



# One-Time Programs from Commodity Hardware

Harry Eldridge<sup>1</sup>, Aarushi Goel<sup>2(✉)</sup>, Matthew Green<sup>1</sup>, Abhishek Jain<sup>1</sup>,  
and Maximilian Zinkus<sup>1</sup>

Johns Hopkins University, Baltimore, USA  
{hmedgreen abhishek zinkus}@cs.jhu.edu  
<sup>2</sup> NTT Research, Sunnyvale, USA  
aarushi.goel@ntt-research.com

**Abstract** One-time programs, originally formulated by Goldwasser et al. [26], are a powerful cryptographic primitive with compelling applications. Known solutions for one-time programs, however, require specialized secure hardware that is not widely available (or, alternatively, access to blockchains and very strong cryptographic tools).

In this work we investigate the possibility of realizing one-time programs from a recent and now more commonly available hardware functionality: the *counter lock ox*. A counter lockbox is a stateful functionality that protects an encryption key under a user-specified password, and enforces a limited number of incorrect guesses. Counter lockboxes have become widely available in consumer devices and cloud platforms.

We show that counter lockboxes can be used to realize one-time programs for general functionalities. We develop a number of techniques to reduce the number of counter lockboxes required for our constructions, that may be of independent interest.

## 1 Introduction

One-time programs, formulated by Goldwasser et al. [26], are a flexible and powerful cryptographic primitive with compelling applications to limited-attempt authentication, fuzzy vaults, limited-query differential-private data analysis, and even autonomous ransomware and beyond. In the standard model, one-time programs are known to be impossible to realize purely in software [13, 26]. To evade this impossibility, prior works have examined the problem of building one-time programs from secure hardware tokens [26, 30], or alternatively, using blockchains [28].

The works of [26, 30] employ tamper-proof hardware that implements *one-time memory* – a simple, stateful functionality that allows anyone to read one location, after which all other locations become inaccessible. While these results are practical and work in a variety of settings, they have mainly garnered theoretical interest. The likely cause is that one-time memory tokens have not been available as a standard feature of popular personal or cloud computing platforms. While it is possible to realize these tokens using programmable smart cards or

HSMs [17, 32, 52], such development typically requires expensive equipment and considerable development effort. Moreover, the few affordable platforms that support custom programming may provide weak or limited security guarantees. If portability is not required, tamper-proof hardware tokens can also be realized through virtualization: *secure enclaves* such as Intel SGX [42] and ARM TrustZone [46] offer tamper-resilience under relatively strong adversarial assumptions such as operating system (OS) compromise. Indeed, if such an enclave platform is considered trusted, it is likely easier to implement an entire one-time functionality within the enclave. However, implicit trust in an enclave provider is unacceptable in some threat models, and the soundness of this trust regardless of threat model has been repeatedly called into question [14, 19, 45]. These execution environments also typically place limitations on end-users’ ability to deploy arbitrary code [6, 33, 50].

**Counter Lockboxes.** Recently, a new generation of device- and cloud-based secure hardware has become available to end users. This includes secure co-processors that are now built into many smartphones and tablets, including the Apple Secure Enclave Processor (SEP) [3] and Google’s Titan M2 [27] co-processor. It also includes specialized Hardware Security Modules (HSMs) that have recently been deployed within the data centers of consumer cloud providers; these can be accessed remotely from consumer devices to implement services such as Apple’s Cloud Key Vault [39], Android Backup [48], WhatsApp backup [38], and Signal Secure Value Recovery [40]. Notably, these systems are not aimed at enterprise customers; they are configured to protect end-user cryptographic keys, even from attacks that might be launched by the device manufacturer or cloud provider themselves. These systems are now being used across billions of devices, making them more broadly accessible to consumers than any prior secure hardware platform.

Unlike secure enclave environments such as TrustZone or SGX, these consumer-oriented hardware devices do not allow end-devices to securely execute arbitrary programs. Instead, they present a limited interface to the device’s application software. Since the primary purpose of these systems is to protect encryption keys under user-selected passwords, the most common interface is a functionality akin to what we describe as a *counter lockbox*.<sup>1</sup> To initialize a lockbox, the application software provides a password to the hardware along with a *maximum attempt limit*. At any later point, the software can retrieve the decryption key by providing the correct password. To protect the key against guessing attacks, the hardware increments a tamper-resistant counter for each incorrect guess: when this counter exceeds the maximum attempt limit, the hardware deletes the stored key. Given that this lockbox functionality has been deployed at massive scale, it represents an attractive building block for constructing more sophisticated cryptographic protocols.

---

The term *counter lock box* was previously introduced by Apple for its SEP [3]. We use it in this work to refer to a broad class of similar functionalities.

**Using Lockboxes to Construct One-Time Programs.** The ubiquity of this basic lockbox functionality motivates us to investigate the following question: *can such a simple functionality be used to achieve general secure computation?* In this work, we answer the question in the affirmative: given access to a sufficient number of lockboxes, we show that it is possible to realize the full power of one-time programs.

This result has important practical implications: since lockboxes are increasingly available to consumer hardware, this approach provides a “backdoor” route to constructing obfuscated software, even on hardware that does not directly support this functionality. This capability facilitates many constructive applications. For example, it can be used to build sophisticated attempt-limiting authentication functionalities. A limited-attempt fuzzy vault [34] can release cryptographic secrets when a user provides an input that satisfies some complex approximate function such as biometric matching or inexact string comparison [15]. Obfuscated software also enables privacy-preserving applications such as differentially-private statistical data analysis, where query limits must be enforced to maintain a privacy budget [22]. This functionality has a dark side as well: one-time programs allow for the creation of *autonomous ransomware* [10, 20, 36], a form of malware with no command-and-control infrastructure: in this paradigm, decryption keys are revealed only when the user provides the malware with proof of payment on a public blockchain. This last concern illustrates how carefully system designers must tread when exposing secure lockbox functionality to users and developers, since as we demonstrate in this work, even this relatively weak primitive can be leveraged into powerful secure computation. The lower bounds for this transformation also raise practical concerns: system designers may wish to know *how many* instances may be safely exposed to users before the power of these constructions can be exploited.

## 1.1 Our Results

In this work, we show that it is possible to construct secure one-time programs (OTP) using multiple instances of the counter lockbox functionality. Our main result is a construction of OTP for general functionalities based on one-way functions that requires a *constant* number of counter lockboxes per input-bit. This asymptotically matches prior constructions of one-time programs [25] in the number of hardware tokens utilized.

**Theorem 1 (Informal).** *Assuming the existence of one-way functions, for any functionality  $F$ , there exists a construction of one-time programs in the lockbox-hybrid model that makes  $\mathcal{O}(1)$  invocations to the lockbox functionality per input bit of  $F$ .*

We present our main result with counter lockboxes that allow exactly one password attempt. In practice, lockboxes may allow more attempts. For example, lockboxes may fix the maximum number of attempts to some system-wide constant (*e.g.*, 10 attempts.) To handle such cases, we demonstrate an extension

of our main construction that supports lockboxes with *any* number of password attempts. The resulting scheme requires the same number of lockboxes as before.

**Reducing The Number of Hardware Tokens.** We observe that at the cost of stronger assumptions, it is possible to achieve an asymptotic reduction in the *total* number of counter lockboxes. In particular, by using laconic oblivious transfer (LOT) [18] with malicious receiver security, we can reduce the total number of lockboxes to be *independent* of the input size and to depend only on the security parameter.

Our transformation is *generic*, and is applicable to any OTP construction (including prior known schemes). As such, this might be of independent interest.

**Theorem 2 Informal).** *Assuming the existence of malicious receiver laconic oblivious transfer, for any functionality  $F$ , there exists a construction of one-time programs that makes  $\mathcal{O}(\lambda)$  total invocations to the lockbox functionality (where  $\lambda$  is the security parameter).*

LOT schemes with malicious receiver security can be generically constructed by compiling the receiver message of existing LOT schemes with succinct arguments of knowledge (SNARKs) [11, 44] either in the random oracle model, or by relying on knowledge assumptions.

**Our Approach.** Our starting point is the observation from the work of Goldwasser et al. [26] that garbled circuits [51] are almost like one-time programs, except the seeming need of interactive oblivious transfer (OT) to transmit the wire labels corresponding to an evaluator’s input. Fortunately, a one-time memory (OTM) token naturally yields the OT functionality, which paves the way for constructing one-time programs from OTM tokens.

Unlike OTMs, however, a natural use of counter lockboxes yields a “leaky” OT functionality, where the receiver is able to learn *both* sender inputs with some constant probability (we elaborate on this in Sect. 2). By applying standard OT combiner techniques [31, 43], the leaky OT functionality can be transformed into secure OT. However, this results in a significant overhead in the number of lockboxes required. Specifically, this approach requires  $O(\lambda)$  lockboxes *per input bit* of the functionality, as opposed to  $O(1)$  OTMs required in prior works.

Towards obtaining our result in Theorem 1, we observe that  $O(1)$  lockboxes per input bit are sufficient to instantiate a leaky “batch” oblivious transfer functionality, where the receiver can learn both sender inputs for an a priori bounded constant fraction of the input bits. We then devise a way to construct a secure (i.e., “non-leaky”) batch-OT from leaky batch-OT via *robust garbling* – a form of garbling where security holds even if the receiver learns both labels for a constant fraction of the input wires – for special functions. The secure batch-OT can then be used together with standard garbled circuits to obtain one-time programs for general functions.

Finally, we demonstrate that using laconic OT, the task of designing OTP for general functions with arbitrary input lengths can be reduced to the task of designing an OTP for functions whose input length is a fixed polynomial

in the security parameter. As a result of this reduction, we are able to “compress” the effective input size, thereby achieving a reduction in the number of required hardware tokens. As we discuss later, this transformation requires an LOT scheme that achieves simulation-based security against malicious receivers.

**Real World Implications.** In order to assess the practical feasibility of our one-time programs, we need to consider several cost factors – number of hardware tokens required, cost of each hardware token, time to generate the OTP, and the size of software component of the OTP.

In our first construction, the main consideration is hardware. Indeed, besides the use of lockboxes to implement leaky batch OT, the rest of our construction comprises of robust garbling for special functions – an efficient, information-theoretic gadget, and regular garbled circuits. The efficiency of state-of-the-art constructions of regular garbled circuits is well-established in prior works [47]. In Sect. 8.1, we evaluate the concrete number of lockboxes required to implement one-time programs in practice and observe that there is a notable (albeit, constant factor) expansion from the input length to the number of total lockboxes required due to the use of binary linear error-correcting codes in our scheme. Overall, our results show that one-time programs may be practical for small to modest-sized inputs using a number of lockboxes that may be practical on today’s systems or systems that will be available in the near future. Because such one-time programs may allow for destructive applications, our concrete bounds on the number of lockboxes can provide safety guidance for system developers who expose such functionalities to application developers.

Given our current understanding of LOT schemes, our second transformation is primarily of theoretical interest at the moment. We first note that recent works [1, 29] have achieved significant improvements in concrete efficiency of LOT by allowing for linear decryption times (as opposed to poly-logarithmic decryption complexity achieved in the initial works). Our transformation only requires the laconic digest property of LOT and is not sensitive to decryption complexity. As such, it can be instantiated using the state-of-the-art LOT schemes with linear decryption complexity. However, the main efficiency bottleneck stems from the fact that our transformation requires a “non-interactive” version of LOT which is obtained by evaluating the LOT sender algorithm *inside a garbled circuit*. For current LOT schemes, this translates to evaluating *public-key* operations inside a garbled circuit for every receiver input bit, which to our current understanding, is quite expensive. Our work, therefore, motivates the design of new LOT schemes (with potentially linear decryption times) with “garbling friendly” sender algorithms.

