



Is ML-Based Cryptanalysis Inherently Limited? Simulating Cryptographic Adversaries via Gradient-Based Methods

Avital Shafran¹(✉), Eran Malach², Thomas Ristenpart³, Gil Segev¹,
and Stefano Tessaro⁴

¹ School of Computer Science and Engineering, Hebrew University of Jerusalem,
Jerusalem 91904, Israel

{avital.shafran,segev}@cs.huji.ac.il

² Kempner Institute for the Study of Natural and Artificial Intelligence,
Harvard University, MA, USA
emalach@fas.harvard.edu

³ Department of Computer Science, Cornell Tech, New York, USA
ristenpart@cornell.edu

⁴ Paul G. Allen School of Computer Science & Engineering,
University of Washington, Seattle, WA, USA
tessaro@cs.washington.edu

Abstract. Given the recent progress in machine learning (ML), the cryptography community has started exploring the applicability of ML methods to the design of new cryptanalytic approaches. While current empirical results show promise, the extent to which such methods may outperform classical cryptanalytic approaches is still somewhat unclear.

In this work, we initiate exploration of the theory of ML-based cryptanalytic techniques, in particular providing new results towards understanding whether they are fundamentally limited compared to traditional approaches. Whereas most classic cryptanalysis crucially relies on directly processing individual samples (e.g., plaintext-ciphertext pairs), modern ML methods thus far only interact with samples via gradient-based computations that average a loss function over all samples. It is, therefore, conceivable that such gradient-based methods are inherently weaker than classical approaches.

We introduce a unifying framework for capturing both “sample-based” adversaries that are provided with direct access to individual samples and “gradient-based” ones that are restricted to issuing gradient-based queries that are averaged over all given samples via a loss function.

This research was partially supported by the Israel Science Foundation (Grant No. 1336/22), the European Union (ERC, FTRC, 101043243), NSF grants CNS-2055169, CNS-2120651, CNS-2026774 and CNS-2154174, a JP Morgan Faculty Award, a CISCO Faculty Award, and a gift from Microsoft. Views and opinions expressed are however those of the author(s) only and do not necessarily reflect those of the European Union or the European Research Council. Neither the European Union nor the granting authority can be held responsible for them. This work was performed while A. Shafran was visiting Cornell Tech.

© International Association for Cryptologic Research 2024

L. Reyzin and D. Stebila (Eds.): CRYPTO 2024, LNCS 14925, pp. 37–71, 2024.

https://doi.org/10.1007/978-3-031-68391-6_2

Within our framework, we establish a general feasibility result showing that any sample-based adversary can be simulated by a seemingly-weaker gradient-based one. Moreover, the simulation exhibits a nearly optimal overhead in terms of the gradient-based simulator’s running time. Finally, we extend and refine our simulation technique to construct a gradient-based simulator that is fully parallelizable (crucial for avoiding an undesirable overhead for parallelizable cryptanalytic tasks), which is then used to construct a gradient-based simulator that executes the particular and highly useful gradient-descent method.

Taken together, although the extent to which ML methods may outperform classical cryptanalytic approaches is still somewhat unclear, our results indicate that such gradient-based methods are not inherently limited by their seemingly restricted access to the provided samples.

1 Introduction

The interplay between cryptography and learning theory has been extensively studied over the years, consistently leading to fundamental insights, both theoretical and practical. Traditionally, the study of this interplay has mostly focused on establishing a fruitful relationship between the hardness of cryptographic problems and that of learning problems (see, for example, [Val84, KV89, Riv91, Kha93, Reg05, KS09, ABW10, SZB21] and the references therein). Quite recently, however, following the tremendous progress in the area of machine learning, the cryptographic community has gradually started revealing an additional exciting avenue: Exploring the applicability of modern machine learning methods (in particular, those based on training deep neural networks, or DNNs) to the design of new cryptanalytic approaches that may hopefully improve upon the classic, somewhat “manual”, cryptanalytic approaches.

ML-Based Cryptanalysis. ML-based cryptanalytic methods have so far been developed mostly for attacking block ciphers, starting with the work of Gohr [Goh19] focusing on differential cryptanalysis of SPECK [BSS+13], and for attacking the LWE problem, starting with the work of Wenger, Chen, Charton, and Lauter [WCC+22]. In both cases, essentially, the same high-level ML-based methodology was followed: First, a large sample set was used to train a DNN using a certain network architecture and learning algorithm. Then, the resulting DNN was challenged with respect to an appropriate key recovery task or an indistinguishability task. For example, the sample set used by Gohr consisted of samples $((C_1, C_2), y)$ corresponding to encryptions under a secret key of either two uniformly random plaintexts ($y = 0$) or uniformly chosen plaintexts that have a fixed difference $\Delta = P \oplus P'$ ($y = 1$).

Similarly, the sample set used by Wenger et al. consisted of a large number of independently generated LWE samples $(x, y = x \cdot s + e \bmod q)$ for a secret s and added noise e .

The work of Gohr and that of Wenger et al. motivated a large and growing amount of follow-up work, as we further discuss in Sect. 1.2. This includes

additional applications of ML-based methods both in the context of block ciphers [JKM20, BGP+21, GLN22, BB22] and in the context of the LWE problem [LSW+23a, LSW+23b, SWL+24]. The extent to which these ML-based techniques outperform classic ones has been the subject of debate [DPS23], and further empiricism is warranted. We suggest complementing empiricism with theoretical treatments. To date, however, no formal frameworks have been suggested.

Sample-Based vs. Gradient-Based Methods. In this work, we initiate the exploration of the theoretical foundations of ML-based cryptanalysis, focusing on the question of whether learning-based approaches are *fundamentally* limited compared to others.

Motivating our exploration is the observation that modern ML-based methods fall into a specific template (as discussed above) that only access samples via gradient-based computations, which are averaged over all samples via a “loss function”. In comparison, traditional approaches use full access to directly process samples. This crucial difference motivates the following question:

Are ML-based cryptanalytic methods inherently limited due to their seemingly restricted access to cryptographic samples?

1.1 Our Contribution

We initiate a foundational exploration of the extent to which machine learning methods may lead to significant cryptanalytic advances. First, we introduce a unifying framework that models both “sample-based” adversaries and “gradient-based” ones, and put forward strong simulation notions capturing a concrete and quantifiable extent to which a sample-based adversary may be simulated by a gradient-based one. Then, within our framework, we provide a somewhat surprising negative answer to the above fundamental question via a sequence of increasingly refined general feasibility results, showing that any sample-based adversary can be simulated by a seemingly weaker gradient-based one with no significant loss in efficiency.

Taken together, our results indicate that the extent to which machine learning methods may lead to significant cryptanalytic advances is not inherently limited by their seemingly restricted access to the provided samples. In what follows, we provide a high-level overview of our main contributions.

Modeling Gradient-Based Methods. Our notions of sample-based adversaries and gradient-based ones model algorithms that are provided with access to a set S of samples $(x, y) \in \mathcal{X} \times \mathcal{Y}$, for some finite sets $\mathcal{X}$ and $\mathcal{Y}$. However, while sample-based adversaries receive the sample set S as an explicit input, gradient-based ones receive only the size $|S|$ of the sample set as an explicit input and are then restricted to accessing the sample set itself by querying a corresponding “gradient oracle” $\mathcal{O}_G$.

Our gradient oracle, which models a gradient-based training process, is provided with the sample set S , and receives queries of the form $(\ell, h, \vec{\theta})$, where:

- $\ell : \mathbb{R} \times \mathcal{Y} \rightarrow \mathbb{R}^+$ is a differentiable *loss function*, and
- $h : \mathbb{R}^p \times \mathcal{X} \rightarrow \mathbb{R}$ is a differentiable *model function* equipped with its *model parameters* $\vec{\theta} \in \mathbb{R}^p$.

At a high level, the model function h represents the gradient-based adversary’s current estimate of a mapping between $\mathcal{X}$ and $\mathcal{Y}$, and the loss function ℓ is used by the oracle to determine the extent to which the estimated mapping is accurate. Specifically, given a sample set S and a query $(\ell, h, \vec{\theta})$, the gradient oracle $\mathcal{O}_G$ averages, over all samples $(x, y) \in S$, the gradient $\nabla_{\theta} \ell(h(\vec{\theta}, x), y)$ of the loss function when composed with the model function $h(\vec{\theta}, \cdot)$. Intuitively, this provides a measure of the way the model parameters $\vec{\theta}$ may be modified in order to minimize the value of the loss function. Note that the gradient oracle still enables general access in terms of adversarial choice of ℓ and h that can change adaptively with each query. Common ML algorithms, including cryptanalytic ones, interact with gradient oracles in a much more limited fashion, most often via gradient descent (GD). Here, algorithms start with some initial set of parameters $\vec{\theta}_0$ and then, over a fixed number of iterations, refine them by (1) querying $(\ell, h, \vec{\theta}_{t-1})$ to obtain a gradient vector $\vec{g}$, and (2) applying a fixed update rule $\vec{\theta}_t = \vec{\theta}_{t-1} - \eta \cdot \vec{g}$. Here, η is a parameter called the learning rate.

As we concretely demonstrate in the full version of this work, this modeling of a gradient-based training process is indeed sufficiently general and expressive for capturing the ML-based cryptanalytic methods that have so far been developed. At the same time, such gradient-based and, even more so, GD-based algorithms clearly utilize very limited access to samples. Intuitively, it would seem that this could limit the power of such adversaries, at least compared to ones that have unfettered sample access.

Notions of Simulation. To explore this intuition, we put forward strong simulation notions capturing a concrete and quantifiable extent to which a sample-based adversary may be simulated by a gradient-based one. Our simulation notions are strong in the sense that they enable replacing any sample-based adversary in a cryptographic experiment with a gradient-based adversary that simulates it. At a high level, we quantify the “ ϵ -closeness” of such a simulation in an information-theoretic manner, and say that a gradient-based adversary $\mathcal{B}$ simulates a sample-based one $\mathcal{A}$ if, for any distribution $\mathcal{D}$ that produces samples sets, it holds that

$$\text{SD}_{S \leftarrow \mathcal{D}} \left((S, \mathcal{A}(S)), \left(S, \mathcal{B}^{\mathcal{O}_G(S, \cdot, \cdot, \cdot)} \left(1^{|S|} \right) \right) \right) \leq \epsilon,$$

where SD denotes the statistical distance between the above two distributions. This guarantees that whenever we replace a sample-based adversary in a cryptographic experiment with a gradient-based one that simulates it within distance ϵ , the results of the two experiments differ by statistical distance at most ϵ .

Our results in this work, in fact, provide an even stronger guarantee that we formalize by asking for a single “universal” gradient-based adversary that can be used to simulate *any* sample-based algorithm. For this purpose, such a universal

black-box gradient-based simulator is provided with oracle access both to the gradient-based oracle and to the sample-based adversary it simulates.

Goal: Gradient-Based Simulation—with Low Overhead. Equipped with our framework, we then start exploring the feasibility of constructing gradient-based simulators for sample-based adversaries. Concretely, our first goal is to understand whether any sample-based adversary may be simulated by a gradient-based one. At this point, we crucially observe that such a feasibility result on its own may be insufficient for cryptographic applications: Although simulating within statistical distance ϵ guarantees that the *success probability* of the resulting gradient-based adversary would be essentially identical to that of the sample-based adversary that it simulates, it does not capture the *efficiency overhead* of the simulation.

Thus, given that the security of cryptographic primitives and schemes is measured via the trade-off between the success probability of attacking them and the efficiency of the attack, we must therefore additionally consider the efficiency overhead of the simulator. We capture this overhead by considering the following two key measures of efficiency for any gradient-based simulator: The internal running time of the simulator and the number of queries that it issues to the gradient oracle.

Black-Box Perfect Simulation via DFS-Based Extraction. As our first construction, we present a gradient-based adversary that black-box simulates any sample-based one. Moreover, the simulation is *perfect* in terms of producing a completely identical output distribution, and is highly efficient in terms of the simulator’s internal running time and query complexity (both are essentially linear in the number of samples).

The main observation underlying this construction is that gradient queries are, somewhat counterintuitively, sufficiently expressive for efficiently extracting the entire set of samples. Specifically, we observe that, for any sample set $S = \{(x_i, y_i)\}_{i \in [s]}$ and for any given prefix z , a single gradient query enables to distinguish between the case in which there are no samples x_i that are prefixed with z , the case in which there is exactly one sample x_i that is prefixed with z , and the case in which there is more than one sample x_i that is prefixed with z . We then rely on this observation to realize a recursive depth-first search (DFS) based exploration of the given sample set by interacting with the gradient oracle. Ultimately, our simulator is rather efficient, using $O(|S| \cdot \log |\mathcal{X}|)$ gradient queries and has time overhead $O(|S| \cdot \log^2 |\mathcal{X}|)$ (see Theorem 4.1).

Our first construction provides a surprising negative answer to our main research question. But it is somewhat unsatisfying because, while on the one hand, it constructively shows that gradient-based adversaries are as strong as any sample-based one, it does so using gradient queries that are fundamentally adaptive and, in particular, are not naturally produced by any GD-based adversary, i.e., an adversary that only runs the GD algorithm on the samples. Therefore, it may still be the case that although gradient-based methods are *generally* not inherently limited, the commonly-used methods (most notably, GD and its various variants) are inherently limited.

Gradient Descent via Parallelizable Gradient Queries. We therefore go on to extend our approach to build a simulator that is GD-based. Such a simulator would have to somehow succeed at simulating a sample-based adversary despite being severely restricted in its interaction with the gradient oracle: it is not allowed to process the intermediate responses of the gradient oracle in any way other than determining the next query based on the fixed GD update rule. Crucially, the simulation can only choose the model function h and the loss function ℓ , along with the initial parameters θ_0 , the learning rate η , and the iteration number T . The output of the simulation then can only rely on the final parameters θ_T produced by the GD algorithm. The techniques underlying our DFS-based simulator above are now inapplicable as they rely on adaptively discarding “non-useful” paths in order to be efficient. We need a different approach.

As an intermediate step on our way to construct a GD-based simulator, we first devise a *non-adaptive* gradient-based simulation technique. We present a black-box gradient-based simulator whose gradient queries are issued in a non-adaptive manner and are thus fully parallelizable (i.e., all of its gradient queries can be issued within a single round of parallel queries). This simulator crucially relies on randomization and thus does not provide a perfect simulation as our first simulation. Instead, it provides simulation within any statistical distance ϵ , while its efficiency scales with just $\log(1/\epsilon)$ (therefore, from a cryptographic perspective, this allows choosing ϵ to be a negligible function of the security parameter). In fact, as we discuss in Sect. 5, this construction does not only present an intermediate simulation technique, but enables us to avoid an undesirable sequential overhead in the simulation of parallelizable cryptanalytic tasks.

