

Targeted Lossy Functions and Applications

Willy Quach*

Brent Waters[†]

Daniel Wichs[‡]

June 29, 2021

Abstract

Lossy trapdoor functions, introduced by Peikert and Waters (STOC '08), can be initialized in one of two indistinguishable modes: in injective mode, the function preserves all information about its input, and can be efficiently inverted given a trapdoor, while in lossy mode, the function loses some information about its input. Such functions have found countless applications in cryptography, and can be constructed from a variety of number-theoretic or algebraic “Cryptomania” assumptions. In this work, we introduce *targeted lossy functions (TLFs)*, which relax lossy trapdoor functions along two orthogonal dimensions. First, they do not require an inversion trapdoor in injective mode. Second, the lossy mode of the function is initialized with some target input, and the function is only required to lose information about this particular target. The injective and lossy modes should be indistinguishable even given the target. We construct TLFs from “Minicrypt” assumptions, namely, injective pseudorandom generators, or even one-way functions under a natural relaxation of injectivity. We then generalize TLFs to incorporate *branches*, and construct *all-injective-but-one* and *all-lossy-but-one* variants. We show a wide variety of applications of targeted lossy functions. In several cases, we get the first Minicrypt constructions of primitives that were previously only known under Cryptomania assumptions. Our applications include:

- *Pseudo-entropy functions* from one-way functions.
- Deterministic leakage-resilient message-authentication codes and improved leakage-resilient symmetric-key encryption from one-way functions.
- Extractors for *extractor-dependent sources* from one-way functions.
- Selective-opening secure symmetric-key encryption from one-way functions.
- A new construction of CCA PKE from (exponentially secure) trapdoor functions and injective pseudorandom generators.

We also discuss a fascinating connection to distributed point functions.

*Northeastern University. Email: quach.w@northeastern.edu

[†]UT Austin and NTT Research Inc. Email: bwaters@cs.utexas.edu. Supported by NSF CNS-1908611, Packard Foundation Fellowship, and Simons Investigator Award.

[‡]Northeastern University and NTT Research Inc. Email: wichs@ccs.neu.edu. Research supported by NSF grant CNS-1750795 and the Alfred P. Sloan Research Fellowship.

Contents

1	Introduction	1
1.1	Our Results	3
1.2	Our Techniques	6
1.3	Relation to Distributed Point Functions	10
2	Preliminaries	11
3	Definitions	12
3.1	Targeted Lossy Functions (TLFs)	12
3.2	Targeted All-Lossy-But-One Functions (T-ALBO)	13
3.3	Targeted All-Injective-But-One Functions (T-AIBO)	14
4	Constructions	14
4.1	Targeted Lossy Functions	14
4.2	T-AIBOs	16
4.3	T-ALBOs	18
5	Applications of T-ALBOs	24
5.1	Pseudo-Entropy Functions.	24
5.2	Extractors for Extractor-Dependent Sources	25
5.3	Deterministic Leakage-Resilient MACs	28
5.4	Leakage-Resilient Symmetric Encryption	29
5.5	Symmetric-Key Encryption Secure against Selective Opening Attacks	30
6	Application of T-AIBOs to CCA Security	36
6.1	Encryption with Randomness Recovery	37
6.2	CCA-secure Encryption from T-AIBOs and Trapdoor Functions	40

1 Introduction

Lossy trapdoor functions, introduced by Peikert and Waters [PW08], are a fundamental cryptographic tool and have found countless applications in many areas of cryptography. They can be constructed from a wide variety of specific number-theoretic and algebraic “Cryptomania” assumptions. In this work, we introduce a relaxation of lossy trapdoor functions that we call *targeted-lossy functions* (TLFs), and show how to instantiate them using “Minicrypt” assumptions, such as injective pseudorandom generators or even one-way functions for a natural variant. We then provide applications of TLFs to a diverse set of problems. For several of these problems, we get the first solutions under one-way functions, where previously only solutions under specific Cryptomania assumptions were known.

Lossy Trapdoor Functions. Lossy trapdoor functions consist of a function family $F_{fk}(\cdot)$ indexed by a public *function key* fk . The function key fk can be generated in one of two modes. In *injective* mode, the function $F_{fk}(\cdot)$ is injective, and therefore each output $y = F_{fk}(x)$ uniquely determines the input x . Furthermore, the public function key fk is generated together with a secret trapdoor td that allows one to efficiently invert the function and recover the input x from the output $y = F_{fk}(x)$. In *lossy* mode, the output $y = F_{fk}(x)$ loses some information about the input x . This is captured by defining a lossiness parameter ℓ and requiring that the size of the image of F_{fk} is at most a $\frac{1}{2^\ell}$ fraction of the size of the domain. In particular, this implies that when x is uniformly random over its domain, then the conditional entropy¹ of x given $F_{fk}(x)$ is at least ℓ bits, meaning that this information about x is lost by $F_{fk}(x)$. The two modes should be *computationally indistinguishable*: given fk , one cannot tell if it was generated in injective mode or lossy mode.

Since their introduction, lossy trapdoor functions have turned out to be incredibly versatile tool and have quickly become an integral part of our cryptographic tool-set. They have found countless and varied applications, including to CCA security, trapdoor functions with many hard-core bits, collision-resistant hash functions, and oblivious transfer [PW08], deterministic encryption [BFO08], analyzing OAEP [KOS10], hedged public-key encryption with bad randomness [BBN⁺09], selective opening security [BHY09], pseudo-entropy functions [BHK11], point-function obfuscation [Zha16], computational extractors [D VW20, GKK20], incompressible encodings [MW20], etc.

Lossy trapdoor functions are known to imply public-key encryption, making them a *Cryptomania* primitive. We currently know how to construct lossy trapdoor functions under most (but not all) concrete Cryptomania assumptions, such as DDH, LWE, Quadratic Residuosity (QR), Decision-Composite Residuosity (DCR), and Phi-hiding [PW08, KOS10, FGK⁺13], but not e.g., factoring, RSA, or the (low noise) LPN assumption.

Targeted-Lossy Functions. In this work, we introduce a relaxation of lossy trapdoor functions, that we call *targeted-lossy functions* (TLFs), with the goal of constructing them under weaker assumptions. TLFs relax the notion of lossy trapdoor functions along two orthogonal dimensions:

- *No inversion trapdoor in injective mode.* When we generate fk in injective mode, we now *only* require that the function $F_{fk}(\cdot)$ is injective, but we no longer require there to be a trapdoor td that allows us to efficiently invert it.

¹Throughout the introduction, entropy refers to min-entropy, and conditional entropy refers to average-case conditional min-entropy [DORS08].

- *Targeted Lossiness.* When we generate fk in lossy mode, we are now also given a target input x^* and only require that $F_{\text{fk}}(x^*)$ loses ℓ bits of information about the particular target x^* . In particular, when the target x^* is chosen uniformly at random and fk is chosen in lossy mode for this target, then the conditional entropy of x^* given $(\text{fk}, F_{\text{fk}}(x^*))$ should be at least ℓ bits.

The two modes should be computationally indistinguishable even given the potential target x^* . In other words, given the pair (fk, x^*) , one cannot distinguish whether fk was chosen in injective mode and independently of x^* or in lossy mode with x^* as the target.

Notice that the first relaxation already appears to take us out of Cryptomania – without a trapdoor, there is no obvious way to use this primitive to construct public-key encryption. This relaxation was considered on its own in prior works (e.g., [BHK11, DVW20]), and is already known to have interesting applications. Unfortunately, there has been no progress towards achieving this relaxation on its own under any Minicrypt assumption, or even under any assumption that doesn't already imply the full notion of lossy trapdoor functions.² This motivates us to consider this relaxation in conjunction with our second relaxation.

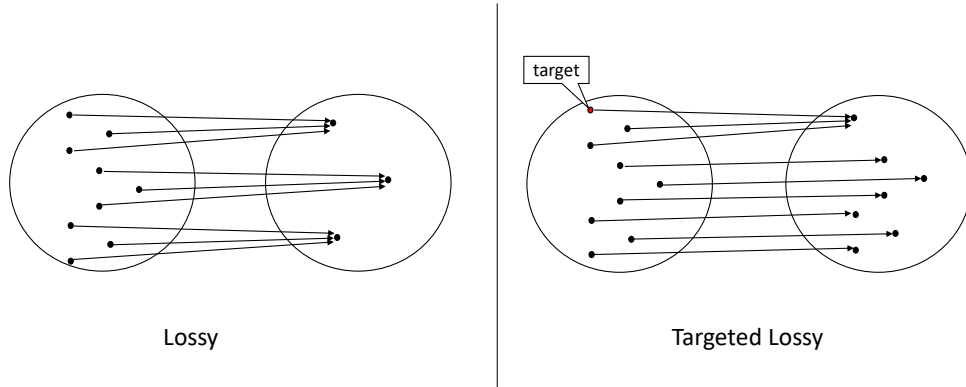


Figure 1: Lossy vs Targeted-Lossy

The second relaxation substantially weakens the lossiness requirement and only asks for targeted lossiness. To highlight the difference between standard vs targeted lossiness, notice that targeted lossiness with parameter ℓ could be achieved by choosing a lossy function key fk for some target x^* , where the output $y = F_{\text{fk}}(x^*)$ has 2^ℓ pre-images in the set $S = \{x : F_{\text{fk}}(x) = y\}$, but for every $x \notin S$, the value $F_{\text{fk}}(x)$ has a unique pre-image x . Such a function would *not* be lossy in the standard sense. For example, if the domain of the function is $\{0, 1\}^n$ with $n = 2\ell$, then, from the point of view of standard lossiness, the function only has negligible information loss $\ell' = O(2^{-\ell})$, even though its targeted-lossiness ℓ can be an arbitrarily large polynomial.³ This example is illustrated in Figure 1. Despite the significant difference between the notions, we show that targeted lossiness suffices in many applications in place of standard lossiness.⁴

²It is also known that, with a sufficiently high lossiness rate, this relaxation on its own would already at least imply collision-resistant hashing [PW08], and therefore is unlikely to follow from one-way functions/permutations.

³The function has an image of size $2^n - 2^\ell + 1$, which we can write as $\frac{1}{2^{\ell'}} 2^n$ for $\ell' = O(2^{-\ell})$.

⁴We could also consider the second relaxation to targeted lossiness on its own, without making the first relaxation (i.e., by still insisting on an inversion trapdoor in injective mode). In that case, the resulting notion would still be a Cryptomania primitive. Interestingly, this notion was considered informally in [GGH19], where it was constructed

TLFs with Branches/Tags. We also consider an augmented notion of TLFs with branches/tags, analogously to prior notions of branches for lossy trapdoor functions [PW08]. In this setting, a single function key fk defines an entire family of functions $F_{\text{fk},\text{tag}}(\cdot)$ with various branches indexed by tag . We can sample the function key fk with a special branch tag^* and a target value x^* , and we require that the pair (fk, x^*) computationally hides tag^* . We define two main variants of this notion, depending on whether the special branch is lossy or injective.

In a targeted *all-injective-but-one* (T-AIBO) family, the special branch tag^* is targeted-lossy and all other branches are injective. In particular, for all $\text{tag} \neq \text{tag}^*$, the function $F_{\text{fk},\text{tag}}$ is injective, while $F_{\text{fk},\text{tag}^*}(x^*)$ loses ℓ bits of information about the target x^* . This notion is most directly analogous to the way branches were defined for standard lossy trapdoor functions of [PW08].

In a targeted *all-lossy-but-one* (T-ALBO) family, the function $F_{\text{fk},\text{tag}^*}$ is injective on the special branch tag^* , while all other branches $F_{\text{fk},\text{tag}}$ are *cumulatively* targeted-lossy. In particular, the cumulative outputs of all lossy branches $(F_{\text{fk},\text{tag}}(x^*))_{\text{tag} \neq \text{tag}^*}$ must lose ℓ -bits of information about the target x^* . An analogous notion of branches for the case of lossy trapdoor functions was previously considered in [CPW20], and a relaxed version without trapdoors (but without the second relaxation to targeted lossiness) was considered implicitly in [BHK11, GKK20] and explicitly in [DVW20]. All prior constructions relied on Cryptomania assumptions.

Relaxing Injectivity. It turns out that we can also relax the injectivity requirement of TLFs, while sufficing for most of our applications. When we choose fk in injective mode, instead of requiring that $F_{\text{fk}}(x)$ uniquely determines x , we only require that it uniquely determines some *property* of x , modeled as an arbitrary function $P(x)$. In this case, we also require that when fk is in lossy mode for the target value x^* , then $F_{\text{fk}}(x^*)$ loses ℓ bits of information about the same property $P(x^*)$. We can define T-AIBOs and T-ALBOs with relaxed injectivity analogously, and even allow the property P to depend on fk . In the case of T-ALBOs, by setting the property $P(x) = F_{\text{fk},\text{tag}^*}(x)$, this relaxed injectivity requirement is equivalent to insisting that the cumulative outputs of all lossy branches $(F_{\text{fk},\text{tag}}(x^*))_{\text{tag} \neq \text{tag}^*}$ should lose ℓ bits of information about the output of the “injective” branch $F_{\text{fk},\text{tag}^*}(x^*)$.

1.1 Our Results

We construct targeted lossy functions (TLFs) from injective pseudorandom generators (PRGs). We also generalize our construction to targeted all-injective-but-one functions (T-AIBOs) under the same assumption. For all-lossy-but-one functions (T-ALBOs), we need a stronger “doubly injective” PRG $G(x) = (y_0, y_1)$, whose output consists of two halves y_0, y_1 , and the PRG is injective on each half individually (i.e., either one of y_0, y_1 uniquely determines x). We also construct TLFs, T-AIBOs and T-ALBOs with relaxed injectivity from just one-way functions. In all cases, we start with a basic construction that only achieves lossiness of $\ell = 1$ bits, but can then amplify lossiness via parallel repetition to get to an arbitrary polynomial ℓ . (However, the lossiness *rate* of our constructions, defined as ℓ/n where n is the domain size, is only $1/\lambda$, where λ is the security parameter. Achieving a higher lossiness rate in Minicrypt is a fascinating open problem.)

under the CDH assumption, which is not known to imply standard lossy trapdoor functions. Not much else is known about this setting.

Application: Pseudo-entropy Functions. The work of Braverman, Hassidim, and Kalai [BHK11] introduced the notion of a *pseudo-entropy function (PEF)*, and constructed them under the DDH assumption. A PEF $f_k(x)$ has a secret key k and takes as inputs values x . The requirement is that for any a-priori chosen input x^* , we can indistinguishably select the key k so that $f_k(x^*)$ has ℓ bits of true statistical entropy even given $f_k(x)$ for all inputs $x \neq x^*$. We observe that PEFs follow almost immediately from T-ALBOs with relaxed injectivity (just by “renaming” the various components), and therefore get a construction of PEFs from one-way functions. The amount of entropy ℓ in our construction can be set to an arbitrarily large polynomial. (On the other hand, the entropy *rate* of our construction, defined as ℓ/n where n is the key size, is stuck at $1/\lambda$. This is in contrast to the construction of [BHK11], which achieved an entropy rate of $1 - o(1)$ under DDH.)

Application: Leakage-Resilience. *Leakage-resilient cryptography* aims to preserve security even if an adversary can get some partial leakage on the secret key. We consider the setting of memory leakage [AGV09, ADW09, NS09, HLWW13], where an adversary can learn any efficiently computable function of the secret key, as long as the number of leaked bits is bounded by some parameter ℓ . As shown by [BHK11], pseudo-entropy functions (PEFs) are useful for leakage-resilient symmetric-key cryptography and were previously used to construct (selectively secure) deterministic leakage-resilient MACs under DDH. By using our new construction of PEFs from one-way functions, we get the first (selectively secure) deterministic leakage-resilient MACs from just one-way functions. The amount of leakage ℓ that we can tolerate can be set to an arbitrarily large polynomial. (However, the length of the key grows depending on ℓ and the leakage *rate* of our construction, defined as ℓ/n where n is the key size, is stuck at $1/\lambda$. This is in contrast to the construction of [BHK11], which achieved a leakage rate of $1 - o(1)$ under DDH.) We can also use a similar technique to construct leakage-resilient CPA-secure symmetric-key encryption from one-way functions.

We note that a prior work of [HLWW13] constructed leakage-resilient symmetric-key primitives, including CPA-secure symmetric-key encryption and (adaptively secure) MACs from one-way functions. The amount of leakage and the leakage rate are the same as in our construction. However, the MACs in the prior work were inherently randomized, while in this work we get deterministic MACs. This is especially crucial in the context of leakage-resilience since, in a randomized construction, leakage that occurs during a computation may also depend on the randomness of the computation in addition to the secret key, but such leakage was not analyzed by the prior work (and indeed, the proof there would fail). Furthermore, our MAC has a smaller signature size: the ratio of leakage to signature size is $(1 - o(1))$ in our construction while it is $1/\lambda$ in the prior work. For the case of CPA-secure symmetric-key encryption, in the prior work the ciphertext size grew linearly with the leakage bound ℓ , while in our work, only the secret key size grows with the leakage bound ℓ , but the ciphertext size just has a minimal $O(\lambda)$ additive overhead on top of the message length. For both MACs and symmetric-key encryption, our constructions are substantially different from those of [HLWW13].

Application: Extractor-Dependent Sources. The work of Dodis, Vaikuntanathan and Wichs [DVW20], which we will refer to as DVW, defined the notion of (computational) extractors for extractor-dependent sources. The goal is to extract nearly uniform randomness R from an arbitrary source of randomness X that has some sufficient entropy. Classical results show this to be possible

using a *seeded* randomness extractor $R = \text{Ext}(X; S)$, which relies on a public random seed S . As long as the source X is independent of the seed S , the output R is nearly uniform even given S . Usually, we think of the source X as coming from nature and therefore consider it to be worst-case but not adversarial – this is used to justify its independence from S . DVW considered a setting where the seed S is repeatedly used to extract randomness from nature and may therefore impact nature itself (e.g., consider using the timing of interrupts to derive entropy, but the interrupts may depend on processes that may themselves rely on extracted outputs). They model this by assuming that the source which produces X can depend on oracle access to the extractor $\text{Ext}(\cdot; S)$, but is independent of the seed S otherwise, and they refer to such sources as *extractor-dependent sources*. DVW showed that extractors for extractor-dependent sources cannot exist unconditionally and at least imply one-way functions. They also distinguished between two scenarios, depending on whether the source can output some additional correlated *auxiliary information* AUX in addition to the sample X , as long as it preserves the entropy of X . The setting with auxiliary information is considered more realistic. As their main results, DVW show how to construct extractors for extractor-dependent sources in the setting without auxiliary information from sub-exponentially secure one-way functions, and in the setting with auxiliary information from a wide range of Cryptomania assumptions such as DDH, DLIN, LWE or DCR. They also gave some evidence that it would be difficult to construct such extractors from simple Minicrypt primitives, by showing that a large class of constructions — ones where seeing the outputs of the extractor on many inputs uniquely determines the seed — cannot be proven secure via a black box reduction.

Despite the above negative result, in this work we construct extractors for extractor-dependent sources, even in the setting with auxiliary information, from standard one-way functions! Our construction does not require sub-exponential security, is entirely black-box in the one-way function, and achieves the same parameters as the prior constructions from Cryptomania assumptions. We circumvent the negative result of DVW by using a construction that lies outside of the class considered there — by relying on lossiness, we ensure that many outputs don’t uniquely determine the seed — yet can still be instantiated in Minicrypt. Our main technique is to adapt a construction of DVW, which relied on all-lossy-but-one functions (without a trapdoor), and adapt it to only rely on targeted all-lossy-but-one functions.

Applications: Selective Opening Security. We also apply TLFs to the problem of selective opening security [DNRS99, BHY09] for symmetric-key encryption. A selective opening attack considers a scenario where an adversary sees a large number of ciphertexts and adaptively asks to “open” some subset of them; we would like to argue that the adversary does not learn anything about the messages encrypted in the remaining ciphertexts. An opening could correspond to seeing the encryption randomness or, if all the ciphertexts are encrypted under different keys, then seeing the corresponding secret keys. Surprisingly, selective opening security does not follow generically from standard encryption security [BDWY12, HR14, HRW16]. On the other hand, we have constructions of selective-opening secure public-key encryption for both randomness-opening and key-opening under many specific public-key assumptions [BHY09, FHKW10, HLOV11, Hof12, HPW15]. However, the problem does not appear to have been studied in the symmetric-key setting. One piece of good news is that symmetric-key encryption schemes are often “public coin”, meaning that the encryption randomness is sent in the clear as part of the ciphertext. Such schemes are automatically secure against selective randomness-opening attacks, since the randomness is available to the adversary for free! Therefore, we focus on constructing a public-coin symmetric-key encryption that achieves

security under *selective key-opening attacks*. We consider a setting where n secret keys $k_1, \dots, k_n$ are chosen uniformly at random and the adversary is given a CPA oracle for each of these keys. In addition, the adversary gets n challenge ciphertexts, one under each key. The adversary gets to adaptively choose to open some arbitrary subset of the n ciphertexts and receive the corresponding secret keys, and we want to argue that the messages encrypted in the remaining ciphertexts stay hidden. Formalizing this requires some care and we naturally adapt the simulation-based definition of selective security from the public-key setting. We show how to construct such selectively secure symmetric-key encryption from one-way functions via our constructions of T-ALBOs/PEFs.

Application: CCA Encryption from Injective Trapdoor Functions The recent work of [HKW20] gave a black-box construction of CCA-secure public-key encryption from any injective trapdoor function. In this work, we give a completely different construction using targeted all-injective-but-one functions (T-AIBOs). As our final result, we get CCA-secure public-key encryption from any injective trapdoor function with a very high (strongly exponential) level of security and an injective pseudorandom generator. While our end-result is strictly worse than [HKW20] in terms of the assumptions, our construction is conceptually simple and we hope it may point to further applications and/or improvements.

1.2 Our Techniques

Basic Construction. Our basic construction of targeted lossy functions (TLFs) with lossiness $\ell = 1$ is extremely simple. Let $G : \{0, 1\}^\lambda \rightarrow \{0, 1\}^{3\lambda+1}$ be an injective pseudorandom generator (PRG), where λ is the security parameter. Let $\mathcal{H} = \{h : \{0, 1\}^{3\lambda+1} \rightarrow \{0, 1\}^{3\lambda}\}$ be a universal hash function family that compresses the input by 1 bit. We define the function family $F_{\text{fk}} : \{0, 1\}^\lambda \rightarrow \{0, 1\}^{3\lambda}$ via $F_{\text{fk}}(x) = h(G(x))$, where $\text{fk} = h \in \mathcal{H}$.

The above parameters ensure that if we choose $\text{fk} = h$ randomly, then the function F_{fk} is injective with overwhelming probability. In particular, for any $x_0 \neq x_1 \in \{0, 1\}^\lambda$, the probability that $G(x_0)$ and $G(x_1)$ are a collision on h is $2^{-3\lambda}$. By taking a union bound over all such pairs x_0, x_1 , the probability of there being any collision is at most $2^{-\lambda}$.

For the lossy mode of the function, we're given some random target that we denote by x_0^* . We choose an additional random input x_1^* and “program” the hash function h so that the values $G(x_0^*)$ and $G(x_1^*)$ collide, which ensures that $F_{\text{fk}}(x_0^*) = F_{\text{fk}}(x_1^*)$. Since x_0^* and x_1^* are treated symmetrically, the tuple $(\text{fk}, y = F_{\text{fk}}(x_0^*) = F_{\text{fk}}(x_1^*))$ does not disambiguate between them, and hence preserves at least $\ell = 1$ bit of entropy in the target.

We can ensure that programming h with a collision in lossy mode is computationally indistinguishable from choosing a random h in injective mode, even given the target x_0^* . For concreteness, we consider the specific universal hash function $h_a(x) = \text{chop}(a \cdot x)$ that performs a field multiplication over $\mathbb{F}_{2^{3\lambda+1}}$ and chops off the least significant bit. Using the standard representation of field elements, this implies that for all y we have $\text{chop}(y) = \text{chop}(y + 1)$.⁵ In that case, programming h to ensure $G(x_0^*)$ collides with $G(x_1^*)$ means choosing $a = (G(x_0^*) - G(x_1^*))^{-1}$. But, even if we're given the target x_0^* , the value $a = (G(x_0^*) - G(x_1^*))^{-1}$ is indistinguishable from uniform by the pseudorandomness of $G(x_1^*)$. Therefore the lossy mode of choosing a for a target x_0^* is indistinguishable from the injective mode of choosing a uniformly at random. The above summarizes the entire construction and proof of security, highlighting its simplicity!