## 2 Technical Overview

We now describe our main ideas for constructing a one-time program using counter lockboxes. We first describe a basic construction that relies on a fairly large number of lockboxes with only one attempt allowed (denoted  $A = 1$ ). This approach requires  $O(\lambda)$  lockboxes *per bit of input* to the one-time program for

security parameter  $\kappa$ . This construction serves as a technical warm-up and highlights the main challenges in building OTPs from counter lockboxes as opposed to one-time memory (OTM) tokens used by Goldwasser et al. [26].

We then describe our key ideas towards constructing OTPs with many fewer lockboxes, even *constant* per input bit. This asymptotically matches prior constructions based on OTM tokens. Finally, we discuss two extensions. First, we describe a generic method using laconic oblivious transfer [18] (LOT) to reduce the *total* number of lockboxes to be independent of the input size, and to depend only on  $\kappa$ . Second, we describe how our constructions can be extended to support counter lockboxes that allow multiple password attempts.

**Initial Ideas.** Goldwasser et al. [26] proposed a construction of one-time programs using one-time memory (OTM) tokens. Their construction relies on the observation that garbled circuits are almost like one-time programs, except that the sender needs to interact with the receiver (via oblivious transfer) to securely hand over input wire labels for the garbled circuit corresponding to the receiver's input. This interaction can be replaced with OTMs for each input wire: the sender can embed both the 0-label and the 1-label for each wire inside an OTM, and send all the OTMs together with the garbled circuit in *one shot*. The security of OTM ensures that the receiver learns at most one label from each OTM, which it can then use to evaluate the garbled circuit.

While the above idea is intuitive, the security proof requires a bit more care due to the fact that the adversary can choose its input in an *adaptive* fashion and query the OTM tokens in an arbitrary order. In particular, the proof of security requires garbling schemes with adaptive security. Efficient solutions for such garbling schemes are known in the random oracle model [8].

In this work, we build OTPs using a different kind of hardware token, the *counter lockbox*. A natural approach is to emulate the OTM functionality using counter lockboxes. However, an immediate challenge arises. Recall that a counter lockbox protects a secret value with a pre-configured password and limited attempts; if the number of incorrect attempts reaches the threshold, the secret value is irrevocably deleted. A natural idea is to store the two wire labels for each input bit in two separate lockboxes and devise a mechanism that allows a receiver to unlock only one of the two lockboxes. This, however, seems to require revealing only one of the two passwords to the user, returning to the problem of emulating OTM.

## 2.1 Basic Protocol

Our first idea is to use the receiver's input bits as passwords to the lockboxes. Concretely, for each input wire, we can use 0 and 1 as the passwords for the lockboxes that hide the 0-label and 1-label, respectively. The two lockboxes for each wire are then shuffled so that the input-to-password mapping is not known to the receiver.

An honest receiver can simply use the same value to attempt to unlock both lockboxes associated with an input wire. This guarantees that they obtain their desired label from one lockbox and consumes the single attempt of the other.

A malicious receiver may attempt to learn both labels by guessing the password for both of the lockboxes. This will give them only a  $\frac{1}{2}$  chance of success: at least one label remains hidden with that probability. This idea can be leveraged to reduce the adversary's chances of learning both values: instead of embedding each label in a single lockbox, we “distribute” each label across additional lockboxes.

We now discuss the baseline construction of OTP that results from using lockboxes in this manner. A reader already familiar with the garbling based OTP approach may want to skip the next two paragraphs and directly go to the analysis of this baseline construction.

**Generating the OTP.** Let  $C$  be a Boolean circuit with input length  $n$ . The sender first garbles  $C$  to obtain a garbled circuit  $\tilde{C}$  along with  $n$  pairs of wire labels  $(\text{label}_0^i, \text{label}_1^i)$ . It then performs the following steps:

1. Sample uniform bits  $b_1, \dots, b_{2\ell}$ , where  $\ell$  counts the number of lockboxes each label is distributed across.
2. For each  $j = 1$  to  $2\ell$ : first, create an independent lockbox  $L_j^i$  using maximum attempt counter  $A = 1$  and password  $P = b_j$ . Receive the corresponding lockbox secret  $K_j$ .
3. Next, compute  $\text{CT}_0^i = \text{label}_0^i \quad \forall j, b_j=0 \quad K_j$  and  $\text{CT}_1^i = \text{label}_1^i \quad \forall j, b_j=1 \quad K_j$ .

Finally, the sender provides the receiver with the garbled circuit  $\tilde{C}$  and the tuples  $(\text{CT}_0^1, \text{CT}_1^1), \dots, (\text{CT}_0^n, \text{CT}_1^n)$  as well as references to the  $2\ell \cdot n$  lockboxes.

**OTP Evaluation.** To evaluate this program on an input  $x = (x_1, \dots, x_n)$ , the receiver performs the following steps for  $i = 1$  to  $n$ :

1. For  $j = 1$  to  $2\ell$ , attempt to open the lockbox  $L_i$  with password  $x_i$  to retrieve either  $K_j$  or an error (in which case, set  $K_j = 0$ .)
2. Compute  $\text{label}_{x_i}^i = \text{CT}_{x_i}^i \quad \bigwedge_{j=1}^{\lambda} K_j$ .

The receiver can now evaluate  $\tilde{C}$  using the labels  $\text{label}_{x_1}^1, \dots, \text{label}_{x_n}^n$  to obtain a circuit output.

**Analysis.** It is easy to verify correctness of the above construction. What remains is to show that the protocol achieves security, *i.e.*, that a malicious receiver has a negligible chance of recovering more than one label for any input wire. The argument here is simple: to recover both  $(\text{label}_0^i, \text{label}_1^i)$  for some wire  $i$ , the attacker must query each of  $2\ell$  lockboxes  $L_j^i$  using exactly the right passwords. However, since the lockboxes do not reveal the password until the attempt to open is made (at which point, the lockbox either reveals the secret or destroys it), the attacker must succeed in distinguishing between the 0 and 1 lockboxes. With an optimal guessing strategy, this happens with probability  $\frac{1}{2} \cdot \frac{n}{2^{\ell}} \approx \frac{1}{2^{\ell}}$ . Therefore, for  $\ell$  bits of security, we need  $\ell = O(\log \frac{1}{\epsilon})$  lockboxes per-input wire.

**Limitations.** While a decent baseline solution, this simple approach has several limitations. First, the number of lockboxes required grows with  $O(\log \frac{1}{\epsilon})$ , which is significantly worse than the one-time program construction of [26] that requires

a constant number of hardware tokens per wire. Moreover, the above solution does not support lockboxes that allow multiple password attempts, and therefore has limited applicability for real-world use. To address these limitations, in the following sections we present techniques to reduce the number of counter lockboxes required. Later, we also describe approaches for supporting lockboxes that allow multiple password attempts.

## 2.2 Reducing the Number of Lockboxes

Our baseline solution can be seen as implicitly building a secure *combiner* for the OTM functionality. Indeed, the secret-sharing-based approach is also used in prior works that build secure combiners for oblivious transfer (OT) (e.g. [31, 43]). It is natural to ask whether one can obtain a reduction in the number of lockboxes by using a more efficient combiner. To the best of our knowledge, however, all existing methods require an overhead of  $O(\lambda)$  – the same as our baseline solution – when each component is only secure with constant probability.

We now discuss our key insights towards reducing the number of lockboxes required for one-time programs. To streamline this discussion, we start by defining an abstract “leaky” OT primitive and show how to obtain a one-time program using this primitive. Later, we discuss how counter lockboxes can be used to instantiate such a primitive and also analyse the total number of the lockboxes required for this instantiation.

**Insight I: Leaky Batch-OT.** Let us assume we have access to a leaky OT functionality, where the receiver can choose to specify: (1) either a choice bit  $b$  and get sender input  $m_b$  as output, (2) or a special “leakage” option. In this case, it learns both sender inputs  $m_0$  and  $m_1$  with some constant probability, and only one of the these inputs with the remaining probability.

This notion can be generalized to a *leaky batch-OT* functionality, where the receiver is allowed to learn both sender inputs for an a priori bounded *constant fraction* of the OTs. Furthermore, it is easy to see that multiple copies of the leaky OT functionality – one for each input bit – can realize leaky batch-OT. We ask whether it is possible to build one-time programs using leaky batch-OT, *without paying the overhead of standard OT combiners*.

At first, this seems highly unlikely. Indeed, the standard approach to one-time programs – as discussed earlier – involves the use of garbled circuits. Using leaky batch-OT would result in leakage of *both* wire labels for several input wires. The security of standard garbled circuits, however, completely breaks down if both wire labels are leaked even for a single wire (let alone multiple wires).

**Insight II: Robust Garbling.** We address this challenge by using a notion of *robust* garbling – one where security of the garbled function is ensured even if the receiver learns both labels for a constant fraction of the input wires. If achievable, such a tool would be clearly helpful for our task at hand. However, while intuitively appealing, it is not immediately apparent how to formally define such a notion.

With leakage, the adversary may obtain labels for multiple different inputs – inputs differing at bit locations where both wire labels were obtained. Should the adversary then be allowed to learn multiple outputs, or only a single output? Clearly the former conflicts with the one-time nature of the required functionality, and thus we would like to enforce the latter. This raises a new question: *which* output? For example, if the function is such that each input corresponds to a different output, it is not clear how we can enforce the single-output requirement in a meaningful way. Indeed, achieving our intuitive notion of robustness seems impossible for general functions. We note that previously, Alimashaqbeh et al. [2], also considered a notion of robustness in garbled circuits (and more generally in non-interactive secure multiparty computation). However, given their application, they consider a slightly weaker setting, where they are able to assume an a priori fixed output for the adversary and hence do not need to deal with the above issue of “which output to reveal”.<sup>2</sup> Since such assumptions are not applicable to our setting, we cannot rely on their definition of robustness.

We therefore weaken our goal and attempt to define robust garbling for a restricted class of functions that have a huge number of collisions, i.e. where inputs have a certain degree of *redundancy*. If we consider functions where multiple inputs with an overlapping subset of input bits have the same output, we could hope to achieve robustness. Even if the receiver learns multiple labels for the remaining (non-overlapping) bits, it will only learn at most one unique output.

As the following example shows, however, we need to be more careful. Consider two  $n$ -bit input strings  $\mathbf{x}_1$  and  $\mathbf{x}_2$  that share the same first  $n/2$  bits, and another input string  $\mathbf{x}_3$  that shares the same last  $n/2$  bits with  $\mathbf{x}_2$ . Toward the above intuitive description of collisions, if  $\mathbf{x}_1$  and  $\mathbf{x}_2$  correspond to the same output, and  $\mathbf{x}_2$  and  $\mathbf{x}_3$  do as well, by transitivity  $\mathbf{x}_1$  and  $\mathbf{x}_3$  (that do not necessarily share a significant fraction of overlapping bits) also have the same output. Without further specification, this can escalate quickly until all inputs have the same output and we end up with a constant function.

In a pursuit to capture more interesting and non-trivial functions, we specify a class of functions that take inputs of length  $n$ , with respect to a parameter  $\gamma$  and try to capture the idea that there is only at most one unique non- $\perp$  output associated with any  $n - \gamma$  input bits. Note that this is different from saying that inputs with the same subset of  $n - \gamma$  input bits have a unique output. We say that a function is *admissible* if for any  $n - \gamma$  input bits, there exists *at most one unique* combination of the remaining  $\gamma$  bits, such that the output of this function on the combined  $n$ -bit input is a non- $\perp$  value. Moreover, if such a unique combination of the remaining  $\gamma$  bits exists, then it is *easy* to find them using a deterministic procedure.<sup>3</sup> In this work, we consider robust garbling for such admissible functions.