Finally, we build off our parallelizable simulator to obtain a black-box GD-based simulator. The efficiency of this simulator crucially relies on our non-adaptive simulation, as each of its gradient oracle queries essentially encodes a number of non-adaptive queries. While encoding multiple queries, we introduce a mechanism to separate between them, such that each query issued by the GD-based simulator only replicates a single query issued by the non-adaptive simulator. When concluding all iterations, the simulator can extract the samples from the final set of parameters.

Summary and Open Problems. The framework we have introduced enables to formally reason on the applicability of gradient-based methods to the design of new cryptanalytic approaches. Focusing on the restricted manner in which such methods process their provided samples, our results indicate that gradient-based methods are not inherently limited when compared to classical cryptanalytic approaches. Moreover, this holds even when further restricting gradient-based methods to the most common gradient-descent learning methodology.

Our exploration gives rise to a variety of interesting problems, with the aim of both improving our concrete techniques and constructions, and extending the reach of this line of research to consider various restrictions on ML-based methods. In what follows, we briefly exemplify some of these potential future directions:

- Although our techniques already lead to highly efficient constructions in terms of their internal time and query complexities, a natural goal is to establish *lower bounds* on the required overhead for black-box simulating any sample-based adversary. Moreover, given that the running time of our GD-based simulator is somewhat higher than that of our other simulators, it would be interesting to identify any inherent additional overhead imposed by restricting the simulation to use gradient descent.
- Our parallelizable simulator crucially relies on randomization, as discussed above, and does not provide a perfect simulation as our DFS-based deterministic one. A foundational research direction is to explore the role of randomness in gradient-based simulation of sample-based adversaries, and whether it is indeed essential for parallelism in this setting.
- Each of our gradient-based simulators utilizes carefully crafted queries, which are based on non-trivial choices of the loss function, model function, and parameters. This setting is somewhat theoretical and might not fully reflect some common practices in ML. An interesting direction for further research would be to explore whether similar results may be obtained even when using a more restricted class of queries (e.g., fixing the loss function to be a commonly used loss such as the square loss).
- A weakened variant of our gradient oracle may be obtained by adding a certain amount of independent noise to each of its responses. This realistically models, for example, training with limited-precision samples, and training in the presence of random labeling errors. For this and other practically motivated variants of our gradient oracle, it would be fascinating to identify the extent to which they enable simulating sample-based adversaries.

1.2 Related Work

ML-Based Cryptanalysis. As discussed above, ML-based cryptanalytic methods have so far been developed in the contexts of block ciphers and the LWE problem. In the context of block ciphers, Gohr [Goh19] focused on round-reduced versions of Speck32/64 [BSS+13], and trained a ResNet-based [HZR+16] classification neural network to distinguish between ciphertext pairs that originate from plaintexts with a predefined difference, and ciphertext pairs that originate from independent plaintexts. Gohr’s neural distinguisher outperforms classical methods for 5 to 8 rounds of Speck32/64. Additionally, a key-recovery attack for 11 rounds of Speck32/64 was proposed, based on the trained model, with an estimated time complexity of 2^{38} , improving upon the best-known attack of Dinur [Din14] with time complexity of 2^{46} , at the cost of higher data complexity of $2^{14.5}$ ciphertext pairs.

Benamira, Gerault, Peyrin, and Tan [BGP+21] revisited Gohr’s work with the goal of obtaining a better understanding of its success, suggesting that his distinguisher was able to effectively approximate the cipher’s differential distribution tables (DDT). Baksi [BB22] and Jain, Kohli, and Mishra [JKM20] proposed a somewhat different approach, and trained neural networks to classify

input differences from a set of predefined differences, instead of distinguishing between pairs with a predefined difference and independent pairs. A significant amount of follow-up work subsequently applied similar ML-based methods to other block ciphers (see, for example, [HRC21a, SZM21, YK21, BB22, ZW22, BB22, CSY+23, ZW22, CSY+23, BB22, CSY+23, ZW22, JKM20, LLS+22, ZWw22, HRC21b, BGL+21, LLS+22, GLN22, BLY+23]).

In the context of the LWE problem, Wenger, Chen, Charton, and Lauter [WCC+22] trained a neural network on LWE samples with a shared secret s , which they then used for a secret-recovery attack. This enables them to extract secrets for very sparse (only 3 or 4 nonzero bits) and low-dimension (up to 128) instances of LWE. Li et al. [LSW+23b] proposed an improvement based on a preprocessing step, in which the LWE samples are processed by BKZ [CN11], a lattice reduction algorithm, for obtaining LWE samples with smaller coordinate variance. For larger moduli, this enabled them to extract secrets for higher dimensions (up to 350) and larger Hamming weights (up to 60). A follow-up work by Li et al. [LSW+23a] was then able to reduce the modulus size as well as to additionally recover ternary and small Gaussian secrets. Stevens et al. [SWL+24] proposed additional improvements that enabled to improve efficiency and reduce sample complexity. The extent to which these methods can outperform classical methods is still unclear, as also noted by Ducas et al. [DPS23].

In the full version of this work we show how the previous ML-based cryptanalytic methods are indeed captured within our framework.

PAC vs. Differential-Based Learning

In the context of machine learning theory, some works focused on studying the power and limitations of learning with GD and stochastic gradient-descent (SGD), a popular variant of GD. Abbe and Sandon [AS20] showed that a neural network trained with SGD with a batch size of 1 (i.e., each gradient update is based on the gradient of a single example) can implement any PAC learning algorithm [Val84]. Namely, any algorithm for learning with examples can be used to generate a neural network, such that running batch-1 SGD over this network simulates the learning algorithm. This result was extended by Abbe et al. [AKM+21], showing that SGD with a larger batch size can also simulate algorithms for learning from examples, as long as the gradients are not noisy, where the noise is modeled by the level of arithmetic precision. These results show that with low noise, i.e., high arithmetic precision, SGD is equivalent to PAC learning from examples, while SGD with high noise, i.e., low precision, is equivalent to learning from Statistical Queries [Kea98], a learning framework known to be weaker than PAC.

These works focused on comparing different learning frameworks (e.g., learning with SGD versus PAC and Statistical Query learning), but did not study the power of gradient-based algorithms in the context of cryptography. From the technical perspective, the work of Abbe et al. focuses on (a variant of) the SQ learning model as an intermediate learning model in between PAC learning and differential-based learning. In the cryptographic context, however, the SQ learning model does not seem to directly capture realistic applications (either

classic ones or recent ML-based ones as those discussed above). Therefore, we focus on directly exploring the interplay between samples-based methods and gradient-based ones.

Specifically, in the SQ setting a learner interacts with the SQ oracle by issuing statistical queries about the sample set, while in the gradient-based setting, the learner interacts with the gradient oracle by issuing queries that consist of an estimated mapping (as a parameterized function) and a loss function (used to estimate the quality of the estimate). This difference in interaction with the oracle requires to structure the oracle queries differently in order to extract the same information, namely the sample set. For this, Abbe et al. constructed “counting” queries, that ask for the number of samples that match some rule, e.g., prefix. In our gradient-based setting, we construct the queries by designing a dedicated loss function and parameterized model function such that the gradient with respect to the samples will directly encode samples that match some rule, e.g., prefix. Most crucially, our direct approach enables us to focus on quantitative measures that are more fundamental in the cryptographic setting, such as the running time and query complexity of attackers, while Abbe et al. focus most of their technical attention on identifying the level of numeric precision that separates differential-based learning from either PAC learning or SQ learning.

Memorization in Machine Learning. Neural networks are known to be able to “memorize” some of their training samples. Empirically it has been shown that large enough deep neural networks can achieve good accuracy on large, randomized training sets [ZBH+17]. A line of previous works explored this phenomenon using neural tangent kernels [JGH18] to show that large enough networks, trained with GD, can, in theory, memorize the label of each sample in the training set [ADH+19, DLL+19, Dan20]. Some works showed how individual training samples can be reconstructed from trained models, both in the vision and in the text domains [CLE+19, HVY+22, BHY+23]. Song et al. [SRS17] demonstrated that by maliciously modifying the training algorithm, an adversary can cause intended memorization and increase its ability to extract samples or information about the training set. Seen in this light, our simulators can be viewed as malicious training algorithms, and our results go beyond prior ones by showing how to extract the *entire* training set (rather than just a few points), and moreover to do so efficiently, as we show in this work. That said, our simulators do not attempt to simultaneously memorize training data while also achieving good performance for some baseline tasks.

1.3 Paper Organization

The remainder of this paper is organized as follows. First, in Sect. 2 we present several standard notions that are used throughout this work. In Sect. 3 we present our framework for modeling sample-based adversaries and gradient-based ones. In Sect. 4 we establish our general feasibility result by presenting a gradient-based adversary that black-box simulates any sample-based one. In Sects. 5 and 6 we then show that our approach extends to providing a gradient-based simulator

whose gradient queries are fully parallelizable, and a gradient-based simulator that follows the particular gradient-descent algorithm, respectively.

Due to space limitations, some of our contributions are provided in the full version of this work. This includes extending our results from the single-bit output case to the multi-bit one, and from the full-batch gradient oracle to the mini-batch case (as used in the popular SGD algorithm), as well as demonstrating the generality of our notion of gradient-based adversaries by exemplifying how previous ML-based cryptanalytic methods are indeed captured by it. The full proof of the gradient-descent simulation described in Sect. 6 is also provided in the full version of this work.

2 Preliminaries

For an integer $n \in \mathbb{N}$ we denote by $[n]$ the set $\{1, \dots, n\}$. For a distribution X we denote by $x \leftarrow X$ the process of sampling a value x from the distribution X . Similarly, for a set $\mathcal{X}$ we denote by $x \leftarrow \mathcal{X}$ the process of sampling a value x from the uniform distribution over $\mathcal{X}$. For a binary string $x \in \{0, 1\}^n$, we interchangeably refer to x both as a binary vector $x = x_1 \cdots x_n \in \{0, 1\}^n$ and as a real-valued one $\vec{x} = (x_1, \dots, x_n) \in \mathbb{R}^n$. For two random variables X and Y over a finite domain Ω we denote by $\text{SD}(X, Y)$ their statistical distance, which is defined as

$$\text{SD}(X, Y) = \frac{1}{2} \sum_{\omega \in \Omega} |\Pr[X = \omega] - \Pr[Y = \omega]|.$$

Let $\mathcal{M}_{n,k}$ be a family of functions mapping from n -bit inputs to k -bit outputs. The family $\mathcal{M}_{n,k}$ is said to be *pairwise-independent* if for every $x_1, x_2 \in \{0, 1\}^n$ such that $x_1 \neq x_2$, and for every $y_1, y_2 \in \{0, 1\}^k$, it holds that

$$\Pr_{M \leftarrow \mathcal{M}_{n,k}} [M(x_1) = y_1 \wedge M(x_2) = y_2] = \frac{1}{2^{2k}}.$$

For a differentiable function $f : \mathbb{R}^n \rightarrow \mathbb{R}$, the gradient $\nabla f : \mathbb{R}^n \rightarrow \mathbb{R}^n$ with respect to $\vec{x} = (x_1, \dots, x_n)$ is defined as the vector of its partial derivatives with respect to each component of $\vec{x}$, i.e.,

$$\nabla f(\vec{x}) = \left(\frac{\partial f(\vec{x})}{\partial x_1}, \dots, \frac{\partial f(\vec{x})}{\partial x_n} \right) \in \mathbb{R}^n.$$

For functions defined over multiple inputs, we use subscript to denote the input with respect to which the gradient is computed. For example, for a differentiable function $f : \mathbb{R}^{n_1} \times \mathbb{R}^{n_2} \rightarrow \mathbb{R}$ with two input variables $x \in \mathbb{R}^{n_1}$ and $z \in \mathbb{R}^{n_2}$, its gradient with respect to x is defined as

$$\nabla_x f(\vec{x}, \vec{z}) = \left(\frac{\partial f(\vec{x}, \vec{z})}{\partial x_1}, \dots, \frac{\partial f(\vec{x}, \vec{z})}{\partial x_n} \right) \in \mathbb{R}^{n_1}.$$

3 Our Framework

In this section, we present our framework for modeling both “sample-based” adversaries that are provided with direct access to individual samples and “gradient-based” ones that are restricted to issuing gradient-based queries that are averaged over all given samples via a loss function. We begin by formally defining our notion of sample-based adversaries and exemplifying the range of fundamental cryptographic applications that they enable us to capture in this context. Then, we formally define our notion of gradient-based adversaries, and provide strong simulation notions capturing a concrete and quantifiable extent to which a sample-based adversary may be simulated by a gradient-based one.

Sample-Based Adversaries. Our notion of sample-based adversaries captures the basic form of an algorithm that explicitly receives as input a sample set, denoted S , and produces an output. Throughout this work, when we refer to a “sample set”, we in fact refer to a *list* that contains pairs of the form $(x, y) \in \mathcal{X} \times \mathcal{Y}$, for some finite sets $\mathcal{X}$ and $\mathcal{Y}$ (the term *set* was chosen for consistency with the common ML terminology of learning from datasets).

Definition 3.1 (Sample-Based Adversary). *Let $\mathcal{X}$ and $\mathcal{Y}$ be finite sets, and let $\mathcal{D}$ be a distribution over $(\mathcal{X} \times \mathcal{Y})^*$. A sample-based adversary $\mathcal{A}$ is a potentially-randomized algorithm that, when given as input a sample set $S \leftarrow \mathcal{D}$, produces an output $\mathcal{A}(S) \in \{0, 1\}^*$.*

Despite its simplicity, our notion of a sample-based adversary captures a wide range of fundamental cryptographic applications, including both unpredictability and indistinguishability ones. Specifically, providing the above distribution $\mathcal{D}$ (which produces the sample sets in Definition 3.1) with access to a typically-keyed cryptographic primitive enables to capture various standard security notions corresponding to the primitive, as we now exemplify.