⁵The addition here is over $\mathbb{F}_{2^{3\lambda+1}}$ which is of characteristic 2.

If we don't have an injective PRG, the same construction above already achieves a relaxed form of injectivity. Namely, the injective mode of the function uniquely determines the property $P(x) = G(x)$, while the lossy mode of the function loses 1-bit of information about the same property $P(x^*) = G(x^*)$ for the target x^* .

The above only achieves lossiness of 1 bit, but can amplify the lossiness arbitrarily via parallel repetition. Given a TLF F_{fk} with 1 bit of lossiness, we define $F'_{\text{fk}}(x_1, \dots, x_\ell) = F_{\text{fk}_1}(x_1) \parallel \dots \parallel F_{\text{fk}_\ell}(x_\ell)$ for $\text{fk}' = (\text{fk}_1, \dots, \text{fk}_\ell)$ to get ℓ bits of lossiness. While the lossiness amount can be made arbitrarily large, the lossiness rate (defined as the ratio of the lossiness ℓ to the input size) is stuck at $\frac{1}{\lambda}$ and is not improved by parallel repetition.⁶

Targeted All-Injective-But-One Functions (T-AIBOs). We can also easily extend the basic construction to get a T-AIBO. For branches $\text{tag} \in \{0, 1\}^t$, we need to define a family of functions $F_{\text{fk}, \text{tag}}(\cdot)$ such that, for a special branch tag^* , the function $F_{\text{fk}, \text{tag}^*}$ is lossy and for all other branches it is injective. We can achieve this generically from any TLF without branches where the injective function key fk is uniformly random, as is the case in our basic construction. We simply set $\text{fk} = h$ to be a pairwise-independent hash function and then apply it to the value tag to derive a function key $\hat{\text{fk}} = h(\text{tag})$ for the basic TLF; the output of $F_{\text{fk}, \text{tag}}(x)$ is then set to $F_{\hat{\text{fk}}}(x)$. We program the hash so that the special branch tag^* maps to a lossy function key $\hat{\text{fk}}^*$ for the target value x^* . The output of the hash on any other $\text{tag} \neq \text{tag}^*$ is random and independent, and therefore the resulting TLF function key $\hat{\text{fk}}$ is injective with overwhelming probability.

Targeted All-Lossy-But-One Functions (T-ALBOs). Getting T-ALBOs is more involved. Recall that we need a family of functions $F_{\text{fk}, \text{tag}}(\cdot)$ such that there is a special branch tag^* on which the function is injective and, on all other branches $\text{tag} \neq \text{tag}^*$, it is cumulatively targeted-lossy for some target x^* , meaning that the entire collection of outputs on *all* the lossy tags $(F_{\text{fk}, \text{tag}}(x^*))_{\text{tag} \neq \text{tag}^*}$ must lose ℓ -bits of information about x^* . We start with an approach that was originally proposed by [BHK11], and later abstracted more explicitly in [DVW20], as a way of converting lossy (trapdoor) functions into all-lossy-but-one (trapdoor) functions in the non-targeted setting. We first describe this approach and then show how to adapt it to the targeted setting.

The basic idea of [BHK11, DVW20] is to rely on function composition. As a first step, assume we have a lossy (trapdoor) function F_{fk} where both the domain and the range are $\{0, 1\}^n$, and in particular are the same. We can use it to construct an all-lossy-but-one (trapdoor) function $\bar{F}_{\text{fk}, \text{tag}}$ with tags in $\{0, 1\}^t$. We define the function key $\text{fk} = ((\text{fk}_1^0, \text{fk}_1^1), \dots, (\text{fk}_t^0, \text{fk}_t^1))$ to consist of $2t$ function keys for the underlying lossy function and we define

$$\bar{F}_{\text{fk}, \text{tag}}(x) = (F_{\text{fk}_t^{\text{tag}_t}} \circ F_{\text{fk}_{t-1}^{\text{tag}_{t-1}}} \circ \dots \circ F_{\text{fk}_1^{\text{tag}_1}})(x)$$

where tag_i denotes the i 'th bit of tag . We set the t function keys $\text{fk}_i^{\text{tag}_i^*}$ corresponding to the injective branch tag^* to be injective and the other t function keys $\text{fk}_i^{1-\text{tag}_i^*}$ to be lossy. Since the

⁶We could slightly improve the lossiness rate of the basic construction to $O(\log(\lambda))/\lambda$ by using a $t = \text{poly}(\lambda)$ -wise independent hash function and programming it to have t collisions instead of just 1 collision. This would come at the cost of a larger function key fk . This slight improvement in lossiness rate would only be of interest if we were to consider exact security. Otherwise, asymptotic polynomial/negligible security is too coarse-grained to capture this improvement since it does not even distinguish between λ and λ^ϵ for $\epsilon > 0$; in other words, in the asymptotic setting we can anyway “cheat” and make the rate as high $1/\lambda^\epsilon$ by changing the security parameter to λ^ϵ and weakening exact security accordingly.

composition of injective functions is injective, it holds that $F_{\text{fk}, \text{tag}^*}$ is an injective function. On the other hand, for any lossy branch $\text{tag} \neq \text{tag}^*$, there exists some i such that $\text{tag}_i \neq \text{tag}_i^*$ and therefore one of the functions $F_{\text{fk}_i, \text{tag}_i}$ applied during the computation of $F_{\text{fk}, \text{tag}}(x)$ will be lossy and lose ℓ bits of information about its input, which is the same as losing ℓ bits of information about x since its input is a permutation of x . This shows that for each lossy branch $\text{tag} \neq \text{tag}^*$, the function $F_{\text{fk}, \text{tag}}$ is individually lossy. But we can even show that the lossy branches are cumulatively lossy. This is because the only information revealed about x by all the $2^t - 1$ lossy branches cumulatively, $(F_{\text{fk}, \text{tag}}(x^*))_{\text{tag} \neq \text{tag}^*}$, can be deduced from the t values one gets by applying the first $i - 1$ injective functions followed by a lossy one in position i , for $i = 1, \dots, t$. Each such output reveals at most $n - \ell$ bits of information about x and hence in they reveal at most $t(n - \ell)$ bits in total. This gives lossiness $\ell' = n - t(n - \ell)$, which can be large if ℓ is very close to n and t is small relative to n ; e.g., if $\ell = n(1 - o(1))$ and $t = o(n)$ then $\ell' = (1 - o(1))n$. Indeed, one can get such parameters from DDH.

Unfortunately, there are several issues with applying the above approach in our case. Firstly, our basic TLF does not have the same domain and range: it maps an input in $\{0, 1\}^\lambda$ to an output in $\{0, 1\}^{3\lambda}$. This makes it difficult to even syntactically rely on the above approach. Fortunately, this is relatively easy to fix. We can redefine our basic TLF with modified parameters $F_{\text{fk}} : \{0, 1\}^{3\lambda} \rightarrow \{0, 1\}^{3\lambda}$ via $F_{\text{fk}}(x) = h(G(x))$ where now we have $x \in \{0, 1\}^{3\lambda}$, the injective PRG is of the form $G : \{0, 1\}^{3\lambda} \rightarrow \{0, 1\}^{3\lambda+1}$, and the universal hash functions are of the form $h : \{0, 1\}^{3\lambda+1} \rightarrow \{0, 1\}^{3\lambda}$. This lets us syntactically use F_{fk} in the above construction to define a family with branches. Unfortunately, now F_{fk} is no longer injective when fk is chosen in “injective mode”. However, we can regain injectivity by first pre-processing a smaller input $x \in \{0, 1\}^\lambda$ via an injective PRG $G' : \{0, 1\}^\lambda \rightarrow \{0, 1\}^{3\lambda}$. We define the overall function with branches $\bar{F}_{\text{fk}, \text{tag}} : \{0, 1\}^\lambda \rightarrow \{0, 1\}^{3\lambda}$ via:

$$\bar{F}_{\text{fk}, \text{tag}}(x) = (F_{\text{fk}_t, \text{tag}_t} \circ F_{\text{fk}_{t-1}, \text{tag}_{t-1}} \circ \dots \circ F_{\text{fk}_1, \text{tag}_1})(G'(x)).$$

This preserves injectivity on the branch tag^* since, with overwhelming probability, each function component $F_{\text{fk}_i, \text{tag}_i^*}$ is injective over the domain of inputs $(F_{\text{fk}_{i-1}, \text{tag}_{i-1}^*} \circ \dots \circ F_{\text{fk}_1, \text{tag}_1^*})(G'(x))$ of size 2^λ .

For targeted lossiness, we can now chose each of the targeted lossy keys $\text{fk}_i^{1-\text{tag}_i^*}$ with the target $x^*[i] = (F_{\text{fk}_{i-1}, \text{tag}_{i-1}^*} \circ \dots \circ F_{\text{fk}_1, \text{tag}_1^*})(G'(x^*))$. This is enough to show that each lossy branch is individually targeted-lossy. Unfortunately, we don't get cumulative lossiness. Recall that the argument we employed in the previous paragraph only gave cumulative lossiness $\ell' = n - t(n - \ell)$, which was only meaningful when the initial lossiness ℓ was a large fraction of the domain size n . But in our case $\ell = 1$ and hence the above does not give us any meaningful bound on ℓ' , even for tag size $t = 2$.

To solve the above issue, we need to control the lossiness more carefully to ensure that the leakages from different lossy tags don't add up. We do so by going under the hood of our basic TLF construction. We set $x_0^* = x^*$ to be the target and choose a uniformly random and independent x_1^* . We then choose all the lossy function keys $\text{fk}_i^{1-\text{tag}_i^*}$ to ensure that the two values x_0^*, x_1^* collide on every lossy branch. We can do so by programming the universal hash function in the i th lossy key to ensure that it has a collision on $G(x_0^*[i]), G(x_1^*[i])$ where $x_b^*[i] = (F_{\text{fk}_{i-1}, \text{tag}_{i-1}^*} \circ \dots \circ F_{\text{fk}_1, \text{tag}_1^*})(G'(x_b^*))$. This guarantees that $F_{\text{fk}_i, 1-\text{tag}_i^*}(x_0^*[i]) = F_{\text{fk}_i, 1-\text{tag}_i^*}(x_1^*[i])$ and so x_0^* and x_1^* collide on every lossy branch. While this ensures cumulative lossiness, we now lose indistinguishability. The reason is that the

randomness of $G(x_1^*[i])$ is used twice: once to define the lossy key $\text{fk}_i^{1-\text{tag}_i^*}$, and once to define $x_1^*[i+1]$, which is used to define the lossy key $\text{fk}_{i+1}^{1-\text{tag}_{i+1}^*}$. Since we're reusing the same randomness, the lossy keys for different values i will appear correlated and can then be distinguished from injective keys. We can fix this issue by using two different PRGs G_0 and G_1 for the 0 functions $F_{\text{fk}_i^0}$ and the 1 functions $F_{\text{fk}_i^1}$ respectively. We need G_0, G_1 to each be injective and also to be mutually pseudorandom so that, for a random x the values $G_0(x), G_1(x)$ look like random and independent values.⁷ With this modification, we preserve indistinguishability. This is because, the lossy function key $\text{fk}_i^{1-\text{tag}_i^*}$ now only relies on $G^{1-\text{tag}^*}(x_1^*[i])$ while the value $x_1^*[i+1]$ used to define $\text{fk}_{i+1}^{1-\text{tag}_{i+1}^*}$ only relies on $G^{\text{tag}^*}(x_1^*[i])$, and hence we can argue that the two values look random and independent.

T-ALBO with Relaxed Injectivity. We can also get a T-ALBO with relaxed injectivity by using the same construction as above with standard PRGs rather than injective PRGs, and therefore from one-way functions. The observation is that, when we program the lossy mode to ensure that the target $x^* = x_0^*$ collides with some random x_1^* on every lossy tag, it's very unlikely that x_0^* and x_1^* would collide on the injective tag, even when the PRG is not injective. Therefore, although the injective output $F_{\text{fk}, \text{tag}^*}(x^*)$ may not uniquely determine x^* , we ensure that $F_{\text{fk}, \text{tag}^*}(x^*)$ has 1-bit of entropy even given $(F_{\text{fk}, \text{tag}}(x^*))_{\text{tag} \neq \text{tag}^*}$. In other words, the injective mode of the function reveals at least 1 bit of information about x^* that was lost by all the lossy evaluations on x^* . In fact, we can even ensure that the above holds if we shorten the output size of the function to just 1 bit (in which case, the function certainly can't be injective). We do so by applying a universal hash function with 1-bit output at the end, and programming it to ensure that $F_{\text{fk}, \text{tag}^*}(x_0^*)$ and $F_{\text{fk}, \text{tag}^*}(x_1^*)$ hash to different bits. This ensures 1 bit of entropy in a 1 bit output, and therefore the output of the injective branch is *uniformly random*, even given the outputs of all the lossy branches. We can amplify from 1 bit to many bits via parallel repetition.

Applications of T-ALBOs. We notice that T-ALBOs with relaxed injectivity directly give us *pseudo-entropy functions*, just by relabeling the components. We define the secret key k of the pseudo-entropy function as $k = (\text{fk}, s)$ to consist of a function key fk for a T-ALBO and a uniformly random input s for it. We then define the pseudo-entropy function $f_k(x) = F_{\text{fk}, \text{tag}=x}(s)$ which interprets its input x as a branch and evaluates the T-ALBO on s . For any input x^* chosen a-priori, we can choose $k = (\text{fk}, s)$ by selecting a random s and choosing fk with the injective branch x^* and the target s . This guarantees the properties of a pseudo-entropy function: the value k chosen this way is indistinguishable from an honestly chosen k that is independent of x^* , but ensures that $f_k(x^*)$ has ℓ bits of statistical entropy even given $f_k(x)$ for all $x \neq x^*$. This shows that T-ALBOs directly give pseudo-entropy functions. In fact, this gives a pseudo-entropy function where the key k consists of a uniformly random secret component s and a public but carefully chosen component fk defined in terms of s and the point x^* on which we want to ensure statistical entropy. Conversely, a pseudo-entropy function of this form also gives a T-ALBO. By using a T-ALBO where lossiness ℓ is equal to the output size (as we showed can be done above), we can even ensure that $f_k(x^*)$ is uniformly random given $f_k(x)$ for all $x \neq x^*$.

Our applications to *leakage-resilient MACs*, *leakage-resilient symmetric-key encryption*, and *extractors for extractor-dependent sources* follow as interesting applications of pseudo-entropy functions.

⁷Equivalently, we can think of $G_0(x), G_1(x)$ as the left/right halves of a single PRG $G(x)$.

For *selective-opening security*, we notice that our pseudo-entropy function has an additional feature. Not only can we ensure that $f_k(x^*)$ is uniformly random given $f_k(x)$ for all $x \neq x^*$, but for any output y we can even efficiently find a key k_y such that $f_{k_y}(x^*) = y$ and $f_{k_y}(x) = f_k(x)$ for all $x \neq x^*$. Intuitively, this additional feature gives us the ability to efficiently “equivocate”, which is used to get selective-opening security. In particular, a simulator can efficiently find a key k_y to open a challenge ciphertext to any value it wants, and k_y looks consistent even to an adversary that got access to a CPA oracle.

Application of T-AIBOs to CCA Security. We also give an application of T-AIBOs to CCA-secure encryption from any trapdoor function with a sufficient high level of security. We describe a simplified version of this result, and the main body gives a more general treatment. Let $F_{\text{fk}, \text{tag}}$ be a T-AIBO with λ -bit input and $\ell = 1$ bits of lossiness, as we constructed from injective pseudorandom generators. Let f_{pk} be a family of trapdoor functions (not necessarily permutations) with input length $n = \lambda^3$. Our CCA encryption public key consists of the pair (pk, fk) and the secret key is the trapdoor of the trapdoor function. The encryption procedure selects a random $r \in \{0, 1\}^n$ and parses it as $r = (r_1, \dots, r_d)$ with $d = \lambda^2$ and $r_i \in \{0, 1\}^\lambda$. It also selects a one-time signature key pair (vk, sk) . It computes $y = f_{\text{pk}}(r)$ and $y_1 = F_{\text{fk}, \text{vk}}(r_1), \dots, y_{\lambda^2} = F_{\text{fk}, \text{vk}}(r_d)$ then uses a Goldreich-Levin hardcore bit of r to one-time pad the message and signs everything under vk . The decryption procedure checks the signature, inverts y to recover r and checks that $y_1, \dots, y_d$ were computed correctly: if so it recovers the hardcore bit and decrypts the message, else it rejects.

To prove CCA security, we select fk to be lossy on the branch $\text{tag} = \text{vk}$ and the target value r that correspond to the challenge ciphertext. We can then simulate the decryption procedure without knowing the trapdoor td by brute-force inverting all the values $y_i = F_{\text{fk}, \text{vk}}(r_i)$ in $\tilde{O}(2^\lambda)$ time. In the challenge ciphertext, the value $r = (r_1, \dots, r_d)$ has $d = \lambda^2$ bits of entropy even given $y_1, \dots, y_d$. We argue that this makes it hard to recover r even given the trapdoor function output $f_{\text{pk}}(r)$ and the ability to run in $\tilde{O}(2^\lambda)$ time. We show that this follows from very strong exponential hardness of the trapdoor function: we need to assume that for input length $n = \lambda^3$ no adversary running in time $\tilde{O}(2^\lambda)$ can invert the function with better than $\frac{2^{\lambda^2}}{2^{\lambda^3}}$ probability. While this is a strong assumption, note that the trivial attack that tries 2^λ random inputs only has success probability $\frac{2^\lambda}{2^{\lambda^3}}$ and generic non-uniform attacks [DGK17] can’t do better than $\frac{2^{\tilde{O}(\lambda)}}{2^{\lambda^3}}$.

1.3 Relation to Distributed Point Functions

We observe an interesting connection between T-ALBOs (with relaxed injectivity), pseudo-entropy functions, and distributed-point functions (DPFs) [GI14, BGI15]. In fact, even though the notions look very different and were introduced with different goals in mind, they are essentially equivalent. We already discussed the connection between T-ALBOs and pseudo-entropy functions, and so we now show the connection to DPFs.

Distributed point functions were defined in the context of 2-server private information retrieval (PIR). They consist of a function family $f_k : [N] \rightarrow \{0, 1\}$. Given some target $x^* \in [N]$, it should be possible to choose two keys k_0, k_1 such that $f_{k_0}(x^*) \neq f_{k_1}(x^*)$ differ on the target point, but for all other points $x \neq x^*$ they are the same $f_{k_0}(x) = f_{k_1}(x)$. Each of the keys k_0, k_1 should individually computationally hide the value x^* . This gives 2-server PIR. When a client wants to retrieve a value $\text{DB}[x^*]$ at the location $x^* \in [N]$ of a database $\text{DB} \in \{0, 1\}^N$, it chooses the two keys k_0, k_1 using a target x^* and sends k_b to server b . Each server b computes $y_b = \bigoplus_{x \in [N]} f_{k_b}(x) \cdot \text{DB}[x]$

and the client computes $y_0 \oplus y_1 = \text{DB}[x^*]$. Neither server individually learns anything about the location x^* by the hiding property of the DPF.

A pseudo-entropy function with 1-bit output and 1-bit entropy almost already gives a DPF. If we select the key k to preserve entropy on x^* , then $f_k(x^*)$ has 1 bit of entropy even given $f_k(x)$ for all $x \neq x^*$. That means that there must be some key k' such that $f_{k'}(x^*) \neq f_k(x^*)$ but $f_{k'}(x) = f_k(x)$ for all $x \neq x^*$. We can define $k_0 = k$ and $k_1 = k'$ to get the two DPF keys for the point x^* . The only difficulty is ensuring that we can kind k' efficiently. Recall that in our construction of pseudo-entropy functions for T-ALBOs, we set $k = (\text{fk}, s)$ where fk is a function key of a T-ALBO with the “injective” branch x^* and the target input s . When we choose fk , our T-ALBO construction in turn sets $s_0 = s$, picks a random s_1 and ensures that the function outputs collide on s_0, s_1 for all branches $x \neq x^*$ but differ on the branch x^* . Therefore, we can efficiently set $k_0 = (\text{fk}, s_0)$ and $k_1 = (\text{fk}, s_1)$.

Interestingly, although our construction of T-ALBOs was initially inspired by the works of [BHK11, DVW20], we observe in retrospect that it is very similar to the construction of DPFs in [BGI15]. Indeed the function composition construction of T-ALBOs using two PRGs G_0, G_1 is similar to the GGM construction of PRFs from PRGs [GGM86], and the use of hash functions h is similar to the use of “correction words” in the adaptation of GGM to DPFs in [BGI15].

We hope that the connections between all these notions help foster a better understanding of each of them. The fact that completely different motivations and construction approaches surreptitiously converged to yield related notions and constructions should perhaps be viewed as a good indication of just how fundamental these ideas are.

2 Preliminaries

Basic Notation. For an integer N , we let $[N] := \{1, 2, \dots, N\}$. For a set S we let $x \leftarrow S$ denote sampling x uniformly at random from S . For a distribution $\mathcal{S}$ we let $x \leftarrow \mathcal{S}$ denote sampling x according to the distribution. We will denote the security parameter by λ . We say a function $f(\lambda)$ is negligible, denoted $f(\lambda) = \text{negl}(\lambda)$, if $f(\lambda) = O(\lambda^{-c})$ for every constant $c > 0$. A function $f(\lambda)$ is polynomial, denoted $f(\lambda) = \text{poly}(\lambda)$, if $f(\lambda) = O(\lambda^c)$ for some constant $c > 0$. We say that an event occurs with overwhelming probability if it holds with probability $1 - \text{negl}(\lambda)$. For a randomized algorithm A , we will sometimes explicitly denote the randomness coins it uses, writing $A(x; \text{coins})$. We will write $D_1 \stackrel{c}{\approx} D_2$ if the (ensembles of) distributions D_1 and D_2 are computationally indistinguishable.

Information Theory. For two random variables X, Y with support $\text{supp}(X)$ and $\text{supp}(Y)$ respectively, we define their statistical distance $\text{SD}(X, Y)$ as

$$\text{SD}(X, Y) := \sum_{u \in \text{supp}(X) \cup \text{supp}(Y)} \frac{1}{2} |\Pr[X = u] - \Pr[Y = u]|.$$

Two ensembles of random variables $X = \{X_\lambda\}_\lambda, Y = \{Y_\lambda\}_\lambda$ are statistically close if $\text{SD}(X_\lambda, Y_\lambda) = \text{negl}(\lambda)$.

The min-entropy $H_\infty(X)$ of a random variable X is defined as

$$H_\infty(X) := -\log\left(\max_{x \in \text{supp}(X)} \Pr[X = x]\right).$$

Following Dodis et al. [DORS08], we define the (average) conditional min-entropy of X given Y as: $H_\infty(X|Y) = -\log \left(\mathbb{E}_{y \leftarrow Y} \left[2^{-H_\infty(X|Y=y)} \right] \right)$. Note that $H_\infty(X|Y) = k$ iff the optimal strategy for guessing X given Y succeeds with probability 2^{-k} .

3 Definitions

3.1 Targeted Lossy Functions (TLFs)

We first define our basic notion of targeted lossy functions (TLF).

Definition 3.1 (Targeted Lossy Functions). *A targeted lossy function (TLF) function family with input length $n = n(\lambda)$, output length $m \geq n$ and lossiness parameter $\ell = \ell(\lambda)$ consists of PPT algorithms $(\text{InjectiveGen}, \text{LossyGen}, F)$ with the following syntax:*

- $\text{fk} \leftarrow \text{InjectiveGen}(1^\lambda)$: generates a function key fk .
- $\text{fk} \leftarrow \text{LossyGen}(1^\lambda, x^*)$: on input a target value $x^* \in \{0, 1\}^n$, generates a function key fk .
- $y = F_{\text{fk}}(x)$: a deterministic algorithm which, on input fk along with a value $x \in \{0, 1\}^n$, outputs $y \in \{0, 1\}^m$.

