upper-bound two positive terms of a sum by the whole sum over $\mathbb{F}$. We conclude the proof by summing $\Delta_{\vec{\ell}}$ over $\mathcal{L}$. $\square$

D Invariance of Inner-Product Masking

In order to prove Proposition 2 and Corollary 1, we leverage the following results. The first one is a direct application of Lemma 9, and allows to convert an $\vec{\omega}$ -encoding into a $\vec{1}$ -encoding, and inversely.

Proposition 11. *For any fixed $y \in \mathbb{F}$, and any $\vec{\omega} \in (\mathbb{F}_p^*)^d$ the vector $\llbracket y \rrbracket_{\vec{1}} = (y_1, \dots, y_d)$ is a $\vec{1}$ -encoding of y if and only if the vector $\llbracket y \rrbracket_{\vec{\omega}} = (\omega_1^{-1} \cdot y_1, \dots, \omega_d^{-1} \cdot y_d)$ is an $\vec{\omega}$ -encoding of y .*

Proposition 11 straightforwardly implies the following corollary.

Corollary 5. *For any $y \in \mathbb{F}$, the random vector $\llbracket y \rrbracket_{\vec{\omega}}$ is uniformly distributed over $y + \text{Span}(\vec{\omega})^\perp$ if and only if $\llbracket y \rrbracket_{\vec{1}}$ is uniformly distributed over $y + \text{Span}(\vec{1})^\perp$.*

We are now ready to prove Proposition 2.

Proof of Proposition 2. Define the pointwise product $\psi_{\vec{\omega}} : \vec{x} \mapsto \vec{\omega} \odot \vec{x}$.¹⁴ Let $\llbracket y \rrbracket_{\vec{\omega}}$ be the uniform random vector over $y + \text{Span}(\vec{\omega})^\perp$. Then, by applying Corollary 5

¹⁴ $\psi_{\vec{\omega}}$ is invertible, and its inverse corresponds to $\psi_{\vec{\omega}^{-1}}$, where $\vec{\omega}^{-1} = (\omega_1^{-1}, \dots, \omega_d^{-1})$.

twice, the random vector $\llbracket y \rrbracket_{\vec{\omega}'} = \psi_{\vec{\omega}'}^{-1} \circ \psi_{\vec{\omega}}(\llbracket y \rrbracket_{\vec{\omega}})$ is uniform over $y + \text{Span}(\vec{\omega}')^\perp$. Moreover, for any function $\tau : \mathbb{F}^d \rightarrow \mathcal{L}$, we have

$$\tau(\llbracket y \rrbracket_{\vec{\omega}}) = \tau(\psi_{\vec{\omega}}^{-1} \circ \psi_{\vec{\omega}'}(\llbracket y \rrbracket_{\vec{\omega}'})) .$$

Define therefore $\tau' : \vec{x} \mapsto \tau \circ \psi_{\vec{\omega}}^{-1} \circ \psi_{\vec{\omega}'}(\vec{x})$. It follows that for any $\vec{\ell} \in \mathcal{L}$,

$$\mathbb{E}_{\llbracket y \rrbracket_{\vec{\omega}} \leftarrow y + \text{Span}(\vec{\omega})^\perp} \left[\Pr(\tau(\llbracket y \rrbracket_{\vec{\omega}}) = \vec{\ell}) \right] = \mathbb{E}_{\llbracket y \rrbracket_{\vec{\omega}'} \leftarrow y + \text{Span}(\vec{\omega}')^\perp} \left[\Pr(\tau'(\llbracket y \rrbracket_{\vec{\omega}'})) = \vec{\ell} \right] .$$

It remains to prove that τ' is m -bounded. This follows from the fact that τ and τ' have the same image space. Interestingly, it is clear from the definition of τ' that if τ verifies the independence assumption, then so does τ' . $\square$

Corollary 1 is a direct consequence of Proposition 2.

Proof of Corollary 1. By definition, an $\vec{\omega}$ encoding is ϵ -leakage resilient against m bounded leakage if and only if $\sup_{\tau \in \text{Bnd}_m} \mathbf{M}_{\vec{\omega}}(\mathbf{L}) \leq \epsilon$, where

$$\mathbf{M}_{\vec{\omega}}(\mathbf{L}) = \sup_{y^{(0)}, y^{(1)} \in \mathbb{F}} \text{TV} \left(\Pr(\tau(\text{Enc}_{\vec{\omega}}(y^{(0)}))) ; \Pr(\tau(\text{Enc}_{\vec{\omega}}(y^{(1)}))) \right) .$$