The most interesting applications in our setting, however, are those in which the produced sample sets correspond to ones that may be realistically used by ML-based cryptanalytic methods. Such sample sets are typically generated by *independently* sampling each pair (x, y) from a given distribution that originates from a *hard-to-compute mapping* between x and y (and, more generally, a *hard-to-compute relation*). For unpredictability applications, this captures, in particular, key recovery for various forms of keyed primitives (e.g., block ciphers, pseudorandom functions and encryption schemes) as well as unforgeability for MACs and signatures. In both cases, sample sets may consist, for example, of pairs $(x, F_{\text{sk}}(x))$, where F is the keyed primitive, and all of the x values are independently sampled from a given distribution (e.g., the uniform distribution, which would correspond to a random message attack). The goal of the adversary is either to extract the key sk (for key recovery) or to produce a new “valid” pair

(x^*, y^*) for a value x^* that was not included in sample set (for unforgeability of MACs and signatures).¹

As somewhat expected, the above applications do not cover adaptive attacks (e.g., ones in which an attacker would choose the x values one by one in an adaptive manner after observing each corresponding output $F_{\text{sk}}(x)$). Indeed, realistic sample sets for ML-based methods are of a rather static flavor, as discussed above, and we concretely exemplify this in the full version of this work using those utilized by Gohr [Goh19] and by Wenger, Chen, Charton, and Lauter [WCC+22]. Nevertheless, it is important to note that our framework and results capture arbitrary sample sets that may be generated by any distribution $\mathcal{D}$, and not only those that may be viewed as realistic for ML-based cryptanalysis.

Gradient-Based Adversaries. Our notion of a gradient-based adversary relies on a corresponding gradient oracle. This oracle, as formally defined in Fig. 1, receives as input a sample set S , two functions denoted ℓ and h , as well as a vector of parameters $\vec{\theta}$ for the function h . The first function, ℓ , typically referred to as the *loss function*, is a differentiable function $\ell : \mathbb{R} \times \mathcal{Y} \rightarrow \mathbb{R}^+$. The second function, h , typically referred to as the *model*, is a parameterized differentiable function $h : \mathbb{R}^p \times \mathcal{X} \rightarrow \mathbb{R}$, where $\vec{\theta} \in \mathbb{R}^p$ is its vector of parameters ($\vec{\theta}$ is typically referred to as the *model parameters*). Generally speaking, the function h represents the adversary’s current estimate of a mapping between $\mathcal{X}$ and $\mathcal{Y}$, and the loss function ℓ is used by the oracle to determine the extent to which the estimated mapping is accurate.²

The gradient oracle $\mathcal{O}_G(S, \ell, h, \vec{\theta})$:

1. Compute and output $\vec{g} \leftarrow \frac{1}{|S|} \sum_{(\vec{x}, y) \in S} \nabla_{\theta} \ell(h(\vec{\theta}, \vec{x}), y)$.

Fig. 1. The gradient oracle $\mathcal{O}_G$.

For a given query $(S, \ell, h, \vec{\theta})$, the oracle evaluates $\hat{y}_i = h(\vec{\theta}, \vec{x}_i)$ for each sample $(\vec{x}_i, y_i) \in S$, and compares it to y_i using ℓ , i.e., $\ell(\hat{y}_i, y_i)$. Then, the oracle

¹ Similarly, for indistinguishability applications, realistic ML sample sets naturally arise in various pseudorandomness experiments (such as those used to define the security of weak pseudorandom functions, or hard-core predicates for a one-way function). In these applications, the goal of the adversary is to distinguish between two distributions from which the sample set S is sampled: In one distribution it consists of pairs (x, y) where y is a pseudorandom value computed from x (e.g., the output of a weak pseudorandom function or a hard-core predicate applied to the inverse of x), and in the other distribution it consists of such pairs where y is uniformly distributed.

² While it may seem redundant to distinguish between h and $\vec{\theta}$ and to provide the oracle both as input (since $\vec{\theta}$ can contain a function description and h can be a universal function computing the function provided by $\vec{\theta}$), we nevertheless distinguish between the two since this would be useful for defining our next class of adversaries, in which both h and $\vec{\theta}$ admit a particular structure.

computes the gradient of the evaluation of the loss function with respect to the parameters $\vec{\theta}$. Denote the individual parameters as $\vec{\theta} = (\theta_1, \dots, \theta_p)$, then the gradient is the vector of partial derivatives with respect to each parameter:

$$\nabla_{\theta} \ell(h(\vec{\theta}, \vec{x}_i), y_i) = \left(\frac{\partial \ell(h(\vec{\theta}, \vec{x}_i), y_i)}{\partial \theta_1}, \dots, \frac{\partial \ell(h(\vec{\theta}, \vec{x}_i), y_i)}{\partial \theta_p} \right).$$

Then, the oracle outputs the average of this gradient over all samples in the set S . Computing the gradient of the loss function, composed with $h(\vec{\theta}, \cdot)$, with respect to the parameters $\vec{\theta}$ provides a measure for the way these parameters may be modified in order to minimize the value of the loss function. In other words, the oracle evaluates the performance of $h(\vec{\theta}, \cdot)$ over the sample set S using the loss function ℓ and then computes the gradient in order to suggest a way to improve this performance. Throughout this work, we do not explicitly require the oracle to verify that all queries consist of differentiable functions. This is since, in all of our results, all issued queries consist of differentiable functions (as explicitly established by our proofs).

Note that although it is not essential for our framework to explicitly force any computational constraints over the functions ℓ and h , we do however assume their computation and derivation can be computed in polynomial time when considering a sufficient level of precision (we would like to avoid, for example, a function h that performs an exhaustive search over an exponentially-large key space). Instead, we include the gradient computation and derivation time in the total running time of a gradient-based adversary $\mathcal{B}$, which includes the time required for expressing and forwarding a certain explicit representation (e.g., an arithmetic circuit) of ℓ , h , and the parameters $\vec{\theta}$ to the oracle $\mathcal{O}_{\mathcal{G}}$.

Definition 3.2 (Gradient-Based Adversary). *Let $\mathcal{X}$ and $\mathcal{Y}$ be finite sets, and let $\mathcal{D}$ be a distribution over $(\mathcal{X} \times \mathcal{Y})^*$. A gradient-based adversary $\mathcal{B}$ is a potentially-randomized algorithm that, when provided with access to the oracle $\mathcal{O}_{\mathcal{G}}(S, \cdot, \cdot, \cdot)$ for a set $S \leftarrow \mathcal{D}$, and given as input $1^{|S|}$, produces an output $\mathcal{B}^{\mathcal{O}_{\mathcal{G}}(S, \cdot, \cdot, \cdot)}(1^{|S|}) \in \{0, 1\}^*$.*

Definition 3.2 allows the adversary $\mathcal{B}$ to perform arbitrary computations and interact with the oracle $\mathcal{O}_{\mathcal{G}}$ in any way, under the only restriction that the queries must consist of differentiable functions. However, in the most common learning methodology, the learner (adversary in our setting) runs a particular iterative optimization algorithm in order to learn an estimate of some function or mapping. The Gradient Descent (GD) algorithm is the base for the most popularly used optimization algorithms. In this algorithm, given a fixed loss function ℓ and function h , the learner chooses the first oracle query using some heuristic, meaning it chooses initial parameters $\vec{\theta}_0$. Then, the following queries are dependent on a computation of a specific form over the oracle's responses. This computation, known as the *GD update rule*, receives the current query and the oracle's response to it, and defines the next one. More specifically, at each time step t , given the current query $(\ell, h, \vec{\theta}_t)$ and the oracle's response $\vec{g}_t$, the parameters of next query $(\ell, h, \vec{\theta}_{t+1})$ are defined as:

$$\vec{\theta}_{t+1} = \vec{\theta}_t - \eta \cdot \vec{g}_t ,$$

where $\eta \in \mathbb{R}$ is known as the step size or learning rate. In other words, the GD algorithm uses the gradients to direct the parameters of h towards a set of parameters that minimizes the loss functions.

Therefore, we additionally consider the following notion of a *GD-based adversary* as a refinement of a gradient-based one: A GD-based adversary $\mathcal{B}_{\text{GD}}[T, \ell, h, \eta, \vec{\theta}_0]$ is a gradient-based adversary that runs a GD algorithm for T iterations, and can then perform any additional computation over the output of the T -th iteration. We denote by `PostProcess` such additional computation, which can be viewed either as part of the internal description of the adversary $\mathcal{B}_{\text{GD}}[T, \ell, h, \eta, \vec{\theta}_0]$ or as a supplied parameter. See Fig. 2 for the formal description.

The GD-based adversary $\mathcal{B}_{\text{GD}}^{\mathcal{O}_{\text{G}}(S, \cdot, \cdot, \cdot)}$:

1. For $t \in \{1, \dots, T\}$:
 - 1.1 Query $\mathcal{O}_{\text{G}}(S, \cdot, \cdot, \cdot)$ with $(\ell, h, \vec{\theta}_{t-1})$ to obtain $\vec{g}_t$.
 - 1.2 Update $\vec{\theta}_t = \vec{\theta}_{t-1} - \eta \cdot \vec{g}_t$.
2. Output $v = \text{PostProcess}(\vec{\theta}_T)$.

Fig. 2. The GD-based adversary $\mathcal{B}_{\text{GD}}$.

In the full version of this work, we demonstrate that our notion of a GD-based adversary enables us to capture the ML-based cryptanalytic methods that have been studied so far. We exemplify this by focusing on the methods developed by Gohr [Goh19] and by Wenger, Chen, Charton, and Lauter [WCC+22] (as discussed in Sect. 1.2 and which have already led to additional similar methods), and showing that they indeed fit into our framework.

We emphasize that GD is not the only optimization algorithm used in learning settings, and that there are many others, such as SGD [RM51], ADAM [KB14], and many more [Rud16]. However, in this work, we focus our attention on the GD algorithm, as it is the base algorithm behind many of the popularly used algorithms. This is not a limitation of our work, and all of our results can be adapted to other algorithms as well. In particular, in the full version of this work, we extend our approach to the commonly used *stochastic mini-batch* setting (which also represents the key difference between the GD and SGD algorithms).

Our Notions of Simulation. We capture the extent to which a gradient-based adversary simulates a sample-based one by measuring the statistical distance between their input-output distributions. This enables to replace any sample-based adversary in a cryptographic experiment with a gradient-based adversary that simulates it, and ensures that the statistical distance between the experiments is bounded by the “closeness” of the simulation.

Definition 3.3 (ϵ -simulation). Let $\mathcal{X}$ and $\mathcal{Y}$ be finite sets, let $\epsilon > 0$, and let $\mathcal{A}$ and $\mathcal{B}$ be a sample-based adversary and a gradient-based adversary, respectively. Then, we say that $\mathcal{B}$ ϵ -simulates $\mathcal{A}$ with respect to $\mathcal{X} \times \mathcal{Y}$ if for any distribution $\mathcal{D}$ over $(\mathcal{X} \times \mathcal{Y})^*$, it holds that

$$\text{SD} \left((S, \mathcal{A}(S)), \left(S, \mathcal{B}^{\mathcal{O}_G(S, \cdot, \cdot, \cdot)} \left(1^{|S|} \right) \right) \right) \leq \epsilon,$$

where $S \leftarrow \mathcal{D}$ in both distributions.

Definition 3.3 presents a highly intuitive notion of simulation, providing a strong information-theoretic guarantee that enables to replace any sample-based adversary with a gradient-based one that simulates it. Our results in this work, in fact, provide an even stronger guarantee obtained by naturally extending Definition 3.3 to ask for a single “universal” gradient-based algorithm that ϵ -simulates for *any* sample-based algorithm. For this purpose, such a universal gradient-based algorithm $\mathcal{B}$ is provided with oracle access both to the gradient-based oracle $\mathcal{O}_G$ and to the sample-based adversary $\mathcal{A}$ that it simulates.

Definition 3.4 (Black-box ϵ -simulation). Let $\mathcal{X}$ and $\mathcal{Y}$ be finite sets and let $\epsilon > 0$. We say that a gradient-based algorithm $\mathcal{B}$ black-box ϵ -simulates all sample-based adversaries with respect to $\mathcal{X} \times \mathcal{Y}$ if for any distribution $\mathcal{D}$ over $(\mathcal{X} \times \mathcal{Y})^*$ and for any sample-based adversary $\mathcal{A}$ it holds that

$$\text{SD} \left((S, \mathcal{A}(S)), \left(S, \mathcal{B}^{\mathcal{O}_G(S, \cdot, \cdot, \cdot), \mathcal{A}(\cdot)} \left(1^{|S|} \right) \right) \right) \leq \epsilon,$$

where $S \leftarrow \mathcal{D}$ in both distributions.

Whenever Definitions 3.3 and 3.4 are satisfied by a gradient-based adversary $\mathcal{B}$ for $\epsilon = 0$ (as in the case with our simulator in Sect. 4), we refer to such an ϵ -simulation as *perfect simulation*.

Measuring the Simulation Overhead. Finally, for identifying the overhead incurred when simulating a sample-based adversary by a gradient-based one, throughout this work we focus on the following main measures of efficiency for such algorithms $\mathcal{A}$ and $\mathcal{B}$:

- We denote by $T_{\mathcal{A}} = T_{\mathcal{A}}(s, |\mathcal{X}|, |\mathcal{Y}|)$ and $T_{\mathcal{B}} = T_{\mathcal{B}}(s, |\mathcal{X}|, |\mathcal{Y}|, T_{\mathcal{A}})$ the running time of $\mathcal{A}$ and $\mathcal{B}$, respectively, while naturally allowing the running time of $\mathcal{B}$ to depend on that of $\mathcal{A}$. However, when considering a *black-box* simulator $\mathcal{B}$, we measure its running time as a function of only s , $|\mathcal{X}|$ and $|\mathcal{Y}|$, while separately accounting for the number of queries that it issues to $\mathcal{A}$. In order to account for the internal runtime of $\mathcal{B}$, we assume as a standard baseline that simple arithmetic operations (e.g., additions and multiplications) over a small constant number of elements from $\mathcal{X}$ and $\mathcal{Y}$ can be executed in unit cost.
- We denote by $Q_{\mathcal{B}} = Q_{\mathcal{B}}(s, |\mathcal{X}|, |\mathcal{Y}|, T_{\mathcal{A}})$ the number of gradient-oracle queries issued by $\mathcal{B}$ (where, again, we allow the number of queries issued by $\mathcal{B}$ to depend on the running time of $\mathcal{A}$ – although this will not be used by our constructions).