**OTPs from Robust Garbling.** Let us now assume that we have robust garbling for this restricted class of functions. We now describe how we can leverage robust garbling to build OTPs for general functions. Let  $F$  be the

<sup>2</sup> we refer the reader to Sect. 2.5 for a more detailed comparison with their work.

<sup>3</sup> The reason why we need this deterministic procedure will be explained shortly.

intended OTP functionality. Then, consider a new functionality  $F'$  such that  $F'(\text{enc}(\mathbf{x})) = F(\mathbf{x})$ , where  $F'$  is an admissible function amenable to robust garbling and  $\text{enc}$  is some mapping function that allows us to map inputs of  $F$  to inputs of  $F'$ . Concretely, we can use an error-correcting code (ECC) as the mapping function  $\text{enc}$  that can introduce redundancy in the mapped input to help ensure that  $F'$  satisfies the above conditions of being an amenable function.

This idea can now be used to design an OTP for  $F$  as follows: (1) The sender garbles  $F$  using a regular garbling scheme. (2) For each input wire  $i$  and bit  $b \in \{0, 1\}$ , it defines  $F'_{i,b}$  such that on input  $\text{enc}(\mathbf{x})$ ,  $F'_{i,b}$  runs the ECC decoding function  $\text{dec}$  to decode  $\mathbf{x}$  and then if  $\mathbf{x}[i] = b$  it outputs the  $b$ -label for the  $i$ -th wire, and otherwise it outputs  $\perp$ . For any ECC with distance  $\gamma + 1$ , there is only one “valid” codeword associated with any  $n - \gamma$ -bit message, hence, it is easy to see that  $\text{dec}$  (and as a result  $F'_{i,b}$ ) is an admissible function. (3) The sender garbles each  $F'_{i,b}$  using robust garbling. An important point to note is that each  $F'_{i,b}$  takes the same input  $\text{enc}(\mathbf{x})$ . (4) The sender uses this observation to concatenate input labels for each  $F'_{i,b}$  and embed them inside the leaky batch-OT.

**Constructing Leaky OT.** We now describe our idea for constructing leaky OT (and consequently leaky batch-OT). Intuitively, our leaky oblivious transfer functionality allows the receiver to obtain *both* sender inputs with some constant probability.

Our construction of leaky OT is quite natural: in fact, we use the same approach as in the base protocol discussed earlier, where the sender prepares  $2\ell$  lockboxes (where  $\ell$  is some constant) and distributes the “0” and “1” message across  $\ell$  lockboxes. As before, in order to learn both sender inputs, the adversary must correctly guess the passwords for each of the  $2\ell$  associated lockboxes. The adversary then succeeds with a constant probability of  $\approx \frac{1}{2^{\sigma-1}}$ .

For leaky batch-OT, when considering a collection of  $n$  such leaky OTs, the probability that an adversary can successfully obtain both sender inputs for a constant fraction of the OTs is  $\approx \frac{1}{2^{\sigma-1}}$ . Now, observe that if  $n$  is sufficiently large (say  $n = O(\quad)$ ), then the probability  $\approx \frac{1}{2^{\sigma-1}}$  is negligible in  $\quad$ , even if  $\ell$  is some constant value. While this analysis is somewhat simplified, it suffices for the purposes of this discussion. More details can be found in the technical sections.

Importantly, the above insight gives us significant improvement in the required number of lockboxes. Specifically, we now only require a constant number of lockboxes per OT (or input wire). However, as discussed before, in order to implement our idea of combining leaky batch-OT with robust garbling, the length of input to this leaky batch-OT is slightly longer than our “real” input. In particular, the input to our leaky batch-OT is an ECC encoding of the receiver’s input. If we use binary linear ECCs with constant rate, then the length of this codeword is  $n + \gamma$  where  $\gamma = O(n)$ , and we need a total of  $\ell \cdot (n + \gamma)$  lockboxes, which in an amortized sense is a constant number of lockboxes per  $n$ -bits.

**Handling Adaptivity.** We now highlight some important subtleties regarding the security definitions of leaky batch-OT and robust garbling.

In our OTP constructions, we use robust garbling in conjunction with leaky batch-OT. Specifically, the receiver obtains labels for a robust-garbled circuit from the leaky batch-OT. From our prior discussion on leaky batch-OT, it is clear that an adversary can obtain both labels for some (e.g.  $\gamma$  out of  $n$ ) of the input wires of this robust garbling. Moreover, recall that in above construction of leaky batch OT, given the entire set of lockboxes, an adversary can query them in *any* order of its choosing. In fact, it can “adaptively” decide an order based on the outcomes of previously queried lockboxes. In other words, the adversary can be “fully adaptive”. Our definition of leaky batch-OT must allow for this flexibility and our robust garbling must also support this “fully adaptive” setting.

Since the adversary can potentially learn both labels for some of the inputs, for simulation, we need a way to predict the output based only on the input bits for which the adversary gets exactly one label. This is why we require that the set of admissible functions admit a deterministic procedure to predict the only (if any) valid associated output.

Finally, we remark that since the adversary can choose to ask for the *second* label of some input wires in any order, the simulator would not know until the last query which  $n - \gamma$  input bits it must consider to predict the output. However, by then it might be “too late” to correctly simulate garbling. To overcome this, we make a crucial observation about our construction of leaky batch-OT from lockboxes: recall that in our construction we have  $2\ell$  lockboxes associated with every index  $i \in [n]$ . If an adversarial receiver successfully opens the relevant lockboxes and learns *one of the sender messages* (say  $\text{msg}_i^b$ ) associated with that index, it is easy to predict if the adversary will also be able to learn the *other sender message* (say  $\text{msg}_i^{1-b}$ ) corresponding to that index. Indeed, if the adversary made any incorrect password attempts for any of the  $\ell$  lockboxes associated with  $\text{msg}_i^{1-b}$ , then the simulator can predict that the adversary will never be able to learn  $\text{msg}_i^{1-b}$ . However, if no incorrect password attempts were made for those  $\ell$  lockboxes, then the adversary can be certain that the remaining (unopened) lockboxes associated with index- $i$  have password  $1 - b$  and can always successfully open them and learn  $\text{msg}_i^{1-b}$ .

Therefore, we model our definition of leaky batch-OT to require the following: whenever the adversary makes a query for a particular index, it must specify whether it plans to query the second message for this index in the future. Moreover, since we only want to allow for some bounded leakage, the number of indices for which the adversary can make this request is bounded by a parameter  $\gamma$ . This observation helps ensure that the simulator of robust garbling does not need to wait until the “last query” to determine which  $n - \gamma$  input bits it must consider to predict the output. Instead, this can be determined once the adversary makes at least one query for each of the  $n$  indices.

**Constructing Robust Garbling.** We now discuss robust garbling for a subclass of admissible functions. As discussed earlier, such a construction for a restricted function class suffices for our use in the construction of OTP. In particular, we consider admissible functions of the form  $f = (\mathbf{M}, \mathbf{u}, \mathbf{z})$ , where

$\mathbf{M} \in \{0, 1\}^{k \times n}$ ,  $\mathbf{u} \in \{0, 1\}^k$  are public and  $\mathbf{z} \in \{0, 1\}^k$  is private, such that on any input  $\mathbf{x} \in \{0, 1\}^n$ ,  $f(\mathbf{x}) = \begin{cases} \mathbf{z} & \text{if } \mathbf{u} = \mathbf{M}\mathbf{x} \\ \mathbf{z}' \xleftarrow{\$} \{0, 1\}^k & \text{otherwise} \end{cases}$

While all “invalid” inputs must lead to a  $\perp$  output in admissible functions, the above function instead outputs a random  $\mathbf{z}'$ . We note that this is not a problem in our setting (and the above function is still admissible). This is because in our OTP construction, the value  $\mathbf{z}$  will correspond to labels of the garbled circuit that garbles the actual function for which we compute the OTP. In the case that the output of the above function is a random unrelated value instead of a valid label, the receiver will be able to detect this while evaluating and demarcate this output as essentially equivalent to  $\perp$ . We elaborate more on this in Sect. 6.2.

Benhamouda et al. [9] design a non-interactive *multi-party* computation (NIMPC) protocol for such functions, but where  $\mathbf{M}, \mathbf{u}, \mathbf{z}$  could be matrices and vectors in any field and where each party contributes one element of  $\mathbf{x}$  as input. This NIMPC protocol can be re-imagined as a robust garbling for such functionalities, when  $\mathbf{M}, \mathbf{u}, \mathbf{z}$  are matrices and vectors over the Boolean field. Previously, Almashaqbeh et al. [2] leveraged a similar observation (of combining this NIMPC protocol with a regular garbled circuit) towards designing a garbling scheme that remains robust in the presence of an adversary who gets access to both labels for a fraction of the input-wires. However, there are some important differences between our definition and theirs; see Sect. 2.5 for a discussion).

The NIMPC protocol in [9] is presented in two phases – (1) an *offline pre-processing phase* that outputs private messages to each party and a broadcast message to all parties, and (2) an *online phase* where each party deterministically computes and broadcasts a single message based on its input and the private message output in the pre-processing phase. We observe that when working over a Boolean field, the broadcast message of the offline phase can be viewed as a garbling of the above function. Since there are only two-possible values for each element of the input vector  $\mathbf{x}$ , we can compute both possible messages corresponding to each element that the parties are expected to send in the online phase, and these may essentially act as the wire labels for the garbled circuit.

More concretely, this robust garbling works as follows: (1) sample a random matrix  $\mathbf{s} \xleftarrow{\$} \{0, 1\}^{k \times k}$  and compute  $\mathbf{s}'_i = \mathbf{s} \cdot \mathbf{M}_{\cdot, i}$  for each  $i \in [n]$ . (2) For input wire  $i \in [n]$ , the 0-label  $\text{label}_{i,0} \in \{0, 1\}^k$  is sampled randomly and the 1-label is computed as  $\text{label}_{i,1} = \text{label}_{i,0} \oplus \mathbf{s}'_i$ . (3) The garbled function is defined as  $\tilde{f} = \mathbf{z} \oplus \mathbf{s} \cdot \mathbf{u} \oplus \bigoplus_{i \in [n]} \text{label}_{i,0}$ . To evaluate, the receiver can simply exclusive-or all the appropriate labels with  $\tilde{f}$ . In Sect. 6.2, we show this construction satisfies the above notion of robust garbling, and that if  $\mathbf{x}$  satisfies  $\mathbf{u} = \mathbf{M}\mathbf{x}$ , then  $\mathbf{z} = \tilde{f} \oplus \bigoplus_{i \in [n]} \text{label}_{i, \mathbf{x}[i]}$ , otherwise, this evaluation will output random  $\mathbf{z}'$ .

### 2.3 Reducing Lockboxes Using Laconic OT

We now describe a generic method for achieving an asymptotic reduction in the total number of counter lockboxes by using laconic oblivious transfer (LOT) [18].

Recall that our previous construction requires a total of  $\mathcal{O}(n)$  lockboxes for  $n$ -bit inputs. Using LOT, we can reduce the number of lockboxes to be independent of the input size and only depend on the security parameter (as determined by the LOT scheme).

An LOT scheme allows a receiver to commit to a large input  $x \in \{0,1\}^n$  via a short *hash* whose size is a fixed polynomial in the security parameter. Subsequently, a sender with inputs  $(m_0, m_1)$  and an index  $i$  sends a short message to the receiver. Using this message, the receiver can recover  $m_{x[i]}$  but  $m_{1-x[i]}$  remains computationally hidden.<sup>4</sup> Moreover, the hash value can be reused by the sender to transmit different messages to the receiver, based on different choices of indices  $i$ .