Let us now take τ maximizing $\mathbf{M}_{\vec{\omega}}(\mathbf{L})$. According to Proposition 2, there exists τ' such that

$$\begin{aligned} \Pr(\tau(\text{Enc}_{\vec{\omega}}(y^{(0)}))) &= \Pr(\tau'(\text{Enc}_{\vec{\Gamma}}(y^{(0)}))) \\ \Pr(\tau(\text{Enc}_{\vec{\omega}}(y^{(1)}))) &= \Pr(\tau'(\text{Enc}_{\vec{\Gamma}}(y^{(1)}))) . \end{aligned}$$

This proves that $\sup_{\tau \in \text{Bnd}_m} \mathbf{M}_{\vec{\omega}}(\mathbf{L}) \leq \sup_{\tau \in \text{Bnd}_m} \mathbf{M}_{\vec{\Gamma}}(\mathbf{L})$. To conclude the proof, it suffices to prove that the inequality holds in the opposite direction too, *i.e.* $\sup_{\tau \in \text{Bnd}_m} \mathbf{M}_{\vec{\omega}}(\mathbf{L}) \geq \sup_{\tau \in \text{Bnd}_m} \mathbf{M}_{\vec{\Gamma}}(\mathbf{L})$. We prove this true via a similar reasoning as above: we take τ maximizing $\mathbf{M}_{\vec{\Gamma}}(\mathbf{L})$, and use Proposition 2 to construct τ' such that

$$\begin{aligned} \Pr(\tau(\text{Enc}_{\vec{\Gamma}}(y^{(0)}))) &= \Pr(\tau'(\text{Enc}_{\vec{\omega}}(y^{(0)}))) \\ \Pr(\tau(\text{Enc}_{\vec{\Gamma}}(y^{(1)}))) &= \Pr(\tau'(\text{Enc}_{\vec{\omega}}(y^{(1)}))) . \end{aligned}$$

This proves that $\sup_{\tau \in \text{Bnd}_m} \mathbf{M}_{\vec{\omega}}(\mathbf{L}) \geq \sup_{\tau \in \text{Bnd}_m} \mathbf{M}_{\vec{\Gamma}}(\mathbf{L})$, and concludes the proof. $\square$

E Adaptive Leakage

Proof sketch of Corollary 3. In what follows, we outline the modifications needed in the proof of Lemma 1. First, we note that if we include the revelation of $\vec{\Omega}$

after the leakage, their metric becomes $\mathbb{E}_{\tilde{\Omega}} [M_{\tilde{\Omega}}(L)]$. When m equals 1, the step in their proof that weakens the strong extractor to a standard one (Lemma 21) becomes unnecessary. As a result, we can use a stronger version of their Lemma 22 where the extractor is strong, and the resulting security continues to hold even when $\tilde{\omega}$ is fully revealed at the end. We use different fields for the coefficients (they use $\mathbb{F}^d \setminus \{0^d\}$ while we have $(\mathbb{F} \setminus \{0\})^d$), but this has no impact on the security bound of Lemma 22, which is independent of it. For these reasons, the bound of Lemma 1 stays the same in this case too. To compare it with our bound, we first rewrite δ as a function of λ , yielding

$$\delta = \frac{1}{2} - \frac{\lambda - \log(\gamma)}{d \log |\mathbb{F}|}$$

$$\epsilon = 2 \cdot \left[|\mathbb{F}|^{\frac{3}{2} - \frac{d}{2}} \cdot \frac{2^\lambda}{\gamma} + \gamma |\mathbb{F}| \right]$$

Next, we minimize over the parameter γ , which is a proof artifact hiding the actual dependencies. To do it, we study derivative of ϵ with respect to γ , which is

$$\frac{d\epsilon}{d\gamma} = 2 \cdot \left[-|\mathbb{F}|^{\frac{3}{2} - \frac{d}{2}} \cdot \frac{2^\lambda}{\gamma^2} + |\mathbb{F}| \right],$$

and observe that, for $\gamma \geq 0$ it is positive if and only if $\gamma \geq |\mathbb{F}|^{\frac{1}{4} - \frac{d}{4}} \cdot 2^{\frac{\lambda}{2}}$. This means that $\gamma_{\min} = |\mathbb{F}|^{\frac{1}{4} - \frac{d}{4}} \cdot 2^{\frac{\lambda}{2}}$ minimizes ϵ . Plugging $\gamma_{\min}$ within ϵ , we get $\epsilon = |\mathbb{F}|^{1 - \frac{d}{4}} \cdot 2^{\frac{\lambda+1}{2}}$. Applying Markov's inequality, the corollary follows. $\square$