Simulation via Neural Networks. Neural networks (NN) are a particularly important class of (typically) non-linear differentiable functions and are the most popular family of model functions h used in the ML literature. As discussed above, in our framework the model functions h are assumed to be computable in polynomial time, i.e., their computation can be implemented using Turing machines (TM) running in polynomial time. As such, they can be represented by a bounded-size NN. In a nutshell, this holds as poly-sized TMs can be implemented using poly-sized Boolean circuits [Coo23, Lev73], which in turn can be represented by bounded-sized NNs [SSBD14]. As a result, although being the most common function family in the ML community, in particular in the ML-based cryptanalysis literature, we do not need to explicitly limit the class of functions h to neural networks as any h can be represented as one.

Distinct vs. Arbitrary Samples. Our gradient-based simulators, which we present throughout the following sections, assume that the sample sets S given as input to the gradient oracle $\mathcal{O}_G(S, \cdot, \cdot, \cdot)$ consist of “distinct” samples. Specifically, letting $S = \{(x_i, y_i)\}_{i \in [s]} \subset \mathcal{X} \times \mathcal{Y}$, our gradient-based simulators assume that $x_i \neq x_j$ for any $i \neq j \in [s]$ (i.e., that the x -values within any sample set are always distinct). This, however, may not be a valid assumption whenever the underlying distribution $\mathcal{D}$ may generate non-distinct samples (say, with some non-negligible probability).

Nevertheless, this assumption does not limit the scope of our framework. This is due to the fact that machine-learning methods naturally allow to preprocess sample sets, where in our case a naive single-pass serialization suffices. That is, any sample set $S = \{(x_i, y_i)\}_{i \in [s]}$ can be easily transformed into a serialized one $S' = \{(i||x_i, y_i)\}_{i \in [s]}$ by concatenating an index to each sample. This increases the bit-length of the x -values from $\log |\mathcal{X}|$ to $\log |\mathcal{X}| + \log s$, which for our results yields only a minor lower-level overhead with no asymptotic effect.³ Thus, in the remainder of this work, we assume that all sample sets which are given as input to the gradient oracle always consist of distinct samples (and we will not explicitly serialize and deserialize them).

4 Perfect Simulation via DFS-Based Extraction

Equipped with our framework for modeling sample-based and gradient-based adversaries, in this section we establish a general feasibility result by presenting a gradient-based adversary that black-box simulates any sample-based one. Moreover, the simulation is *perfect*, ensuring a completely identical output distribution (see Definition 3.4), and exhibits a *nearly optimal overhead* in terms of the gradient-based adversary’s running time (when compared to that of the simulated sample-based adversary).

³ Looking ahead, our gradient-based simulators would in fact extract the serialized sample set S' via their oracle queries, and then invoke the sample-based adversary $\mathcal{A}$ providing it with the original sample set S as input. This naturally requires deserializing each sample, which again yields only a minor lower-level overhead with no asymptotic effect.

For simplicity, here we state and prove our result for samples over $\mathcal{X} \times \mathcal{Y}$ where $\mathcal{X} = \{0, 1\}^n$ for some integer $n \in \mathbb{N}$ and $\mathcal{Y} = \{0, 1\}$. In the full version of this work we then extend our result to capture the more general setting in which $\mathcal{Y} = \{0, 1\}^m$ for some integer $m \in \mathbb{N}$. We prove the following theorem:

Theorem 4.1. *Let $\mathcal{X} = \{0, 1\}^n$ for some $n \geq 1$, and let $\mathcal{Y} = \{0, 1\}$. There exists a gradient-based adversary $\mathcal{B}$ that black-box perfectly-simulates all sample-based adversaries with respect to $\mathcal{X} \times \mathcal{Y}$, where:*

- $\mathcal{B}$ runs in time $T_{\mathcal{B}}(|S|, |\mathcal{X}|, |\mathcal{Y}|) = O(|S| \cdot \log^2 |\mathcal{X}|)$.
- $\mathcal{B}$ issues $Q_{\mathcal{B}}(|S|, |\mathcal{X}|, |\mathcal{Y}|) = O(|S| \cdot \log |\mathcal{X}|)$ queries to the gradient oracle, and a single query to the simulated sample-based adversary.

For the more general setting in which $\log |\mathcal{Y}| = m > 1$, for some integer $m \in \mathbb{N}$, our simulator $\mathcal{B}$ runs in time $T_{\mathcal{B}}(|S|, |\mathcal{X}|, |\mathcal{Y}|) = O(|S| \cdot \log |\mathcal{X}| \cdot (\log |\mathcal{X}| + \log |\mathcal{Y}|))$ and issues the same number of queries $Q_{\mathcal{B}}(|S|, |\mathcal{X}|, |\mathcal{Y}|) = O(|S| \cdot \log |\mathcal{X}|)$. We refer the reader to the full version of this work for the extended result. In what follows, we first discuss in more detail the overhead of our gradient-based simulator provided by Theorem 4.1. Next, we provide a high-level technical overview of our proof of Theorem 4.1, and then provide its formal proof.

The Overhead of our Gradient-Based Simulation. In terms of running time, as our gradient-based simulator issues only a single query to the simulated sample-based adversary, then the running time overhead is simply the simulator’s internal running time $\Omega(|S| \cdot \log^2 |\mathcal{X}|)$. Note that any sample-based adversary that merely examines the entire sample set S would run in time $\Omega(|S| \cdot \log |\mathcal{X}|)$, and therefore compared to all such sample-based adversaries our overhead is asymptotically optimal within a multiplicative factor of at most $O(\log |\mathcal{X}|)$.

Proof Overview. Inspired by the work of Abbe et al. [AKM+21] (as discussed in Sect. 1.2), our main observation underlying the proof of Theorem 4.1 is that gradient queries are sufficiently expressive for efficiently extracting the entire set of samples. That is, we show that there exists a gradient based-adversary $\mathcal{B}$ that, for any set $S \subseteq \mathcal{X} \times \mathcal{Y}$, can issue $O(|S| \cdot \log |\mathcal{X}|)$ queries to the gradient oracle $\mathcal{O}_{\mathcal{G}}(S, \cdot, \cdot, \cdot)$ and extract the set S . Then, a single query to the simulated sample-based adversary $\mathcal{A}$ provides the output $\mathcal{A}(S)$.

More specifically, our gradient-based simulator (which is formally described in Fig. 3 below), is based on the following main observation: For any sample set $S = \{(x_i, y_i)\}_{i \in [s]}$ and for any given prefix z , we can issue a query that enables to distinguish between the case in which there are no samples x_i that are prefixed with z , the case in which there is exactly one sample x_i that is prefixed with z , and the case in which there is more than one sample x_i that is prefixed with z .⁴ Thus, beginning with empty string $z = \varepsilon$, this enables our simulator to realize a recursive DFS-based exploration of the sample set S . Specifically, when exploring a path corresponding to a certain prefix z , if no samples are prefixed with z then

⁴ Recall, as discussed in Sect. 3, that throughout this work we assume that sample sets always consist of distinct samples, as otherwise straightforward single-pass serialization may be applied.

this path is terminated, if there is exactly one sample prefixed with z then the gradient oracle's response in fact enables to extract the sample, and if there is more than one such sample then the simulator continues recursively with the exploration paths corresponding to the prefixes $z0$ and $z1$. The gradient queries issued by $\mathcal{B}$ enable, in particular, to always correctly distinguish between these three cases, and this ensures that all samples are eventually extracted (and that no false samples are "extracted").

From a more technical perspective, the simulator $\mathcal{B}$ uses the same loss function ℓ across all queries, which is defined as $\ell(\hat{y}, y) = (2 + y) \cdot \hat{y}$. As for our choice of the constant 2, since each $y \in \{0, 1\}$ is a single bit, we have that $y < 2$. Looking ahead, our proof relies on this property to distinguish between the case where a single sample is prefixed by some z to the case where multiple samples are. In the full version of this work we extend our results to the multi-bit label setting, and show how to adjust the loss function accordingly. Letting $\ell'(h(\vec{\theta}, \vec{x}), y)$ denote the partial derivative of ℓ with respect to $h(\vec{\theta}, \vec{x})$, it holds that

$$\ell'(h(\vec{\theta}, \vec{x}), y) = 2 + y,$$

and therefore the response $\vec{g}$ to any query $(\ell, h, \vec{\theta})$ issued by $\mathcal{B}$ to the gradient oracle $\mathcal{O}_G(S, \cdot, \cdot, \cdot)$ is of the following form (recall Fig. 1 for the definition of the gradient oracle):

$$\begin{aligned} \vec{g} &= \frac{1}{|S|} \sum_{(\vec{x}, y) \in S} \ell'(h(\vec{\theta}, \vec{x}), y) \cdot \nabla_{\theta} h(\vec{\theta}, \vec{x}) \\ &= \frac{1}{|S|} \sum_{(\vec{x}, y) \in S} (2 + y) \cdot \nabla_{\theta} h(\vec{\theta}, \vec{x}). \end{aligned}$$

At each step of the exploration $\mathcal{B}$ can now tailor the function h , whose gradient $\nabla_{\theta} h(\vec{\theta}, \vec{x})$ depends on the current prefix z , so that it enables to count the number of samples prefixed with z (and, in case of a single match, to extract it). Specifically, $\mathcal{B}$ chooses h and $\vec{\theta}$ such that $\nabla_{\theta} h(\vec{\theta}, \vec{x})$ returns $(\vec{x}, 1)$ if $\vec{x}$ is prefixed by z (the additional entry 1 enables counting), and otherwise returns the all-zeros vector. That is, for a current prefix $z = z_1 \dots z_k \in \{0, 1\}^k$, and for some vector of parameters $\vec{w} \in \mathbb{R}^{n+1}$, we set $\vec{\theta} = (\vec{w}, z) = ((\vec{w})_1, \dots, (\vec{w})_{n+1}, z_1, \dots, z_k)$ and define z to be a non-differentiable parameter. Note that the exact value of $\vec{w}$ is not important as it is only used for its gradient behaviour, however for consistency we will define it as the all-zeros vector $(0, \dots, 0) \in \mathbb{R}^{n+1}$. We can obtain the desired gradient behavior by defining the function h as follows:

$$\begin{aligned} h(\vec{\theta}, \vec{x}) &= \vec{w} \cdot ((\vec{x})_1, \dots, (\vec{x})_n, 1) \cdot \mathbb{1}_{x_1, \dots, x_k = z} \\ &= (w_1(\vec{x})_1, \dots, w_n(\vec{x})_n, w_{n+1}) \cdot \mathbb{1}_{x_1, \dots, x_k = z}, \end{aligned}$$

and we get:

$$\nabla_{\theta} h(\vec{\theta}, \vec{x}) = \left(\frac{\partial h(\vec{\theta}, \vec{x})}{\partial w_1}, \dots, \frac{\partial h(\vec{\theta}, \vec{x})}{\partial w_n}, \frac{\partial h(\vec{\theta}, \vec{x})}{\partial w_{n+1}} \right)$$

$$\begin{aligned}
&= ((\vec{x})_1, \dots, (\vec{x})_n, 1) \cdot \mathbb{1}_{x_1, \dots, k=z} \\
&= \begin{cases} (0, \dots, 0) \in \mathbb{R}^{n+1} & \text{if } x_{1\dots k} \neq z \\ ((\vec{x})_1, \dots, (\vec{x})_n, 1) \in \mathbb{R}^{n+1} & \text{otherwise} \end{cases}
\end{aligned}$$

Note that although $\vec{\theta}$ contains both the differentiable parameters $\vec{w}$ and the non-differentiable parameters z , for clarity we denote by ∇_{θ} the gradient with respect to all differentiable parameters (i.e., ∇_{θ} represents ∇_w). Using the function h and parameters $\vec{\theta}$, as described above, enables us to realize a recursive DFS-based exploration for extracting any sample set S .

The black-box gradient-based simulator $\mathcal{B}^{\mathcal{O}_G(S, \cdot, \cdot, \cdot), \mathcal{A}(\cdot)}(1^s)$:

1. Set $\ell(\hat{y}, y) = (2 + y) \cdot \hat{y}$.
2. Set $z = \varepsilon$. // empty bit-string
3. Compute $D = \text{CHECK-MATCH}(z)$.
4. Obtain $v \leftarrow \mathcal{A}(D)$ and output v .

The recursive procedure CHECK-MATCH(z):

1. Define $k = |z|$ // bit-length of z
2. Define $\vec{w} = (0, \dots, 0) \in \mathbb{R}^{n+1}$ and set $\vec{\theta} = (\vec{w}, z)$.
3. Define $h(\vec{\theta}, \vec{x}) = \vec{w} \cdot ((\vec{x})_1, \dots, (\vec{x})_n, 1) \cdot \mathbb{1}_{x_1, \dots, k=z}$.
4. Obtain $\vec{g} = \mathcal{O}_G(S, \ell, h, \vec{\theta}) \in \mathbb{R}^{n+1}$
5. If $(\vec{g})_{n+1} = 0$: // no matches
Return $\emptyset$.
6. If $0 < |S| \cdot (\vec{g})_{n+1} < 4$: // single match
Define $\hat{x} = (1/(\vec{g})_{n+1}) \cdot (\vec{g})_{1\dots n}$ and $\hat{y} = |S| \cdot ((\vec{g})_{n+1} - 2)$
Return $(\hat{x}, \hat{y})$
7. Else:
Define $z_0 = z_0 z_1 \dots z_k 0$ and $z_1 = z_0 z_1 \dots z_k 1$ // concat 0 and 1 to z
Return $\text{CHECK-MATCH}(z_0) \cup \text{CHECK-MATCH}(z_1)$

Fig. 3. The black-box gradient-based simulator $\mathcal{B}$.

Proof of Theorem 4.1. Let $\mathcal{X} = \{0, 1\}^n$ for some $n \geq 1$ and $\mathcal{Y} = \{0, 1\}$. We show that the gradient-based simulator $\mathcal{B}$ described in Fig. 3 black-box perfectly simulates all sample-based adversaries with respect to $\mathcal{X} \times \mathcal{Y}$. In what follows we first prove the correctness of the simulation, and then analyze the running time and query complexity of $\mathcal{B}$.