At a high-level, we can use LOT to “compress” the effective input size, thereby achieving an asymptotic reduction in the number of lockboxes. More specifically, let  $C$  be a circuit with  $n$ -bit inputs. We can build a one-time program for  $C$  using the following two-step approach:

1. First, we compute an adaptively secure garbled circuit  $\tilde{C}$  for  $C$  together with a set of wire labels.
2. Now let **Send** be the next-message sender function in an LOT scheme. Let us consider  $n$  different copies ( $\text{Send}_1, \dots, \text{Send}_n$ ) of **Send**, where the  $i$ -th copy is hardwired with an index  $i \in [n]$  and a pair of labels  $(\text{lab}_i^0, \text{lab}_i^1)$ . Here,  $\text{lab}_i^b$  is the  $b$ -th label corresponding to the  $i$ -th input bit computed in the first step. Now, consider a new circuit **Send** that computes all of the functions  $\text{Send}_1, \dots, \text{Send}_n$  (in parallel). The input to this circuit is the LOT receiver message  $H$  – namely, the hash of an input  $x$  (to the original circuit  $C$ ). We now create a one-time program **OTP** for **Send** with  $\mathcal{O}(|H|)$  counter lockboxes using the scheme described in the previous sub-section. The final one-time program **OTP** for circuit  $C$  consists of **OTP** and the garbled circuit  $\tilde{C}$  computed in the first step.

To evaluate the one-time program **OTP** on an input  $x$ , a receiver first computes an LOT hash  $H$  of  $x$  and evaluates **OTP** on input  $H$ . Using the output values, it evaluates the garbled circuit  $\tilde{C}$  and returns its output.

It is easy to verify that the above construction achieves correctness. In order to prove security, we need to be able to *extract* the input of the receiver. However, from the security of **OTP**, we can only hope to extract the input to **OTP**, namely,  $H$ , which is presumably the LOT hash of some input  $x$ . In order to extract the actual  $x$ , we therefore require an LOT scheme that achieves simulation-based security against malicious receivers.

It is well known that such an LOT scheme cannot be constructed using standard black-box simulation techniques [21]. However, if we rely on random oracles or knowledge assumptions, then such a scheme can be constructed by compiling

<sup>4</sup> We emphasize that LOT is non-trivial even without privacy for receivers. While receiver privacy can be generically added [18], we do not require it for our transformation.

an LOT scheme with a succinct argument of knowledge (SNARK) [11, 44]. Due to space constraints we defer the formal description of our OTP construction using Laconic OT to the full version of the paper.

## 2.4 Counter Lockboxes with Multiple Password Attempts

Up to this point we have only considered counter lockboxes that allow for a *single* attempt to guess the password. For some real-world instantiations of counter lockboxes e.g. [39, 40], this may not be a valid assumption. We now discuss how our construction of leaky batch-OT can be adapted to support counter lockboxes that allow for *any number* of password attempts.

A natural approach is that the sender may simply “burn” all but one attempt from each lockbox they configure. However, this may be undesirable, especially in a cloud-based lockbox setting or if the sender does not wish to track the state of each lockbox. Therefore, we also provide a subtler approach described in this section and more fully examined in the full version of the paper.

Let  $z$  be the number of password attempts allowed by a counter lockbox functionality. We modify the previous construction as follows: once the sender decides that a particular lockbox should be a  $b$ -lockbox for a choice bit  $b$ , they do not simply set its password to  $b$ . Instead, they create  $z$  distinct strings  $\text{bin}(1)\|b, \dots, \text{bin}(z)\|b$  – each ending with bit  $b$ , where  $\text{bin}(i)$  denotes the binary representation of  $i$ . The sender then selects one of these  $z$  at random and sets it as the password for the counter lockbox.

For any choice bit  $b$ , an honest receiver can simply generate and try all of the  $z$  potential passwords for any lockbox. This guarantees that it can open all of the required lockboxes to reconstruct the desired label for its choice bit. On the other hand, the adversary gains no new advantage from having  $z$  attempts since there are  $2z$  potential password choices for any lockbox. In particular, an adversary can do no better in determining whether a lockbox is a  $b$ -lockbox than by “committing” to some  $b$  and trying  $b$  concatenated with each possible prefix string. We can therefore achieve the same parameters for the multiple password attempt case as in the single attempt case. Due to space constraints we defer a formal treatment of this topic to the full version of the paper.

## 2.5 Related Work

Chaum and Pederson [16] were the first to propose the use of tamper-proof hardware for cryptography purposes, and Goldreich and Ostrovsky [24] explored its application to software protection. Goldwasser, Kalai and Rothblum [26] introduced the notion of one-time programs as well as one-time memory tokens. Further improvements to their construction were investigated by Goyal et al. [30] and Bellare et al. [8]. More recently, Goyal and Goyal [28] investigated the use of blockchains to construct one-time programs.

Prior to our work, Almashaqbeh et al. [2] also leveraged the techniques from [9] to achieve a form of robustness in non-interactive secure computation using garbled circuits in a different context. There are some key differences between

|                                                                                                                                                                                                                                                                                                  |
|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Functionality $\mathcal{F}_f^{\text{OTP}}$                                                                                                                                                                                                                                                       |
| <b>Create:</b> Upon receiving <code>create, sid, <math>P_i, P_j, x</math></code> from $P_i$ where $x$ is a string do: <ol style="list-style-type: none"> <li>1. Send <code>create, sid, <math>P_i, P_j</math></code> to <math>P_j</math>.</li> <li>2. Store <math>P_i, P_j, x</math>.</li> </ol> |
| <b>Execute:</b> On receiving <code>run, sid, <math>P_i, y</math></code> from party $P_j$ , find the stored tuple $P_i, P_j, x$ if no such tuple exists, do nothing.) Send <code><math>f(x, y)</math></code> to $P_j$ and delete tuple $P_i, P_j, x$ .                                            |

**Fig 1** Ideal functionality for a one-time program (OTP), parameterized with a specific function  $f$ , quoted from [30].

our work and theirs: we provide a general definition of robust garbling that accounts for the challenges involved in determining the adversary’s input (and output) in our setting involving “leakage”. In particular, as discussed earlier, since it is unclear how to define robust garbling for general functions, we define a class of admissible functions and robust garbling for such functions (as discussed in Sect. 2.2). In contrast, their definitions assume an a priori fixed input (and output) for the adversary, and are not applicable to our setting. Further, our definitions (unlike theirs) account for *fully adaptive* adversaries, which is crucial to our setting where the adversary can query the lockboxes in arbitrary order.

### 3 Preliminaries

We include preliminary definitions and discussion for *computational indistinguishability*, the *UC-Framework*, *adaptive projective garbling schemes*, *linear error-correcting codes* and *succinct non-interactive arguments of knowledge (SNARKs)* in the full version of the paper.

#### 3.1 One-Time Programs

One-time Programs (OTP) were introduced by [26]. At a high level, a one-time program for a function  $f$  enables a party to evaluate  $f$  on any one input of its choice. The security of a one-time program dictates that no efficient adversary should be able to learn anything about the function  $f$ , beyond what can be inferred from its output  $f(x)$  on any one input  $x$  of its choice.

Similar to Goyal et al. [30], we model one-time programs as a two-party non-interactive protocol that is secure against malicious receivers. We define the ideal functionality for a one-time program in Fig. 1.

### 4 Counter Lockboxes

In this section, we formalize our notion of *counter lockboxes*. A *counter lockbox*, or just “lockbox,” is a stateful abstraction for securely storing cryptographic secrets such that they are protected by a human-memorable password. To create a new

lockbox, a requester provides a password  $P$  and a maximum attempt counter  $A$ . The lockbox then generates random value  $K$  and returns  $K$  to the requester. The lockbox also stores internally  $A$ , some data with which it can re-compute  $K$  given  $P$ , and some information it can use to check if a future password guess matches  $P$ .

At a later point, a requester can provide some password  $P'$  to the lockbox, which will use its internal state to check if  $P'$  produces a match. If so, the lockbox recomputes and returns  $K$  to the requester. If the password does not produce a match, the lockbox decrements  $A$ . After  $A$  incorrect guesses the lockbox completely erases its internal content, preventing the value of  $K$  from ever being retrieved.

We model the lockbox functionality as  $\mathcal{F}_\lambda^{\text{Lockbox}}$ , described in Fig. 2. In this work, we study cryptography in the  $\mathcal{F}_\lambda^{\text{Lockbox}}$ -hybrid model.

| Functionality $\mathcal{F}^{\text{Lockbox}}$                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                             |
|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <p><b>Create:</b> On input <code>create, <math>P_i, P_j, sid, id, P, A</math></code> from party <math>P_i</math> where <math>A &gt; 0</math>, send <code>create, <math>P_i, P_j, sid, id</math></code> to <math>P_j</math>. Sample <math>K \in \{0, 1\}^*</math>, store the tuple <math>(P_i, P_j, sid, id, P, A, K, 0)</math>, and send <math>K</math> to <math>P_i</math>.</p> <p><b>Open:</b> On input <code>open, <math>P_i, sid, id, P</math></code> from party <math>P_j</math>:</p> <ul style="list-style-type: none"> <li>– If a tuple <math>(P_i, P_j, sid, id, P, A, K, N)</math> does not exist, then do nothing.</li> <li>– If <math>N = A</math> then delete the tuple and return <code>expired</code>.</li> <li>– Otherwise if <math>P = P</math> then delete and replace the tuple with <math>(P_i, P_j, sid, id, P, A, K, 0)</math> and return <math>K</math>.</li> <li>– If <math>P \neq P</math> then delete and replace the tuple with <math>(P_i, P_j, sid, id, P, A, K, N + 1)</math> and return <code>bad-guess</code>.</li> </ul> |

**Fig 2** Ideal functionality for a counter lockbox. This simplified interface assumes that the lockbox “secret”  $K$  is a random string of length  $\kappa$  and that the password guess is directly compared to a stored password.

**On the Communication Model.** Previous works using secure hardware tokens [26] assume a two-party model in which a *sender* provisions stateful tokens and sends them to the *receiver*, who then uses them to evaluate a one-time program. This model can be directly adapted to cloud-based lockbox functionalities by simply forwarding references to the appropriate online locations. Lockboxes on a fixed device, however, may require adapted usage. For example, unlike hardware tokens, lockboxes implemented within the Apple SEP are an integral component of the device and cannot easily be removed or replaced. Hence our results can rely on the following different usage scenarios:

1. In a cloud-based scenario, the sender provisions a series of lockboxes on a shared (accessible to both parties) server, such as an Apple Cloud Key Vault [39] HSM or Google Titan [48] HSM in a remote data center. The sender then provides the location (IP address or URL) of these lockboxes along with some auxiliary data to the receiver. The receiver accesses these lockboxes to evaluate the one-time program.

2. In a device-based scenario, the sender provisions a device (such as an Apple iOS device with a SEP) with lockboxes and then physically delivers the device to the receiver. Given the physical security of the SEP [3], these lockboxes are designed to resist device forensics. Auxiliary data can be transmitted within the regular device memory, and evaluation could even be facilitated by custom on-device software such as an iOS app if deemed acceptable to the evaluator.
3. In a further device-centric instantiation, the sender and receiver may not be physically co-located. To provision lockboxes on the receiver’s secure hardware, the sender employs a cryptographic protocol that enables secure message transmission to the receiver’s secure hardware, while entirely bypassing the receiver’s ability to observe this provisioning. For example, Apple’s SEP supports a cryptographic protocol for communications between the SEP and application processor within a single device. With appropriate key management, this could be repurposed to allow a remote party to communicate securely with a receiver’s SEP.

In all three settings, we assume that the hardware itself is secure against logical and physical attacks: this means that the only way to access lockbox secrets is through the password interface the hardware exposes. By contrast we assume that, at least at program execution time, the receiver has full control of the remaining portions of the device processor and can query the lockbox interface arbitrarily.