We require the following properties:

Injectivity: With overwhelming probability over the choice of $\text{fk} \leftarrow \text{InjectiveGen}(1^\lambda)$, the function F_{fk} is injective over its domain $\{0, 1\}^n$.

ℓ -Lossiness: For random variables $x^* \leftarrow \{0, 1\}^n$, $\text{fk} \leftarrow \text{LossyGen}(1^\lambda, x^*)$, we have

$$H_\infty(x^* \mid \text{fk}, F_{\text{fk}}(x^*)) \geq \ell.$$

Indistinguishability: For all $x^* \in \{0, 1\}^n$, we have the computational indistinguishability

$$(x^*, \text{fk}_{\text{inj}}) \stackrel{c}{\approx} (x^*, \text{fk}_{\text{loss}})$$

where $\text{fk}_{\text{inj}} \leftarrow \text{InjectiveGen}(1^\lambda)$ and $\text{fk}_{\text{loss}} \leftarrow \text{LossyGen}(1^\lambda, x^*)$.

Relaxing Injectivity. We can also define a variant of TLFs with *relaxed injectivity*, where we require the injective mode to only uniquely determine some property $P(x)$ while lossy mode loses ℓ -bits of information on $P(x^*)$ for the target x^* . In particular, we require that there exists some function $P : \{0, 1\}^n \rightarrow \{0, 1\}^*$ for which the following holds:

- *Relaxed Injectivity:* With overwhelming probability over the choice of $\text{fk} \leftarrow \text{InjectiveGen}(1^\lambda)$ it holds that for all $x, x' \in \{0, 1\}^n$ if $F_{\text{fk}}(x) = F_{\text{fk}}(x')$ then $P(x) = P(x')$.
- *ℓ -Lossiness:* For random variables $x^* \leftarrow \{0, 1\}^n$, $\text{fk} \leftarrow \text{LossyGen}(1^\lambda, x^*)$, we have

$$H_\infty(P(x^*) \mid \text{fk}, F_{\text{fk}}(x^*)) \geq \ell.$$

3.2 Targeted All-Lossy-But-One Functions (T-ALBO)

Next, we define a tagged version of TLFs, that we name targeted all-lossy-but-one functions (T-ALBO).

Definition 3.2 (T-ALBO). *A targeted all-lossy-but-one (T-ALBO) function family with input length $n = n(\lambda)$, output length $m \geq n$, tag length $t = t(\lambda)$ and lossiness parameter $\ell = \ell(\lambda)$, consists of PPT algorithms (InjectiveGen, LossyGen, F) with the following syntax:*

- $\text{fk} \leftarrow \text{InjectiveGen}(1^\lambda)$: generates a function key fk .
- $\text{fk} \leftarrow \text{LossyGen}(1^\lambda, \text{tag}^*, x^*)$: on input $\text{tag}^* \in \{0, 1\}^t, x^* \in \{0, 1\}^n$, generates a function key fk .
- $y = F_{\text{fk}, \text{tag}}(x)$: a deterministic algorithm, which, on input fk, tag along with a value $x \in \{0, 1\}^n$, outputs $y \in \{0, 1\}^m$.

We require the following properties:

Injectivity: *With overwhelming probability over the choice of $\text{fk} \leftarrow \text{InjectiveGen}(1^\lambda)$, for all $\text{tag} \in \{0, 1\}^t$, the function $F_{\text{fk}, \text{tag}}$ is injective over the domain $\{0, 1\}^n$. Moreover, for any tag^*, x^* , with overwhelming probability over the choice of $\text{fk} \leftarrow \text{LossyGen}(1^\lambda, \text{tag}^*, x^*)$, the function $F_{\text{fk}, \text{tag}^*}$ on tag tag^* is injective.*

ℓ -Lossiness: *For all $\text{tag}^* \in \{0, 1\}^t$ and random variables $x^* \leftarrow \{0, 1\}^n, \text{fk} \leftarrow \text{LossyGen}(1^\lambda, \text{tag}^*, x^*)$, we have $H_\infty(x^* \mid \text{fk}, (F_{\text{fk}, \text{tag}}(x^*))_{\text{tag} \neq \text{tag}^*}) \geq \ell$. We use $(F_{\text{fk}, \text{tag}}(x^*))_{\text{tag} \neq \text{tag}^*}$ to denote the (ordered) list of outputs of the function $F_{\text{fk}, \text{tag}}(x^*)$ on all $2^t - 1$ possible lossy tags $\text{tag} \neq \text{tag}^*$.*

Indistinguishability: *For all $\text{tag}^* \in \{0, 1\}^t$ and all $x^* \in \{0, 1\}^n$, we have*

$$(\text{tag}^*, x^*, \text{fk}_{\text{inj}}) \stackrel{c}{\approx} (\text{tag}^*, x^*, \text{fk}_{\text{loss}})$$

where $\text{fk}_{\text{inj}} \leftarrow \text{InjectiveGen}(1^\lambda)$ and $\text{fk}_{\text{loss}} \leftarrow \text{LossyGen}(1^\lambda, \text{tag}^, x^*)$.*

Relaxing Injectivity: Entropy-Preserving T-ALBOs. We can also relax injectivity in much the same way as we did for TLFs, by requiring that injective mode uniquely determines some property $P(x)$ while lossy mode loses information on $P(x^*)$. Here, we can even allow the property P to depend on fk, tag^* . In this case, without loss of generality, we can set $P(x) = F_{\text{fk}, \text{tag}^*}(x)$ to be the output on the “injective” tag, and therefore it tautologically holds that $F_{\text{fk}, \text{tag}^*}(x)$ determines $P(x)$. Hence this notion just requires that seeing the output of the function on input x^* over all lossy branches $\text{tag} \neq \text{tag}^*$ preserves some entropy of the output $F_{\text{fk}, \text{tag}^*}(x^*)$ on the “injective” branch tag^* . We call this notion *entropy preserving*. This notion also meaningfully allows us to make the output much smaller than the input size, and potentially just 1-bit.

Definition 3.3 (Entropy-Preserving T-ALBO). *An entropy-preserving T-ALBO with input length $n = n(\lambda)$, output length $m = m(\lambda)$, tag length $t = t(\lambda)$ and lossiness parameter $\ell = \ell(\lambda)$, consists of algorithms (InjectiveGen, LossyGen, F). We require that they satisfy the same indistinguishability property as in Definition 3.2. However, we replace the injectivity and lossiness properties with the following entropy-preserving property. For any fixed $\text{tag}^* \in \{0, 1\}^t$ we have:*

$$H_\infty(F_{\text{fk}}(\text{tag}^*, x^*) \mid \text{fk}, (F_{\text{fk}, \text{tag}}(x^*))_{\text{tag} \neq \text{tag}^*}) \geq \ell,$$

where we define the random variables $x^ \leftarrow \{0, 1\}^n, \text{fk} \leftarrow \text{LossyGen}(1^\lambda, \text{tag}^*, x^*)$.*

We say that (InjectiveGen, LossyGen, F) is maximally entropy-preserving if $\ell = m$, where m is the output size of the T-ALBO.

3.3 Targeted All-Injective-But-One Functions (T-AIBO)

Last, we define another tagged variant of TLFs that we call *targeted all-injective-but-one lossy functions* (T-AIBO). In a T-AIBO, the branches $\text{tag} \neq \text{tag}^*$ are injective, whereas only tag^* corresponds to a lossy branch.

Definition 3.4 (T-AIBO). *A targeted all-injective-but-one T-AIBO function family with input length $n = n(\lambda)$, output length $m \geq n$, tag length $t = t(\lambda)$ and lossiness parameter $\ell = \ell(\lambda)$, consists of PPT algorithms $(\text{InjectiveGen}, \text{LossyGen}, F)$ with the following syntax:*

- $\text{fk} \leftarrow \text{InjectiveGen}(1^\lambda)$: generates a function key fk .
- $\text{fk} \leftarrow \text{LossyGen}(1^\lambda, \text{tag}^*, x^*)$: on input $\text{tag}^* \in \{0, 1\}^t, x^* \in \{0, 1\}^n$, generates a function key fk .
- $y = F_{\text{fk}, \text{tag}}(x)$: a deterministic polynomial time algorithm, which, on input fk, tag along with a value $x \in \{0, 1\}^n$ outputs $y \in \{0, 1\}^m$.

We require the following properties:

Injectivity on injective branches: *With overwhelming probability over the choice of $\text{fk} \leftarrow \text{InjectiveGen}(1^\lambda)$, we have that for all tags $\text{tag} \in \{0, 1\}^t$, the function $F_{\text{fk}, \text{tag}}$ is injective over the domain $\{0, 1\}^n$. Moreover, for any tag^*, x^* , with overwhelming probability over $\text{fk} \leftarrow \text{LossyGen}(1^\lambda, \text{tag}^*, x^*)$, we have that for all tags $\text{tag} \neq \text{tag}^*$, the function $F_{\text{fk}, \text{tag}}$ is injective.*

ℓ -Lossiness: *For any $\text{tag}^* \in \{0, 1\}^t$ and random variables $x^* \leftarrow \{0, 1\}^n, \text{fk} \leftarrow \text{LossyGen}(1^\lambda, \text{tag}^*, x^*)$, we have $H_\infty(x^* \mid \text{fk}, F_{\text{fk}, \text{tag}^*}(x^*)) \geq \ell$.*

Indistinguishability: *For any $\text{tag}^* \in \{0, 1\}^t$ and $x^* \in \{0, 1\}^n$, we have the computational indistinguishability*

$$(\text{tag}^*, x^*, \text{fk}_{\text{inj}}) \stackrel{c}{\approx} (\text{tag}^*, x^*, \text{fk}_{\text{loss}})$$

where $x^* \leftarrow \{0, 1\}^n, \text{fk}_{\text{inj}} \leftarrow \text{InjectiveGen}(1^\lambda)$ and $\text{fk}_{\text{loss}} \leftarrow \text{LossyGen}(1^\lambda, \text{tag}^*, x^*)$.

We could also relax injectivity as we did for TLFs and T-ALBOs. Since we do not consider this notion in the paper, we omit it for simplicity.

4 Constructions

In this section we present our constructions of TLFs and its variants. In Section 4.1, we give the construction of basic TLFs (Theorem 4.1). Then, we show in Section 4.2 our construction of a T-AIBO (Theorem 4.2). Finally, in Section 4.3, we build both a T-ALBO (Theorem 4.3) and a maximally entropy-preserving T-ALBO (Theorem 4.4).

4.1 Targeted Lossy Functions

We start with our base construction of TLF. We prove the following:

Theorem 4.1 (TLFs from Injective PRGs). *Let $\ell = \ell(\lambda)$ be a polynomial. Assuming the existence of injective PRGs, there exists a TLF with input length $n = \ell\lambda$, output length $m = 3\ell\lambda$ and lossiness ℓ .*

Let $G : \{0, 1\}^\lambda \rightarrow \{0, 1\}^{3\lambda+1}$ be an injective PRG. Let $\mathbb{F} = \mathbb{F}_{2^{3\lambda+1}}$ be the field of size $2^{3\lambda+1}$. We represent elements of $\mathbb{F}$ in the standard manner as $3\lambda + 1$ coefficients of polynomials in $\mathbb{F}_2[X]/(f)$ for some appropriate polynomial f , so that adding elements in $\mathbb{F}$ can be performed by adding their coefficients component wise. Throughout this section we will identify $\{0, 1\}^{3\lambda+1}$ with $\mathbb{F}$. For $a \in \mathbb{F}$, represented as a bit string of length $3\lambda + 1$, define $\text{chop}(a)$ as the first 3λ bits of the representation of a (i.e., the last bit, corresponding to the constant term of the polynomial, is chopped off). Let $e = 1_{\mathbb{F}} \in \mathbb{F}$ be the field element that has 0s in the first 3λ positions and 1 in the last position. Note that for all $x \in \mathbb{F}$ we have $\text{chop}(x + e) = \text{chop}(x)$, and hence for all x_1, x_2 , $\text{chop}(x_1) = \text{chop}(x_2)$ if and only if $x_1 = x_2$ or $x_1 = x_2 + e$.

We first construct an LTF (InjectiveGen, LossyGen, F) with input length $n = \lambda$ and output length $m = 3\lambda$, and lossiness 1 as follows.

- **InjectiveGen**(1^λ): Sample $a \leftarrow \mathbb{F}$ and output $\text{fk} = a$.
- **LossyGen**($1^\lambda, x^*$): On input $x^* \in \{0, 1\}^\lambda$, set $x_0^* := x^*$ and sample $x_1^* \leftarrow \{0, 1\}^\lambda \setminus \{x^*\}$. Set $a = e \cdot (G(x_0^*) - G(x_1^*))^{-1}$, and output $\text{fk} = a$.
- $F_{\text{fk}}(x) = \text{chop}(a \cdot G(x)) \in \{0, 1\}^{3\lambda}$.

Claim 4.1.1. *Suppose $G : \{0, 1\}^\lambda \rightarrow \{0, 1\}^{3\lambda+1}$ is an injective PRG. Then (InjectiveGen, LossyGen, F) is a TLF with input length $n = \lambda$, output length $m = 3\lambda$ and lossiness $\ell = 1$.*

Proof. Note that in lossy mode, a is well-defined by injectivity of G . We prove injectivity, ℓ -lossiness and indistinguishability.

Injectivity. Fix any $x_0 \neq x_1 \in \mathbb{F}$. By injectivity of G , we have $G(x_0) \neq G(x_1)$. Therefore, $F_{\text{fk}}(x_0) = F_{\text{fk}}(x_1)$ iff $\text{chop}(a \cdot G(x_0)) = \text{chop}(a \cdot G(x_1))$, which occurs iff $a = e \cdot (G(x_0) - G(x_1))^{-1}$ or $a = 0$. This happens with probability $2/|\mathbb{F}|$ over the randomness of $\text{fk} \leftarrow \text{InjectiveGen}(1^\lambda)$. By taking a union bound over all pairs of distinct inputs $x_0, x_1 \in \{0, 1\}^\lambda$, we obtain that the probability over $\text{fk} \leftarrow \text{InjectiveGen}(1^\lambda)$ that F_{fk} is not injective is at most $\frac{2^{2\lambda+1}}{|\mathbb{F}|} = \frac{1}{2^\lambda}$.

($\ell = 1$)-Lossiness. Let $x_0^* = x^* \leftarrow \{0, 1\}^\lambda$ be a random target, and let $x_1^* \leftarrow \{0, 1\}^\lambda \setminus \{x_0^*\}$ denotes the random value sampled during the execution of $\text{fk} \leftarrow \text{LossyGen}(1^\lambda, x_0^*)$; we will denote such an execution via $\text{fk} = \text{LossyGen}(1^\lambda, x_0^*; x_1^*)$. We think of x_0^*, x_1^* as random variables, which in turn define the random variables $\text{fk} = \text{LossyGen}(1^\lambda, x_0^*; x_1^*)$, $y = F_{\text{fk}}(x_0^*) = F_{\text{fk}}(x_1^*)$. We observe that the resulting distribution of fk, y does not reveal anything about x_0^*, x_1^* beyond the (unordered) set $\{x_0^*, x_1^*\}$, due to the symmetry of how x_0^*, x_1^* are treated by **LossyGen**. In other words, $x_0^* \rightarrow \{x_0^*, x_1^*\} \rightarrow (\text{fk}, y)$ forms a Markov chain. A data processing inequality gives:

$$H_\infty(x_0^* \mid \text{fk}, y) \geq H_\infty(x_0^* \mid \{x_0^*, x_1^*\}) \geq 1,$$

where the last inequality follows since one cannot predict x_0^* given the (unordered) set $\{x_0^*, x_1^*\}$ with probability better than $1/2$. Note that x_0^*, x_1^* are uniformly random over $\{0, 1\}^\lambda$ conditioned on $x_0^* \neq x_1^*$.

Indistinguishability. We define a hybrid experiment where **LossyGen** is modified as follows:

- **LossyGen**($1^\lambda, x^*$): Set $x_0^* := x^*$ and select $u_1^* \leftarrow \mathbb{F}$. If $u_1^* = G(x_0^*)$, output $\text{fk} = 0$. Otherwise output $a = e \cdot (G(x_0^*) - u_1^*)^{-1}$.

The output of $\widetilde{\text{LossyGen}}$ is indistinguishable from the output of LossyGen by PRG security of G , noting that $u_1^* = G(x_0^*)$ with negligible probability $1/2^{3\lambda+1}$ over the randomness of u_1^* alone.

Moreover, even given x^* , the output of $\widetilde{\text{LossyGen}}$ is uniformly random in $\mathbb{F}$ over the choice of u_1^* alone, and is therefore identically distributed as the output of InjectiveGen . $\square$

Amplifying (absolute) lossiness. The construction above only have lossiness 1. We note here that we can amplify this lossiness, which gives Theorem 4.1.

The idea is that (absolute) lossiness can be amplified by partitioning a (longer) input into blocks and applying an independent TLF on each chunk. Suppose TLF $(\text{InjectiveGen}, \text{LossyGen}, F)$ is a TLF with input size n , output size m and lossiness ℓ . Let $k = k(\lambda)$ be a polynomial. The modified scheme is defined as follows:

- $\overline{\text{InjectiveGen}}(1^\lambda)$: For all $i \in [k]$, compute $\text{fk}_i \leftarrow \text{InjectiveGen}(1^\lambda)$. Output $\{\text{fk}_i\}_{i \in [k]}$.
- $\overline{\text{LossyGen}}(1^\lambda, x^*)$: On input $x^* \in \{0, 1\}^{kn}$, parse $x^* = x_1 \| \dots \| x_k \in \{0, 1\}^{kn}$ where $x_i \in \{0, 1\}^n$ for all $i \in [k]$. For all $i \in [k]$, compute $\text{fk}_i \leftarrow \text{LossyGen}(1^\lambda, x_i)$. Output $\{\text{fk}_i\}_{i \in [k]}$.
- $\overline{F}_{\text{fk}}(x)$: On input $x \in \{0, 1\}^{kn}$, parse $x = x_1 \| \dots \| x_k \in \{0, 1\}^{kn}$ where $x_i \in \{0, 1\}^n$ for all $i \in [k]$. For all $i \in [k]$, compute $y_i = F_{\text{fk}_i}(x_i)$. Output $(y_1 \| \dots \| y_k)$.

The resulting scheme is a TLF with input size kn , output size km and lossiness $k\ell$. This is at the cost of making the input longer, and therefore doesn't affect the lossiness rate. Applying the above to our construction from Claim 4.1.1, we obtain Theorem 4.1.

Remark: Relaxed Injectivity. If we take our construction of TLFs above but remove the requirement that the PRG is injective, we get relaxed injectivity with the property $P(x) = G(x)$. The proof is otherwise identical.

4.2 T-AIBOs

We describe our construction of T-AIBO. We prove the following:

Theorem 4.2 (T-AIBOs from Injective PRGs). *Let $\ell = \ell(\lambda)$ and $t = t(\lambda)$ be polynomials. Assuming the existence of injective PRGs, there exists a T-AIBO with input length $n = \ell\lambda$, tag length t , output length $m = \ell \cdot (3\lambda + t)$ and lossiness ℓ .*

We build on our construction of TLF from Section 4.1. Recall that we built our TLF as $F_{\text{fk}}(s) = \text{chop}(a \cdot G(s))$, where $a \in \mathbb{F}$ forms the key fk . In order to build a T-AIBO, we now compute $a_{\text{tag}} = h_k(\text{tag})$ where h is a pairwise independent hash function.

More formally, let $t = t(\lambda)$ be the tag length. Let $n = 3\lambda + t + 1$, and let $\mathbb{F} = \mathbb{F}_{2^n}$. We consider elements $\text{tag} \in \{0, 1\}^t$ as elements of $\mathbb{F}$ (for instance by considering any injection induced by the coefficient embedding of $\mathbb{F}$ as in Section 4.1, setting the remaining $3\lambda + 1$ entries as 0), over which we define the pairwise independent hash function

$$h_{u,v}(\text{tag}) = u \cdot \text{tag} + v \in \mathbb{F},$$

where $u, v \in \mathbb{F}$. We will use the following useful algorithm related to h :

- **Equivocate**(tag, y): on input tag $\in \mathbb{F}$ and $y \in \mathbb{F}$, sample $u \leftarrow \mathbb{F}$, compute $v = y - u \cdot \text{tag}$, and output (u, v) .

Namely, **Equivocate**(tag, y) samples a random key (u, v) conditioned on $h_{u,v}(\text{tag}) = y$. Note that for any fixed tag $\in \mathbb{F}$, we have that $h_{u,v}(\text{tag})$ is uniform over $\mathbb{F}$ over the randomness of (u, v) . As a result, for all tag $\in \mathbb{F}$, we have:

$$(\text{tag}, (u, v) \leftarrow \mathbb{F} \times \mathbb{F}, y = h_{u,v}(\text{tag})) \equiv (\text{tag}, (u, v) \leftarrow \text{Equivocate}(\text{tag}, y), y \leftarrow \mathbb{F}) \quad (1)$$

We now describe our first construction of a T-AIBO with lossiness $\ell = 1$. Let $t = t(\lambda)$ and $n = 3\lambda + t + 1$, $G : \{0, 1\}^\lambda \rightarrow \{0, 1\}^{3\lambda+t+1}$ be an injective PRG and h be the pairwise independent hash function from above. We define the following algorithms:

- **InjectiveGen**(1^λ): Sample $(u, v) \leftarrow \mathbb{F} \times \mathbb{F}$ and set $\text{fk} = (u, v)$.
- **LossyGen**($1^\lambda, \text{tag}^*, x^*$): Set $x_0^* = x^*$ and sample $x_1^* \leftarrow \{0, 1\}^\lambda \setminus \{x_0^*\}$. Let $a = e \cdot (G(x_0^*) - G(x_1^*))^{-1}$. Compute $(u, v) \leftarrow \text{Equivocate}(\text{tag}^*, a)$ and output $\text{fk} = (u, v)$.
- $F_{\text{fk}}(\text{tag}, x)$: Output $\text{chop}(h_{u,v}(\text{tag}) \cdot G(x))$.

Claim 4.2.1. Suppose $G : \{0, 1\}^\lambda \rightarrow \{0, 1\}^{3\lambda+t+1}$ is an injective PRG. Then $(\text{InjectiveGen}, \text{LossyGen}, F)$ is a T-AIBO with input length λ , tag length t , output length $3\lambda + t$, and lossiness $\ell = 1$.

Proof. We prove injectivity, 1-lossiness and indistinguishability.

Injectivity follows by the same argument as the TLF of Claim 4.1.1. In particular, an identical argument shows that, by injectivity of G , for any fixed tag $\in \{0, 1\}^t$, the probability that $F_{\text{fk}, \text{tag}}$ is not injective is at most $\frac{2^{2\lambda+1}}{|\mathbb{F}|} = \frac{1}{2^{t+\lambda}}$. An union bound over the 2^t possible tags shows that the probability that $F_{\text{fk}, \text{tag}}$ is not injective for some tag is at most $\frac{1}{2^\lambda}$.

The proof of 1-lossiness is identical to the proof of Claim 4.1.1. In particular, we define random variables $x_0^* = x^* \leftarrow \{0, 1\}^\lambda$ denoting the target, $x_1^* \leftarrow \{0, 1\}^\lambda \setminus \{x_0^*\}$ denoting the random value sampled during the execution of $\text{fk} \leftarrow \text{LossyGen}(1^\lambda, \text{tag}^*, x_0^*)$, and $y = F_{\text{fk}, \text{tag}^*}(x_0^*) = F_{\text{fk}, \text{tag}^*}(x_1^*)$. We observe that the resulting distribution of fk, y does not reveal anything about x_0^*, x_1^* beyond the (unordered) set $\{x_0^*, x_1^*\}$, due to the symmetry of how x_0^*, x_1^* are treated by **LossyGen**. In other words, $x_0^* \rightarrow \{x_0^*, x_1^*\} \rightarrow (\text{fk}, y)$ forms a Markov chain. A data-processing inequality then shows:

$$H_\infty(x_0^* \mid \text{fk}, y) \geq H_\infty(x_0^* \mid \{x_0^*, x_1^*\}) \geq 1.$$

For indistinguishability, we first argue that for any tag $\in \{0, 1\}^t$ and any $x^* \in \{0, 1\}^n$, the value a sampled during **LossyGen**($1^\lambda, \text{tag}^*, x^*$) is computationally indistinguishable from uniformly random by PRG security of G , over the randomness of x_1^* . Then, indistinguishability follows by Eq. (1). $\square$

As for TLFs and T-ALBOs one can amplify lossiness by concatenating T-AIBOs evaluations on blocks of the input, which gives Theorem 4.2.