Let $\mathcal{D}$ be a distribution over $(\mathcal{X} \times \mathcal{Y})^*$. For proving that

$$\text{SD} \left((S, \mathcal{A}(S)), \left(S, \mathcal{B}^{\mathcal{O}_G(S, \cdot, \cdot, \cdot), \mathcal{A}(\cdot)} \left(1^{|S|} \right) \right) \right) = 0,$$

where $S \leftarrow \mathcal{D}$ in both distributions, we in fact prove a stronger statement. We prove that for any $s \geq 1$ and for any set $S \subseteq (\mathcal{X} \times \mathcal{Y})^s$, the distribution of

the output produced by the computation $\mathcal{B}^{\mathcal{O}_G(S, \cdot, \cdot, \cdot), \mathcal{A}(\cdot)}(1^s)$ is identical to the distribution of the output produced by the computation $\mathcal{A}(S)$. For this purpose, it suffices to show that, for any such set S , the algorithm $\mathcal{B}$ always extracts the sample set S via its queries to the gradient oracle (i.e., that in Step 3 of $\mathcal{B}$'s description it holds that $D = S$).

At each invocation of the procedure CHECK-MATCH with some prefix z , the algorithm $\mathcal{B}$ queries the oracle $\mathcal{O}_G(S, \cdot, \cdot, \cdot)$ with $(\ell, h, \vec{\theta})$, and obtains a response $\vec{g}$ computed as follows:

$$\vec{g} = \frac{1}{|S|} \sum_{(\vec{x}, y) \in S} \nabla_{\theta} \ell(h(\vec{\theta}, \vec{x}), y) = \frac{1}{|S|} \sum_{(\vec{x}, y) \in S} \ell'(h(\vec{\theta}, \vec{x}), y) \cdot \nabla_{\theta} h(\vec{\theta}, \vec{x}).$$

Denoting by $S_z \subseteq S$ the set of samples $(\vec{x}, y) \in S$ for which $\vec{x}$ is prefixed by z , we obtain:

$$\begin{aligned} & \frac{1}{|S|} \sum_{(\vec{x}, y) \in S} (2 + y) \cdot \nabla_{\theta} h(\vec{\theta}, \vec{x}) \\ &= \frac{1}{|S|} \left(\sum_{(\vec{x}, y) \in S_z} (2 + y) \cdot \nabla_{\theta} h(\vec{\theta}, \vec{x}) + \sum_{(\vec{x}, y) \in S \setminus S_z} (2 + y) \cdot \nabla_{\theta} h(\vec{\theta}, \vec{x}) \right) \\ &= \frac{1}{|S|} \left(\sum_{(\vec{x}, y) \in S_z} (2 + y) \cdot ((\vec{x})_1, \dots, (\vec{x})_n, 1) \right. \\ & \quad \left. + \sum_{(\vec{x}, y) \in S \setminus S_z} (2 + y) \cdot (0, \dots, 0) \right) \\ &= \frac{1}{|S|} \sum_{(\vec{x}, y) \in S_z} (2 + y) \cdot ((\vec{x})_1, \dots, (\vec{x})_n, 1) \\ &= \frac{1}{|S|} \sum_{(\vec{x}, y) \in S_z} ((2 + y) \cdot (\vec{x})_1, \dots, (2 + y) \cdot (\vec{x})_n, (2 + y)), \end{aligned}$$

and thus for the response $\vec{g}$ provided by the oracle it holds that

$$\vec{g} = \frac{1}{|S|} \sum_{(\vec{x}, y) \in S_z} ((2 + y) \cdot (\vec{x})_1, \dots, (2 + y) \cdot (\vec{x})_n, (2 + y)).$$

We now distinguish between the following three cases depending on the size of the set S_z :

- **Case I:** $|S_z| = 0$. In this case $\vec{g} = (0, \dots, 0) \in \mathbb{R}^{n+1}$, and therefore the procedure CHECK-MATCH returns $\perp$ at Step 5.
- **Case II:** $|S_z| = 1$. In this case, for $S_z = \{(\vec{x}', y')\}$ we get:

$$\vec{g} = \frac{1}{|S|} \sum_{(\vec{x}, y) \in S_z} ((2 + y) \cdot (\vec{x})_1, \dots, (2 + y) \cdot (\vec{x})_n, (2 + y))$$

$$= \frac{1}{|S|} \cdot ((2 + y') \cdot (\vec{x}')_1, \dots, (2 + y') \cdot (\vec{x}')_n, (2 + y')) ,$$

thus the $(n + 1)$ -th component of the vector $\vec{g}$ will contain $\frac{1}{|S|} \cdot (2 + y')$. As $y' \in \{0, 1\}$ we get that $y' < 2$, and in particular $0 < 2 + y' < 2 + 2 = 4$. Therefore, the procedure CHECK-MATCH extracts and returns the pair $(\vec{x}', y')$ at Step 6.

- For the extraction of y' , we multiply the $(n + 1)$ -th component of $\vec{g}$ by $|S|$ and subtract 2, and we get

$$\hat{y} = |S| \cdot \left(\frac{1}{|S|} \cdot (2 + y') \right) - 2 = y' .$$

- For the extraction of $\vec{x}'$ we use the first n components of the vector $\vec{g}$. By multiplying element-wise by $\frac{|S|}{(2 + y')}$ we get

$$\hat{x} = \frac{|S|}{(2 + y')} \cdot \left(\frac{1}{|S|} \cdot (2 + y') \cdot \vec{x}' \right) = \vec{x}' .$$

- **Case III:** $|S_z| \geq 2$. Without loss of generality, we will analyze this case for $|S_z| = 2$. Let $S_z = \{(\vec{x}_1, y_1), (\vec{x}_2, y_2)\}$. We have:

$$\begin{aligned} \vec{g} &= \frac{1}{|S|} \sum_{(\vec{x}, y) \in S_z} ((2 + y) \cdot (\vec{x})_1, \dots, (2 + y)) \\ &= \frac{1}{|S|} (((2 + y_1) \cdot (\vec{x}_1)_1 + (2 + y_2) \cdot (\vec{x}_2)_1), \dots, (4 + y_1 + y_2)) \end{aligned}$$

We can see that the $(n + 1)$ -th component of $\vec{g}$ now contains $4 + y_1 + y_2 \geq 4$ and therefore the procedure CHECK-MATCH will continue the search over longer prefix values at Step 7.

Thus, we have shown that in each call to CHECK-MATCH we either extract and return one sample, return $\perp$, or return nothing and continue the search recursively. By starting from an empty string z , that matches all samples in S , we extend it bit-by-bit until we reach all prefixes of the sample set. This ensures that we extract the entire set S , thus obtaining $D = S$ and proving the correctness of our algorithm.

After establishing correctness, we now bound the number of queries issued by $\mathcal{B}$ as well as $\mathcal{B}$'s running time. As there are $|S|$ distinct samples, the search will reach at most $2|S|$ leaves of the resulting tree: $|S|$ leaves corresponding to the samples, and an additional $|S|$ leaves corresponding to prefixes that differ in their last bit from a sample. The maximal depth of the DFS tree is at most the bit-length n of the samples, and therefore each time we reach a leaf, we visit n nodes on the way. Although some of these nodes may be shared between different samples (meaning they correspond to prefixes that match more than one sample), there are at most $2|S| \cdot n$ visited nodes.

Each one of these node visits corresponds to one run of CHECK-MATCH, in which the simulator $\mathcal{B}$ issues one oracle query and performs some basic arithmetic computations over the output $\vec{g}$. For issuing the oracle query the simulator encodes a parameter vector $\vec{\theta} = (\vec{w}, z) \in \mathbb{R}^{(n+1)+k}$, where k ranges from 0 to n . Therefore, we bound the number of parameters by $O(n)$. The output $\vec{g}$ is of size $n + 1$ and thus in total we get that we perform $O(n)$ basic arithmetic computations in each run of CHECK-MATCH. We conclude that by running CHECK-MATCH for $O(|S| \cdot n)$ times the simulator issues at most $Q_{\mathcal{B}} = O(|S| \cdot n) = O(|S| \cdot \log |\mathcal{X}|)$ queries and runs in time at most $T_{\mathcal{B}} = O((|S| \cdot n) \cdot n) = O(|S| \cdot n^2) = O(|S| \cdot \log^2 |\mathcal{X}|)$ (assuming that basic arithmetic operations over n -bits are values are counted at unit cost). ■

5 Statistical Simulation with Fully-Parallelizable Gradient Queries

In Sect. 4 we established a general feasibility result by presenting a gradient-based adversary that black-box simulates any sample-based one. Although the presented simulator is highly efficient in terms of both its running time and query complexity, it issues its gradient queries in a sequential manner. Over the years, however, parallelization techniques have played an instrumental role in speeding up cryptanalytic tasks. Such techniques have provided parallel algorithms for a variety of cryptanalytic tasks, with notable examples ranging from classic ones such as discrete logarithm algorithms via parallel variants of Pollard’s rho method [Pol78, vOW99, BKN+10, BCC+10, BKK+12] and collision finding algorithms via parallel variants of Wagner’s generalized birthday attack [Wag02, Ber07, BLN+09] to a wide variety of parallel algorithms for lattice problems [Mic] (for additional examples see also [Nie12, Bos15] and the references therein). Thus, when simulating a sample-based adversary for a cryptanalytic task that may highly benefit from parallelization techniques, a sequential query pattern as exhibited by our gradient-based adversary presented in Sect. 4 may introduce an undesirable simulation bottleneck.

In this section we show that our sequential simulation approach can be modified to result in a simulator whose gradient queries are fully parallelizable (i.e., all of its gradient queries can be issued within a single round of parallel queries). This comes at a rather minor cost of statistical simulation instead of perfect simulation (i.e., simulation within statistical distance ϵ for any fixed $\epsilon > 0$ as captured by Definition 3.4) together with a slight increase in the running time and query complexity of the simulator for some ranges of the parameters (depending logarithmically on $1/\epsilon$ and thus enabling to efficiently support negligible values of ϵ). As we show below, for most natural choices of the parameters, the multiplicative overhead compared to our sequential simulator is in fact constant.

As in Sect. 4, for simplicity here we state and prove our result for samples over $\mathcal{X} \times \mathcal{Y}$ where $\mathcal{X} = \{0, 1\}^n$ for some integer $n \in \mathbb{N}$ and $\mathcal{Y} = \{0, 1\}$. We prove the following theorem:

Theorem 5.1. *Let $\mathcal{X} = \{0, 1\}^n$ for some $n \geq 1$, and let $\mathcal{Y} = \{0, 1\}$. For any $\epsilon > 0$ there exists a gradient-based adversary $\mathcal{B}$ that black-box ϵ -simulates all sample-based adversaries with respect to $\mathcal{X} \times \mathcal{Y}$, where:*

- $\mathcal{B}$ runs in time $T_{\mathcal{B}}(|S|, |\mathcal{X}|, |\mathcal{Y}|) = O(|S| \cdot \log(|S|/\epsilon) \cdot \log |\mathcal{X}|)$.
- $\mathcal{B}$ issues $Q_{\mathcal{B}}(|S|, |\mathcal{X}|, |\mathcal{Y}|) = O(|S| \cdot \log(|S|/\epsilon))$ parallel queries to the gradient oracle, followed by a single query to the simulated sample-based adversary.

Recall that our sequential simulator provided by Theorem 4.1 runs in time $O(|S| \cdot \log^2 |\mathcal{X}|)$ and issues $O(|S| \cdot \log |\mathcal{X}|)$ gradient queries. Thus, in terms of both the running time and query complexity, the performance of our parallel simulator provided by Theorem 5.1 matches that of our sequential simulator within a multiplicative gap of $\log(|S|/\epsilon)/\log |\mathcal{X}|$. Recall that $|S| \leq |\mathcal{X}| = \{0, 1\}^n$, and therefore even for an exponentially-small $\epsilon = \exp(-\Omega(n))$ this multiplicative gap is constant.

For the more general setting in which $|\mathcal{Y}| = m > 1$ for some integer $m \in \mathbb{N}$, our simulator $\mathcal{B}$ runs in time $T_{\mathcal{B}}(|S|, |\mathcal{X}|, |\mathcal{Y}|) = O(|S| \cdot \log(|S|/\epsilon) \cdot (\log |\mathcal{X}| + \log |\mathcal{Y}|))$ and issues the same number of queries $Q_{\mathcal{B}}(|S|, |\mathcal{X}|, |\mathcal{Y}|) = O(|S| \cdot \log(|S|/\epsilon))$. We refer the reader to the full version of this work for the extended result.

In what follows we first provide a high-level technical overview of our proof of Theorem 5.1 when compared to that of Theorem 4.1, and then provide its formal proof.

Proof Overview. Recall that our sequential simulation approach, as detailed in Sect. 4, relies on the observation that gradient queries are sufficiently expressive for efficiently extracting the entire set of samples. Specifically, we observed that, for any sample set $S = \{(x_i, y_i)\}_{i \in [s]}$ and for any given prefix z , a single gradient query enables to distinguish between the case in which there are no samples x_i that are prefixed with z , the case in which there is exactly one sample x_i that is prefixed with z , and the case in which there is more than one sample x_i that is prefixed with z . We then used this observation to realize a recursive and *completely deterministic* DFS-based exploration of the sample set S , leading to a highly sequential process for extracting all samples (e.g., whenever two samples share a prefix of length ℓ , then our simulator must issue at least ℓ sequential queries in order to extract them).

A natural approach for obtaining a less sequential process is to rely on randomization, and a naive attempt would be to guess a prefix z , and hope that it matches exactly one sample (this is obviously fully parallelizable). Unfortunately, on the one hand, for n -bit samples x_i , guessing a rather short prefix (say, of length $O(\log n)$ bits) may not enable to isolate any sample (e.g., consider a sample set S in which all samples x_i share the same $n/2$ -bit prefix). On the other hand, guessing a rather long prefix (say, of length $\Omega(n)$ bits) may lead to an exponentially small probability of even hitting the prefix of any sample. Thus, this naive attempt completely fails.

Our approach relies on a more subtle form of randomization, where instead of guessing a prefix, we guess an output of a hash function $M : \{0, 1\}^n \rightarrow$

$\{0, 1\}^k$ that is independently sampled for each gradient query from a pairwise-independent function family. Recall that the pairwise independence guarantee provided by such a function family is that for any $x_i \neq x_j$ it holds that the random variables $M(x_i)$ and $M(x_j)$ are independent and uniformly distributed over the choice of the function M from the given function family. This guarantees that, even for a short output length $k = O(\log n)$, each sample x_i would be eventually isolated from all other samples, except any some pre-specified failure probability ϵ that would determine the number of queries. Moreover, since the output length is short, we can guess the output $z = M(x_i)$ of the hash function with a sufficiently high probability and include it in the gradient query to enable the extraction of the sample.