**Discussion.** In all prior hardware-token models, the sender physically transmits the device to the receiver and it is assumed that there is no “backward communication channel” to the sender. Indeed, such a channel can lead to privacy loss for the receiver.

However, one could consider a stronger model, where the sender does in fact have the ability to inspect lockboxes after the receiver is done querying them. In such a model, to prevent the sender from learning receiver’s input bits, it is important to ensure that the following three states of lockboxes remain indistinguishable – (1) lockboxes with leftover password attempts, (2) lockboxes that were “destroyed” because of failed password attempts and (3) ones that are still presumably “functional” because they were opened using the correct password. For the first kind, we can use a simple defense and ask the receiver to consume all password attempts on each.

For the remaining two forms, our ideal lockbox functionality implicitly assumes that an adversary cannot distinguish between hardware that outputs the secret and one where the secret was destroyed because of failed password attempts. In the above stronger model, hardware that matches this ideal functionality clearly will not “leak” extra information once its attempts have been expired. It simply outputs  $\perp$ , and there is no way to distinguish between “expired during evaluation without producing a secret” and “did output the secret but expired later as a defensive cleanup measure.” While in general, real hardware may not behave like an ideal function, our definition of this ideal functionality is inspired by precise technical specifications from vendors such as Apple (see e.g. Apple iOS Security Guide), and there seems to be strong evidence that hardware

| Functionality $\mathcal{F}_{(n,\gamma)}^{\text{OT}}$                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                              |
|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <p><b>Initialize:</b> Upon receiving <code>init, sid, id, sen, rec, <math>\{m_{i,0}, m_{i,1}\}_{i \in [n]}</math></code> from the sender <code>sen</code>, where <math>\{m_{i,0}, m_{i,1}\}_{i \in [n]} \in \mathcal{M}^{2n}</math>, send <code>init, sid, id, sen, rec</code> to the receiver <code>rec</code> and store the tuple <math>(sid, id, sen, rec, \{m_{i,0}, m_{i,1}\}_{i \in [n]}, \mathcal{S}_1, \mathcal{S}_2, \text{counter})</math>, where <math>\mathcal{S}_1 = \mathcal{S}_2 = \emptyset</math> and <math>\text{counter} = 0</math>.</p> <p><b>Open:</b> Upon receiving <code>open, sid, id, sen, rec, i, b, choice</code> from party <code>rec</code>, where <math>\text{choice} \in \{\text{both}, \text{single}\}</math>, find the stored tuple <math>(sid, id, sen, rec, \{m_{i,0}, m_{i,1}\}_{i \in [n]}, \mathcal{S}_1, \mathcal{S}_2, \text{counter})</math> if no such tuple exists, do nothing).</p> <ul style="list-style-type: none"> <li>– If <math>i \in \mathcal{S}_1</math>, do nothing.</li> <li>– Else if <math>i \in \mathcal{S}_2</math>, send <math>m_{i,b}</math> to <code>rec</code>.</li> <li>– Else, do the following: <ul style="list-style-type: none"> <li>• If <math>\text{choice} = \text{single}</math>, send <math>m_{i,b}</math> to <code>rec</code>, then delete and replace the tuple with <math>(sid, id, sen, rec, \{m_{i,0}, m_{i,1}\}_{i \in [n]}, \mathcal{S}_1 \cup \{i\}, \mathcal{S}_2, \text{counter})</math></li> <li>• else, if <math>\text{choice} = \text{both}</math> and <math>\text{counter} = \text{return forbidden}</math>. Else if <math>\text{counter} &lt; \gamma</math> send <math>m_{i,b}</math> to <code>rec</code>, then delete and replace the tuple with <math>(sid, id, sen, rec, \{m_{i,0}, m_{i,1}\}_{i \in [n]}, \mathcal{S}_1, \mathcal{S}_2 \cup \{i\}, \text{counter} + 1)</math>.</li> </ul> </li> </ul> |

Fig 3 Ideal functionality for leaky batch-OT

will satisfy it. As a result, our constructions remain secure in this stronger model as long as the hardware behaves similarly to the ideal functionality.

## 5 Leaky Batch-OT

In this section, we present and formalize a notion of *leaky batch-OT* and show how it can be realised using counter lockboxes.

### 5.1 Definition

Leaky batch oblivious transfer is a two-party functionality between a sender and receiver, where the sender initially inputs  $n$  pairs of messages  $\{(m_{i,0}, m_{i,1})\}_{i \in [n]}$  where each  $m_{i,b}$  is in some message domain  $\mathcal{M}$ . For each  $i \in [n]$ , the receiver inputs a single bit  $b \in \{0, 1\}$  and obtains  $m_{i,b}$ . Additionally, at most  $\gamma$  times, the receiver is allowed to input  $i$  and obtain  $m_{i,1-b}$ , assuming they have previously received  $m_{i,b}$ . Our specific formulation is more nuanced. We give a formal definition of this *reactive* functionality in Fig. 3.

### 5.2 Construction

In this section, we construct a protocol for leaky batch-OT using counter lockboxes. Recall that our definition of leaky batch-OT only allows the receiver to obtain both messages for at most  $\gamma$  indices  $i \in [n]$ . Therefore, we show that if  $\ell$  is set to  $\lceil -\log_2(\frac{\gamma}{n}) \rceil + 1$  then except with some negligible probability in  $n$  a malicious receiver can successfully obtain keys of all  $2\ell$  lockboxes for at most  $\gamma$  indices  $i \in [n]$ . We give a formal description of this protocol in the  $\mathcal{F}_{\lambda}^{\text{Lockbox-}}$  hybrid model in Fig. 4. Due to space constraints we defer the formal proof of security to the full version of the paper and include a proof sketch here.

**Theorem 3.** *There exists a protocol for securely realizing the leaky batch-OT functionality  $\mathcal{F}_{(n,\gamma)}^{\text{OT}}$  (Fig. 3) against a malicious sender and receiver, in the  $\mathcal{F}_{\lambda}^{\text{Lockbox}}$ -hybrid model, where the sender only sends a single message to receiver, while the receiver does not to send any messages to the sender.*

**Proof Sketch.** We first present a simulator that in the ideal world simulates sending lockboxes to the adversary as in the protocol. We then show that as long as the adversary is not able to successfully open all  $2\ell$  lockboxes associated with more than  $\gamma$  input wires, with BAD denoting the event that the adversary succeeds in doing so, the transcript output by the simulator is indistinguishable from that computed in the real world.

We then proceed to show that the probability that the BAD event happens is negligible when  $\ell = \lceil -\log_2(\frac{\gamma}{n}) \rceil + 1$ . For this, we first show that the probability of the adversary successfully opening all  $2\ell$  lockboxes for one wire is at most  $p = (\frac{1}{2}) \cdot \frac{1}{(2-\gamma)}$ . As intuition, the adversary can do no better when guessing passwords than just guessing whatever password is in the majority among the remaining lockboxes. For  $\ell$  of the guesses this gives them a 50% chance of success, producing the  $(\frac{1}{2})$  term. The second term follows from a similar but more involved calculation. Following our derivation of  $p$ , we use a Chernoff bound to show that the overall probability is negligible in  $n$ . We then conclude that when  $n$  is large, i.e.  $O(\frac{1}{p})$ , the protocol is secure.

- **Sender:** Let  $\ell := \lceil -\log_2(\frac{\gamma}{n}) \rceil + 1$ . Given inputs  $\{m_{i,0}, m_{i,1}\}_{i \in [n]}$ , the sender `sen` samples a fresh `sid`. For each  $i \in [n]$ , do the following:
  1. Sample a random permutation  $\pi_i : [2\ell] \rightarrow [2\ell]$ .
  2. Sample  $2\ell$  unique ids  $\{id_{i,j}\}_{j \in [2\ell]}$ .
  3. For each  $j \in [2\ell]$ ,
    - If  $\pi_i(j) \leq \ell$ , invoke  $\mathcal{F}^{\text{Lockbox}}$  on arguments `create, sen, rec, sid, idi,πi(j), 0, 1` and obtain  $K_{i,0}^{\pi_i(j)}$  in return.
    - Else, invoke  $\mathcal{F}^{\text{Lockbox}}$  on arguments `create, sen, rec, sid, idi,πi(j), 1, 1` and get  $K_{i,1}^{\pi_i(j)}$  in return.
  4. Compute  $C_{i,0} := m_{i,0} \oplus \bigoplus_{j=1}^{\ell} K_{i,0}^j$ .
  5. Compute  $C_{i,1} := m_{i,1} \oplus \bigoplus_{j=\ell+1}^{2\ell} K_{i,1}^j$ .
  6. Send  $\{C_{i,0}, C_{i,1}\}_{i \in [n]}$  to the receiver `rec`.
- **Receiver.** Given a set of input bits  $\{b_i\}_{i \in [n]}$  and upon receiving  $\{(sid, id_{i,\pi_i(j)}, \text{sen}, \text{rec})\}_{j \in [2\ell], i \in [n]}$  from the  $\mathcal{F}^{\text{Lockbox}}$  functionalities and  $\{C_{i,0}, C_{i,1}\}_{i \in [n]}$  from the sender, the receiver proceeds as follows for each  $i \in [n]$ :
  1. For each  $j \in [2\ell]$ , invoke  $\mathcal{F}^{\text{Lockbox}}$  on arguments `open, sen, sid, idi,πi(j), bi` to receive either  $K_{i,b_i}^{\pi_i(j)}$  or `bad_guess`, in which case set  $K_{i,b_i}^{\pi_i(j)} = 0$ .
  2. Compute  $m_{i,b_i} = C_{i,b_i} \oplus \bigoplus_{j=1}^{2\ell} K_{i,b_i}^{\pi_i(j)}$ .

**Fig 4** Protocol for leaky batch OT

## 6 Robust Garbling

In this section, we formalize the notion of robust garbling for a class of admissible functions. We then present a robust garbling scheme for a sub-class of such functions, with fully adaptive, information-theoretic security.

### 6.1 Definitions

In a robust garbling scheme, we want to capture the requirement that even if the receiver obtains both labels for some of the input wires, it should only be able to learn exactly one output. However, this poses the following conundrum: on the one hand, we are allowing the receiver to obtain labels for *multiple* inputs. On the other hand, we do not want it to learn more than *one* output. How do we reconcile these requirements?

While achieving a reconciliation seems impossible for general functions, we can hope to do so for functions where the inputs have some level of *redundancy*. In other words, if only a subset of the input bits are sufficient to determine the output of the function, we can hope to construct a garbling scheme where even if the receiver learns multiple labels for the remaining bits, it will only learn at most one uniquely defined output.

We now give a formal definition of such a class of functions.

**Definition 1 Function Class  $\mathcal{F}^{n,\gamma}$ .**  $\mathcal{F}^{n,\gamma}$  contains all functions  $f : \{0,1\}^n \rightarrow \{0,1\}^* \cup \{\perp\}$  such that for any set  $\mathcal{S} \subset [n]$  of size  $(n - \gamma)$  and any set of bits  $\{x_i\}_{i \in \mathcal{S}}$ , there exists at most one “valid”  $\{x_i\}_{i \in \bar{\mathcal{S}}}$  such that  $f(x_1, \dots, x_n) \neq \perp$ .

Further, there is an associated function  $\text{Expand} : \{0,1\}^{(n-\gamma)} \rightarrow \{0,1\}^n$  such that for every  $\{x_i\}_{i \in \mathcal{S}}$ :

1. If  $\exists \{x_i\}_{i \in \bar{\mathcal{S}}}$ , such that  $f(x_1, \dots, x_n) \neq \perp$ , then  $\text{Expand}(\{x_i\}_{i \in \mathcal{S}}) = (x_1, \dots, x_n)$ .
2. Else,  $f(\text{Expand}(\{x_i\}_{i \in \mathcal{S}})) = \perp$ .