Remark 4.1 (Weaker forms of injectivity). For our applications, we only need a weaker form of injectivity that ensures that any tag $\in \{0, 1\}^t$ induces an injective function with overwhelming probability. This is in contrast to all of the tags being simultaneously injective with overwhelming probability. This allows to have a slightly more efficient construction, by considering a PRG $G : \{0, 1\}^\lambda \rightarrow \{0, 1\}^{3\lambda+1}$, which in turn gives a T-AIBO with output size 3λ (as opposed to $3\lambda + t$ where t is the tag space).

Remark 4.2 (Generic Construction of T-AIBOs from TLFs). In the above, we build a T-AIBO starting from the particular TLF from Section 4.1. We note that our transformation above can be made semi-generic, by assuming that the injective function keys fk generated by the `InjectiveGen` procedure of the base TLF are (computationally indistinguishable from) uniformly random (as is the case in our construction). In that case, by mapping branches tag to function keys fk via a programmable pairwise independent hash function, we generically obtain a T-AIBO from such a TLF, in the same way as above.

4.3 T-ALBOs

We now describe our construction of T-ALBOs. We prove the following theorems:

Theorem 4.3 (T-ALBOs from Injective PRGs). *Let $\ell = \ell(\lambda)$ and $t = t(\lambda)$ be any polynomials. Assuming the existence of injective PRGs, there exists a T-ALBO with input length $n = \ell\lambda$, tag length t , output length $m = \ell \cdot (3\lambda + t)$ and lossiness ℓ .*

Theorem 4.4 (Entropy-Preserving T-ALBOs from OWFs). *Let $\ell = \ell(\lambda)$ and $t = t(\lambda)$ be any polynomials. Assuming the existence of one-way functions, there exists an entropy-preserving T-ALBO with input length $n = \ell\lambda$, tag length t , output length $m = \ell$ and lossiness ℓ . In particular, such a T-ALBO is maximally entropy preserving.*

Again, we build on our construction of TLF from Section 4.1. We refer to our technical overview for a high level overview of the following construction. We begin with our construction of a standard T-ALBO (satisfying injectivity) from any injective PRG.

Building blocks. Let $n = 3\lambda + t$. Let $\mathbb{F} = \mathbb{F}_{2^{3\lambda+t+1}}$. Let $G : \{0, 1\}^n \rightarrow \{0, 1\}^{2(n+1)}$ be a PRG. We will write $G(x) = (G^0(x), G^1(x))$. In particular, G^0 and G^1 are PRGs with input size n and output size $n + 1$; we will furthermore assume that each of the functions G^0 and G^1 is injective. We define $(\text{InjectiveGen}^0, \text{LossyGen}^0, F^0)$ and $(\text{InjectiveGen}^1, \text{LossyGen}^1, F^1)$ as follows:

- $\text{InjectiveGen}^b(1^\lambda)$: Sample $a \leftarrow \mathbb{F}$ and output $\text{fk} = a$.
- $\text{LossyGen}^b(1^\lambda, x^*, x_1^*)$: On input $x^*, x_1^* \in \{0, 1\}^n$, set $x_0^* = x^*$. If $G^b(x_0^*) = G^b(x_1^*)$, output $\text{fk} = 0$. Otherwise compute $a = e \cdot (G^b(x_0^*) - G^b(x_1^*))^{-1}$, and output $\text{fk} = a$.
- $F_{\text{fk}}^b(x) = \text{chop}(a \cdot G^b(x)) \in \{0, 1\}^n$.

Let $G' : \{0, 1\}^\lambda \rightarrow \{0, 1\}^n$ be an injective PRG.

Construction. We define a T-ALBO $(\overline{\text{InjectiveGen}}, \overline{\text{LossyGen}}, \overline{F})$ as follows.

- $\overline{\text{InjectiveGen}}(1^\lambda)$: For all $i \leq t$ and $b \in \{0, 1\}$, sample $\text{fk}_i^b \leftarrow \text{InjectiveGen}^b(1^\lambda)$, and output $\text{fk} = \{\text{fk}_i^b\}_{i \in [t], b \in \{0, 1\}}$.
- $\overline{\text{LossyGen}}(1^\lambda, \text{tag}^*, x^*)$: Sample $x_1^* \leftarrow \{0, 1\}^\lambda \setminus \{x^*\}$, and set $x_0^{(0)} = G'(x^*)$, and $x_1^{(0)} = G'(x_1^*)$. For $i = 1$ to t , sample

$$\text{fk}_i^{\text{tag}_i^*} \leftarrow \text{InjectiveGen}^{\text{tag}_i^*}(1^\lambda).$$

Then set

$$\begin{aligned} x_0^{(i)} &= F_{\text{fk}_i^{\text{tag}_i^*}}^{\text{tag}_i^*}(x_0^{(i-1)}) = F_{\text{fk}_i^{\text{tag}_i^*}}^{\text{tag}_i^*} \circ \dots \circ F_{\text{fk}_1^{\text{tag}_1^*}}^{\text{tag}_1^*}(x_0^{(0)}) \\ x_1^{(i)} &= F_{\text{fk}_i^{\text{tag}_i^*}}^{\text{tag}_i^*}(x_1^{(i-1)}) = F_{\text{fk}_i^{\text{tag}_i^*}}^{\text{tag}_i^*} \circ \dots \circ F_{\text{fk}_1^{\text{tag}_1^*}}^{\text{tag}_1^*}(x_1^{(0)}). \end{aligned}$$

Sample

$$\text{fk}_i^{1-\text{tag}_i^*} \leftarrow \text{LossyGen}^{1-\text{tag}_i^*}(1^\lambda, x_0^{(i-1)}, x_1^{(i-1)}).$$

The output is

$$\overline{\text{fk}} = (\text{fk}_i^{(b)})_{i \in [t], b \in \{0,1\}}.$$

- $\overline{F}_{\overline{\text{fk}}, \text{tag}}(x)$: Output

$$y = F_{\text{fk}_t^{\text{tag}_t}}^{\text{tag}_t} \circ \dots \circ F_{\text{fk}_1^{\text{tag}_1}}^{\text{tag}_1}(G'(x)).$$

Claim 4.4.1. Suppose G^0 , G^1 and G' are injective PRGs, and $G = (G^0, G^1)$ is a PRG. Then $(\text{InjectiveGen}, \text{LossyGen}, \overline{F})$ is a T-ALBO with input length λ , tag length t , output length $3\lambda + t + 1$, and lossiness 1.

Proof. We first show a useful property of $(\text{InjectiveGen}^0, \text{LossyGen}^0, F^0)$ and $(\text{InjectiveGen}^1, \text{LossyGen}^1, F^1)$.

- For all $x^* \in \{0, 1\}^n$, we have the following computational indistinguishability:

$$(\text{fk}_{inj}^0, F_{\text{fk}_{inj}^0}^0(x_1^*), x^*, \text{fk}_{los}^1) \stackrel{c}{\approx} (\text{fk}_{inj}^0, u, x^*, \text{fk}_{inj}^1), \quad (2)$$

where $\text{fk}_{inj}^0 \leftarrow \text{InjectiveGen}^0(1^\lambda)$, $\text{fk}_{inj}^1 \leftarrow \text{InjectiveGen}^1(1^\lambda)$, $x_1^* \leftarrow \{0, 1\}^n$ and $\text{fk}_{los}^1 \leftarrow \text{LossyGen}^1(1^\lambda, \text{tag}^*, x^*, x_1^*)$.

Symmetrically, we have:

$$(\text{fk}_{inj}^1, F_{\text{fk}_{inj}^1}^1(x_1^*), x^*, \text{fk}_{los}^0) \stackrel{c}{\approx} (\text{fk}_{inj}^1, u, x^*, \text{fk}_{inj}^0), \quad (3)$$

where $\text{fk}_{inj}^1 \leftarrow \text{InjectiveGen}^1(1^\lambda)$, $\text{fk}_{inj}^0 \leftarrow \text{InjectiveGen}^0(1^\lambda)$, $x_1^* \leftarrow \{0, 1\}^n$ and $\text{fk}_{los}^0 \leftarrow \text{LossyGen}^0(1^\lambda, \text{tag}^*, x^*, x_1^*)$.

These properties follow by PRG security of $G = (G^0, G^1)$, so that $F_{\text{fk}_{inj}^b}^b(x_1^*)$ is computationally indistinguishable from uniformly random over the randomness of $G^b(x_1^*)$, while indistinguishability of fk_{inj}^{1-b} and fk_{los}^{1-b} follows over the (independent) randomness of $G^{1-b}(x_1^*)$.

We now prove that the construction above is a T-ALBO.

Injectivity. Fix $x_0 \neq x_1 \in \{0, 1\}^\lambda$. By injectivity of G' , we have $G'(x_0) \neq G'(x_1)$. Let $i \in [t]$ and $b \in \{0, 1\}$, and let $x_0^{(i)} \neq x_1^{(i)} \in \{0, 1\}^n$. By injectivity of G^b , the probability over $\text{fk}_i^b \leftarrow \text{InjectiveGen}^b(1^\lambda)$ that $F_{\text{fk}_i^b}^b(x_0^{(i)}) = F_{\text{fk}_i^b}^b(x_1^{(i)})$ is $2/|\mathbb{F}|$ (which correspond to either $a = e \cdot (G^b(x_0^{(i)} - G^b(x_1^{(i)}))^{-1}$ or $a = 0$). In particular, for any fixed $\text{tag} \in \{0, 1\}^t$, we have, by taking an union bound over $i \in [t]$, that the probability over $\text{fk} \leftarrow \text{InjectiveGen}(1^\lambda)$ that $\overline{F}_{\overline{\text{fk}}, \text{tag}}(G'(x_0)) = \overline{F}_{\overline{\text{fk}}, \text{tag}}(G'(x_1))$ is

at most $2t/|\mathbb{F}|$. Then, by union bound over all $\text{tag} \in \{0, 1\}^t$ and pairs of input $x_0 \neq x_1 \in \{0, 1\}^\lambda$, the probability that there exists a tag tag and two inputs $x_0 \neq x_1$ such that $\overline{F}_{\overline{\text{fk}}, \text{tag}}(G'(x_0)) = \overline{F}_{\overline{\text{fk}}, \text{tag}}(G'(x_1))$ is at most $\frac{2t \cdot 2^t \cdot 2^{2\lambda}}{|\mathbb{F}|} = \frac{t}{2^\lambda}$, which is negligible.

An almost identical argument (without the union bound over all tags) shows that the branch tag^* is injective in lossy mode.

1-Lossiness. Fix $\text{tag}^* \in \{0, 1\}^t$. Let $x_0^* \leftarrow \{0, 1\}^\lambda$ be the target, and let $x_1^* \leftarrow \{0, 1\}^\lambda \setminus \{x_0^*\}$ be the value sampled during $\overline{\text{fk}} \leftarrow \overline{\text{LossyGen}}(1^\lambda, \text{tag}^*, x_0^*)$.

First, we claim that, for all $\text{tag} \neq \text{tag}^*$, $\overline{F}_{\overline{\text{fk}}, \text{tag}}(x_0^*) = \overline{F}_{\overline{\text{fk}}, \text{tag}}(x_1^*)$. To see this, fix any $\text{tag} \neq \text{tag}^*$, and let i denote the smallest index in $[t]$ such that $\text{tag}_i \neq \text{tag}_i^*$. Recall that we have

$$\text{fk}_i^{1-\text{tag}_i^*} \leftarrow \text{LossyGen}^{1-\text{tag}_i^*}(1^\lambda, x_0^{(i-1)}, x_1^{(i-1)}),$$

and in particular, by construction of LossyGen and F :

$$F_{\text{fk}_i^{\text{tag}_i}}^{\text{tag}_i}(x_0^{(i-1)}) = F_{\text{fk}_i^{\text{tag}_i}}^{\text{tag}_i}(x_1^{(i-1)}).$$

As $\text{tag}_j = \text{tag}_j^*$ for all $j < i$, we have by construction of $x_0^{(i-1)}$ and $x_1^{(i-1)}$:

$$y_i := F_{\text{fk}_i^{\text{tag}_i}}^{\text{tag}_i} \circ \dots \circ F_{\text{fk}_1^{\text{tag}_1}}^{\text{tag}_1}(G'(x_0^*)) = F_{\text{fk}_i^{\text{tag}_i}}^{\text{tag}_i} \circ \dots \circ F_{\text{fk}_1^{\text{tag}_1}}^{\text{tag}_1}(G'(x_1^*)),$$

and in particular

$$\begin{aligned} \overline{F}_{\overline{\text{fk}}, \text{tag}}(x_0^*) &= F_{\text{fk}_t^{\text{tag}_t}}^{\text{tag}_t} \circ \dots \circ F_{\text{fk}_{i+1}^{\text{tag}_{i+1}}}^{\text{tag}_{i+1}}(y_i) \\ &= \overline{F}_{\overline{\text{fk}}, \text{tag}}(x_1^*) \end{aligned}$$

(where by convention we consider the composition to be empty if $i = t$).

Then, we observe, similarly to the proof of Claim 4.1.1, that the resulting distribution $(\overline{\text{fk}}, (\overline{F}_{\overline{\text{fk}}, \text{tag}}(x^*))_{\text{tag} \neq \text{tag}^*})$ does not reveal anything about x_0^*, x_1^* beyond the (unordered) set $\{x_0^*, x_1^*\}$. This follows since x_0^*, x_1^* are treated symmetrically in the generation of $\overline{\text{fk}}$ and the fact that $\overline{F}_{\overline{\text{fk}}, \text{tag}}(x_0^*) = \overline{F}_{\overline{\text{fk}}, \text{tag}}(x_1^*)$ for all $\text{tag} \neq \text{tag}^*$. In other words, $x_0^* \rightarrow \{x_0^*, x_1^*\} \rightarrow (\overline{\text{fk}}, (\overline{F}_{\overline{\text{fk}}, \text{tag}}(x^*))_{\text{tag} \neq \text{tag}^*})$ forms a Markov chain. A data-processing inequality then shows:

$$H_\infty(x_0^* \mid \overline{\text{fk}}, (\overline{F}_{\overline{\text{fk}}, \text{tag}}(x_0^*))_{\text{tag} \neq \text{tag}^*}) \geq H_\infty(x_0^* \mid \{x_0^*, x_1^*\}) \geq 1,$$

where the last inequality follows since one cannot predict x_0^* given the (unordered) set $\{x_0^*, x_1^*\}$ with probability better than $1/2$. Note that x_0^*, x_1^* are uniformly random over $\{0, 1\}^\lambda$ conditioned on $x_0^* \neq x_1^*$.

Indistinguishability. On a high level, $x_1^{(0)}$ looks pseudorandom by security of G' . Then, we switch lossy keys to injective keys, one by one, using our special TLF property (Eq. (2), Eq. (3)), simultaneously switching $x_1^{(j)}$ to uniform and $\text{fk}_i^{1-\text{tag}_i^*}$ to injective.

Fix $\text{tag}^* \in \{0, 1\}^t$ and $x^* \in \{0, 1\}^\lambda$. We define the following hybrids:

Hybrid H_0 : This is the distribution induced by the lossy mode, namely, the output distribution is:

$$(\text{tag}^*, x^*, \overline{\text{fk}}),$$

where $\overline{\text{fk}} \leftarrow \overline{\text{LossyGen}}(1^\lambda, \text{tag}^*, x^*)$.

Hybrid $H_{1,j}$, $0 \leq j \leq t$: We switch how we generate $\overline{\text{fk}}$.

For all $i \leq t$, sample

$$\text{fk}_i^{\text{tag}^*} \leftarrow \text{InjectiveGen}^{\text{tag}^*}(1^\lambda).$$

Sample $x_1^{(j)} \leftarrow \{0, 1\}^n$, and set, for all $i > j$:

$$x_1^{(i)} = F_{\text{fk}_i^{\text{tag}^*}}^{\text{tag}^*} \circ \dots \circ F_{\text{fk}_1^{\text{tag}^*}}^{\text{tag}^*}(x_1^{(j)}),$$

and set for all $i \geq j + 1$:

$$\text{fk}_i^{1-\text{tag}^*} \leftarrow \text{LossyGen}^{1-\text{tag}^*}(1^\lambda, x_0^{(i-1)}, x_1^{(i-1)}).$$

For all $i < j$, set:

$$\text{fk}_i^{1-\text{tag}^*} \leftarrow \text{InjectiveGen}^{1-\text{tag}^*}(1^\lambda),$$

and set if $j \geq 1$:

$$\text{fk}_j^{1-\text{tag}^*} \leftarrow \text{InjectiveGen}^{1-\text{tag}^*}(1^\lambda).$$

Output $(\text{tag}^*, x^*, \overline{\text{fk}})$ where

$$\overline{\text{fk}} = (\text{fk}_i^b)_{i \in [t], b \in \{0,1\}}.$$

Note that the distribution output by Hybrid $H_{2,t}$ is identical to the one output by $\overline{\text{InjectiveGen}}(1^\lambda)$. Therefore it suffices to prove that the hybrid distributions above are indistinguishable.

Claim 4.4.2. *Assuming G' is a PRG, for all $\text{tag}^* \in \{0, 1\}^t$ and all $x^* \in \{0, 1\}^n$, the distributions*

$$(\text{tag}^*, x^*, \overline{\text{fk}})$$

generated in Hybrids H_0 and $H_{1,0}$ are computationally indistinguishable.

Proof. The only difference between these two hybrids is how $x_1^{(0)}$ is distributed. In H_0 , it is computed as $G'(x_1^*)$ where $x_1^* \leftarrow \{0, 1\}^\lambda$, and in $H_{1,0}$, as uniformly random in $\{0, 1\}^n$. Indistinguishability follows by PRG security of G' . $\square$

Claim 4.4.3. *For all $\text{tag}^* \in \{0, 1\}^t$, all $x^* \in \{0, 1\}^n$, and for all $j \in [t]$, the distributions*

$$(\text{tag}^*, x^*, \overline{\text{fk}})$$

generated in Hybrids $H_{1,j-1}$ and $H_{1,j}$ are computationally indistinguishable.

Proof. Fix $\text{tag}^* \in \{0, 1\}^t$, $x^* \in \{0, 1\}^n$, $j \in [t]$. The only differences between hybrids $H_{1,j-1}$ and $H_{1,j}$, are how $x_1^{(j)}$ and $\text{fk}_j^{1-\text{tag}_j^*}$ are generated.

To argue indistinguishability, we use our special joint indistinguishability property of $(\text{InjectiveGen}^0, \text{LossyGen}^0, F^0)$ and $(\text{InjectiveGen}^1, \text{LossyGen}^1, F^1)$ (Eq. (2), Eq. (3)).

Suppose $\text{tag}_j^* = 0$. We use the output distribution $(\text{fk}^0, u, \bar{x}^*, \text{fk}^1)$ from Eq. (2), setting $\bar{x}^* = x_0^{(j-1)} = F_{\text{fk}_{j-1}^*}^{\text{tag}_{j-1}^*} \circ \dots \circ F_{\text{fk}_1^*}^{\text{tag}_1^*}(x_0^{(0)})$ where $x_0^{(0)} = G'(x^*)$. We set $\text{fk}_j^0 = \text{fk}^0$, $x_1^{(j)} = u$ and $\text{fk}_j^1 = \text{fk}^1$, where we implicitly set x_1^* from Eq. (2) as $x_1^* = x_1^{(j-1)}$ (note that $x_1^{(j-1)}$ is only used to define $x_1^{(j)}$ and $\text{fk}_j^{1-\text{tag}_j^*}$ in $H_{1,j-1}$, and is not used in $H_{1,j}$). We compute all the other components as in Hybrid $H_{1,j}$.

If the distribution of Eq. (2) comes in lossy mode (meaning that fk^0 is in lossy mode), then we obtain the output distribution of Hybrid $H_{1,j-1}$. If the distribution comes in injective mode (meaning that fk^0 is in injective mode), then we obtain the output distribution of Hybrid $H_{1,j}$.

The case $\text{tag}_j^* = 1$ is almost identical, where we use Eq. (3) instead of Eq. (2). $\square$

This overall shows that for any $\text{tag}^* \in \{0, 1\}^t$ and $x^* \in \{0, 1\}^\lambda$, the distributions $(\text{tag}^*, x^*, \bar{\text{fk}})$ induced by Hybrids H_0 and $H_{1,t}$ are computationally indistinguishable, which concludes the proof. $\square$

Amplifying Lossiness. As in our construction of TLF, the construction above of T-ALBO only has lossiness 1. We note that we can also amplify lossiness for T-ALBOs, which gives Theorem 4.3.

We now move on to our construction of (maximally) entropy-preserving T-ALBO.

Entropy-Preserving T-ALBOs. Next, we describe how to obtain a (maximally) *entropy-preserving* T-ALBOs with output size m . We first start by giving a construction of an entropy-preserving T-ALBOs with output size 1 and lossiness 1.

Our starting point is our previous construction of a T-ALBO. Let $\overline{\text{chop}} : \{0, 1\}^n \rightarrow \{0, 1\}$ be the function that outputs the *last* bit of its input. In particular $\overline{\text{chop}}(x + e) \neq \overline{\text{chop}}(x)$ for all $x \in \{0, 1\}^n$. We slightly modify our previous construction as follows. First, we will further chop the final output of the previous construction to match the output size $m = 1$. Second, we slightly modify the LossyGen algorithm to resample a fresh key if x^* and x_1^* collide. The reason is that entropy-preserving would otherwise only hold with overwhelming probability in lossy mode (over the choice of the function key).

Let $(\overline{\text{InjectiveGen}}, \overline{\text{LossyGen}}, \overline{F})$ denote the T-ALBO of Claim 4.4.1. We define our entropy-preserving T-ALBO as follows.

- $\widetilde{\text{InjectiveGen}}(1^\lambda)$: Sample $\text{fk} \leftarrow \overline{\text{InjectiveGen}}(1^\lambda)$. Sample $a \leftarrow \mathbb{F}$ and output (fk, a) .
- $\widetilde{\text{LossyGen}}(1^\lambda, \text{tag}^*, x^*)$: Sample $\text{fk} \leftarrow \overline{\text{LossyGen}}(1^\lambda, \text{tag}^*, x^*)$, which also internally samples $x_1^* \leftarrow \{0, 1\}^\lambda$.

If $\overline{F}_{\text{fk}, \text{tag}^*}(x^*) = \overline{F}_{\text{fk}, \text{tag}^*}(x_1^*)$, sample $b \leftarrow \{0, 1\}$. If $b = 0$, output $\tilde{\text{fk}} = (\perp, \text{tag}^*, x^*, x_1^*)$, and if $b = 1$, output $\tilde{\text{fk}} = (\perp, \text{tag}^*, x_1^*, x^*)$.

Otherwise, compute $y^* = \overline{F}_{\text{fk}, \text{tag}^*}(x^*)$ and $y_1 = \overline{F}_{\text{fk}, \text{tag}^*}(x_1^*)$. Set $a = e \cdot (y^* - y_1)^{-1}$. Output $\tilde{\text{fk}} = (\text{fk}, a)$.

- $\tilde{F}_{\text{fk}, \text{tag}}(x)$: If $\tilde{\text{fk}} = (\perp, \text{tag}^*, x_0, x_1)$, output 1 if $\text{tag} = \text{tag}^*$ and $x = x_1$, and 0 otherwise.

Otherwise output

$$y = \overline{\text{chop}}(a \cdot \overline{F}_{\text{fk}, \text{tag}}(x)).$$

Claim 4.4.4. Let $(\widetilde{\text{InjectiveGen}}, \widetilde{\text{LossyGen}}, \overline{F})$ be the T-ALBO of Claim 4.4.1. Then $(\widetilde{\text{InjectiveGen}}, \widetilde{\text{LossyGen}}, \tilde{F})$ is an entropy-preserving T-AIBO with input length λ , tag length t , output length 1, and lossiness $\ell = 1$.