From a more technical perspective, for each gradient query our simulator samples a string $z \in \{0, 1\}^k$ and a function $M : \{0, 1\}^n \rightarrow \{0, 1\}^k$ from a pairwise independent function family $\mathcal{M}_{n,k}$. For a vector of parameters $\vec{w} \in \mathbb{R}^{n+1}$, we define $\vec{\theta} = (\vec{w}, z, M)$, where we assume M can be encoded using $\ell(n, k) = O(n + k)$ bits, and define z and M to be non-differentiable parameters. We remind that although the exact value of $\vec{w}$ is not important, for consistency we will define it as the all-zeros vector $(0, \dots, 0) \in \mathbb{R}^{n+1}$. The function h is defined as follows:

$$\begin{aligned} h(\vec{\theta}, \vec{x}) &= \vec{w} \cdot ((\vec{x})_1, \dots, (\vec{x})_n, 1) \cdot \mathbb{1}_{M(x)=z} \\ &= (w_1(\vec{x})_1, \dots, w_n(\vec{x})_n, w_{n+1}) \cdot \mathbb{1}_{M(x)=z}. \end{aligned}$$

The simulator then issues an oracle query $(\ell, h, \vec{\theta})$, for which it holds that

$$\nabla_{\theta} h(\vec{\theta}, \vec{x}) = \begin{cases} (0, \dots, 0) \in \mathbb{R}^{n+1} & \text{if } M(x) \neq z \\ ((\vec{x})_1, \dots, (\vec{x})_n, 1) \in \mathbb{R}^{n+1} & \text{otherwise} \end{cases}$$

As we noted in Sect. 4, we denote by ∇_{θ} the gradient with respect to all differentiable parameters in $\vec{\theta}$, in this case this means $\vec{w}$ only. As these queries are now independent of each other, they can all be issued in parallel. Our analysis below shows that, for any $\epsilon > 0$, issuing $T = O(|S| \cdot \log(|S|/\epsilon))$ such queries guarantees that all samples are extracted except with probability ϵ .

Proof of Theorem 5.1. Let $\mathcal{X} = \{0, 1\}^n$ for some $n \geq 1$, $\mathcal{Y} = \{0, 1\}$, and let $\epsilon > 0$. We show that the gradient-based simulator $\mathcal{B}$ described in Fig. 4 black-box ϵ -simulates all sample-based adversaries with respect to $\mathcal{X} \times \mathcal{Y}$. In what follows we first prove the correctness of the simulation, and then analyze the running time and query complexity of $\mathcal{B}$.

Let $\mathcal{D}$ be a distribution over $(\mathcal{X} \times \mathcal{Y})^*$. For proving that

$$\text{SD} \left((S, \mathcal{A}(S)), \left(S, \mathcal{B}^{\mathcal{O}_G(S, \cdot, \cdot, \cdot), \mathcal{A}(\cdot)} \left(1^{|S|} \right) \right) \right) \leq \epsilon,$$

where $S \leftarrow \mathcal{D}$ in both distributions, we in fact prove a stronger statement. We prove that for any $s \geq 1$ and for any set $S \subseteq (\mathcal{X} \times \mathcal{Y})^s$, the output produced by the computation $\mathcal{B}^{\mathcal{O}_G(S, \cdot, \cdot, \cdot), \mathcal{A}(\cdot)}(1^s)$ is distributed as follows: With probability at most ϵ over the internal randomness of $\mathcal{B}$ it outputs $\perp$, and otherwise it is

The black-box gradient-based simulator $\mathcal{B}^{\mathcal{O}_G(S, \cdot, \cdot, \cdot), \mathcal{A}(\cdot)}(1^s)$:

1. Set $k = \lceil \log_2(2s) \rceil$ and $T = \lceil 16s \cdot \ln(s/\epsilon) \rceil$.
2. Let $\mathcal{M}_{n,k}$ be a pairwise-independent function family, where $M : \{0, 1\}^n \rightarrow \{0, 1\}^k$ for every $M \in \mathcal{M}_{n,k}$.
3. Set $\ell(\hat{y}, y) = (2 + y) \cdot \hat{y}$.
4. Set $D = \emptyset$.
5. For $t \in \{1, \dots, T\}$: // can be performed in parallel
 - 5.1. Sample $z \xleftarrow{\$} \{0, 1\}^k$ and $M \xleftarrow{\$} \mathcal{M}_{n,k}$.
 - 5.2. Define $\vec{w}(0, \dots, 0) \in \mathbb{R}^{n+1}$ and set $\vec{\theta} = (w, z, M)$.
 - 5.3. Define $h(\vec{\theta}, \vec{x}) = \vec{w} \cdot \psi(\vec{x})_1, \dots, (\vec{x})_n, 1) \cdot \mathbb{1}_{M(x)=z}$.
 - 5.4. Obtain $\vec{g} = \mathcal{O}_G(S, \ell, h, \vec{\theta}_t) \in \mathbb{R}^{n+1}$.
 - 5.5. If $0 < |S| \cdot (\vec{g})_{n+1} < 4$: // single match
 - i. Define $\hat{x} = (1/(\vec{g})_{n+1}) \cdot (\vec{g})_{1 \dots n}$ and $\hat{y} = |S| \cdot ((\vec{g})_{n+1} - 2)$.
 - ii. Update $D = D \cup \{(\hat{x}, \hat{y})\}$.
6. If $|D| < s$ then output $\perp$, and otherwise obtain $v \leftarrow \mathcal{A}(D)$ and output v .

Fig. 4. The black-box gradient-based simulator $\mathcal{B}$.

identical to the distribution of the output produced by the computation $\mathcal{A}(S)$. For this purpose, it suffices to show that, for any such set S , with probability at least $1 - \epsilon$ over the internal randomness of $\mathcal{B}$, the algorithm $\mathcal{B}$ is able to extract the sample set S via its queries to the gradient oracle.

For every $t \in [T]$ we denote by $z_t \in \{0, 1\}^k$ and $M_t \in \mathcal{M}_{n,k}$ the random variables corresponding to the bit-string and the function sampled at iteration t , respectively. For every $i \in [s]$ and for each iteration $t \in [T]$, we say that the sample $(\vec{x}_i, y_i)$ is *isolated* by the t -th iteration if $M_t(\vec{x}_i) = z_t$ and for every $j \in [s] \setminus \{i\}$ it holds that $M_t(x_i) \neq M_t(x_j)$. For every $i \in [s]$ we let Failure_i denote the event in which the sample $(\vec{x}_i, y_i)$ is not isolated by any iteration, and we let $\text{Failure} = \bigcup_{i \in [s]} \text{Failure}_i$.

We will now show that if the event Failure does not occur, then the gradient-based simulator extracts the entire set S (and otherwise outputs $\perp$). For any iteration $t \in [T]$, the response of the gradient oracle is as follows:

$$\begin{aligned}
 \vec{g} &= \frac{1}{|S|} \sum_{(\vec{x}, y) \in S} \nabla_{\theta} \ell(h(\vec{\theta}, \vec{x}), y) \\
 &= \frac{1}{|S|} \sum_{(\vec{x}, y) \in S} \ell'(h(\vec{\theta}, \vec{x}), y) \cdot \nabla_{\theta} h(\vec{\theta}, \vec{x}) \\
 &= \frac{1}{|S|} \sum_{(\vec{x}, y) \in S} (2 + y) \cdot \nabla_{\theta} h(\vec{\theta}, \vec{x}) .
 \end{aligned}$$

Denoting by $S_{M,z} \subseteq S$ the set of samples $(\vec{x}, y) \in S$ for which it holds that $M(x) = z$, we obtain

$$\begin{aligned}
& \frac{1}{|S|} \sum_{(\vec{x}, y) \in S} (2 + y) \cdot \nabla_{\theta} h(\vec{\theta}, \vec{x}) \\
&= \frac{1}{|S|} \left(\sum_{(\vec{x}, y) \in S_{M,z}} (2 + y) \cdot \nabla_{\theta} h(\vec{\theta}, \vec{x}) + \sum_{(\vec{x}, y) \in S \setminus S_{M,z}} (2 + y) \cdot \nabla_{\theta} h(\vec{\theta}, \vec{x}) \right) \\
&= \frac{1}{|S|} \left(\sum_{(\vec{x}, y) \in S_{M,z}} (2 + y) \cdot ((\vec{x})_1, \dots, (\vec{x})_n, 1) \right. \\
&\quad \left. + \sum_{(\vec{x}, y) \in S \setminus S_{M,z}} (2 + y) \cdot (0, \dots, 0) \right) \\
&= \frac{1}{|S|} \sum_{(\vec{x}, y) \in S_{M,z}} (2 + y) \cdot ((\vec{x})_1, \dots, (\vec{x})_n, 1) \\
&= \frac{1}{|S|} \sum_{(\vec{x}, y) \in S_{M,z}} ((2 + y) \cdot (\vec{x})_1, \dots, (2 + y) \cdot (\vec{x})_n, (2 + y)) ,
\end{aligned}$$

and thus for the response $\vec{g}$ provided by the oracle it holds that

$$\vec{g} = \frac{1}{|S|} \sum_{(\vec{x}, y) \in S_{M,z}} ((2 + y) \cdot (\vec{x})_1, \dots, (2 + y) \cdot (\vec{x})_n, (2 + y)) .$$

We now distinguish between the following three cases depending on the size of the set $S_{M,z}$:

- **Case I:** $|S_{M,z}| = 0$. In this case $\vec{g} = (0, \dots, 0) \in \mathbb{R}^{n+1}$ and therefore the condition in Step 5.5 is not satisfied, and the iteration will end with no samples being extracted.
- **Case II:** $|S_{M,z}| = 1$. In this case it holds that

$$\begin{aligned}
\vec{g} &= \frac{1}{|S|} \sum_{(\vec{x}, y) \in S_{M,z}} ((2 + y) \cdot (\vec{x})_1, \dots, (2 + y) \cdot (\vec{x})_n, (2 + y)) \\
&= \frac{1}{|S|} \cdot ((2 + y') \cdot (\vec{x}')_1, \dots, (2 + y') \cdot (\vec{x}')_n, (2 + y')) ,
\end{aligned}$$

for $S_{M,z} = \{(\vec{x}', y')\}$. The $(n + 1)$ -th component of the vector $\vec{g}$ contains $\frac{1}{|S|} \cdot (2 + y')$. As $y' \in \{0, 1\}$ we get that $y' < 2$, and in particular $0 < 2 + y' < 2 + 2 = 4$. Therefore, the condition in Step 5.5 is satisfied and we will extract the pair $(\vec{x}', y')$ as follows:

- For the extraction of y' , we multiply the $(n + 1)$ -th component of $\vec{g}$ by $|S|$ and subtract 2, and we get

$$\hat{y} = |S| \cdot \left(\frac{1}{|S|} \cdot (2 + y') \right) - 2 = y' .$$

- For the extraction of $\vec{x}'$ we use the first n components of the vector $\vec{g}$ which contains $\frac{1}{|S|} \cdot (2 + y') \cdot \vec{x}'$. By multiplying element-wise by $\frac{|S|}{(2+y')}$ we get

$$\hat{x} = \frac{|S|}{(2+y')} \cdot \left(\frac{1}{|S|} \cdot (2 + y') \cdot \vec{x}' \right) = \vec{x}' .$$

- **Case III:** $|S_{M,z}| \geq 2$. Without loss of generality, we analyze this case for $|S_{M,z}| = 2$. In this case,

$$\begin{aligned} \vec{g} &= \frac{1}{|S|} \sum_{(\vec{x}, y) \in M} ((2 + y)(\vec{x})_1, \dots, (2 + y)) \\ &= \frac{1}{|S|} (((2 + y_1) \cdot (\vec{x}_1)_1 + (2 + y_2) \cdot (\vec{x}_2)_1), \dots, (4 + y_1 + y_2)) , \end{aligned}$$

for $S_{M,z} = \{(\vec{x}_1, y_1), (\vec{x}_2, y_2)\}$. We can see that the $(n + 1)$ -th component of $\vec{g}$ now contains $4 + y_1 + y_2 \geq 4$ and therefore the condition in Step 5.5 is not satisfied and the iteration will end with no samples being extracted.