At a high level, the above definition implies that it is possible to determine the unique output associated with any  $(n - \gamma)$  bits of input.

Next, we formalize the notion of *robust garbling* for this class of functions. In addition to the robustness property discussed above, we also want this garbling scheme to be “fully adaptive”. That is, upon receiving the garbled circuit, the adversary should be allowed to choose its input bit-by-bit, depending on the labels received thus far. We note that this is stronger than the standard notion of adaptivity for garbled circuits [7, 8, 23], where the adversary must specify its *entire input* in one go, after receiving the garbled circuit.

Moreover, as discussed previously, we allow the adversary to receive both labels for some of the input wires. However, in case it plans to obtain the second label for any index, it must specify that at the time of making the first query for that index. This way, once the adversary has received at least one label for each input position, the simulator can determine the output based on the ones

for which the adversary is guaranteed to not make a second query and simulate accordingly. Therefore, we model our simulator for robust garbling to essentially consist of three algorithms ( $\text{SimFunc}$ ,  $\text{SimIn}$ ,  $\text{SimInLast}$ ), where  $\text{SimFunc}$  simulates the garbled circuit using only “public-information” about the circuit (e.g., the size of the circuit).  $\text{SimIn}$  and  $\text{SimInLast}$  are used for simulating the input wire labels, where  $\text{SimInLast}$  is used specifically once the adversary has obtained at least one label for each input wire.

We now present a definition of robust garbling.

**Definition 2 Robust Garbling)** *A robust garbling scheme for functions  $f \in \mathcal{F}^{n,\gamma}$  consists of a tuple of PPT algorithms ( $\text{RobGarble}$ ,  $\text{RobGarbleInp}$ ,  $\text{RobEval}$ ) such that:*

- $(\tilde{f}, \text{st}) \leftarrow \text{RobGarble}(1^\lambda, f)$ : *This is a PPT algorithm that takes as input the security parameter  $1^\lambda$  and a function  $f \in \mathcal{F}^{n,\gamma}$  and outputs a garbling  $\tilde{f}$  and some private state information  $\text{st}$ .*
- $\text{lab}_{i,x_i} \leftarrow \text{RobGarbleInp}(\text{st}, i, x_i)$ : *This is a PPT algorithm that takes as input the state information  $\text{st}$ , an index  $i \in [n]$  and an input bit  $x_i$ , and outputs the corresponding input label  $\text{lab}_{i,x_i}$ .*
- $y = \text{RobEval}(\tilde{f}, \{\text{lab}_{i,x_i}\}_{i \in [n]})$ : *Given a garbling  $\tilde{f}$  and a set of labels  $\{\text{lab}_{i,x_i}\}_{i \in [n]}$  it outputs a value  $y \in \{0, 1\}^k$ .*

**Correctness.** *For every  $n \in \mathbb{N}$ ,  $f \in \mathcal{F}^{n,\gamma}$ , and for each  $\mathbf{x} \in \{0, 1\}^n$ , it holds that:  $\Pr[\text{RobEval}(\tilde{f}, \{\text{lab}_{i,x_i}\}_{i \in [n]}) = f(\mathbf{x})] = 1$ , where  $(\tilde{f}, \text{st}) \leftarrow \text{RobGarble}(1^\lambda, f)$  and  $\forall i \in [n]$ ,  $\text{lab}_{i,x_i} \leftarrow \text{RobGarbleInp}(\text{st}, i, x_i)$ .*

**$\gamma$ -Robust Adaptive Security.** *There exists a PPT simulator  $\text{Sim} = (\text{SimFunc}, \text{SimIn}, \text{SimInLast})$  such that, for any non-uniform PPT adversary  $\mathcal{A}$  there exists a negligible function  $v$  such that:*

$$|\Pr[\text{Exp}_{\mathcal{A}, \text{GC}, \text{Sim}}^{\text{Rob dp}}(1^\lambda, 0) = 1] - \Pr[\text{Exp}_{\mathcal{A}, \text{GC}, \text{Sim}}^{\text{Rob dp}}(1^\lambda, 1) = 1]| \leq v(\cdot)$$

where the experiment  $\text{Exp}_{\mathcal{A}, \text{GC}, \text{Sim}}^{\text{Rob dp}}$  is defined in Fig. 5

## 6.2 Construction

In this section, we present an information-theoretically secure construction of robust garbling for functions of the form  $f = (\mathbf{M}, \mathbf{u}, \mathbf{z}) \in \mathcal{F}^{n,\gamma}$ , where  $\mathbf{M} \in \{0, 1\}^{k \times n}$ ,  $\mathbf{u} \in \{0, 1\}^k$  are public and  $\mathbf{z} \in \{0, 1\}^k$  is private, such that on any

$$\text{input } \mathbf{x} \in \{0, 1\}^n, f(\mathbf{x}) = \begin{cases} \mathbf{z} & \text{if } \mathbf{u} = \mathbf{M}\mathbf{x} \\ \mathbf{z}' \xleftarrow{\$} \{0, 1\}^k & \text{otherwise} \end{cases}.$$

We use  $\mathcal{F}_{\text{linear}}^{n,\gamma}$  to denote this subclass of  $\mathcal{F}^{n,\gamma}$ . While all *invalid* inputs must to lead to a  $\perp$  output in any  $f \in \mathcal{F}^{n,\gamma}$ , functions in  $\mathcal{F}_{\text{linear}}^{n,\gamma}$  instead output a random  $\mathbf{z}'$ . We note that depending on the context, this may not be a problem (and the above function can still be admissible), if the receiver can distinguish a *valid* output  $\mathbf{z}$  from an *invalid* random  $\mathbf{z}'$  potentially using some “additional

| Experiment                                                                                                                                                                                                                                           | $\text{Exp}_{\mathcal{A}, \text{GC}, \text{Sim}}^{\text{Rob dp}}$                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                           |
|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| 1. The adversary specifies a function $f \in \mathcal{F}^{n, \gamma}$ and obtains $\tilde{f}$ , where $\tilde{f}$ is created as follows:                                                                                                             | <ul style="list-style-type: none"> <li>– If <math>b = 0</math>: <math>(\tilde{f}, \text{st}) \leftarrow \text{RobGarble}(1^\lambda, f)</math></li> <li>– If <math>b = 1</math>: <math>(\tilde{f}, \text{st}) \leftarrow \text{SimFunc}(1^\lambda, 1^{ f })</math> Here, we implicitly assume that this simulator can get any public information about <math>f</math>, not just its size.)</li> </ul>                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                        |
| 2. Initialize $\mathcal{S}_1 = \mathcal{S}_2 = \emptyset$ and counter = 0. For each $j \in [n + 1]$ , the adversary $\mathcal{A}$ specifies a tuple $(i_j, x_{i_j}, \text{choice}_i)$ , where $\text{choice}_i \in \{\text{single}, \text{both}\}$ . | <ul style="list-style-type: none"> <li>– If <math>\text{choice}_i = \text{single}</math> and <math>i_j \notin \mathcal{S}_1</math>, update <math>\mathcal{S}_1 = \mathcal{S}_1 \cup \{i_j, x_{i_j}\}</math>. Else if <math>\text{choice}_i = \text{both}</math>, <math>i_j \notin \mathcal{S}_1 \cup \mathcal{S}_2</math> and counter &lt; <math>\gamma</math>, update <math>\mathcal{S}_2 = \mathcal{S}_2 \cup \{i_j\}</math> and set counter = counter + 1. In both cases do the following: <ul style="list-style-type: none"> <li>• If <math>b = 0</math>, output <math>\text{lab}_{i_j, x_{i_j}} \leftarrow \text{RobGarbleInp}(\text{st}, i_j, x_{i_j})</math>.</li> <li>• If <math>b = 1</math> and <math> \mathcal{S}_1 \cup \mathcal{S}_2  &lt; n</math>, output <math>\text{lab}_{i_j, x_{i_j}} \leftarrow \text{SimIn}(\text{st}, i_j, x_{i_j})</math>.</li> <li>• If <math>b = 1</math> and <math> \mathcal{S}_1 \cup \mathcal{S}_2  = n</math>, output <math>\text{lab}_{i_j, x_{i_j}} \leftarrow \text{SimInLast}(\text{st}, i_j, x_{i_j}, \mathcal{S}, \text{out})</math>, where <math>\mathcal{S} \subset [n]</math> is the set of indices <math>i \in [n]</math> such that <math>(i, \cdot) \in \mathcal{S}_1</math> and <math>\text{out} = f _{\text{expand}(\{x_i\}_{i \in \mathcal{S}})}</math>.</li> </ul> </li> <li>– Else if <math>\text{choice}_i = \text{both}</math>, <math>i_j \in \mathcal{S}_1</math> and <math>i_j \in \mathcal{S}_2</math>, do the following. <ul style="list-style-type: none"> <li>• If <math>b = 0</math>, output <math>\text{lab}_{i_j, x_{i_j}} \leftarrow \text{RobGarbleInp}(\text{st}, i_j, x_{i_j})</math>.</li> <li>• If <math>b = 1</math>, output <math>\text{lab}_{i_j, x_{i_j}} \leftarrow \text{SimIn}(\text{st}, i_j, x_{i_j})</math>.</li> </ul> </li> </ul> |

Finally, the adversary outputs a bit  $b$ , which is the output of the experiment.

Fig 5  $\gamma$ -robust adaptivity experiment

information.” In our OTP construction, the value  $\mathbf{z}$  will correspond to labels of the garbled circuit that garbles the actual function for which we compute the OTP. While these labels are also random vectors in  $\{0, 1\}^k$ , the receiver gets “additional information” in the form of the garbled circuit where  $\mathbf{z}$  is used as an input wire label. In case the output of the above function is a random unrelated value instead of a valid label, while evaluating, the receiver will be able to detect this and demarcate this output as essentially equivalent to  $\perp$ .

**Garbling Scheme.** We now present a construction of robust garbling scheme for the above class of functions. As discussed previously, this is adapted from the non-interactive *multi-party* computation (NIMPC) protocol for such functions proposed by Benhamouda et al. [9].

- $\text{RobGarble}(1^\lambda, f)$ :
  1. Sample a random  $\mathbf{s} \xleftarrow{\$} \{0, 1\}^{k \times k}$ .
  2. For each  $i \in [n]$ , sample a random  $\mathbf{r}_i \in \{0, 1\}^k$ .
  3. Set  $\text{st} = \mathbf{s}, \{\mathbf{r}_i\}_{i \in [n]}$ .
  4. Output garbling  $\tilde{f} = \mathbf{z} = \mathbf{s} \cdot \mathbf{u} \quad \mathbf{u} = \sum_{i \in [n]} \mathbf{r}_i$ .
- $\text{RobGarbleInp}(\text{st}, \mathcal{I}, \{x_i\}_{i \in [n] \setminus \mathcal{I}})$ :
  1. Parse  $\text{st} = \mathbf{s}, \{\mathbf{r}_i\}_{i \in [n]}$ .
  2. For each  $i \in [n]$ , compute  $\mathbf{s}'_i = \mathbf{s} \cdot \mathbf{M}_{\cdot, i}$ , where  $\mathbf{M}_{\cdot, i}$  denotes the  $i$ -th column of  $\mathbf{M}$ .
  3. For each  $i \in \mathcal{I}$ , compute and output  $\text{lab}_{i, 0} = \mathbf{r}_i$  and  $\text{lab}_{i, 1} = \mathbf{r}_i \oplus \mathbf{s}'_i$ .
  4. For each  $i \in [n] \setminus \mathcal{I}$ , output  $\text{lab}_{i, x_i} = \mathbf{r}_i \oplus \mathbf{s}'_i \cdot x_i$ .
- $\text{RobEval}(\tilde{f}, \{\text{lab}_{i, x_i}\}_{i \in [n]})$ : Compute and output  $\tilde{f}(\sum_{i \in [n]} \text{lab}_{i, x_i})$ .