We first show that the event $\overline{F}_{\text{fk}, \text{tag}^*}(x^*) = \overline{F}_{\text{fk}, \text{tag}^*}(x_1^*)$ in $\widetilde{\text{LossyGen}}$ only occurs with negligible probability. Indeed, this occurs only if $x^* = x_1^*$, or if there exists some $i \in [t]$ such that $F_{\text{fk}_i, \text{tag}_i^*}^{\text{tag}_i^*}(x_0^{(i-1)}) = F_{\text{fk}_i, \text{tag}_i^*}^{\text{tag}_i^*}(x_1^{(i-1)})$ (using the notation of Claim 4.4.1).

If i^* is the smallest such i , then $x_1^{(i^*-1)}$ is computationally indistinguishable from freshly uniform in $\{0, 1\}^n$ even given x^* and fk (which we showed in the proof of Claim 4.4.1 to argue 1-lossiness). Therefore $F_{\text{fk}_{i^*}, \text{tag}_{i^*}^*}^{\text{tag}_{i^*}^*}(x_1^{(i^*-1)})$ is also computationally indistinguishable from uniformly random in $\{0, 1\}^n$ even given x^* and fk : this follows by our special properties of F^0 and F^1 (Eq. (2), Eq. (3)), using $x_1^* = x_1^{(i^*-1)}$. In particular $F_{\text{fk}_{i^*}, \text{tag}_{i^*}^*}^{\text{tag}_{i^*}^*}(x_0^{(i-1)}) = F_{\text{fk}_{i^*}, \text{tag}_{i^*}^*}^{\text{tag}_{i^*}^*}(x_1^{(i-1)})$ with at most negligible probability. An union bound over $i \in [t]$ gives that the probability that $\tilde{\text{fk}} = (\perp, \text{tag}^*, x_0, x_1)$ for some x_0, x_1 is negligible.

To argue indistinguishability, we note that the previous argument shows that conditioned on $\tilde{\text{fk}} \neq (\perp, \text{tag}^*, x_0, x_1)$, we have that $y_1 = \overline{F}_{\text{fk}, \text{tag}^*}(x_1^*)$, where $(\text{fk}, a) \leftarrow \widetilde{\text{LossyGen}}(1^\lambda, \text{tag}^*, x^*)$, is computationally indistinguishable from uniform even given x^* and fk , and therefore so is $a = e \cdot (y^* - y_1)^{-1}$. Given that $\tilde{\text{fk}} \neq (\perp, \text{tag}^*, x_0, x_1)$ only occurs with negligible probability, indistinguishability then follows by indistinguishability from Claim 4.4.1.

For lossiness, let $x^* = x_0^* \leftarrow \{0, 1\}^\lambda$ be the target, and let $x_1^* \leftarrow \{0, 1\}^\lambda \setminus \{x_0^*\}$ be the value sampled during the generation of $\tilde{\text{fk}}$. We observe that $(\tilde{\text{fk}}, (\tilde{F}_{\text{fk}, \text{tag}}^{\text{tag}}(x^*))_{\text{tag} \neq \text{tag}^*})$ does not reveal anything about the values x_0^*, x_1^* beyond the (unordered) set $\{x_0^*, x_1^*\}$, since they are treated symmetrically during the generation of $\tilde{\text{fk}}$ and $\tilde{F}_{\text{fk}, \text{tag}}^{\text{tag}}(x_0^*) = \tilde{F}_{\text{fk}, \text{tag}}^{\text{tag}}(x_1^*)$ for all $\text{tag} \neq \text{tag}^*$. In other words, we have a Markov chain

$$x_0^* \rightarrow \{x_0^*, x_1^*\} \rightarrow (\tilde{\text{fk}}, (\tilde{F}_{\text{fk}, \text{tag}}^{\text{tag}}(x_0^*))_{\text{tag} \neq \text{tag}^*}).$$

On the other hand, given $\tilde{\text{fk}}$ and $\{x_0^*, x_1^*\}$, the value $\tilde{F}_{\text{fk}, \text{tag}^*}^{\text{tag}^*}(x_0^*)$ completely determines x_0^* since $\tilde{F}_{\text{fk}, \text{tag}^*}^{\text{tag}^*}(x_0^*) \neq \tilde{F}_{\text{fk}, \text{tag}^*}^{\text{tag}^*}(x_1^*)$. Therefore,

$$\begin{aligned} & H_\infty(\tilde{F}_{\text{fk}, \text{tag}^*}^{\text{tag}^*}(x_0^*) \mid \tilde{\text{fk}}, (\tilde{F}_{\text{fk}, \text{tag}}^{\text{tag}}(x_0^*))_{\text{tag} \neq \text{tag}^*}) \\ & \geq H_\infty(\tilde{F}_{\text{fk}, \text{tag}^*}^{\text{tag}^*}(x_0^*) \mid \tilde{\text{fk}}, (\tilde{F}_{\text{fk}, \text{tag}}^{\text{tag}}(x_0^*))_{\text{tag} \neq \text{tag}^*}, \{x_0^*, x_1^*\}) \\ & \geq H_\infty(x_0^* \mid \tilde{\text{fk}}, (\tilde{F}_{\text{fk}, \text{tag}}^{\text{tag}}(x_0^*))_{\text{tag} \neq \text{tag}^*}, \{x_0^*, x_1^*\}) \\ & \geq H_\infty(x_0^* \mid \{x_0^*, x_1^*\}) \\ & \geq 1. \end{aligned}$$

Note that the only property that used the injectivity of the PRGs in 4.4.1 is the injectivity of the T-ALBO. As we do not require injectivity for entropy-preserving T-ALBOs (the output size

being shorter than the input size anyway), we can instantiate the PRGs G and G' with standard (non-necessarily injective) PRGs. In particular, we obtain an entropy-preserving T-ALBO from any one-way function.

Amplifying Lossiness. Again, we can amplify the lossiness ℓ by concatenating images of the previous entropy-preserving T-ALBO on blocks of the input, which results in Theorem 4.4.

5 Applications of T-ALBOs

We first recall in Section 5.1 the definition of pseudo-entropy functions (PEF) introduced by [BHK11]. We note that any entropy-preserving T-ALBO directly gives such a PEF, thus giving a construction from one-way functions (Theorem 5.2). Then, we describe applications of PEF, by constructing (1) extractor-dependent extractors (Theorem 5.4), (2) leakage-resilient, deterministic MACs (Theorem 5.8), and (3) leakage-resilient, public-coin symmetric encryption schemes (Theorem 5.9).

5.1 Pseudo-Entropy Functions.

Definition 5.1 (Pseudo-Entropy Functions.). *A pseudo-entropy function family with input length $n = n(\lambda)$, output length $m = m(\lambda)$, and lossiness parameter $\ell = \ell(\lambda)$ consists of the following PPT algorithms $(\text{Gen}, \text{LossyGen}, f)$:*

- $k \leftarrow \text{Gen}(1^\lambda)$: generates a key k .
- $k \leftarrow \text{LossyGen}(1^\lambda, x^*)$: on input $x^* \in \{0, 1\}^n$, outputs a key k .
- $y = f_k(x)$: on input $x \in \{0, 1\}^n$, deterministically outputs $y \in \{0, 1\}^m$.

We require the following properties:

ℓ -Lossiness: For all $x^* \in \{0, 1\}^n$:

$$H_\infty(f_k(x^*)) | (f_k(x))_{x \neq x^*} \geq \ell$$

over the randomness of $k \leftarrow \text{LossyGen}(1^\lambda, x^*)$. We use $(f_k(x))_{x \neq x^*}$ to denote the (ordered) list of outputs of the function $(f_k(x))$ on all $2^n - 1$ possible inputs $x \neq x^*$.

Indistinguishability: For all $x^* \in \{0, 1\}^n$:

$$\text{Gen}(1^\lambda) \stackrel{c}{\approx} \text{LossyGen}(1^\lambda, x^*).$$

Next, we show the following:

Theorem 5.2 (PEFs from OWFs). *Let $\ell = \ell(\lambda)$ and $t = t(\lambda)$ be polynomials. Assuming the existence of one-way functions, there exists a PEF with input length t , output length ℓ , lossiness ℓ and key length $\ell\lambda$.*

Next, we note that PEFs are directly implied by any T-ALBO.

Claim 5.2.1. *Suppose $(\text{InjectiveGen}, \text{LossyGen}, F)$ is an entropy-preserving T-ALBO with input length n , output length m , tag length t and lossiness parameter ℓ . Then there exists a PEF with input length t , output length m , key size n and lossiness parameter ℓ .*

Proof. Let $(\text{InjectiveGen}, \text{LossyGen}, F)$ be such a T-ALBO. We define a PEF $(\text{Gen}, \text{LossyGen}, f)$ as follows:

- $\text{Gen}(1^\lambda)$: Sample a uniformly random $s \in \{0, 1\}^n$, and compute $\text{fk} \leftarrow \text{InjectiveGen}$. Output $k = (\text{fk}, s)$.
- $\text{LossyGen}(1^\lambda, x^*)$: Sample $s^* \leftarrow \{0, 1\}^n$, and compute $\text{fk} \leftarrow \text{LossyGen}(x^*, s^*)$, interpreting x^* as a tag for the T-ALBO and s^* in its input space. Output $k = (\text{fk}, s^*)$.
- $f_k(x)$: On input $x \in \{0, 1\}^t$ and a key $k = (\text{fk}, s)$, output $F_{\text{fk}, x}(s)$ (again interpreting x as a tag for the T-ALBO, and s as an input to the T-ALBO).

Indistinguishability follows immediately by indistinguishability of the T-ALBO. For lossiness, if $k^* \leftarrow \text{LossyGen}(x^*)$, then

$$\begin{aligned} H_\infty(f_k(x^*) \mid \{f_k(x)\}_{x \neq x^*}) &= H_\infty(F_{\text{fk}, x^*}(s^*) \mid (F_{\text{fk}, x^*}(s^*))_{x \neq x^*}) \\ &\geq H_\infty(F_{\text{fk}, x^*}(s^*) \mid \text{fk}, (F_{\text{fk}, x^*}(s^*))_{x \neq x^*}) \\ &\geq \ell, \end{aligned}$$

where the last equality is due to the entropy-preserving property of the T-ALBO. $\square$

Theorem 5.2 follows by combining Claim 5.2.1 with Theorem 4.4.

5.2 Extractors for Extractor-Dependent Sources

We first recall the definition of extractor-dependent source extractors (ED-extractors), defined in [DVW20]. The following is taken verbatim from [DVW20].

Definition 5.3 (Extractor-Dependent Source Extraction). *An extractor for α -entropy extractor-dependent sources (α -ED-Extractor) consists of two polynomial-time algorithms $(\text{SeedGen}, \text{EExt})$ with the following syntax:*

- $\text{seed} \leftarrow \text{SeedGen}(1^\lambda)$ is a randomized algorithm that generates seed .
- $\text{EExt}(s, \text{seed})$ is a deterministic algorithm that takes a sample $s \in \{0, 1\}^n$, together with seed and outputs a value $y \in \{0, 1\}^m$ for some polynomial length parameters $n = n(\lambda), m = m(\lambda)$.

Consider an adversarial source/distinguisher pair $(\mathcal{D}, \mathcal{S})$ and define the following extraction experiment $\text{EDGame}^{\mathcal{D}, \mathcal{S}}(1^\lambda)$:

- Sample a random bit $b \leftarrow \{0, 1\}$ and a random $\text{seed} \leftarrow \text{SeedGen}(1^\lambda)$.
- Run $(s, \text{aux}) \leftarrow \mathcal{D}^{\text{EExt}(\cdot, \text{seed})}(1^\lambda)$.
- If $b = 0$ set $r = \text{EExt}(s, \text{seed})$ else if $b = 1$ set $r \leftarrow \{0, 1\}^m$.
- Let $b' = \mathcal{S}(1^\lambda, \text{seed}, \text{aux}, r)$.

We say that $\mathcal{D}$ is an α -legal extractor-dependent source if the following conditions hold:

1. The probability that $\mathcal{D}$ queries its oracle on the value x that it outputs is negligible.
2. $H_\infty(X|\text{AUX}, \text{SEED}) \geq \alpha(\lambda)$, where $X, \text{SEED}, \text{AUX}$ denotes the joint distribution of the values $x, \text{seed}, \text{aux}$ in the above experiment.

An α -ED-Extractor is secure if for all α -legal polynomial-time sources $\mathcal{D}$ and all polynomial-time distinguishers $\mathcal{S}$, the above experiment satisfies

$$\left| \Pr[b = b'] - \frac{1}{2} \right| = \text{negl}(\lambda).$$

The rest of the section is dedicated to prove the following theorem:

Theorem 5.4 (ED-Extractors from OWFs). *Assuming the existence of one-way functions there exists an ED-extractor for α -entropy sources with auxiliary information, where $\alpha = \lambda^{\Omega(1)}$.*

We recall the definition of a 2-source extractor, and a construction due to Raz [Raz05].

Definition 5.5 ((Strong, Average-Case) Two-Source Extractor [CG88]). *We say that an efficient function $2\text{Ext} : \{0, 1\}^n \times \{0, 1\}^n \rightarrow \{0, 1\}^m$ is an (e_1, e_2, δ) -strong 2-source extractor if for all random variables (X_1, X_2, Z) such that X_1, X_2 are independent conditioned on Z and $H_\infty(X_1|Z) \geq e_1, H_\infty(X_2|Z) \geq e_2$ we have $\text{SD}((Z, X_2, 2\text{Ext}(X_1; X_2)), (Z, X_2, U_m)) \leq \delta$ where U_m is a uniformly string of length m .*

Theorem 5.6 ([Raz05]). *For any polynomial input length $n = \text{poly}(\lambda)$, any $e_1 = \lambda^{\Omega(1)}$ and any $e_2 = (1/2 + \Omega(1))n$, there exist (e_1, e_2, δ) -extractor with input length n , output length $m = \lambda^{\Omega(1)}$ and error $\delta = 2^{-\lambda^{\Omega(1)}}$.*

Next, we present a construction of an ED-extractor, starting from PEFs. Our construction essentially just abstracts out the construction of [DVW20], which relied on (non-targeted) all-lossy-but-one functions, in terms of PEFs. Doing so essentially shows that the construction [DVW20] only needs a *targeted* form of all-lossy-but-one functions.

Construction Outline. The basic idea is to set the seed $\text{seed} = k$ to consist of a PEF key k for a PEF f_k . The extractor on input x “hashes” x to some much smaller value z and computes $y = f_k(z)$ to be the PEF output. We then think of y as a seed to a standard seeded randomness extractor and output the extracted randomness $r = \text{Ext}(x; y)$. The idea behind the proof of security is to “guess” the value z that the sample x chosen by the source will hash to and select the PEF key K to preserve entropy on z . This is indistinguishable even if the adversary sees the seed later. The above change ensures that $y = f_k(z)$ is uniformly random and independent of all the values $y_i = f_k(z_i)$ that were used to compute prior extractor outputs. Therefore, this essentially corresponds to using a completely fresh and independent seed y when extracting from the challenge sample x , and guarantees that the extracted output looks random.

The above relies on hashing, but we can replace the hash by a standard PRF whose key is part of the seed, and rely on the fact that the source does not see the seed and therefore will not be able to cause a collision. A bigger issue is that the above argument relies on “guessing” and therefore has a super-polynomial security loss. This is fixed by using a series of PEFs with progressively larger

input sizes $i = 1, \dots, i_{\max} = \omega(\log \lambda)$ and targeting our guessing strategy for a particular input size i which is just slightly larger than the number of queries q made by the source and therefore ensures we guess correctly with inverse polynomial security. Doing this correctly, requires a more careful construction and proof.

Construction. We now describe our construction more formally.

Let $i_{\max}, j_{\max} = \omega(\lambda)$, and $w = j_{\max} i_{\max} (i_{\max} + 1)/2$. For $i \in [i_{\max}]$, let (Gen^i, f^i) be a PEF with input size i and lossiness ℓ . Let $F : \{0, 1\}^\lambda \times \{0, 1\}^n \rightarrow \{0, 1\}^v$, $G : \{0, 1\}^v \rightarrow \{0, 1\}^w$. Let $2\text{Ext} : \{0, 1\}^n \times \{0, 1\}^n \rightarrow \{0, 1\}^m$ be a strong 2-source extractor. We construct an ED-extractor $(\text{SeedGen}, \text{EExt})$ as follows:

- $\text{SeedGen}(1^\lambda)$: For $i \in [i_{\max}], j \in [j_{\max}]$, set $k_{i,j} \leftarrow \text{Gen}^i(1^\lambda)$, $K \leftarrow \{0, 1\}^\lambda$. Output $\text{seed} = (\{k_{i,j}\}, K)$.
- $\text{EExt}(s, \text{seed})$: Compute $z = F_K(s)$. Parse $G(z) = \{x_{i,j}\}_{i \in [i_{\max}], j \in [j_{\max}]}$ where $x_{i,j} \in \{0, 1\}^i$ for all $i \in [i_{\max}], j \in [j_{\max}]$. Compute, for all $i \in [i_{\max}], j \in [j_{\max}]$, $y_{i,j} = f_{k_{i,j}}^i(x_{i,j})$, and set $y = \bigoplus_{i,j} y_{i,j}$. Output $2\text{Ext}(x, y)$.

Theorem 5.7. Let $e_1 = \alpha - v - 1$, $e_2 = \ell - v - 1$. Assume 2Ext is a $(e_1, e_2, \text{negl}(\lambda))$ -extractor, that F is a PRF, G is a PRG, and the PEFs (Gen^i, f^i) satisfy ℓ -lossiness. Then $(\text{SeedGen}, \text{EExt})$ is an α -entropy secure ED-extractor.

Proof. For completeness, we describe the hybrid games which we directly adapt from [DVW20], and highlight any differences that occur.

Let $\mathcal{D}, \mathcal{S}$ be a source/distinguisher pair. Let q be the number of extractor queries made by $\mathcal{D}$, and let $i^* = \lceil \log q \rceil + 1$. We consider the following sequence of hybrids.

Hybrid H_0 . This is the ED-Extractor game with a source/distinguisher $\mathcal{D}, \mathcal{S}$. The game proceeds as follows:

- Sample a random bit $b \leftarrow \{0, 1\}$ and a random $\text{seed} \leftarrow \text{SeedGen}(1^\lambda)$.
- Run $(s, \text{aux}) \leftarrow \mathcal{D}^{\text{EExt}(\cdot, \text{seed})}(1^\lambda)$.
- If $b = 0$ set $r = \text{EExt}(s, \text{seed})$ else if $b = 1$ set $r \leftarrow \{0, 1\}^m$.
- Let $b' = \mathcal{S}(1^\lambda, \text{seed}, \text{aux}, r)$.

Hybrid H_1 . We change the way b' is computed. Let $x_{i,j}^* \in \{0, 1\}^i$ denote the values internally computed by the call $\text{EExt}(s, \text{seed})$ in the experiment. Let **BAD** denote the event that, any oracle call from $\mathcal{S}$ to $\text{EExt}(\cdot, \text{seed})$ sampled the internal value $x_{i^*,j} = x_{i^*,j}^*$ for all $j \in [j_{\max}]$.

We now **set $b' \leftarrow \{0, 1\}$ if **BAD** occurs** (and no modification otherwise).

Hybrids 0 and 1 are indistinguishable. This is because the probability that **BAD** occurs is negligible, by PRF security of F and PRG security of G , and where we crucially use $2^{i^*} \geq 2q$.

Hybrid H_2 . We now **sample** $j^* \leftarrow [j_{\max}]$ and $x^* \leftarrow \{0, 1\}^{i^*}$ in the beginning of the experiment, where we recall $i^* = i^* = \lceil \log q \rceil + 1$. Let **GUESS** denote the event that (1) **BAD** does not occur and (2) j^* is the least element in $[j_{\max}]$ such that no oracle call from $\mathcal{S}$ to $\text{EExt}(\cdot, \text{seed})$ sampled the internal value $x_{i^*, j^*} = x_{i^*, j^*}^*$ (which is well-defined if **BAD** does not occur) and (3) $x^* = x_{i^*, j^*}$.

We now **set** $b' \leftarrow \{0, 1\}$ if **¬GUESS** occurs (and no modification otherwise).

GUESS occurs with probability at least $\frac{1}{2 \cdot q \cdot j_{\max}} - \text{negl}(\lambda)$, and thus any distinguisher in H_1 with advantage Adv induces a distinguisher in H_2 with advantage $\text{Adv}/p - \text{negl}(\lambda)$, where $p = 2 \cdot q \cdot j_{\max}$. Note that we crucially use the fact that **GUESS** occurring is independent of the view of the adversary.

Hybrid H_3 . We switch how the seed is generated. After sampling $j^* \leftarrow [j_{\max}]$ and setting $i^* = \lceil \log q \rceil + 1$ and sampling $x^* \leftarrow \{0, 1\}^{i^*}$, we now **generate** $k_{i^*, j^*} \leftarrow \text{LossyGen}(1^\lambda, x^*)$.

Hybrids 2 and 3 are indistinguishable by security of the PEF.

Hybrid H_4 . We now **sample** $r \leftarrow \{0, 1\}^m$ regardless of b .

This is indistinguishable by security of 2Ext . The proof is almost identical to the one of [DVW20, Claims 5.11.1 and 5.11.2]. The only difference is how we argue that $Y_{i^*, j^*} = f_k(x^*)$ has min-entropy conditioned on $L = (f_k(x))_{x \neq x^*}$, $Z = F_K(s)$ the seed of the PRG G , and E which denotes whether **¬BAD** and **GUESS** hold simultaneously.

We have:

$$\begin{aligned} H_\infty(Y_{i^*, j^*} \mid L, Z, E) &\geq H_\infty(Y_{i^*, j^*} \mid L) - v - 1 \\ &\geq \ell - v - 1, \end{aligned}$$

where the first inequality follows from the fact that Z and E have size v and 1 respectively, and the second from ℓ -lossiness of the PEF. $\square$

Wrapping up. We show how to setup the parameters to obtain Theorem 5.4. We can set $\alpha = \lambda^{\Omega(1)}$, $n = \max(\alpha, \lambda)$, $i_{\max} = j_{\max} = \lambda^1$, $v = \min\{\alpha/2, \lambda^1\}$, and use PEFs with input size i , where $i \in [i_{\max}]$, and lossiness and output lengths $k = \ell = n$ (Theorem 5.2).

This makes use of the Raz 2-source extractor with $e_1 = \alpha - v - 1 \geq \alpha/2 - 1 = \lambda^{\Omega(1)}$ and $e_2 = \ell - v - 1 = (1 - o(n)) \cdot n$, $\delta = 2^{-\lambda^{\Omega(1)}}$ and output size $m = \lambda^{\Omega(1)}$ (Theorem 5.6). This overall proves Theorem 5.4.

5.3 Deterministic Leakage-Resilient MACs

We note here, as already observed by [BHK11], that PEFs with super-logarithmic lossiness ℓ directly give deterministic, leakage-resilient MACs with selective unforgeability security. Leakage resilience here denotes the fact that the adversary can learn any efficiently computable function of the secret key, as long as the output size of the function is bounded by some leakage-size parameter L . Selective unforgeability security means that the adversary chooses the message m^* on which it will produce a forgery ahead of time, before seeing the leakage or seeing any authentication tags for chosen messages. We refer the reader to [HLWW13] for a formal definition of leakage-resilient MACs.

Theorem 5.8 (Deterministic Leakage-Resilient MACs from OWFs). *Let $m = \text{poly}(\lambda)$ be any message length and $t \geq \omega(\log \lambda)$ be any tag length. Assuming one-way functions, there exists a*

deterministic MAC with message length m and tag length t that satisfies selective unforgeability even given leakage of size L , as long as $t - L = \omega(\log \lambda)$. The key length of the MAC is $t \cdot \lambda$.

Note that the amount of tolerated leakage in the above theorem is optimal relative to the tag size: if $t - L = O(\log \lambda)$ then the adversary can just leak the first L bits of the tag for some message m^* and guess the remaining $t - L$ bits with inverse polynomial probability. On the other hand, the tolerated leakage is not optimal relative to the key size since the ratio of leakage L to key size is $\leq 1/\lambda$. Under the DDH assumption, [BHK11] shows how to get leakage to key size ratio of $(1 - o(1))$.