We now show that the event **Failure** occurs with probability at most ϵ , and this finalizes the correctness of the simulation. From the definition of the events $\{\text{Failure}_i\}_{i \in [s]}$, for every $i \in [s]$ it holds that

$$\begin{aligned} \Pr[\text{Failure}_i] &= (\Pr[(\vec{x}_i, y_i) \text{ is not isolated by the first iteration}])^T \\ &= \left(\Pr_{z_1, M_1} [\exists j \in [s] \setminus \{i\} : M_1(x_i) = M_1(x_j) \wedge z_1 \neq M_1(x_i)] \right)^T \\ &= \left(1 - \Pr_{z_1, M_1} [\nexists j \in [s] \setminus \{i\} : M_1(x_i) = M_1(x_j) \vee z_1 = M_1(x_i)] \right)^T , \end{aligned}$$

where the first equality follows from the independence across all iterations. Using the union bound, we obtain

$$\begin{aligned} \Pr_{M_1} [\nexists j \in [s] \setminus \{i\} : M_1(x_i) = M_1(x_j)] &= 1 - \Pr_{M_1} [\exists j \in [s] \setminus \{i\} : M_1(x_i) = M_1(x_j)] \\ &\geq 1 - \sum_{j \in [s] \setminus \{i\}} \Pr_{M_1} [M_1(x_i) = M_1(x_j)] \\ &= 1 - \sum_{j \in [s] \setminus \{i\}} \frac{1}{2^k} = 1 - (s - 1) \cdot \frac{1}{2^k} , \end{aligned}$$

where the last equality follows from the pairwise-independence of the function family $\mathcal{M}_{n,k}$. Since z_1 is uniformly sampled and is independent of M_1 , the probability of a match between z_1 and $M_1(x_i)$, given that there are no collisions can be computed as

$$\Pr_{z_1, M_1} [z_1 = M_1(x_i) \mid \nexists j \in [s] \setminus \{i\} : M_1(x_i) = M_1(x_j)] = \frac{1}{2^k} .$$

Therefore,

$$\Pr[\text{Failure}_i] \leq \left(1 - \left(1 - \frac{s-1}{2^k}\right) \cdot \frac{1}{2^k}\right)^T = \left(1 - \left(\frac{1}{2^k} - \frac{s-1}{2^{2k}}\right)\right)^T,$$

and our choice of $k = \lceil \log_2(2s) \rceil$ guarantees that

$$\frac{1}{2^k} - \frac{s-1}{2^{2k}} \geq \frac{s}{2^{2k}} \geq \frac{1}{16s},$$

implying

$$\Pr[\text{Failure}_i] \leq \left(1 - \frac{1}{16s}\right)^T \leq e^{-T/16s}.$$

We conclude by observing that our choice of $T = \lceil 16s \cdot \ln(s/\epsilon) \rceil$ now implies

$$\Pr[\text{Failure}] = \Pr\left[\bigcup_{i \in [s]} \text{Failure}_i\right] \leq s \cdot e^{-T/16s} \leq \epsilon.$$

After establishing correctness, we now bound the number of queries issued by $\mathcal{B}$ as well as $\mathcal{B}$'s running time. In each iteration the simulator issues a single oracle query and performs some basic arithmetic computations over the output $\vec{g} \in \mathbb{R}^{n+1}$. For issuing the oracle query, the simulator encodes the parameter vector $\vec{\theta} = (\vec{w}, z, M) \in \mathbb{R}^{(n+1)+k+\ell(n,k)}$ with $\ell(n, k) = O(n + k)$ being the description length of the pairwise functions. As we defined $k = O(\log |S|)$, we can bound the number of parameters, and correspondingly the runtime of a single iteration, by $O(\log |\mathcal{X}| + \log |S|) = O(\log |\mathcal{X}|)$. By running for $T = O(|S| \cdot \log(|S|/\epsilon))$ iterations we get a total runtime of $T_{\mathcal{B}} = O(|S| \cdot \log(|S|/\epsilon) \cdot \log |\mathcal{X}|)$ and query complexity of $Q_{\mathcal{B}} = O(|S| \cdot \log(|S|/\epsilon))$. $\blacksquare$

6 Statistical Simulation via Gradient Descent

In this section we show that all sample-based adversaries can be simulated not only using general gradient-based methods (as established by Theorems 4.1 and 5.1), but in fact using the particular gradient-descent (GD) algorithm. Recall that whereas our notion of a gradient-based adversary allows querying the gradient oracle in a rather arbitrary manner, our notion of a GD-based adversary forces the application of a specific update rule for each query based on the response provided to the previous one. As discussed in Sect. 3, this restriction models the GD algorithm, which is considered the most common learning methodology.

Specifically, as detailed in Fig. 2, a GD-based algorithm starts by initializing the parameters θ according to some predefined distribution. Then, in each iteration, the algorithm updates these parameters based on the current gradient information $\vec{g}$ (which, in our case, is received from the gradient oracle) for

“directing” the update towards parameters that minimize the loss function (and thus correspond to a better estimate). In other words, at every iteration $t \in [T]$, given the current parameters $\bar{\theta}^{(t-1)}$ and the gradient vector $\vec{g}$, the algorithm updates the parameters by computing

$$\bar{\theta}^{(t)} = \bar{\theta}^{(t-1)} - \eta \cdot \vec{g},$$

where η is the algorithm’s “step size” (also known as the “learning rate”) which determines the size of the steps taken in the direction of the gradient.

Our result in this section shows that, even in this significantly more restricted setting, we can construct a GD-based adversary $\mathcal{B}_{\text{GD}}$ that black-box ϵ -simulates any sample-based one. As in Sects. 4 and 5, for simplicity here we state and prove our result for samples over $\mathcal{X} \times \mathcal{Y}$ where $\mathcal{X} = \{0, 1\}^n$ for some integer $n \in \mathbb{N}$ and $\mathcal{Y} = \{0, 1\}$. We prove the following theorem:

Theorem 6.1. *Let $\mathcal{X} = \{0, 1\}^n$ for some $n \geq 1$, and let $\mathcal{Y} = \{0, 1\}$. For any $\epsilon > 0$ there exists a GD-based adversary $\mathcal{B}_{\text{GD}}$ that black-box ϵ -simulates all sample-based adversaries with respect to $\mathcal{X} \times \mathcal{Y}$, where:*

- $\mathcal{B}_{\text{GD}}$ runs in time $T_{\mathcal{B}_{\text{GD}}}(|S|, |\mathcal{X}|, |\mathcal{Y}|) = O((|S| \cdot \log(|S|/\epsilon))^2 \cdot \log |\mathcal{X}|)$.
- $\mathcal{B}_{\text{GD}}$ issues $Q_{\mathcal{B}_{\text{GD}}}(|S|, |\mathcal{X}|, |\mathcal{Y}|) = O(|S| \cdot \log(|S|/\epsilon))$ queries to the gradient oracle, followed by a single query to the simulated sample-based adversary.

For the more general setting in which $|\mathcal{Y}| = m > 1$ for some integer $m \in \mathbb{N}$, our simulator $\mathcal{B}$ runs in time $T_{\mathcal{B}}(|S|, |\mathcal{X}|, |\mathcal{Y}|) = O((|S| \cdot \log(|S|/\epsilon))^2 \cdot (\log |\mathcal{X}| + \log |\mathcal{Y}|))$ and issues the same number of queries $Q_{\mathcal{B}}(|S|, |\mathcal{X}|, |\mathcal{Y}|) = O(|S| \cdot \log(|S|/\epsilon))$. We refer the reader to the full version of this work for the extended result.

Due to space limitations, the formal proof of Theorem 6.1 is provided as supplementary material in the full version of this work. In the remainder of this section, we provide an overview of this proof and explaining the main ideas underlying our GD-based simulator described in Fig. 5.

Proof Overview. The main challenge underlying the proof of Theorem 6.1 stems from the “non-adaptive” nature of the GD updates. In Sects. 4 and 5, we extracted samples after each query, for queries that isolated a single sample. Here, we are allowed to extract samples only after completing the entire execution of the GD algorithm, meaning concluding all T iterations. Towards this goal, we extend our parameters so that they can store all of the gradient oracle’s responses, and we extract all samples at once. This requires the size of the parameters to grow from $O(\log |\mathcal{X}|)$ to $O(T \cdot \log |\mathcal{X}|)$. In addition, in our previous constructions we defined the parameters $\bar{\theta}$ to include a vector of additional parameters $\vec{w} \in \mathbb{R}^{n+1}$ that enabled to obtain the desired gradient behavior. As we only used $\vec{w}$ for its gradient vector, any specific choice of $\vec{w}$ could be used, and for consistency we used the all-zeros vector $\vec{w} = (0, \dots, 0) \in \mathbb{R}^{n+1}$. In the GD setting, we can no longer extract the samples directly from the value of $\vec{g}$, and must write the value of $\vec{g}$ to that of $\vec{w}$ in each iterations. Therefore, here we require $\vec{w}$ to be initialized as the all-zeros vector.

The black-box GD-based simulator $\mathcal{B}_{\text{GD}}^{\mathcal{O}_{\text{G}}(S, \cdot, \cdot, \cdot), \mathcal{A}(\cdot)}(1^s)$:

1. Set $k = \lceil \log_2(2s) \rceil$, $T = \lceil 16s \cdot \ln(s/\epsilon) \rceil$, and $\eta = -s$.
2. Set $\ell(\hat{y}, y) = (2 + y) \cdot \hat{y}$.
3. Initialize $\vec{\theta}^{(0)} = (\vec{\theta}_1^{(0)}, \dots, \vec{\theta}_T^{(0)})$, such that for all $t \in [T]$ we denote by $\vec{\theta}_t^{(0)}$ the set of parameters $(\vec{w}_t^{(0)}, z_t, M_t, \kappa_t^{(0)})$ initialized as follows:
 - $\vec{w}_t^{(0)} = (0, \dots, 0) \in \mathbb{R}^{n+1}$
 - $z_t \in \{0, 1\}^k$
 - $M_t \in \mathcal{M}_{n,k}$
 - $\kappa_t^{(0)} = 0$
4. Define $h((\vec{w}, z, M, \kappa), \vec{x}) = (\vec{w} \cdot ((\vec{x})_1, \dots, (\vec{x})_n, 1)) \cdot \mathbb{1}_{M(x)=z} + \frac{\kappa}{s}$.
5. Define $\tilde{h}(\vec{\theta}, \vec{x}) = \sum_{j=1}^T \phi(\kappa_{j-1}, \kappa_j, h(\vec{\theta}_j, \vec{x}))$.
6. For $t \in \{1, \dots, T\}$: // Running GD for T iterations
 - 6.1 Obtain $\vec{g} = \mathcal{O}_{\text{G}}(S, \ell, \tilde{h}, \vec{\theta}^{(t-1)}) \in \mathbb{R}^{|\vec{\theta}^{(t-1)}|}$
 - 6.2 Update $\vec{\theta}^{(t)} = \vec{\theta}^{(t-1)} - \eta \cdot \vec{g}$.
7. Set $D = \emptyset$.
8. For $t \in \{1, \dots, T\}$: // decode extracted samples from parameters $\vec{w}_1^{(T)}, \dots, \vec{w}_T^{(T)}$
 - 8.1 If $0 < \vec{w}_t^{(T)} < 4$:
 - i. Define $\hat{x} = (1/(\vec{w}_t^{(T)})_{n+1}) \cdot (\vec{w}_t^{(T)})_{1..n}$ and $\hat{y} = (\vec{w}_t^{(T)})_{n+1} - 2$.
 - ii. Update $D = D \cup \{(\hat{x}, \hat{y})\}$.
9. Compute $v \leftarrow \mathcal{A}(D)$ and output v .

Fig. 5. The black-box gradient-based simulator $\mathcal{B}_{\text{GD}}$.

Another main difference compared to Sect. 5 lies in the selection of the string z and the pairwise independent function M in each of the T iterations. In the gradient-based construction the adversary samples a fresh pair of (z, M) in each iteration and encodes them in the parameters $\vec{\theta}$ so they can be used by the gradient oracle $\mathcal{O}_{\text{G}}$, i.e., $\vec{\theta} = (\vec{w}, z, M)$. In the GD-based setting this is no longer possible, as the GD update rule does not allow the adversary to encode arbitrary new values into the parameters $\vec{\theta}$ in each iteration. As such, the sampling of the z 's and M 's for all iterations must be made in advance. Namely, our GD-based adversary running for T iterations will sample in advance T pairs (z, M) and encode all of them at once in the initial parameters $\vec{\theta}$. However, in order to use a different set of parameters $(z, M, \vec{w})$ in each iteration, a “clock” mechanism must be introduced, to direct the computation to use only the relevant parameters at each iteration. For this, following Abbe et al. [AKM+21], we include in each set of parameters $(z, M, \vec{w})$ an additional bit κ indicating whether this set was already used in some iteration. We initialize $\kappa = 0$, and after using the corresponding set of parameters, the value of κ is updated to 1. At each iteration, we only use the *first* set of parameters for which $\kappa = 0$.

More formally, we construct the model function and model parameters as follows. For any parameter p , we denote by $p^{(t)}$ the value of the parameter at the end of iteration t . Then, for every $t \in [T]$ the adversary initializes a set of parameters $\vec{\theta}_t^{(0)} = (\vec{w}_t^{(0)}, z_t, M_t, \kappa_t^{(0)})$ such that $\vec{w}_t^{(0)} = (0, \dots, 0) \in \mathbb{R}^{n+1}$, $z_t \in \{0, 1\}^k$,

$M_t \in \mathcal{M}_{n,k}$, and $\kappa_t^{(0)} = 0$. Similarly to the construction in Fig. 4, we define the parameters z and M to be non-differentiable, and as such their value does not change between different iterations. We get that the entire set of parameters at initialization time is of the form:

$$\vec{\theta}^{(0)} = \left(\vec{\theta}_1^{(0)}, \dots, \vec{\theta}_T^{(0)} \right) .$$

As the GD update rule requires the subtraction of the gradient vector $\vec{g}$ from the vector of parameters $\vec{\theta}$, we modify our notation of the gradient. In our previous results we denoted by ∇_{θ} the gradient with respect to only the differentiable parameters in $\vec{\theta}$, so that when $\vec{\theta}$ is composed of p_{diff} differentiable parameters and $p_{\text{non-diff}}$ non-differentiable ones, we get that $|\vec{g}| = p_{\text{diff}}$. Here we will denote by ∇_{θ} the gradient with respect to the differentiable parameters, and in positions corresponding to non-differentiable parameters we will return 0, and we get $|\vec{g}| = p_{\text{diff}} + p_{\text{non-diff}}$.

Let $\phi : \mathbb{R}^3 \rightarrow \mathbb{R}$ be a differentiable function such that

$$\phi(\alpha_1, \alpha_2, \alpha_3) = \begin{cases} \alpha_3 & \alpha_1 = 1 \wedge \alpha_2 = 0 \\ 0 & \alpha_1 = \alpha_2 = 0 \\ 0 & \alpha_1 = \alpha_2 = 1 \\ * & \text{otherwise} \end{cases} ,$$

where the “*” symbol may correspond to any arbitrary values that ensure the continuity of the function, required for it to be differentiable. Then, we define the differentiable model $\tilde{h}$ as follows:

$$\tilde{h}(\vec{\theta}, \vec{x}) = \sum_{j=1}^T \phi \left(\kappa_{j-1}, \kappa_j, h(\vec{\theta}_j, \vec{x}) \right) ,$$

where h is the differentiable function defined as in Fig. 4, and we set $\kappa_0 = 1$. We prove that this construction indeed provides the clock mechanism described above, such that when running GD for T iterations we get that in each iteration we initiate an independent evaluation of h on a fresh set of parameters $\vec{\theta}_t$.

It is important to note that while we extended the construction of our fully-parallelizable gradient-based adversary, the construction here is not necessarily parallelizable. This is due to GD being an *iterative* algorithm in its nature, where a fully-parallelizable solution is less intuitive. The construction in Fig. 4 can be viewed as a one-step parallelizable variant of GD. The DFS-based construction described in Fig. 3 does not fit the GD-based construction as all “guesses” must be made in advance and in a non-adaptive manner.