We prove the following theorem in the full-version of our paper.

**Theorem 4** *There exists an information-theoretically secure robust adaptive garbling scheme for each every function  $f \in \mathcal{F}_{\text{linear}}^{n,\gamma}$ .*

## 7 One-Time Program

In this section we use the tools built in previous sections to construct a one-time program. In addition to leaky batch-OT and robust garbling for  $\mathcal{F}_{\text{linear}}^{n,\gamma}$ , we make use of a standard adaptive, projective garbled circuit and linear error-correcting codes over  $\mathbb{F}_2$ .

We instantiate our one-time program construction using a  $[n, k, \gamma+1]_2$ -binary linear error-correcting code, where  $k$  is the message length,  $n$  is the code-word length, and  $\gamma+1$  is the distance. We give a formal description of this construction in the  $\mathcal{F}_{(n,\gamma)}^{\text{OT}}$ -hybrid model. While an honest receiver does not use the “leaky” aspect of our leaky batch-OT to receive both  $(\text{lab}'_{j,0}, \text{lab}'_{j,1})$  for any index  $j$ , a malicious receiver can certainly try to exploit it. However, since the number of “double-labels” that they can obtain is capped at  $\gamma$  (and our robust garbling is secure as long as double-labels for at most  $\gamma$  input wires are revealed), they will never receive enough to successfully obtain both labels for any input wire of the adaptive garbled circuit. As a result, even a malicious receiver will only be able to learn the output for a single input.

**Protocol.** We give a formal description of the OTP protocol in Fig. 6, in the  $\mathcal{F}_{(n,\gamma)}^{\text{OT}}$ -hybrid model, using  $[n, k, \gamma+1]$ -binary linear error-correcting codes, an adaptive projective garbled circuit ( `daGarbleCkt`, `daGarbleInp`, `daEvalCkt`) and a robust function garbling scheme (`RobGarble`, `RobGarbleInp`, `RobEval`) for  $\mathcal{F}_{\text{linear}}^{n,\gamma}$ .

Note that for all  $i \in [k]$  and  $b \in \{0,1\}$ , the function  $F_{i,b}$  belongs to the  $\mathcal{F}_{\text{linear}}^{n,\gamma}$  class of functions described in Sect. 6.2. It is easy to identify when the output is  $\perp$ , as the output of each function is an input wire label for a garbled circuit. The use of an error-correcting code grants the properties required by  $\mathcal{F}_{\text{linear}}^{n,\gamma}$ . By the definition of minimum distance, any set of  $(n - \gamma)$  fixed bits will define only a single valid codeword, and the `Expand` function is simply a lookup for the codeword uniquely defined by those bits. Finally, each  $F_{i,b}$  clearly meets the linear construction requirement of  $\mathcal{F}_{\text{linear}}^{n,\gamma}$ . We prove the following theorem in the full-version of our paper.

**Theorem 5** *Assuming the existence of one-way functions, there exists a non-interactive protocol for securely realizing  $\mathcal{F}_f^{\text{OTP}}$  against a semi-honest sender and malicious receiver in the  $\mathcal{F}_{(n,\gamma)}^{\text{OT}}$ -hybrid model.*

## 8 Concrete Analysis

In this section, we present a concrete analysis to investigate the suitability of our schemes for real-world applications. In Sect. 8.1, we estimate the number of lockboxes required for different input lengths. In Sect. 8.2, we discuss how lockboxes

- **Sender:** Given an input  $f$ , the sender **sen** proceeds as follows:
  1. Express  $f$  as a circuit  $C$ , then compute  $\tilde{C}, \{\text{lab}_{i,b}\}_{i \in [k], b \in \{0,1\}} \leftarrow \text{daGarbleCkt } 1, C$ .
  2. Instantiate a linear error-correcting code with length  $n$ , rank  $k$ , minimum distance  $\geq 1$  and generating matrix  $\mathbf{G}$ .
  3. For each  $i \in [k]$ , and each  $b \in \{0,1\}$ , compute a matrix  $\mathbf{M}_{i,b}$  and vector  $\mathbf{u}_{i,b}$  such that  $\mathbf{u}_{i,b} = \mathbf{M}_{i,b} \cdot \mathbf{y}$  if and only if  $\mathbf{y}$  is a valid codeword generated using  $\mathbf{G}$  and its corresponding word has bit  $b$  at position  $i$ , i.e.  $\mathbf{u}_{i,b} = \mathbf{M}_{i,b} \cdot \mathbf{y} \iff \exists \mathbf{x} \in \{0,1\}^k, \mathbf{y}^\top = \mathbf{x}^\top \cdot \mathbf{G} \wedge \mathbf{x}_i = b$ . Then, define the following function:
 
$$F_{i,b}(\mathbf{y}) = \begin{cases} \text{lab}_{i,b} & \text{if } \mathbf{u}_{i,b} = \mathbf{M}_{i,b} \cdot \mathbf{y} \\ \mathbf{z} \leftarrow \{0,1\}^k & \text{otherwise} \end{cases}$$
- Next, compute  $\tilde{F}_{i,b}, \{\text{robustLab}_{j,b}^{i,b}\}_{j \in [n], b \in \{0,1\}} \leftarrow \text{RobGarble } 1, F_{i,b}$ .
- 4. Define  $\overrightarrow{\text{robustLab}}_{j,b} := \{\text{robustLab}_{j,b}^{i,b}\}_{i \in [k], b \in \{0,1\}}$  for all  $j \in [n], b \in \{0,1\}$ .
- 5. Sample a fresh  $\text{sid}$  and  $\text{id}$  and invoke  $\mathcal{F}_{(n,\gamma)}^{\text{OT}}$  on arguments  $\text{init}, \text{sid}, \text{id}, \text{send}, \text{rec}, \{\overrightarrow{\text{robustLab}}_{j,0}, \overrightarrow{\text{robustLab}}_{j,1}\}_{j \in [n], b \in \{0,1\}}$ .
- 6. Send  $\tilde{C}, \{\tilde{F}_{i,b}\}_{i \in [k], b \in \{0,1\}}$  to the receiver **rec**.
- **Receiver:** Given an input  $\mathbf{x}$  and upon receiving  $\text{sid}, \text{id}, \text{sen}, \text{rec}$  from  $\mathcal{F}_{(n,\gamma)}^{\text{OT}}$  and  $\tilde{C}, \{\tilde{F}_{i,b}\}_{i \in [k], b \in \{0,1\}}$  from the sender, the receiver proceeds as follows:
  1. Compute  $\mathbf{y} := \mathbf{x}^\top \cdot \mathbf{G}$ .
  2. For each  $j \in [n]$ , invoke  $\mathcal{F}_{(n,\gamma)}^{\text{OT}}$  on arguments  $\text{open}, \text{sid}, \text{id}, \text{sen}, \text{rec}, j, \mathbf{y}[j]$  and get  $\overrightarrow{\text{robustLab}}_{j,\mathbf{y}[j]} = \{\text{robustLab}_{j,\mathbf{y}[j]}^{i,b}\}_{i \in [k], b \in \{0,1\}}$  in return.
  3. For each  $i \in [k]$ , compute  $\text{lab}_{i,x_i} = \text{RobEval } \tilde{F}_{i, [i]}, \overrightarrow{\text{robustLab}}_{j,\mathbf{y}[j]}^{i, [i]}_{j \in [n]}$ .
  4. Compute and output  $\mathbf{z} \leftarrow \text{daEvalCkt } \tilde{C}, \{\text{lab}_{i, [i]}\}_{i \in [k]}$ .

Fig 6 OTP protocol

can be instantiated using commodity hardware and the associated costs and finally in Sect. 8.3, we discuss some potential applications of our construction.

### 8.1 Number of Lockboxes

We use lockboxes to implement the leaky batch-OT functionality and the input to this functionality is an encoding of the “real” input of the receiver. For encoding, we require linear binary ECC with a constant rate. More often than not, finding optimal binary ECC for specific input lengths  $k$  typically requires iterating over all possible alphabets in the domain. In our case, the problem of choosing optimal codes, is made worse by the fact that we don’t necessarily require codes with optimal distance  $\gamma$  or the smallest codeword length  $n$ . Instead, we want a code that gives the smallest value of  $2n\ell$ , while ensuring that  $\left(\frac{e^{\epsilon/p} - 1}{(\epsilon/p)^{\epsilon/p}}\right)^{np} < \frac{1}{2^{O(\lambda)}}$ , where  $p = (\frac{1}{2}) \cdot \frac{1}{(2-1)}$  and  $\epsilon = \gamma/n$  (See Sect. 5.2 for details). To simplify this problem and to get an estimate of how many lockboxes are required, we pick a particular binary ECC with constant rate and find values

of  $n$ ,  $\gamma$  and  $\ell$  that give the smallest value of  $2n\ell$  withing this encoding scheme. In particular, we use Justesen codes [35].

**Table 1** Lockboxes required for various input lengths with statistical security parameter 50

| Input Length ( $k$ ) | Codeword Length ( $n$ ) | $(n, k, m, \gamma)$       | $\ell$ | Total LB ( $2n\ell$ ) | LB/Bit ( $2n\ell/k$ ) |
|----------------------|-------------------------|---------------------------|--------|-----------------------|-----------------------|
| 192                  | 496                     | (43, 32, 6, 12)           | 7      | 7224                  | 37.625                |
| 256                  | 752                     | (47, 32, 8, 16)           | 7      | 10528                 | 41.125                |
| 560                  | 1302                    | (93, 80, 7, 14)           | 7      | 18228                 | 32.55                 |
| 5000                 | 14180                   | (709, 500, 10, 400)       | 4      | 113440                | 22.688                |
| 300000               | 735720                  | (24524, 20000, 15, 13080) | 4      | 5885760               | 19.6192               |

**Encoding with Justesen Codes.** Justesen codes are derived as the code concatenation of a Reed-Solomon code and the Wozencraft ensemble. The encoding algorithm works as follows – the given binary input string of length  $k$  is divided into  $k'$  blocks of length  $m$  each. This new vector of length  $k'$  is encoded using the Reed Solomon code  $(n', k', n' - k' + 1)$  over field  $GF(2^m)$ . Finally, the resulting  $n'$  blocks of length  $m$  each are encoded using Wozencraft ensemble. We use a particular Wozencraft ensemble [41], that yields a final codeword of length  $2mn'$ . The minimum distance  $\gamma$  of the resulting code is  $\sum_{i \in [g]} i \cdot \binom{2m}{i}$ , where  $g$  is the smallest integer such that  $\sum_{i \in [g]} \binom{2m}{i} \leq n' - k' + 1$ .

**Estimating the Optimal no. of Lockboxes.** Since,  $n'$  here can potentially take any value  $< 2^m$  (and  $m \in [1, k]$ ), a bruteforce approach to find optimal values even within Justesen code will result in an exponential search. To reduce the search space, we observe that for any given input length  $k$  and distance  $\gamma$ , it suffices to only look at the smallest admissible value of  $n'$ . Greater values of  $n'$  for the same  $k$  and  $\gamma$  yield worse security and larger values of  $2n\ell$ . We use this observation to deploy the following strategy – for any input length  $k$ , iterate over all possible values of  $m \in [1, k]$ , compute all corresponding admissible values of  $g, \gamma$  and set  $n' = k' + \left(\sum_{i \in [g]} \binom{2m}{i}\right) - 1$  (this significantly reduces potential domain for  $n'$ ). For each such combination of  $(m, n', \gamma)$ , we calculate security for reasonable values of  $\ell$  and find the combination of  $(n', k', m, \gamma, \ell)$  that results in the fewest total number of lockboxes, while ensuring that the security is at least  $2^{-50}$ .