Note that, by using complexity leveraging, one can obtain a fully-secure, leakage-resilient MAC with the same parameters, if we suppose the underlying PEF is sub-exponentially secure. This, in turn, can be based on the existence of sub-exponentially secure one-way functions.

We can build the MAC of Theorem 5.8 from any PEF, by viewing the PEF key k as the MAC key (generated using `Gen` for the actual scheme) and $f_k(x)$ as the MAC of message x . More formally, if $(\text{Gen}, \text{LossyGen}, f)$ is a PEF, we define the MAC as follows:

- `KeyGen`(1^λ): Output $k \leftarrow \text{Gen}(1^\lambda)$;
- `Sign`(k, x): Output $f_k(x)$;
- `Verify`(k, t, x): Output 1 if $t = f_k(x)$, and 0 otherwise.

For security, if k were generated as $k \leftarrow \text{LossyGen}(x^*)$, then $f_k(x^*)$ still would have (at least) $\ell - L$ bits of (min-)entropy given the MAC of *all* other messages $x \neq x^*$ and any leakage on k of size L , and would in particular be unpredictable whenever $\ell - L \geq \omega(\log \lambda)$. Moreover such a k is indistinguishable from the MAC key by indistinguishability of the PEF. Note here that we crucially need the reduction to know the target message x^* during the key generation, so that the above only yields a *selectively secure* MAC, where the adversary declares ahead of time the message over which he wants to produce a forgery. One can generically achieve full security using complexity leveraging, thus relying on stronger sub-exponential security of the PEF (which follows by sub-exponentially secure PRGs).

Theorem 5.8 follows by combining the construction above with Theorem 5.2.

5.4 Leakage-Resilient Symmetric Encryption

Next, we remark that any PEF also yields a leakage-resilient symmetric encryption scheme. In a leakage-resilient scheme, the adversary can get some leakage of the secret key, potentially after making some CPA-encryption queries but before seeing the challenge ciphertext. As in the case of MACs, we assume the leakage consists of an arbitrary efficiently computable function applied to the secret key, as long as the output size of the function is at most L bits for some leakage parameter L . Our scheme is public-coin, meaning that all the internal randomness of the encryption procedure is explicitly contained in the ciphertext. We refer the reader to [HLWW13] for a formal definition of leakage-resilient symmetric-key encryption.

Theorem 5.9 (Leakage-Resilient Symmetric Encryption from OWFs). *For any polynomial leakage parameter $L = L(\lambda)$ and message length $m = m(\lambda)$, assuming one-way functions exists, there exists a public-coin, leakage-resilient symmetric encryption scheme with message length m , ciphertext size $m + O(\lambda)$, and key size $O((L + \lambda)\lambda)$.*

We recall the definition of a (strong) seeded extractor.

Definition 5.10 (Strong Seeded Extractors). *An efficient function $\text{Ext} : \{0, 1\}^n \times \{0, 1\}^d \rightarrow \{0, 1\}^m$ is a strong (k, ϵ) -extractor if for every (n, k) -source X ,*

$$\text{SD}((U_d, \text{Ext}(X, U_d)), (U_d, U_m)) \leq \epsilon.$$

Theorem 5.11 ([GUV07]). *For every constant $\alpha > 0$, and all positive integers n, k and all $\epsilon > 0$, there is an explicit construction of a (k, ϵ) -extractor with input length n , output length $m = (1 - \alpha)k - O(\log n + \log(1/\epsilon))$ and seed length $d = O(\log n + \log(1/\epsilon))$.*

Construction. Fix a polynomial $L = L(\lambda)$ for the amount of tolerated leakage. Let $\ell = L + 3\lambda$ and let $(\text{Gen}, \text{LossyGen}, f)$ be a PEF with input length λ , lossiness parameter ℓ , output size ℓ and key length $L \cdot \lambda$. Let Ext be a $(2\lambda, \text{negl}(\lambda))$ -extractor with input length ℓ output length λ as guaranteed by Theorem 5.11, and let $G : \{0, 1\}^\lambda \rightarrow \{0, 1\}^m$ be a PRG. The encryption scheme is defined as:

- $\text{KeyGen}(1^\lambda)$: Output $k \leftarrow \text{Gen}(1^\lambda)$,
- $\text{Enc}(k, \text{msg})$: Sample a uniform seed for Ext , a uniform PEF input $x \leftarrow \{0, 1\}^\lambda$ and output
$$\text{ct} = (\text{seed}, x, G(\text{Ext}(f_k(x); \text{seed}))) \oplus \text{msg},$$
- $\text{Dec}(k, \text{ct} = (\text{seed}, x, y))$: Output $\text{msg} = G(\text{Ext}(f_k(x); \text{seed})) \oplus y$.

Claim 5.11.1. *Assuming the security of the underlying building blocks, the above construction gives a public-coin, symmetric encryption scheme with message length m and leakage bound L .*

Proof. We switch to an indistinguishable hybrid where we select the value x^* for the challenge ciphertext at the very beginning and set $k \leftarrow \text{LossyGen}(x^*)$. This is indistinguishable even given k in full and therefore certainly given leakage on k . We also choose the values x in all other encryption queries uniformly at random from $\{0, 1\}^\lambda \setminus \{x^*\}$, which is statistically indistinguishable. Now we argue that $f_k(x^*)$ has (at least) $\ell - L$ bits of (min-)entropy conditioned on the outputs of all the encryption queries and any leakage on k of size L . As a result, $\text{Ext}(\text{seed}, f_k(x))$ is statistically close to uniformly random. Therefore, we can rely on PRG security to switch $G(\text{Ext}(\text{seed}, f_k(x^*)))$ in the challenge ciphertext to a uniformly random value. In this case the adversary does not learn anything about msg which concludes the proof. $\square$

Combining the above with Theorem 5.2 and Theorem 5.11 gives the parameters of Theorem 5.9.

5.5 Symmetric-Key Encryption Secure against Selective Opening Attacks

The goal of this section is to prove the following:

Theorem 5.12 (Selective Opening Security from OWFs). *Assuming one-way functions exist, there exists a symmetric-key encryption scheme that achieves simulation-security against selective opening of keys and randomness (Definition 5.13).*

In Section 5.5.1, we recall the notion of security against selective opening attacks, and provide a simulation-security definition in the context of symmetric-key encryption (Definition 5.13). In Section 5.5.2, we consider a strengthened version of PEFs featuring some equivocability, and provide a construction based on one-way functions (Theorem 5.12).

5.5.1 Selective Opening Security.

We first define selective opening security in the symmetric-key setting. In the public-key encryption setting, an adversary, given ciphertexts encrypting correlated messages under different public keys, can adaptively request so-called *openings* of a subset of these ciphertexts. This allows the adversary to obtain either the randomness used to encrypt the subset of ciphertexts (sometimes referred to as sender corruptions) or subset of secret keys allowing to decrypt (sometimes referred to as receiver corruptions).

In the symmetric-key setting, an adversary receives ciphertexts of correlated messages encrypted under different secret keys. We focus on the setting where openings correspond to secret keys, which allow the adversary to decrypt the associated subset of ciphertexts. In fact, we will also handle randomness opening for free, as our scheme is public-coin, namely, all the randomness of the encryption is included as part of the ciphertext. In other words, we are able to provide security against both kinds of openings.

In terms of definitions, there are mainly two ways of formalizing selective opening security, either through a weaker indistinguishability-based definition, or a stronger simulation-based definition. The main drawback of the indistinguishability-based definition is that it significantly restricts the possible correlations of the encrypted messages. Fortunately, we manage to construct the stronger notion of simulation-based selective opening security against openings of both keys and randomness, which the notion we focus on in this paper.

One drawback of our simulation-based definition is that we only allow a single opening query per secret key in the experiment. This is essentially inherent: a secret key sk cannot explain more than sk ciphertexts to any arbitrary messages, or would otherwise violate entropy lower bounds.

Definition 5.13 (Selective Opening Simulation Security in the Symmetric-Key Setting). *Let $d = d(\lambda)$ be a polynomial. Let $(\text{KeyGen}, \text{Enc}, \text{Dec})$ be a symmetric-key encryption scheme. For PPT algorithms $\mathcal{A} = (\mathcal{A}_1, \mathcal{A}_2, \mathcal{A}_3)$ and $\mathcal{S} = (\mathcal{S}_1, \mathcal{S}_2, \mathcal{S}_3)$ and $\mathcal{P}$, we consider the following experiments:*

Experiment $\text{Exp}_{\text{Sim-SOA}}^{\text{real}}(\mathcal{A}, 1^\lambda)$:

1. Sample, for all $i \in [d]$, $\text{sk}_i \leftarrow \text{KeyGen}(1^\lambda)$.
2. The adversary computes $(\mathcal{D}, \text{state}_1) \leftarrow \mathcal{A}_1^{\text{Enc}(\text{sk}_1, \cdot), \dots, \text{Enc}(\text{sk}_d, \cdot)}(1^\lambda)$, where $\mathcal{D}$ is a distribution over $\mathcal{M}^n$, where $\mathcal{M}$ is the message space of the encryption scheme.
3. Sample $\vec{\text{msg}}^* \leftarrow \mathcal{D}$.
4. Compute for all $i \in [d]$: $\text{ct}_i^* = \text{Enc}(\text{sk}_i, \text{msg}_i^*; r_i^*)$ where r_i^* are the random coins used to encrypt.
5. The adversary computes $(I, \text{state}_2) \leftarrow \mathcal{A}_2^{\text{Enc}(\text{sk}_1, \cdot), \dots, \text{Enc}(\text{sk}_d, \cdot)}(\text{state}_1, \text{ct}_1^*, \dots, \text{ct}_d^*)$, where $I \subseteq [d]$.
6. The adversary computes $\text{output} \leftarrow \mathcal{A}_3^{\text{Enc}(\text{sk}_1, \cdot), \dots, \text{Enc}(\text{sk}_d, \cdot)}(\text{state}_2, (\text{msg}_i^*, r_i^*, \text{sk}_i)_{i \in I})$.
7. The output of the experiment is the bit $\mathcal{P}(\mathcal{D}, \vec{\text{msg}}^*, I, \text{output})$.

Experiment $\text{Exp}_{\text{Sim-SOA}}^{\text{ideal}}(\mathcal{S}, 1^\lambda)$:

1. The simulator samples $(\mathcal{D}, \text{state}_1) \leftarrow \mathcal{S}_1(1^\lambda)$.
2. Sample $\vec{\text{msg}}^* \leftarrow \mathcal{D}$.

3. The simulator computes $(I, \text{state}_2) \leftarrow \mathcal{S}_2(\text{state}_1, 1^{|\text{msg}_1^*|}, \dots, 1^{|\text{msg}_d^*|})$.
4. The simulator computes $\text{output} \leftarrow \mathcal{S}_3(\text{state}_2, (\text{msg}_i^*)_{i \in I})$.
5. The output of the experiment is the bit $\mathcal{P}(\mathcal{D}, \overrightarrow{\text{msg}}^*, I, \text{output})$.

We say that $(\text{KeyGen}, \text{Enc}, \text{Dec})$ is $\text{Sim} - \text{SOA}$ -secure if for all polynomial $d = d(\lambda)$, for all PPT adversary $\mathcal{A}$ in $\text{Exp}_{\text{Sim-SOA}}^{\text{real}}(\mathcal{A}, 1^\lambda)$ and for all distinguisher $\mathcal{P}$, there exists a simulator $\mathcal{D}$ in $\text{Exp}_{\text{Sim-SOA}}^{\text{ideal}}(\mathcal{S}, 1^\lambda)$ such that:

$$\left| \Pr[\text{Exp}_{\text{Sim-SOA}}^{\text{real}}(\mathcal{A}, 1^\lambda) = 1] - \Pr[\text{Exp}_{\text{Sim-SOA}}^{\text{ideal}}(\mathcal{S}, 1^\lambda) = 1] \right| \leq \text{negl}(\lambda).$$

5.5.2 Equivocal T-ALBOs and PEFs

Our main building block to achieve selective opening security is a stronger notion of PEFs that we call *equivocal* PEFs.

Definition 5.14 (Equivocal Pseudo-Entropy Functions). *The syntax of an equivocal PEF with input size n and output size m is similar to the one of a PEF (Definition 5.1), except that $\text{LossyGen}(1^\lambda, x^*)$ outputs (k, state) where state is some secret state. We require an additional PPT algorithm*

- $\text{Equivocate}(\text{state}, y) \rightarrow k_{\text{sim}}^*$: on input $y \in \{0, 1\}^m$ (in the output space of the PEF) and state , output k_{sim}^* (in the key of the PEF).

We require the following properties:

- *Consistency*: For any $x^* \in \{0, 1\}^t, y \in \{0, 1\}^m$: if $(k^*, \text{state}) \leftarrow \text{LossyGen}(1^\lambda, \text{tag}^*, x^*)$ and $k_{\text{sim}}^* \leftarrow \text{Equivocate}(\text{state}, y)$ then $f_k(x) = f_{k_{\text{sim}}^*}(x)$ for all $x \neq x^*$. Furthermore, $f_{k_{\text{sim}}^*}(x^*) = y$.
- *Indistinguishability*: The following distributions are computationally indistinguishable for all $x^* \in \{0, 1\}^n$:

$$(x^*, k) \stackrel{c}{\approx} (x^*, k_{\text{sim}}^*),$$

where $x^* \leftarrow \{0, 1\}^n, k \leftarrow \text{Gen}(1^\lambda), (k^*, \text{state}) \leftarrow \text{LossyGen}(1^\lambda, \text{tag}^*, x^*), y \leftarrow \{0, 1\}^m$, and $k_{\text{sim}}^* \leftarrow \text{Equivocate}(\text{state}, y)$.

Remark 5.1 (Relation with Somewhere Equivocal PRFs ([HJO⁺16])). We note, perhaps surprisingly, that the notion of equivocal PEFs exactly matches the notion of (1-point equivocal) *somewhere equivocal PRFs* (SEPRFs) [HJO⁺16]. In fact, [HJO⁺16] builds SEPRFs from one-way functions using techniques very similar to the distributed point function of [BGI15], which, as previously discussed, are connected to our notion of T-ALBO and therefore PEFs. Interestingly, [HJO⁺16] introduced SEPRFs in a yet different context from our other applications (but not entirely unrelated to selective opening security), that is, adaptive security of Yao's garbled circuits. Indeed, it was used to construct a form of (somewhere) non-committing symmetric-key encryption, and non-committing public-key encryption is known to imply selective key-opening security in the public-key setting [HPW15].

Next, we show the following:

Theorem 5.15 (Equivocal PEFs from OWFs). *Let $\ell = \ell(\lambda)$ and $t = t(\lambda)$ be any polynomials, and let $m \leq 3\lambda + t$. Assuming the existence of one-way functions, there exists an equivocal PEF with input length t , output length ℓ and lossiness ℓ .*

Construction. We build an equivocal PEF by building an associate form of *equivocal T-ALBO*, and seeing the resulting T-ALBO as a PEF (as in Claim 5.2.1).

Definition 5.16 (Equivocal T-ALBO). *The syntax of an equivocal T-ALBO is similar to the one of a T-ALBO (Definition 3.2), except that $\text{LossyGen}(1^\lambda, \text{tag}^*, x^*)$ outputs $(\text{fk}, \text{state})$ where state is some secret state. We require an additional PPT algorithm*

- $\text{Equivocate}(\text{state}, y) \rightarrow x_{\text{sim}}^*$: on input $y \in \{0, 1\}^m$ (that is, the output space of the T-ALBO) and state , output $x_{\text{sim}}^* \in \{0, 1\}^n$ (that is, the input space of the T-ALBO).

We require the following properties:

- *Consistency:* For any $x^* \in \{0, 1\}^n, \text{tag}^* \in \{0, 1\}^t, y \in \{0, 1\}^m$: if $(\text{fk}, \text{state}) \leftarrow \text{LossyGen}(1^\lambda, \text{tag}^*, x^*)$ and $x_{\text{sim}}^* \leftarrow \text{Equivocate}(\text{state}, y)$ then $F_{\text{fk}, \text{tag}}(x^*) = F_{\text{fk}, \text{tag}}(x_{\text{sim}}^*)$ for all $\text{tag} \neq \text{tag}^*$. Furthermore, $F_{\text{fk}, \text{tag}^*}(x_{\text{sim}}^*) = y$.
- *Indistinguishability:* The following distributions are computationally indistinguishable for all $\text{tag}^* \in \{0, 1\}^t$:

$$(\text{tag}^*, \text{fk}_{\text{inj}}, x^*) \stackrel{c}{\approx} (\text{tag}^*, \text{fk}_{\text{los}}, x_{\text{sim}}^*),$$

where $x^* \leftarrow \{0, 1\}^n, \text{fk}_{\text{inj}} \leftarrow \text{InjectiveGen}(1^\lambda), (\text{fk}_{\text{los}}, \text{state}) \leftarrow \text{LossyGen}(1^\lambda, \text{tag}^*, x^*), y \leftarrow \{0, 1\}^m$, and $x_{\text{sim}}^* \leftarrow \text{Equivocate}(\text{state}, y)$.

Our construction of an equivocal T-ALBO with one-bit output follows directly from our construction of entropy-preserving T-ALBO with one-bit output (Theorem 4.4). We only slightly modify the key generation in lossy mode to output a state, and define our new equivocation algorithm:

- $\widetilde{\text{LossyGen}}(1^\lambda, \text{tag}^*, x^*)$: Sample $\text{fk} \leftarrow \widetilde{\text{LossyGen}}(1^\lambda, \text{tag}^*, x^*)$, which also internally samples $x_1^* \leftarrow \{0, 1\}^\lambda$. If $F_{\text{fk}, \text{tag}^*}(x^*) = F_{\text{fk}, \text{tag}^*}(x_1^*)$, repeat the above. Otherwise, compute $y^* = \widetilde{F}_{\text{fk}, \text{tag}^*}(x^*)$ and $y_1 = \widetilde{F}_{\text{fk}, \text{tag}^*}(x_1^*)$. Set $a = e \cdot (y^* - y_1)^{-1}$ and **state** = $(\text{tag}^*, x^*, x_1^*)$. Output $(\text{fk}, a, \text{state})$.
- $\text{Equivocate}(\text{state}, y)$: On input $\text{state} = (\text{tag}^*, x^*, x_1^*)$ and $y \in \{0, 1\}$, compute $b_0 = \widetilde{F}_{\text{fk}, \text{tag}^*}(x^*)$, $b_1 = \widetilde{F}_{\text{fk}, \text{tag}^*}(x_1^*)$. If $b_0 = y$, output x^* , and output x_1^* otherwise.

We keep the algorithms $\widetilde{\text{InjectiveGen}}$ and $\widetilde{F}$ unchanged, as in the construction of Claim 4.4.1. Note that Equivocate is well-defined as $b_0 \neq b_1$ by construction, so that either $y = b_0$ or $y = b_1$.

Consistency follows by the proof of lossiness of Claim 4.4.1, where we proved that $F_{\text{fk}, \text{tag}}(x^*) = F_{\text{fk}, \text{tag}}(x_1^*)$ for all $\text{tag} \neq \text{tag}^*$. The equality $F_{\text{fk}, \text{tag}^*}(x_{\text{sim}}^*) = y$ is by construction.

For indistinguishability, we have by indistinguishability of the entropy-preserving T-ALBO that for all $\text{tag}^* \in \{0, 1\}^t$, $(\text{tag}^*, \text{fk}_{\text{inj}}, x^*) \stackrel{c}{\approx} (\text{tag}^*, \text{fk}_{\text{los}}, x^*)$ where $x^* \leftarrow \{0, 1\}^n, \text{fk}_{\text{inj}} \leftarrow \text{InjectiveGen}(1^\lambda)$, and $\text{fk}_{\text{los}} \leftarrow \text{LossyGen}(1^\lambda, \text{tag}^*, x^*)$. Then, we showed in the proof of Claim 4.4.1 that the views $(\text{tag}^*, \text{fk}_{\text{los}}, x^*)$ and $(\text{tag}^*, \text{fk}_{\text{los}}, x_1^*)$ are identically distributed. Therefore, the distribution $(\text{tag}^*, \text{fk}, x_{\text{sim}}^*)$, where $x^* \leftarrow \{0, 1\}^n, (\text{fk}, \text{state}) \leftarrow \text{LossyGen}(1^\lambda, \text{tag}^*, x^*), y \leftarrow \{0, 1\}^m, x_{\text{sim}}^* \leftarrow \text{Equivocate}(\text{state}, y)$, which picks a random element of $\{x^*, x_1^*\}$, is also indistinguishable from $(\text{tag}^*, \text{fk}_{\text{inj}}, x^*)$.

Finally, output size can be amplified by concatenating the output of many such equivocal T-ALBOs with 1-bit output. Combined with the above, this gives the following:

Theorem 5.17. *Let $k = k(\lambda)$ and $t = t(\lambda)$ be any polynomials. Assuming the existence of one-way functions, there exists an equivocal T-ALBO with input size $\{0, 1\}^{k\lambda}$, tag space $\{0, 1\}^t$, and output space $\{0, 1\}^k$.*

As in Claim 5.2.1, equivocal PEFs can be readily constructed from any equivocal T-ALBO by defining inputs of the PEF as T-ALBO tags, and PEF keys as T-ALBO inputs. This yields Theorem 5.15.

5.5.3 Selective Opening Security from Equivocal PEFs

We now describe our construction of a symmetric-key encryption scheme which is secure against selective opening of keys and randomness, thus proving Theorem 5.12.

Our main idea will be to use an equivocal PEF to construct a symmetric-key encryption scheme which enjoys a (weak) key-equivocation property. Achieving security against selective opening of randomness will be for free as our scheme is public-coin.

Construction. Let $(\text{Gen}, \text{LossyGen}, f)$ be an equivocal PEF (Definition 5.14) with input size n and output size m . We define the following encryption scheme:

- $\text{KeyGen}(1^\lambda)$: Sample $k \leftarrow \text{Gen}(1^\lambda)$, and output $\text{sk} = k$.
- $\text{Enc}(\text{sk}, \text{msg})$: Sample $x \leftarrow \{0, 1\}^n$ and output
$$\text{ct} = (x, f_k(x) \oplus \text{msg}).$$
- $\text{Dec}(\text{sk}, \text{ct})$: Parse ct as $\text{ct} = (x, y)$, compute

$$\text{msg} = y \oplus f_k(x).$$

Theorem 5.18. *Suppose $(\text{Gen}, \text{LossyGen}, f)$ is an equivocal PEF (Definition 5.14). Then $(\text{KeyGen}, \text{Enc}, \text{Dec})$ is simulation-secure against selective openings of keys and randomness (Definition 5.13).*

Proof. Let $\mathcal{A} = (\mathcal{A}_1, \mathcal{A}_2, \mathcal{A}_3)$ be any PPT adversary in $\text{Exp}_{\text{Sim-SOA}}^{\text{real}}$, and $\mathcal{P}$ be a distinguisher. We define our simulator $\mathcal{S} = (\mathcal{S}_1, \mathcal{S}_2, \mathcal{S}_3)$ for $\text{Exp}_{\text{Sim-SOA}}^{\text{ideal}}$ as follows.

- $\mathcal{S}_1(1^\lambda)$: For all $i \in [d]$, sample $x_i^* \leftarrow \{0, 1\}^n$, and compute $(k_i^*, \text{state}_i^{\text{PEF}}) \leftarrow \text{LossyGen}(1^\lambda, x_i^*)$. Compute $(\mathcal{D}, \text{state}_1^A) \leftarrow \mathcal{A}_1^{\text{Enc}(\text{sk}_1, \cdot), \dots, \text{Enc}(\text{sk}_d, \cdot)}(1^\lambda)$ where encryption queries are answered using k_i^* , as:

$$\text{Enc}(\text{sk}_i, m) = (x, f_{k_i^*}(x) \oplus \text{msg}),$$

where $x \leftarrow \{0, 1\}^n$. Set $\text{state}_1 = (\text{state}_1^A, (\text{state}_i^{\text{PEF}}, x_i^*)_{i \in [d]})$. Output $(\mathcal{D}, \text{state}_1)$.