References

- [ABW10] Applebaum, B., Barak, B., Wigderson, A.: Public-key cryptography from different assumptions. In: Proceedings of the Forty-second ACM Symposium on Theory of Computing, pp. 171–180(2010)
- [ADH+19] Arora, S., Du, S., Hu, W., Li, Z., Wang, R.: Fine-grained analysis of optimization and generalization for overparameterized two-layer neural networks. In: International Conference on Machine Learning, pp. 322–332. PMLR (2019)
- [AKM+21] Abbe, E., Kamath, P., Malach, E., Sandon, C., Srebro, N.: On the power of differentiable learning versus PAC and SQ learning. *Adv. Neural. Inf. Process. Syst.* **34**, 24340–24351 (2021)
- [AS20] Abbe, E., Sandon, C.: Poly-time universality and limitations of deep learning. arXiv preprint [arXiv:2001.02992](https://arxiv.org/abs/2001.02992) (2020)
- [BB22] Baksı, A., Baksı, A.: Machine learning-assisted differential distinguishers for lightweight ciphers. In: Classical and Physical Security of Symmetric Key Cryptographic Algorithms, pp. 141–162(2022)
- [BCC+10] Bernstein, D.J., Chen, H.-C., Cheng, C.-M., Lange, T., Niederhagen, R., Schwabe, P., Yang, B.-Y.: ECC2K-130 on NVIDIA GPUs. In: Gong, G., Gupta, K.C. (eds.) INDOCRYPT 2010. LNCS, vol. 6498, pp. 328–346. Springer, Heidelberg (2010). https://doi.org/10.1007/978-3-642-17401-8_23
- [Ber07] Bernstein, D.J.: Better price-performance ratios for generalized birthday attacks. Workshop Record of SHARCS 2007: Special-purpose Hardware for Attacking Cryptographic Systems (2007)
- [BGL+21] Bao, Z., Guo, J., Liu, M., Ma, L., Tu, Y.: Conditional differential-neural cryptanalysis. *IACR Cryptol. ePrint Arch.* **2021**, 719 (2021)
- [BGP+21] Benamira, A., Gerault, D., Peyrin, T., Tan, Q.Q.: A deeper look at machine learning-based cryptanalysis. In: Canteaut, A., Standaert, F.-X. (eds.) EUROCRYPT 2021. LNCS, vol. 12696, pp. 805–835. Springer, Cham (2021). https://doi.org/10.1007/978-3-030-77870-5_28
- [BHY+23] Buzaglo, G., et al.: Deconstructing data reconstruction: Multiclass, weight decay and general losses. arXiv preprint [arXiv:2307.01827](https://arxiv.org/abs/2307.01827) (2023)
- [BKK+12] Bos, J.W., Kaihara, M.E., Kleinjung, T., Lenstra, A.K., Montgomery, P.L.: Solving a 112-bit prime elliptic curve discrete logarithm problem on game consoles using sloppy reduction. *Int. J. Appl. Cryptogr.* **2**(3), 212–228 (2012)
- [BKN+10] Bos, J.W., Kleinjung, T., Niederhagen, R., Schwabe, P.: ECC2K-130 on cell CPUs. In: Bernstein, D.J., Lange, T. (eds.) AFRICACRYPT 2010. LNCS, vol. 6055, pp. 225–242. Springer, Heidelberg (2010). https://doi.org/10.1007/978-3-642-12678-9_14
- [BLN+09] Bernstein, D.J., Lange, T., Niederhagen, R., Peters, C., Schwabe, P.: FSB-day: Implementing Wagner’s generalized birthday attack against the SHA-3 round-1 candidate FSB. In: Roy, B., Sendrier, N. (eds.) INDOCRYPT 2009. LNCS, vol. 5922, pp. 18–38. Springer, Heidelberg (2009). https://doi.org/10.1007/978-3-642-10628-6_2
- [BLY+23] Bao, Z., Lu, J., Yao, Y., Zhang, L.: More insight on deep learning-aided cryptanalysis. In: International Conference on the Theory and Application of Cryptology and Information Security, pp. 436–467. Springer (2023). https://doi.org/10.1007/978-981-99-8727-6_15

- [Bos15] Bos, J.W.: Parallel cryptanalysis. Summer school on real-world crypto and privacy, Croatia (2015). Slides available at <https://summerschool-croatia.cs.ru.nl/2015/ParallelCryptanalysis.pdf>
- [BSS+13] Beaulieu, R., Shors, D., Smith, J., Treatman-Clark, S., Weeks, B., Wingers, L.: The simon and speck families of lightweight block ciphers. Cryptology ePrint Archive, Paper 2013/404 (2013). <https://eprint.iacr.org/2013/404>
- [CLE+19] Carlini, N., Liu, C., Erlingsson, Ú., Kos, J., Song, D.: The secret sharer: Evaluating and testing unintended memorization in neural networks. In: 28th USENIX Security Symposium (USENIX Security 19), pp. 267–284 (2019)
- [CN11] Chen, Y., Nguyen, P.Q.: BKZ 2.0: better lattice security estimates. In: Lee, D.H., Wang, X. (eds.) ASIACRYPT 2011. LNCS, vol. 7073, pp. 1–20. Springer, Heidelberg (2011). https://doi.org/10.1007/978-3-642-25385-0_1
- [Coo23] Cook, S.A.: The complexity of theorem-proving procedures. In: Logic, Automata, and Computational Complexity: The Works of Stephen A. Cook, pp. 143–152 (2023)
- [CSY+23] Chen, Y., Shen, Y., Yu, H., Yuan, S.: A new neural distinguisher considering features derived from multiple ciphertext pairs. Comput. J. **66**(6), 1419–1433 (2023)
- [Dan20] Daniely, A.: Neural networks learning and memorization with (almost) no over-parameterization. Adv. Neural. Inf. Process. Syst. **33**, 9007–9016 (2020)
- [Din14] Dinur, I.: Improved differential cryptanalysis of round-reduced speck. In: Joux, A., Youssef, A. (eds.) SAC 2014. LNCS, vol. 8781, pp. 147–164. Springer, Cham (2014). https://doi.org/10.1007/978-3-319-13051-4_9
- [DLL+19] Du, S., Lee, J., Li, H., Wang, L., Zhai, X.: Gradient descent finds global minima of deep neural networks. In: International Conference on Machine Learning, pp. 1675–1685. PMLR (2019)
- [DPS23] Ducas, L., Postlethwaite, E., Sotáková, J.: Salsa Verde versus the actual state of the art. CRYPTO 2023 Rump Session Talk. <https://crypto.iacr.org/2023/rump/crypto2023rump-paper13.pdf> (2023)
- [GLN22] Gohr, A., Leander, G., Neumann, P.: An assessment of differential-neural distinguishers. Cryptology ePrint Archive, Paper 2022/1521 (2022). <https://eprint.iacr.org/2022/1521>
- [Goh19] Gohr, A.: Improving attacks on round-reduced Speck32/64 using deep learning. In: Boldyreva, A., Micciancio, D. (eds.) CRYPTO 2019. LNCS, vol. 11693, pp. 150–179. Springer, Cham (2019). https://doi.org/10.1007/978-3-030-26951-7_6
- [HRC21a] Hou, Z., Ren, J., Chen, S.: Cryptanalysis of round-reduced simon32 based on deep learning. Cryptology ePrint Archive, Paper 2021/362 (2021). <https://eprint.iacr.org/2021/362>
- [HRC21b] Hou, Z., Ren, J., Chen, S.: Improve neural distinguisher for cryptanalysis. Cryptology ePrint Archive, Paper 2021/1017 (2021). <https://eprint.iacr.org/2021/1017>
- [HVV+22] Haim, N., Vardi, G., Yehudai, G., Shamir, O., Irani, M.: Reconstructing training data from trained neural networks. Adv. Neural. Inf. Process. Syst. **35**, 22911–22924 (2022)
- [HZR+16] He, K., Zhang, X., Ren, S., Sun, J.: Deep residual learning for image recognition. In: Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, pp. 770–778 (2016)

- [JGH18] Jacot, A., Gabriel, F., Hongler, C.: Neural tangent kernel: convergence and generalization in neural networks. *Adv. Neural Inform. Process. Syst.* **31** (2018)
- [JKM20] Jain, A., Kohli, V., Mishra, G.: Deep learning based differential distinguisher for lightweight cipher present. *Cryptology ePrint Archive*, Paper 2020/846 (2020). <https://eprint.iacr.org/2020/846>
- [KB14] Kingma, D.P., Ba., J.: Adam: a method for stochastic optimization. *arXiv preprint* [arXiv:1412.6980](https://arxiv.org/abs/1412.6980) (2014)
- [Kea98] Kearns, M.: Efficient noise-tolerant learning from statistical queries. *J. ACM (JACM)* **45**(6), 983–1006 (1998)
- [Kha93] Kharitonov, M.: Cryptographic hardness of distribution-specific learning. In: *Proceedings of the Twenty-fifth Annual ACM Symposium on Theory of Computing*, pp. 372–381 (1993)
- [KS09] Klivans, A.R., Sherstov, A.A.: Cryptographic hardness for learning intersections of halfspaces. *J. Comput. Syst. Sci.* **75**(1), 2–12 (2009)
- [KV89] Kearns, M.J., Valiant, L.G.: Cryptographic limitations on learning boolean formulae and finite automata. In: *Proceedings of the 21st Annual ACM Symposium on Theory of Computing*, pp. 433–444 (1989)
- [Lev73] Levin, L.A.: Universal sequential search problems. *Problemy peredachi informatsii* **9**(3), 115–116 (1973)
- [LLS+22] Lu, J., Liu, G., Sun, B., Li, C., Liu, L.: Improved (related-key) differential-based neural distinguishers for simon and simeck block ciphers. *arXiv preprint* [arXiv:2201.03767](https://arxiv.org/abs/2201.03767) (2022)
- [LSW+23a] Li, C., Sotakova, J., Wenger, E., Allen-Zhu, Z., Charton, F., Lauter, K.: Salsa verde: a machine learning attack on learning with errors with sparse small secrets. *arXiv preprint* [arXiv:2306.11641](https://arxiv.org/abs/2306.11641) (2023)
- [LSW+23b] Li, C., Sotáková, J., Wenger, E., Malhou, M., Garcelon, E., Charton, F., Lauter, K.: Salsa picante: a machine learning attack on lwe with binary secrets. *arXiv preprint* [arXiv:2303.04178](https://arxiv.org/abs/2303.04178) (2023)
- [Mic] Micciancio, D.: Parallel algorithms for lattice problems. <https://cseweb.ucsd.edu/~daniele/LatticeLinks/Parallel.html>
- [Nie12] Niederhagen, R.: Parallel Cryptanalysis. PhD thesis, Eindhoven University of Technology (2012). <http://polycephaly.org/thesis/niederhagen-thesis-printed.pdf>
- [Pol78] Pollard, J.M.: Monte Carlo methods for index computation (mod p). *Math. Comput.* **32**, 918–924 (1978)
- [Reg05] Regev, O.: On lattices, learning with errors, random linear codes, and cryptography. In: *Proceedings of the 37th Annual ACM Symposium on Theory of Computing*, pp. 84–93 (2005)
- [Riv91] Rivest, R.L.: Cryptography and machine learning. In: Imai, H., Rivest, R.L., Matsumoto, T. (eds.) *ASIACRYPT 1991*. LNCS, vol. 739, pp. 427–439. Springer, Heidelberg (1993). https://doi.org/10.1007/3-540-57332-1_36
- [RM51] Robbins, H., Monro, S.: A stochastic approximation method. *Annals math. stat.*, 400–407 (1951)
- [Rud16] Ruder, S.: An overview of gradient descent optimization algorithms. *arXiv preprint* [arXiv:1609.04747](https://arxiv.org/abs/1609.04747) (2016)
- [SRS17] Song, C., Ristenpart, T., Shmatikov, V.: Machine learning models that remember too much. In: *Proceedings of the 2017 ACM SIGSAC Conference on Computer and Communications Security*, pp. 587–601 (2017)

- [SSBD14] Shalev-Shwartz, S., Ben-David, S.: Understanding machine learning: From theory to algorithms. Cambridge university press (2014)
- [SWL+24] Stevens, S., et al.: SALSA FRESCA: Angular embeddings and pre-training for ML attacks on learning with errors. Cryptology ePrint Archive, Paper 2024/150 (2024)
- [SZB21] Song, M.J., Zadik, I., Bruna, J.: On the cryptographic hardness of learning single periodic neurons. Adv. Neural. Inf. Process. Syst. **34**, 29602–29615 (2021)
- [SZM21] Su, H.-C., Zhu, X.-Y., Ming, D.: Polytopic attack on round-reduced Simon32/64 using deep learning. In: Wu, Y., Yung, M. (eds.) Inscrypt 2020. LNCS, vol. 12612, pp. 3–20. Springer, Cham (2021). https://doi.org/10.1007/978-3-030-71852-7_1
- [Val84] Valiant, L.G.: A theory of the learnable. Commun. ACM **27**(11), 1134–1142 (1984)
- [vOW99] van Oorschot, P.C., Wiener, M.J.: Parallel collision search with cryptanalytic applications. J. Cryptol. **12**(1), 1–28 (1999)
- [Wag02] Wagner, D.: A generalized birthday problem. In: Yung, M. (ed.) CRYPTO 2002. LNCS, vol. 2442, pp. 288–304. Springer, Heidelberg (2002). https://doi.org/10.1007/3-540-45708-9_19
- [WCC+22] Wenger, E., Chen, M., Charton, F., Lauter, K.E.: Salsa: attacking lattice cryptography with transformers. Adv. Neural. Inf. Process. Syst. **35**, 34981–34994 (2022)
- [YK21] Yadav, T., Kumar, M.: Differential-ML distinguisher: machine learning based generic extension for differential cryptanalysis. In: Longa, P., Ràfols, C. (eds.) LATINCRYPT 2021. LNCS, vol. 12912, pp. 191–212. Springer, Cham (2021). https://doi.org/10.1007/978-3-030-88238-9_10
- [ZBH+17] Zhang, C., Bengio, S., Hardt, M., Recht, B., Vinyals, O.: Understanding deep learning requires rethinking generalization. In: International Conference on Learning Representations (2017)
- [ZW22] Zhang, L., Wang, Z.: Improving differential-neural distinguisher model for des, chaskey, and present. arXiv preprint [arXiv:2204.06341](https://arxiv.org/abs/2204.06341) (2022)
- [ZWw22] Zhang, L., Wang, Z., Wang, B.: Improving differential-neural cryptanalysis. Cryptology ePrint Archive, Paper 2022/183 (2022). <https://eprint.iacr.org/2022/183>