- $\mathcal{S}_2(\text{state}_1, 1^{|\text{msg}_1|}, \dots, 1^{|\text{msg}_d|})$: Compute, for all $i \in [d]$:

$$\text{ct}_i^* = (x_i^*, y_i^*),$$

where $y_i^* \leftarrow \{0, 1\}^m$.

Compute $(I, \text{state}_2^A) \leftarrow \mathcal{A}_2^{\text{Enc}(\text{sk}_1, \cdot), \dots, \text{Enc}(\text{sk}_d, \cdot)}(\text{state}_1^A, \text{ct}_1^*, \dots, \text{ct}_d^*)$ where again encryption queries are answered using k_i^* , as:

$$\text{Enc}(\text{sk}_i, m) = (x, f_{k_i^*}(x) \oplus \text{msg}).$$

Output $(I, \text{state}_2 = (\text{state}_2^A, \text{state}_1))$.

- $\mathcal{S}_3(\text{state}_2, \{\text{msg}_i^*\}_{i \in I})$: Compute, for all $i \in I$, $k_{i,\text{sim}}^* \leftarrow \text{Equivocate}(\text{state}_i^{\text{PEF}}, y_i^* \oplus \text{msg}_i^*)$.
Compute and output $\text{output} \leftarrow \mathcal{A}_3^{\text{Enc}(\text{sk}_1, \cdot), \dots, \text{Enc}(\text{sk}_d, \cdot)}(\text{state}_2, (\text{msg}_i^*, x_i^*, k_{i,\text{sim}}^*)_{i \in I})$.

We show that the distribution of $\text{Exp}_{\text{Sim-SOA}}^{\text{real}}(\mathcal{A}, 1^\lambda)$ and $\text{Exp}_{\text{Sim-SOA}}^{\text{ideal}}(\mathcal{S}, 1^\lambda)$ are computationally indistinguishable by considering a series of hybrid distributions.

Hybrid H_0 . This is the distribution induced by $\text{Exp}_{\text{Sim-SOA}}^{\text{real}}(\mathcal{A}, 1^\lambda)$.

Hybrid H_1 . We make the following changes:

- We change the way the randomness used to compute the challenge ciphertexts is sampled. We now **sample, for all $i \in [d]$, $x_i^* \leftarrow \{0, 1\}^n$ before sampling secret keys.**
- We change how we answer encryption queries. We **abort the experiment if any encryption query for secret key i samples x_i^* as randomness.**

Hybrid $H_{2,i}$, $i \in [d]$. We make the following changes:

- We change the way the setup of the experiment is performed. We now **sample $x_i^* \leftarrow \{0, 1\}^n$ and set $(k_i^*, \text{state}_i^{\text{PEF}}) \leftarrow \text{LossyGen}(1^\lambda, x_i^*)$.** All the encryption queries throughout the experiment are now answered using k_i^* .
- We change the way the challenge ciphertext for index i is computed. We now compute

$$\text{ct}_i^* = (x_i^*, y_i^* \oplus \text{msg}_i^*),$$

where $y_i^* \leftarrow \{0, 1\}^m$.

- We change the way the values r_i^*, sk_i given to $\mathcal{A}_3$ are computed. Given $(\text{msg}_i^*)_{i \in I}$, we now compute, if $i \in I$:

$$k_{i,\text{sim}}^* \leftarrow \text{Equivocate}(\text{state}_i^{\text{PEF}}, y_i^*),$$

and $\mathcal{A}_3$ is now given as input $(x_i^*, k_{i,\text{sim}}^*)$.

Hybrid H_3 . We make the following changes:

- We do not abort anymore if an encryption query for secret key i samples $x = x_i^*$.
- We change again the way the challenge ciphertexts are computed. We now compute for all $i \in [d]$

$$\text{ct}_i^* = (x_i^*, y_i^*),$$

where $y_i^* \leftarrow \{0, 1\}^m$.

- We change the way the values sk_i given to $\mathcal{A}_3$ are computed. Given $(\text{msg}_i^*)_{i \in I}$, we now compute for all $i \in I$:

$$k_{i,\text{sim}}^* \leftarrow \text{Equivocate}(\text{state}_i^{\text{PEF}}, y_i^* \oplus \text{msg}_i^*).$$

Observe that the view generated by hybrid H_3 now corresponds to the view generated by the simulator $\mathcal{S} = (\mathcal{S}_1, \mathcal{S}_2, \mathcal{S}_3)$ in $\text{Exp}_{\text{Sim-SOA}}^{\text{ideal}}(\mathcal{S}, 1^\lambda)$.

We now prove that the successive hybrids are indistinguishable.

Claim 5.18.1. *The views generated in hybrids H_0 and H_1 are within negligible statistical distance.*

Proof. The only difference between the two hybrids occurs whenever an encryption query for index i samples as randomness $x = x_i^*$. This happens with probability $1/2^n$, and by union bound over the polynomial number Q of encryption queries made by $\mathcal{A}$ in $\text{Exp}_{\text{Sim-SOA}}^{\text{real}}(\mathcal{A}, 1^\lambda)$, this occurs with probability at most $Q/2^n = \text{negl}(\lambda)$. $\square$

Claim 5.18.2. *Suppose the equivocal PEF $(\text{Gen}, \text{LossyGen}, f)$ satisfies consistency and indistinguishability. Then, the views generated in hybrids H_1 and $H_{2,1}$ are indistinguishable. Similarly, the views generated in hybrids $H_{2,i}$ and $H_{2,i+1}$ are indistinguishable for all $i \in \{1, \dots, d-1\}$.*

Proof. Let us argue that H_1 and $H_{2,1}$ are indistinguishable: the proof for hybrids $H_{2,i}$ and $H_{2,i+1}$, $i \in [d-1]$ is identical.

By indistinguishability of $(\text{Gen}, \text{LossyGen}, f)$, we have:

$$(x_1^*, k_1, f_{k_1}(x_1^*)) \stackrel{c}{\approx} (x^*, k_{1,\text{sim}}^*, y_1)$$

where $x_1^* \leftarrow \{0, 1\}^n$, $k_1 \leftarrow \text{KeyGen}(1^\lambda)$, $(k_1^*, \text{state}_1^{\text{PEF}}) \leftarrow \text{LossyGen}(1^\lambda, x_1^*)$ and $k_{1,\text{sim}}^* \leftarrow \text{Equivocate}(\text{state}_1^{\text{PEF}}, y)$. This is because one can compute the third element of the distributions (x, k, y) as $y = f_k(x)$, by consistency in lossy mode.

It remains to argue that the distribution of the answers to the encryption queries are indistinguishable. This follows by consistency, which ensures that $(x, f_{k_1^*}(x))$ is identically distributed to $(x, f_{k_{1,\text{sim}}^*}(x))$ over the randomness of $x \leftarrow \{0, 1\}^n$, unless $x = x_1^*$, but both hybrids abort the experiment in that case. $\square$

Claim 5.18.3. *The views generated in hybrids $H_{2,d}$ and H_3 are identically distributed.*

Proof. For all $i \in [d]$ and all $\text{msg}_i^* \in \{0, 1\}^m$, the distributions $u_i \oplus \text{msg}_i^*$ and u_i^* are uniform over $\{0, 1\}^m$ over the randomness of $u_i^* \leftarrow \{0, 1\}^m$ alone. Setting $y_i^* = u_i \oplus \text{msg}_i^*$ yields hybrid $H_{2,d}$, and $y_i^* = u_i$ hybrid H_3 .

Finally, the probability of aborting because an encryption query samples $x = x_i^*$ is $Q/2^n$ (where Q is the number of encryption queries made by $\mathcal{A}$), which is negligible. $\square$

$\square$

Combined with Theorem 5.15, we obtain Theorem 5.12.

6 Application of T-AIBOs to CCA Security

We prove in this section the following theorem:

Theorem 6.1 (CCA Encryption from Strong Trapdoor Functions). *Let $d = d(\lambda)$, $n = n(\lambda) = \omega(\log \lambda)$, $\rho = d \cdot n$, and $m = \max(n + 1, \lambda)$. Let TDF be a trapdoor function with input length ρ . Suppose that no time $T = 2^n \cdot \text{poly}(\lambda)$ adversary can invert TDF with probability $\frac{2^d}{2^\rho} \cdot \epsilon$ for any non-negligible ϵ . Assume furthermore the existence of an injective PRG $G : \{0, 1\}^n \rightarrow \{0, 1\}^m$.*

Then there exists a CCA-secure (public-key) encryption scheme.

In Section 6.1, we recall the notion of encryption schemes with randomness recovery, and show that a standard construction from trapdoor functions satisfies a special form of leakage resilience Lemma 6.4. In Section 6.2, we show how to build a public-key CCA-secure encryption scheme assuming the existence of strongly secure trapdoor functions, proving Theorem 6.1.

6.1 Encryption with Randomness Recovery

Another core component of our construction is a public-key encryption scheme with randomness recovery, which we base on trapdoor functions (TDF). Informally, such an encryption has two additional properties over standard public-key encryption, namely (1) the decryption algorithm also recovers the random coins used to encrypt, and (2) one can alternatively decrypt ciphertexts using the random coins used to encrypt (as opposed to traditionally with the secret key).

Definition 6.2 (Public-Key Encryption with Randomness Recovery). *A public-key encryption scheme with randomness recovery $(\text{KeyGen}, \text{Enc}, \text{Dec}, \text{Recover})$ has the following syntax:*

- $(\text{pk}, \text{sk}) \leftarrow \text{KeyGen}(1^\lambda)$: *On input the security parameter, output a public pk and a secret key sk .*
- $\text{ct} \leftarrow \text{Enc}(\text{pk}, m)$: *On input a public key pk and a message m , samples some random coins r and output a ciphertext $\text{ct} = \text{Enc}(\text{pk}, m; r)$.*
- $(m, r) \leftarrow \text{Dec}(\text{sk}, \text{ct})$: *On input a secret key sk and a ciphertext ct , output a message m and random coins r .*
- $m \leftarrow \text{Recover}(\text{pk}, \text{ct}, r)$: *On input a public key pk , a ciphertext ct and random coins r , output a message m .*

We require the following correctness properties:

Correctness. *We require that correctness holds perfectly, except with negligible probability over $(\text{pk}, \text{sk}) \leftarrow \text{KeyGen}(1^\lambda)$, namely:*

$$\Pr[\exists m, r, \text{Dec}(\text{sk}, \text{Enc}(\text{pk}, m; r)) \neq (m, r) \vee \text{Recover}(\text{pk}, \text{Enc}(\text{pk}, m; r)) \neq m] \leq \text{negl}(\lambda),$$

over the probability of $(\text{pk}, \text{sk}) \leftarrow \text{KeyGen}(1^\lambda)$.

Security. *We require standard CPA security.*

For our constructions, we will additionally need some form of leakage resilience, namely that security holds even given some particular leakage on the *random coins* used to encrypt. Note that even though one can generically add leakage resilience against randomness leakage by using strong seeded extractors, this does not preserve randomness recovery.

Instead, we will directly show that the particular construction of public-key encryption with randomness recovery from trapdoor functions of [HKW20] is resilient against a particular form of leakage generated by a T-AIBO, assuming strong security of the underlying trapdoor function. We first recall the definition of an injective trapdoor function.

Definition 6.3 (Injective Trapdoor Functions (TDF)). *An injective trapdoor function family $\text{TDF} = (\text{TDF.Setup}, \text{TDF.Eval}, \text{TDF.Invert})$ with input length $\rho = \text{poly}(\lambda)$ and output length $m = \text{poly}(\lambda)$ has the following syntax:*

- $\text{TDF.Setup}(1^\lambda)$: on input the security parameter, output a public key pk and a secret key sk .
- $\text{TDF.Eval}(\text{pk}, x)$: a deterministic algorithm, which, on input a public key pk and an input $x \in \{0, 1\}^\rho$, output $y \in \{0, 1\}^m$.
- $\text{TDF.Invert}(\text{sk}, y)$: on input a secret key sk and an output $y \in \{0, 1\}^m$, output $x \in \{0, 1\}^\rho$.

We require the following properties:

Correctness. We require that:

$$\Pr[\exists x, \text{TDF.Invert}(\text{sk}, \text{TDF.Eval}(\text{pk}, x)) \neq x] \leq \text{negl}(\lambda),$$

over the randomness of $(\text{pk}, \text{sk}) \leftarrow \text{TDF.Setup}(1^\lambda)$, namely, that with overwhelming probability over the setup alone, inversion is perfectly correct (which also implies that $\text{TDF.Eval}(\text{pk}, \cdot)$ is injective).

Security. We say that TDF is hard to invert if for any PPT adversary $\mathcal{A}$, there exists a negligible function ϵ such that:

$$\Pr[x \leftarrow \mathcal{A}(\text{pk}, y)] \leq \epsilon,$$

over the randomness of $(\text{pk}, \text{sk}) \leftarrow \text{TDF.Setup}$, $x \leftarrow \{0, 1\}^\rho$, and where $y = \text{TDF.Eval}(\text{pk}, x)$.

We will also consider strengthened forms of security where $\mathcal{A}$ is allowed to run in specific, potentially super-polynomial time T , and require his success probability to be at most ϵ for some specific function ϵ .

Next, we recall the construction of a randomness-recoverable PKE from any TDF [HKW20].

Let $\text{TDF} = (\text{TDF.Setup}, \text{TDF.Eval}, \text{TDF.Invert})$ be a trapdoor function with input space $\{0, 1\}^\rho$. Let hc denote the Goldreich-Levin hard-core bit [GL89]. We define the following encryption scheme:

- $\text{CPA.KeyGen}(1^\lambda)$: Sample $(\text{pk}, \text{sk}) \leftarrow \text{TDF.Setup}$. Sample some randomness coins for hc , and output $(\text{CPA.pk} = (\text{pk}, \text{coins}), \text{CPA.sk} = (\text{sk}, \text{coins}))$.
- $\text{CPA.Enc}(\text{CPA.pk}, m \in \{0, 1\})$: Sample $r \leftarrow \{0, 1\}^\rho$, and output

$$\text{ct} = (\text{TDF.Eval}(\text{pk}, r), \text{hc}(r; \text{coins}) \oplus m).$$

- $\text{CPA.Dec}(\text{CPA.sk}, \text{ct})$: On input $\text{ct} = (z, b)$, compute $r = \text{TDF.Invert}(\text{sk}, z)$, and output $m = \text{hc}(r; \text{coins}) \oplus b$.
- $\text{CPA.Recover}(\text{pk}, \text{ct}, r)$: On input $\text{ct} = (z, b)$, compute $z' = \text{TDF.Eval}(\text{pk}, r)$. Abort if $z \neq z'$. Otherwise output $m = \text{hc}(r; \text{coins}) \oplus b$.

Note that the random coins used by CPA.Enc correspond exactly to a random input to TDF .

Next, we show that this particular encryption scheme satisfies a specific form of leakage-resilience, provided the trapdoor function TDF is secure enough.

Lemma 6.4 (Leakage-Resilience of CPA). *Let $n = n(\lambda)$, $d = d(\lambda)$ and $t = t(\lambda)$ and $\ell = \ell(\lambda)$ be polynomials. Let $\rho = d \cdot n$. Let $\text{TDF} = (\text{TDF.Setup}, \text{TDF.Eval}, \text{TDF.Invert})$ be a trapdoor function with input space $\{0, 1\}^\rho$. Let $(\text{InjectiveGen}, \text{LossyGen}, F)$ be the specific T-AIBO of Section 4.2 with input space $\{0, 1\}^n$ and tag space $\{0, 1\}^t$ and lossiness ℓ .*

For all $\text{tag}^ \in \{0, 1\}^t$, consider the two following distributions:*

$$\begin{aligned} &((\text{fk}_i)_{i \in [d]}, \text{pk}, \text{TDF.Eval}_{\text{pk}}(r), (F_{\text{fk}_i, \text{tag}^*}(r_i))_{i \in [d]}, \text{coins}, \text{hc}(r; \text{coins})) \\ &((\text{fk}_i)_{i \in [d]}, \text{pk}, \text{TDF.Eval}_{\text{pk}}(r), (F_{\text{fk}_i, \text{tag}^*}(r_i))_{i \in [d]}, \text{coins}, \text{hc}(r; \text{coins}) \oplus 1), \end{aligned}$$

where $r = (r_1 \| \dots \| r_d) \leftarrow \{0, 1\}^\rho$, $\text{fk}_i \leftarrow \text{LossyGen}(\text{tag}^, r_i)$, $\text{pk} \leftarrow \text{TDF.Setup}$, and hc denotes the Goldreich-Levin hard-core bit [GL89] and coins uniformly sampled randomness for hc .*

Suppose that no time T adversary can invert TDF with probability $\frac{2^{\ell d}}{2^\rho} \cdot \epsilon$ for any non-negligible ϵ , and suppose that the T-AIBO satisfies ℓ -lossiness. Then no time $T \cdot \text{poly}(\lambda)$ adversary can distinguish the two distributions above with non-negligible success probability.

Proof. Suppose there exists an PPT adversary with runtime T that distinguishes the two distributions above with non-negligible success probability ϵ . Then there exists a predictor $\mathcal{P}$ with runtime $\text{poly}(\rho) \cdot T$, which, on input $(\text{TDF.Eval}_{\text{pk}}(r), (F_{\text{fk}_i, \text{tag}^*}(r_i))_{i \in [d]})$, outputs r with non-negligible probability $\Omega(\epsilon)$.

Given such a predictor $\mathcal{P}$, we build an inverter for TDF as follows.

- $\mathcal{I}(\text{TDF.Eval}_{\text{pk}}(r))$: Sample, for $i \in [d]$, $r_i^* \leftarrow \{0, 1\}^n$, and set

$$r^* = \mathcal{P}((\text{TDF.Eval}_{\text{pk}}(r), (F_{\text{fk}_i, \text{tag}^*}(r_i^*))_{i \in [d]}).$$

Output r^* .

First, for any fixed $i \in [d]$, the probability over $r_i^* \leftarrow \{0, 1\}^n$ that $F_{\text{fk}_i, \text{tag}^*}(r_i^*) = F_{\text{fk}_i, \text{tag}^*}(r_i)$ is at least $\frac{2^\ell}{2^n}$. This is because in the construction of Theorem 4.2, each output $F_{\text{fk}_i, \text{tag}^*}(r_i)$ has at least 2^ℓ preimages for all $r_i \in \{0, 1\}^n$ and all $\text{fk}_i \leftarrow \text{LossyGen}(\text{tag}^*, r_i)$.

As a result, the probability that this happens for all $i \in [d]$ is at least $\frac{2^{\ell \cdot d}}{2^{n \cdot d}} = \frac{2^{\ell \cdot d}}{2^\rho}$. In other words, $\mathcal{I}$ correctly guesses all values $(F_{\text{fk}_i, \text{tag}^*}(r_i^*))_{i \in [d]}$ with probability $\frac{2^{\ell \cdot d}}{2^\rho}$, and therefore, by correctness of $\mathcal{P}$, $\mathcal{I}$ correctly outputs r with probability at least $\frac{2^{\ell \cdot d}}{2^\rho} \cdot \epsilon$, contradicting the security of TDF. $\square$

Remark 6.1. One can interpret Lemma 6.4 as showing that the encryption scheme with randomness recovery CPA of [HKW20] is resilient to some particular form of leakage, namely, leakage generated by the T-AIBO F of Section 4.2 (in lossy mode). Most notably, we point out that (1) the leakage is over the encryption randomness, which we can usually fix generically using strong seeded extractors (but does not work here as it would not preserve randomness recovery) and (2) it is a particular form of *entropy-bounded* leakage ([NS09, DHLW10]) in the sense that the size of the leakage is quite bigger than the entropy loss, namely, $|F_{\text{fk}, \text{tag}}(r)| > \ell$. This is as opposed to the perhaps more common form of *bounded-size* leakage.

Remark 6.2 (Usage of a Specific T-AIBO). Lemma 6.4 is stated as using the specific leakage function, defined by the specific T-AIBO of Section 4.2. In fact, the only property we use is that, for all fixed input, the leakage function can be guessed with slightly better probability than trivial. As a result, any T-AIBO satisfying this property can be used in Lemma 6.4, and consequently in our construction Section 6.2.

Remark 6.3 (Running Time VS Success Probability). Lemma 6.4 and its proof involve algorithms that have low running time but small success probability, namely, we build a time $T \cdot \text{poly}(\lambda)$ inverter with success probability $\frac{2^{\ell-d}}{2^\rho} \cdot \epsilon$. We note here that we can also build a time $T \cdot \frac{2^\rho}{2^{\ell-d}} \cdot \text{poly}(\lambda)$ inverter that succeeds with probability $\Omega(\epsilon)$. This is by having the inverter attempt $\frac{2^\rho}{2^{\ell-d}}$ samples for the values r_i^* . Then, with constant probability (over the randomness of these samples alone), one of them is right, conditioned on which the predictor $\mathcal{P}$ succeeds in inverting with probability ϵ .

This, in turn allows to alternatively instantiate Theorem 6.1 with a TDF such that no time $T = 2^{\rho-(\ell d-n)}$ adversary succeeds in inverting TDF with non-negligible probability.

6.2 CCA-secure Encryption from T-AIBOs and Trapdoor Functions

We now describe our construction of a CCA-secure encryption scheme. Let $d = d(\lambda)$, $n = n(\lambda)$, $t = t(\lambda)$, $\rho = d \cdot n$. We will use the following components:

- the T-AIBO $(\text{InjectiveGen}, \text{LossyGen}, F)$ of Section 4.2 with input length n and tag space $\{0, 1\}^t$, satisfying ℓ -lossiness from Theorem 4.2;
- the encryption scheme $(\text{CPA.KeyGen}, \text{CPA.Enc}, \text{CPA.Dec}, \text{CPA.Recover})$ of Section 6.1, with encryption randomness space $\{0, 1\}^\rho$ and ciphertext size $\{0, 1\}^\tau$, which can be built from any trapdoor function TDF with input length ρ . In particular it satisfies *randomness recovery* (Definition 6.2), and satisfies some special form of leakage resilience with respect to the T-AIBO above (Lemma 6.4). More precisely, we will assume that no time $T = 2^n \cdot \text{poly}(\lambda)$ adversary can invert the underlying TDF with probability $\frac{2^d}{2^\rho} \cdot \epsilon$ for any non-negligible ϵ .
- a one-time strongly unforgeable signature $(\text{Sig.KeyGen}, \text{Sig.Sign}, \text{Sig.Verify})$ with message space $\{0, 1\}^{\tau+d \cdot m}$, where τ is the size of ciphertexts from CPA, and with verification keys of size (at most) t .

We now describe our CCA-secure encryption scheme:

- $\text{KeyGen}(1^\lambda)$: Set $(\text{CPA.pk}, \text{CPA.sk}) \leftarrow \text{CPA.KeyGen}$ and compute, for all $i \in [d]$: $\text{fk}_i \leftarrow \text{InjectiveGen}$. Set $\text{pk} = (\text{CPA.pk}, (\text{fk})_{i \in [d]})$ and $\text{sk} = \text{CPA.sk}$ and output (pk, sk) .
- $\text{Enc}(\text{pk}, \text{msg})$: Sample $r \leftarrow \{0, 1\}^\rho$. Compute

$$C = \text{CPA.Enc}(\text{CPA.pk}, \text{msg}; r).$$

Sample $(\text{Sig.vk}, \text{Sig.sk}) \leftarrow \text{Sig.KeyGen}$. Parse r as $r = (r_1 \| \dots \| r_d)$ where $r_i \in \{0, 1\}^n$. Parsing the verification key of the signature as a tag for the T-AIBO, compute, for $i \in [d]$:

$$y_i = F_{\text{fk}_i, \text{Sig.vk}}(r_i).$$

Compute

$$\sigma = \text{Sig.Sign}(\text{Sig.sk}, (C \| y_1 \| \dots \| y_d)),$$

and output

$$\text{ct} = (\text{Sig.vk}, C, (y_i)_{i \leq d}, \sigma).$$

- $\text{Dec}(\text{sk}, \text{ct})$: Parse ct as $(\text{Sig.vk}, C, (y_i)_{i \leq d}, \sigma)$. Output $\perp$ if $\text{Sig.Verify}(\text{Sig.vk}, (C \| y_1 \| \dots \| y_d), \sigma) = 0$.

Else compute $(\text{msg}, r) \leftarrow \text{CPA.Dec}(\text{ct})$, using the randomness recovery property of the CPA-encryption scheme. Parse r as $(r_1 \| \dots \| r_d)$ where $r_i \in \{0, 1\}^n$ for all $i \in [d]$.

If $\text{CPA.Enc}(\text{msg}, r) = C$, and if $F_{\text{fk}_i, \text{Sig.vk}}(r_i) = y_i$, output msg . Otherwise output $\perp$.

Claim 6.4.1. *Suppose that no time $T = 2^n \cdot \text{poly}(\lambda)$ adversary can invert the TDF underlying CPA with probability $\frac{2^d}{2^p} \cdot \epsilon$ for any non-negligible ϵ , that T-AIBO is the one constructed in Section 4.2, and $(\text{Sig.KeyGen}, \text{Sig.Sign}, \text{Sig.Verify})$ is one-time strongly unforgeable. Then $(\text{KeyGen}, \text{Enc}, \text{Dec})$ is CCA-secure.*

As noted in Remark 6.2, one could replace the specific T-AIBO of Section 4.2 with one satisfying a generic non-trivial guessing property in lossy mode.

Proof. We consider a sequence of hybrids.

Hybrid H_0 . This is the standard CCA experiment.

Hybrid H_1 . The challenger changes how it answers decryption queries and the challenge ciphertext. It now picks $(\text{Sig.vk}^*, \text{Sig.sk}^*) \leftarrow \text{Sig.KeyGen}$ at the beginning of the experiment, and **answers $\perp$ to any decryption query whose first component is Sig.vk^* .**

The challenge ciphertext is generated using $(\text{Sig.vk}^*, \text{Sig.sk}^*)$, namely, the other components of the challenge ciphertext are computed as:

$$\begin{aligned} C^* &= \text{CPA.Enc}(\text{CPA.pk}, \text{msg}_b^*; r^*), \\ y_i^* &= F_{\text{fk}_i, \text{Sig.vk}^*}(r_i^*), \\ \sigma^* &= \text{Sig.Sign}(\text{Sig.sk}^*, (C \| y_1^* \| \dots \| y_d^*)), \end{aligned}$$

where $r^* = (r_1^* \| \dots \| r_d^*) \leftarrow \{0, 1\}^\rho$, and $\text{msg}_0^*, \text{msg}_1^*$ are the challenge messages sent by the adversary. The challenge ciphertext is $\text{ct}^* = (\text{Sig.vk}^*, C^*, (y_i^*)_{i \in [d]}, \sigma^*)$.

Hybrid H_2 . The challenger changes how it generates the public key of the scheme, and more precisely how it samples the T-AIBO keys fk_i . It now samples $r^* = (r_1^* \| \dots \| r_d^*) \leftarrow \{0, 1\}^\rho$ and computes fk_i as $\text{fk}_i \leftarrow \text{LossyGen}(\text{Sig.vk}^*, r_i^*)$.

Hybrid H_3 . The challenger changes how it answers decryption queries. Given a ciphertext $\text{ct} = (\text{Sig.vk}, C, (y_i)_{i \leq d}, \sigma)$, the challenger:

1. Checks that $\text{Sig.Verify}(\text{Sig.vk}, (C \| y_1 \| \dots \| y_d), \sigma) = 1$, and outputs $\perp$ otherwise;
2. **For all $i \in [d]$, it enumerates over all $r_i \in \{0, 1\}^n$ the values $F_{\text{fk}_i, \text{Sig.vk}}(r_i)$, and outputs the first such r_i such that $F_{\text{fk}_i, \text{Sig.vk}}(r_i) = y_i$. Otherwise it outputs $\perp$;**
3. Using the randomness recovery property of $(\text{CPA.KeyGen}, \text{CPA.Enc}, \text{CPA.Dec})$, it recovers $m \leftarrow \text{Recover}(\text{pk}, C, r)$ where $r = (r_1 \| \dots \| r_d)$, and checks that $\text{ct} = \text{CPA.Enc}(\text{CPA.pk}, \text{msg}; r)$, and outputs $\perp$ otherwise;
4. Outputs msg .

Hybrid H_4 . The challenger changes how it produces the challenge ciphertext. It now computes $C = \text{CPA.Enc}(\text{CPA.pk}, 0; r^*)$.

We now prove that successive hybrids are indistinguishable.

Claim 6.4.2. *Suppose Sig is one-time strongly unforgeable. Then the views of the adversary in H_0 and H_1 are indistinguishable.*

Proof. The views in H_0 and H_1 differ exactly when the adversary makes a decryption query with respect to the verification key of the challenge ciphertext Sig.vk^* such that the associated signature σ verifies. We distinguish several cases.

- The ciphertext of such a query uses the same signature σ^* and message $(C^* \| y_1^* \| \dots \| y_t^*)$ as the challenge ciphertext, namely queries the challenge ciphertext. By definition of the CCA experiment, this corresponds to the case where the adversary guesses the challenge ciphertext and queries it before receiving the actual challenge ciphertext from the challenger. This only happens with negligible probability.
- Otherwise, the ciphertext of such a query uses a different signature $\sigma \neq \sigma^*$. We argue that such an adversary then induces an adversary against the one-time strongly unforgeability of Sig . A reduction samples the parameters (pk, sk) of the real scheme, and uses them to answer decryption queries. It then receives $\text{msg}_0, \text{msg}_1$ from the adversary, and chooses a random bit $b \leftarrow \{0, 1\}$. It interacts with the unforgeability experiment challenger to receive a verification key vk^* . As in the real scheme, computes $C^* = \text{CPA.Enc}(\text{CPA.pk}, m_b; r^*)$, and $y_i^* = F_{\text{fk}_i, \text{Sig.vk}}(r_i^*)$ where $r^* = (r_1^* \| \dots \| r_d^*) \leftarrow \{0, 1\}^\rho$. It queries the challenger for the unforgeability experiment with message $m = (C^* \| y_1^* \| \dots \| y_d^*)$, and receives a signature σ^* . It sets the challenge ciphertext as $\text{ct}^* = (\text{Sig.vk}^*, C^*, (y_i)_{i \in [d]}, \sigma^*)$.

Now whenever the adversary for the CCA game decryption query $\text{ct} \neq \text{ct}^*$ containing Sig.vk^* such that the associated signature verifies, either the message or the signature differs from the challenge ciphertext (by analysis of the previous case above). In either case, such a decryption query then induces a forgery for Sig .

□

Claim 6.4.3. *Suppose $(\text{InjectiveGen}, \text{LossyGen}, F)$ satisfies indistinguishability. Then the views of the adversary in H_1 and H_2 are indistinguishable.*

Proof. This follows by switching how the keys fk_i are generated, one by one, from InjectiveGen to $\text{LossyGen}(\text{vk}^*, r_i^*)$, over $i \in [d]$. The reduction between two consecutive sub-hybrids picks $\text{tag}^* = \text{Sig.vk}^*$, receives $s^* = r_i^* \in \{0, 1\}^n$ and fk , and uses fk in the public key and r_i^* as the appropriate randomness block to generate the challenge ciphertext. □

Claim 6.4.4. *Suppose $(\text{InjectiveGen}, \text{LossyGen}, F)$ satisfies injectivity on injective branches, and that CPA is statistically correct. Then the views of the adversary in H_2 and H_3 are distributed within negligible statistical distance.*

Proof. First, note that in both hybrids, all queries of the form $\text{ct} = (\text{Sig.vk}^*, C, (y_i)_{i \in [d]}, \sigma)$ are answered with $\perp$. By the injectivity on injective branches of the T-AIBO, one can check that in both hybrids, a decryption query on $\text{ct} = (\text{Sig.vk}, C, (y_i)_{i \in [d]}, \sigma)$ where $\text{Sig.vk} \neq \text{Sig.vk}^*$ does not output $\perp$ if and only if:

- $\text{Sig.Verify}(\text{vk}, (C\|y_1\| \cdots \|y_d), \sigma) = 1$;
- r_i is the only input such that $F_{\text{fk}_i, \text{Sig.vk}}(r_i) = y_i$ (where $\text{Sig.vk} \neq \text{Sig.vk}^*$ ensures injectivity of $F_{\text{fk}_i, \text{Sig.vk}}(\cdot)$ by injectivity on injective branches);
- msg is correctly recovered as $\text{msg} \leftarrow \text{Recover}(\text{pk}, C, m)$;
- $C = \text{CPA.Enc}(\text{CPA.pk}, \text{msg}; r)$.

In particular, the hybrids differ whenever $F_{\text{fk}_i, \text{Sig.vk}}$ is not injective or Recover is not correct, which both happen with negligible probability by injectivity on injective branches of $\text{fk}_i \leftarrow \text{LossyGen}(\text{Sig.vk}^*, r_i^*)$ and correctness of $(\text{CPA.pk}, \text{CPA.sk}) \leftarrow \text{CPA.KeyGen}(1^\lambda)$, respectively. $\square$

Claim 6.4.5. *Suppose $\text{CPA} = (\text{CPA.KeyGen}, \text{CPA.Enc}, \text{CPA.Dec})$ is instantiated with a trapdoor function TDF such that no time $T = 2^n \cdot \text{poly}(\lambda)$ adversary can invert TDF with probability $\frac{2^d}{2^\rho} \cdot \epsilon$ for any non-negligible ϵ . Then the views of the adversary in H_3 and H_4 are indistinguishable.*

Proof. We show that any polynomial-time distinguisher between H_3 and H_4 induces a time $d \cdot 2^n \cdot \text{poly}(\lambda)$ against the distribution of Lemma 6.4. The intuition is that Lemma 6.4 ensures that CPA is leakage resilient when the leakage is computed using the T-AIBO F (in lossy mode).

Our reduction samples $(\text{Sig.vk}^*, \text{Sig.sk}^*)$ and sets $\text{tag}^* = \text{Sig.vk}^*$. It receives a sample from the distribution of Lemma 6.4

$$((\text{fk}_i)_{i \in [d]}, \text{CPA.pk}, \text{TDF.Eval}_{\text{pk}}(r^*), (F_{\text{fk}_i, \text{tag}^*}(r_i^*))_{i \in [d]}, b),$$

where $r = (r_1 \| \dots \| r_d) \leftarrow \{0, 1\}^\rho$, $\text{fk}_i \leftarrow \text{LossyGen}(\text{tag}^*, r_i)$, $\text{pk} \leftarrow \text{TDF.Setup}$, and $b \in \{0, 1\}$.

The reduction sets the public key as $\text{pk} = (\text{CPA.pk}, (\text{fk}_i)_{i \in [d]})$, and sets the challenge ciphertext as

$$\begin{aligned} C^* &= \text{CPA.Enc}(\text{CPA.pk}, \text{msg}^*; r^*) = (\text{TDF.Eval}_{\text{pk}}(r^*), b), \\ y_i^* &= F_{\text{fk}_i, \text{Sig.vk}^*}(r_i^*), \\ \sigma^* &= \text{Sig.Sign}(\text{Sig.sk}^*, (C\|y_1^*\| \cdots \|y_d^*\|), \end{aligned}$$

It answers decryption queries as in Hybrid H_3 , so that each decryption query is answered in time $d \cdot 2^n$. This simulates perfectly Hybrid H_3 if $b = \text{hc}(r) \oplus m$, and Hybrid H_4 if $b = \text{hc}(r)$, and in particular any distinguisher between H_3 and H_4 induce a time $d \cdot 2^n$ distinguisher for the distribution of Lemma 6.4. Then, Lemma 6.4 implies that such a distinguisher implies a time $2^n \cdot \text{poly}(\lambda)$ inverter against TDF with success probability $\frac{2^d}{2^\rho} \cdot \epsilon$ for some non-negligible ϵ . $\square$

This finishes the proof of CCA-security of $(\text{KeyGen}, \text{Enc}, \text{Dec})$. $\square$

Theorem 6.1 follows from Claim 6.4.1, using the fact that one-time strongly unforgeable signatures are implied by one-way functions (which are implied by either PRGs or trapdoor functions).

References

- [ADW09] Joël Alwen, Yevgeniy Dodis, and Daniel Wichs. Leakage-resilient public-key cryptography in the bounded-retrieval model. In Shai Halevi, editor, *Advances in Cryptology – CRYPTO 2009*, volume 5677 of *Lecture Notes in Computer Science*, pages 36–54. Springer, Heidelberg, August 2009.
- [AGV09] Adi Akavia, Shafi Goldwasser, and Vinod Vaikuntanathan. Simultaneous hardcore bits and cryptography against memory attacks. In Omer Reingold, editor, *TCC 2009: 6th Theory of Cryptography Conference*, volume 5444 of *Lecture Notes in Computer Science*, pages 474–495. Springer, Heidelberg, March 2009.
- [BBN⁺09] Mihir Bellare, Zvika Brakerski, Moni Naor, Thomas Ristenpart, Gil Segev, Hovav Shacham, and Scott Yilek. Hedged public-key encryption: How to protect against bad randomness. In Mitsuru Matsui, editor, *Advances in Cryptology – ASIACRYPT 2009*, volume 5912 of *Lecture Notes in Computer Science*, pages 232–249. Springer, Heidelberg, December 2009.
- [BDWY12] Mihir Bellare, Rafael Dowsley, Brent Waters, and Scott Yilek. Standard security does not imply security against selective-opening. In David Pointcheval and Thomas Johansson, editors, *Advances in Cryptology – EUROCRYPT 2012*, volume 7237 of *Lecture Notes in Computer Science*, pages 645–662. Springer, Heidelberg, April 2012.
- [BFO08] Alexandra Boldyreva, Serge Fehr, and Adam O’Neill. On notions of security for deterministic encryption, and efficient constructions without random oracles. In David Wagner, editor, *Advances in Cryptology – CRYPTO 2008*, volume 5157 of *Lecture Notes in Computer Science*, pages 335–359. Springer, Heidelberg, August 2008.
- [BGI15] Elette Boyle, Niv Gilboa, and Yuval Ishai. Function secret sharing. In Elisabeth Oswald and Marc Fischlin, editors, *Advances in Cryptology – EUROCRYPT 2015, Part II*, volume 9057 of *Lecture Notes in Computer Science*, pages 337–367. Springer, Heidelberg, April 2015.
- [BHK11] Mark Braverman, Avinatan Hassidim, and Yael Tauman Kalai. Leaky pseudo-entropy functions. In Bernard Chazelle, editor, *ICS 2011: 2nd Innovations in Computer Science*, pages 353–366. Tsinghua University Press, January 2011.
- [BHY09] Mihir Bellare, Dennis Hofheinz, and Scott Yilek. Possibility and impossibility results for encryption and commitment secure under selective opening. In Antoine Joux, editor, *Advances in Cryptology – EUROCRYPT 2009*, volume 5479 of *Lecture Notes in Computer Science*, pages 1–35. Springer, Heidelberg, April 2009.
- [CG88] Benny Chor and Oded Goldreich. Unbiased bits from sources of weak randomness and probabilistic communication complexity. *SIAM J. Comput.*, 17(2):230–261, 1988.
- [CPW20] Suvaradip Chakraborty, Manoj Prabhakaran, and Daniel Wichs. Witness maps and applications. In Aggelos Kiayias, Markulf Kohlweiss, Petros Wallden, and Vassilis Zikas, editors, *PKC 2020: 23rd International Conference on Theory and Practice of Public Key Cryptography, Part I*, volume 12110 of *Lecture Notes in Computer Science*, pages 220–246. Springer, Heidelberg, May 2020.

- [DGK17] Yevgeniy Dodis, Siyao Guo, and Jonathan Katz. Fixing cracks in the concrete: Random oracles with auxiliary input, revisited. In Jean-Sébastien Coron and Jesper Buus Nielsen, editors, *Advances in Cryptology – EUROCRYPT 2017, Part II*, volume 10211 of *Lecture Notes in Computer Science*, pages 473–495. Springer, Heidelberg, April / May 2017.
- [DHLW10] Yevgeniy Dodis, Kristiyan Haralambiev, Adriana López-Alt, and Daniel Wichs. Cryptography against continuous memory attacks. In *51st Annual Symposium on Foundations of Computer Science*, pages 511–520. IEEE Computer Society Press, October 2010.
- [DNRS99] Cynthia Dwork, Moni Naor, Omer Reingold, and Larry J. Stockmeyer. Magic functions. In *40th Annual Symposium on Foundations of Computer Science*, pages 523–534. IEEE Computer Society Press, October 1999.
- [DORS08] Yevgeniy Dodis, Rafail Ostrovsky, Leonid Reyzin, and Adam D. Smith. Fuzzy extractors: How to generate strong keys from biometrics and other noisy data. *SIAM J. Comput.*, 38(1):97–139, 2008.
- [DVW20] Yevgeniy Dodis, Vinod Vaikuntanathan, and Daniel Wichs. Extracting randomness from extractor-dependent sources. In Anne Canteaut and Yuval Ishai, editors, *Advances in Cryptology – EUROCRYPT 2020, Part I*, volume 12105 of *Lecture Notes in Computer Science*, pages 313–342. Springer, Heidelberg, May 2020.
- [FGK⁺13] David Mandell Freeman, Oded Goldreich, Eike Kiltz, Alon Rosen, and Gil Segev. More constructions of lossy and correlation-secure trapdoor functions. *Journal of Cryptology*, 26(1):39–74, January 2013.
- [FHKW10] Serge Fehr, Dennis Hofheinz, Eike Kiltz, and Hoeteck Wee. Encryption schemes secure against chosen-ciphertext selective opening attacks. In Henri Gilbert, editor, *Advances in Cryptology – EUROCRYPT 2010*, volume 6110 of *Lecture Notes in Computer Science*, pages 381–402. Springer, Heidelberg, May / June 2010.
- [GGH19] Sanjam Garg, Romain Gay, and Mohammad Hajiabadi. New techniques for efficient trapdoor functions and applications. In Yuval Ishai and Vincent Rijmen, editors, *Advances in Cryptology – EUROCRYPT 2019, Part III*, volume 11478 of *Lecture Notes in Computer Science*, pages 33–63. Springer, Heidelberg, May 2019.
- [GGM86] Oded Goldreich, Shafi Goldwasser, and Silvio Micali. How to construct random functions. *Journal of the ACM*, 33(4):792–807, October 1986.
- [GI14] Niv Gilboa and Yuval Ishai. Distributed point functions and their applications. In Phong Q. Nguyen and Elisabeth Oswald, editors, *Advances in Cryptology – EUROCRYPT 2014*, volume 8441 of *Lecture Notes in Computer Science*, pages 640–658. Springer, Heidelberg, May 2014.
- [GKK20] Ankit Garg, Yael Tauman Kalai, and Dakshita Khurana. Low error efficient computational extractors in the CRS model. In Anne Canteaut and Yuval Ishai, editors, *Advances in Cryptology – EUROCRYPT 2020, Part I*, volume 12105 of *Lecture Notes in Computer Science*, pages 373–402. Springer, Heidelberg, May 2020.

- [GL89] Oded Goldreich and Leonid A. Levin. A hard-core predicate for all one-way functions. In *21st Annual ACM Symposium on Theory of Computing*, pages 25–32. ACM Press, May 1989.
- [GUV07] V. Guruswami, C. Umans, and S. Vadhan. Unbalanced expanders and randomness extractors from parvaresh-vardy codes. In *Twenty-Second Annual IEEE Conference on Computational Complexity (CCC’07)*, pages 96–108, 2007.
- [HJO⁺16] Brett Hemenway, Zahra Jafargholi, Rafail Ostrovsky, Alessandra Scafuro, and Daniel Wichs. Adaptively secure garbled circuits from one-way functions. In Matthew Robshaw and Jonathan Katz, editors, *Advances in Cryptology – CRYPTO 2016, Part III*, volume 9816 of *Lecture Notes in Computer Science*, pages 149–178. Springer, Heidelberg, August 2016.
- [HKW20] Susan Hohenberger, Venkata Koppula, and Brent Waters. Chosen ciphertext security from injective trapdoor functions. In Daniele Micciancio and Thomas Ristenpart, editors, *Advances in Cryptology – CRYPTO 2020, Part I*, volume 12170 of *Lecture Notes in Computer Science*, pages 836–866. Springer, Heidelberg, August 2020.
- [HLOV11] Brett Hemenway, Benoît Libert, Rafail Ostrovsky, and Damien Vergnaud. Lossy encryption: Constructions from general assumptions and efficient selective opening chosen ciphertext security. In Dong Hoon Lee and Xiaoyun Wang, editors, *Advances in Cryptology – ASIACRYPT 2011*, volume 7073 of *Lecture Notes in Computer Science*, pages 70–88. Springer, Heidelberg, December 2011.
- [HLWW13] Carmit Hazay, Adriana López-Alt, Hoeteck Wee, and Daniel Wichs. Leakage-resilient cryptography from minimal assumptions. In Thomas Johansson and Phong Q. Nguyen, editors, *Advances in Cryptology – EUROCRYPT 2013*, volume 7881 of *Lecture Notes in Computer Science*, pages 160–176. Springer, Heidelberg, May 2013.
- [Hof12] Dennis Hofheinz. All-but-many lossy trapdoor functions. In David Pointcheval and Thomas Johansson, editors, *Advances in Cryptology – EUROCRYPT 2012*, volume 7237 of *Lecture Notes in Computer Science*, pages 209–227. Springer, Heidelberg, April 2012.
- [HPW15] Carmit Hazay, Arpita Patra, and Bogdan Warinschi. Selective opening security for receivers. In Tetsu Iwata and Jung Hee Cheon, editors, *Advances in Cryptology – ASIACRYPT 2015, Part I*, volume 9452 of *Lecture Notes in Computer Science*, pages 443–469. Springer, Heidelberg, November / December 2015.
- [HR14] Dennis Hofheinz and Andy Rupp. Standard versus selective opening security: Separation and equivalence results. In Yehuda Lindell, editor, *TCC 2014: 11th Theory of Cryptography Conference*, volume 8349 of *Lecture Notes in Computer Science*, pages 591–615. Springer, Heidelberg, February 2014.
- [HRW16] Dennis Hofheinz, Vanishree Rao, and Daniel Wichs. Standard security does not imply indistinguishability under selective opening. In Martin Hirt and Adam D. Smith, editors, *TCC 2016-B: 14th Theory of Cryptography Conference, Part II*, volume 9986

- of *Lecture Notes in Computer Science*, pages 121–145. Springer, Heidelberg, October / November 2016.
- [KOS10] Eike Kiltz, Adam O’Neill, and Adam Smith. Instantiability of RSA-OAEP under chosen-plaintext attack. In Tal Rabin, editor, *Advances in Cryptology – CRYPTO 2010*, volume 6223 of *Lecture Notes in Computer Science*, pages 295–313. Springer, Heidelberg, August 2010.
 - [MW20] Tal Moran and Daniel Wichs. Incompressible encodings. In Daniele Micciancio and Thomas Ristenpart, editors, *Advances in Cryptology – CRYPTO 2020, Part I*, volume 12170 of *Lecture Notes in Computer Science*, pages 494–523. Springer, Heidelberg, August 2020.
 - [NS09] Moni Naor and Gil Segev. Public-key cryptosystems resilient to key leakage. In Shai Halevi, editor, *Advances in Cryptology – CRYPTO 2009*, volume 5677 of *Lecture Notes in Computer Science*, pages 18–35. Springer, Heidelberg, August 2009.
 - [PW08] Chris Peikert and Brent Waters. Lossy trapdoor functions and their applications. In Richard E. Ladner and Cynthia Dwork, editors, *40th Annual ACM Symposium on Theory of Computing*, pages 187–196. ACM Press, May 2008.
 - [Raz05] Ran Raz. Extractors with weak random seeds. In Harold N. Gabow and Ronald Fagin, editors, *37th Annual ACM Symposium on Theory of Computing*, pages 11–20. ACM Press, May 2005.
 - [Zha16] Mark Zhandry. The magic of ELFs. In Matthew Robshaw and Jonathan Katz, editors, *Advances in Cryptology – CRYPTO 2016, Part I*, volume 9814 of *Lecture Notes in Computer Science*, pages 479–508. Springer, Heidelberg, August 2016.