We report the number of lockboxes required for some input lengths in Table 1. As expected, the number of lockboxes per input wire decreases as the number of inputs increase. By replacing Wozencraft ensemble with BCH codes [12], we can hope to get small improvements for larger input lengths; however, for smaller inputs, BCH codes are unlikely to help. Overall, due to the lack of efficient binary linear ECC, the number of required lockboxes are unlikely to be significantly better than the ones computed using Justesen codes. Our laconic OT-based

construction offers some relief in this regard: for instance, if the length of digest output by the receiver is 256 bits, we require 10,528 total lockboxes for *any* input length.

## 8.2 Instantiating Lockboxes

To realize counter lockboxes from the widely-available device- and cloud-based hardware, some implementation considerations arise. In this section, we provide brief background on each candidate lockbox and the practical considerations involving their use.

**Cloud-Based Backup Services.** Apple’s Cloud Key Vault was introduced in 2016 when Apple added functionality to encrypt and store user-controlled encryption keys within hardware security modules (HSM) to remove Apple’s own ability to access them. Each iCloud account (registered email address) has access to a Cloud Key Vault record, which corresponds to a password-protected HSM entry which allows up to ten attempts<sup>5</sup> via the Secure Remote Password [4, 5] (SRP) protocol. Notably, this requires one email address per lockbox, as Apple allocates a single Cloud Key Vault entry to each user account.

Similar to Apple’s Cloud Key Vault, Google introduced HSM-based user-controlled encryption to protect backups even from insider threats [37]. Their system relies on the Titan [48] HSM hardware, and similarly implements a password-based attempt-limited authentication service which can naturally be viewed as a counter lockbox. Akin to Apple’s Cloud Key Vault, Google allocates a single backup service instance per user account, and so each lockbox requires a registered Google account (email address) to be deployed. Both iCloud and Google accounts can be acquired for free, but acquiring multiple accounts can require evading anti-spam measures.

Signal, the secure messaging platform, offers users a backup method relying on user-controlled encryption inaccessible to Signal’s servers. This service is called Secure Value Recovery, or SVR. SVR allows users to set a PIN, and gives them ten attempts to authenticate to an Intel SGX enclave to retrieve their backup data. As a secure enclave, SGX itself is capable of running one-time programs. However, to end users only a basic API is exposed which allows authentication attempts over a secure connection. Rather than email-based registration, Signal requires phone numbers, specifically to receive a confirmation SMS. Therefore, each SVR lockbox requires a phone number able to receive an SMS; such numbers cost \$0.50 USD/month each at scale with a service like Twilio [49].

**iOS Devices.** Apple also offers the eponymous counter lockbox as hardware within modern iOS devices (smartphones and tablets) available since Fall 2020. This component emerged with the second-generation Secure Enclave Processor [3] (SEP) and was designed to prevent forensic attacks against the passcode

---

<sup>5</sup> In Sect. 2.4, we discuss generic techniques to convert a multiple-attempt (e.g. 10) lockbox into a single-attempt, including simply “burning”  $n - 1$  attempts of each  $n$ -attempt lockbox before transmitting their locations to the receiver.

attempt counter which moderates access to a device and its filesystem. Although there are few official documents, initial exploration seems to imply that iOS devices are able to support up to 1024 counter lockbox instances simultaneously. Since counter lockboxes are intended for use by iOS itself, third-party developers must interact with them directly on jailbroken devices. Finally, a note on monetary costs: currently, iPad air 4th generation can be purchased for about \$300. Thus, the average cost of each lockbox can be estimated to be about \$0.30 USD.

### 8.3 Applications

Given the cost of each lockbox and the notable expansion between input length and total lockboxes as seen in Table 1, at present the real-world applicability of our constructions is somewhat limited. However, compelling applications involving small input lengths are within reach: Bitcoin addresses are 160-bit hashes, which could be input into a delegated signature one-time program. Down-sampled biometric measurements could be input to fuzzy matching algorithms, or passwords into client-side key derivations for user authentication. Compressed descriptions of aggregations could be input to an offline differentially-private database service to maintain privacy budgets. As lockbox availability grows, these domains will only expand.

**Acknowledgements** The first, second and fourth authors were supported in part by NSF CNS-1814919, NSF CAREER 1942789 and Johns Hopkins University Catalyst award. The fourth author was additionally supported in part by AFOSR Award FA9550-19-1-0200 and the Office of Naval Research Grant N00014-19-1-2294. The first, third and fifth authors were supported by the National Science Foundation under awards CNS-1653110 and CNS-1801479 and by DARPA under Agreements No. HR00112020021 and Agreements No. HR001120C0084. The fifth author was additionally supported by a Google Security & Privacy Award. This work was done in part while the second author was a student at Johns Hopkins University and while the second and fourth authors were visiting University of California, Berkeley. Any opinions, findings and conclusions or recommendations expressed in this material are those of the author(s) and do not necessarily reflect the views of the United States Government or DARPA.

### References

1. Alamati, N., Branco, P., Döttling, N., Garg, S., Hajiabadi, M., Pu, S.: Laconic private set intersection and applications. Cryptology ePrint Archive, Report 2021/728 (2021). <https://eprint.iacr.org/2021/728>
2. Almashaqbeh, G., et al.: Gage MPC: bypassing residual function leakage for non-interactive MPC. Proc. Priv. Enhanc. Technol. **2021**(4), 528–548 (2021)
3. Apple Inc., Secure Enclave. <https://support.apple.com/guide/security/secure-enclave-sec59b0b31ff/web>
4. Apple Inc., Escrow security for iCloud Keychain (2021). <https://support.apple.com/guide/security/escrow-security-for-icloud-keychain-sec3e341e75d/web>

5. Apple Inc., HomeKit communication security (2021). <https://support.apple.com/guide/security/homekit-communication-security-sec3a881ccb1/web>
6. ARM Holdings. Trusted Base System Architecture Documents. <https://www.arm.com/technologies/trustzone-for-cortex-a/tee-reference-documentation>. Subject to Non-Disclosure Agreement
7. Backes, M., Gerling, R.W., Gerling, S., Nürnberger, S., Schröder, D., Simkin, M.: WebTrust – a comprehensive authenticity and integrity framework for HTTP. In: Boureanu, I., Owesarski, P., Vaudenay, S. (eds.) ACNS 2014. LNCS, vol. 8479, pp. 401–418. Springer, Cham (2014). [https://doi.org/10.1007/978-3-319-07536-5\\_24](https://doi.org/10.1007/978-3-319-07536-5_24)
8. Bellare, M., Hoang, V.T., Rogaway, P.: Adaptively secure garbling with applications to one-time programs and secure outsourcing. In: Wang, X., Sako, K. (eds.) ASIACRYPT 2012. LNCS, vol. 7658, pp. 134–153. Springer, Heidelberg (2012). [https://doi.org/10.1007/978-3-642-34961-4\\_10](https://doi.org/10.1007/978-3-642-34961-4_10)
9. Benhamouda, F., Krawczyk, H., Rabin, T.: Robust non-interactive multiparty computation against constant-size collusion. In: Katz, J., Shacham, H. (eds.) CRYPTO 2017, Part I. LNCS, vol. 10401, pp. 391–419. Springer, Cham (2017). [https://doi.org/10.1007/978-3-319-63688-7\\_13](https://doi.org/10.1007/978-3-319-63688-7_13)
10. Bhudia, A., O’Keeffe, D., Sgandurra, D., Hurley-Smith, D.: RansomClave: ransomware key management using SGX. In: The 16th International Conference on Availability, Reliability and Security (2021)
11. Bitansky, N., et al.: The hunting of the SNARK. *J. Cryptol.* **30**(4), 989–1066 (2017)
12. Bose, R.C., Ray-Chaudhuri, D.K.: On a class of error correcting binary group codes. *Inf. Control* **3**(1), 68–79 (1960)
13. Broadbent, A., Gutoski, G., Stebila, D.: Quantum one-time programs. In: Canetti, R., Garay, J.A. (eds.) CRYPTO 2013, Part II. LNCS, vol. 8043, pp. 344–360. Springer, Heidelberg (2013). [https://doi.org/10.1007/978-3-642-40084-1\\_20](https://doi.org/10.1007/978-3-642-40084-1_20)
14. Van Bulck, J., et al.: Breaking virtual memory protection and the SGX ecosystem with foreshadow. *IEEE Micro* **39**(3), 66–74 (2019)
15. Chatterjee, R., Athayle, A., Akhawe, D., Juels, A., Ristenpart, T.: Password typos and how to correct them securely. In: S&P 2016. IEEE (2016)
16. Chaum, D., Pedersen, T.P.: Wallet databases with observers. In: Brickell, E.F. (ed.) CRYPTO 1992. LNCS, vol. 740, pp. 89–105. Springer, Heidelberg (1993). [https://doi.org/10.1007/3-540-48071-4\\_7](https://doi.org/10.1007/3-540-48071-4_7)
17. Chen, Z.: Java Card Technology for Smart Cards: Architecture and Programmer’s Guide. Addison-Wesley Longman Publishing Co., Inc (2000)
18. Cho, C., Döttling, N., Garg, S., Gupta, D., Miao, P., Polychroniadou, A.: Laconic oblivious transfer and its applications. In: Katz, J., Shacham, H. (eds.) CRYPTO 2017, Part II. LNCS, vol. 10402, pp. 33–65. Springer, Cham (2017). [https://doi.org/10.1007/978-3-319-63715-0\\_2](https://doi.org/10.1007/978-3-319-63715-0_2)
19. Dall, F., et al.: CacheQuote: efficiently recovering long-term secrets of SGX EPID via cache attacks. *IACR Trans. Cryptogr. Hardw. Embed. Syst.* **2018**(2), 171–191 (2018)
20. Delgado-Mohatar, O., Sierra-Cámara, J.M., Anguiano, E.: Blockchain-based semi-autonomous ransomware. *Future Gener. Comput. Syst.* **112**, 589–603 (2020)
21. Döttling, N., Garg, S., Goyal, V., Malavolta, G.: Laconic conditional disclosure of secrets and applications. In: Zuckerman, D., (eds.) 60th FOCS, pages 661–685. IEEE Computer Society Press, November 2019
22. Dwork, C., McSherry, F., Nissim, K., Smith, A.: Calibrating noise to sensitivity in private data analysis. In: Theory of Cryptography Conference (2006)

23. Garg, S., Srinivasan, A.: Adaptively secure garbling with near optimal online complexity. In: Nielsen, J.B., Rijmen, V. (eds.) EUROCRYPT 2018, Part II. LNCS, vol. 10821, pp. 535–565. Springer, Cham (2018). [https://doi.org/10.1007/978-3-319-78375-8\\_18](https://doi.org/10.1007/978-3-319-78375-8_18)
24. Goldreich, O., Ostrovsky, R.: Software protection and simulation on oblivious rams. *J. ACM (JACM)* **43**(3), 431–473 (1996)
25. Goldwasser, S., Kalai, Y.T., Rothblum, G.N.: Delegating computation: interactive proofs for muggles. In: Ladner, R.E., Dwork, C. (eds.) 40th ACM STOC, pp. 113–122. ACM Press, May 2008
26. Goldwasser, S., Kalai, Y.T., Rothblum, G.N.: One-time programs. In: Annual International Cryptology Conference, pp. 39–56 (2008)
27. Google. Google Tensor debuts on the new Pixel 6 this fall (2021). <https://blog.google/products/pixel/google-tensor-debuts-new-pixel-6-fall/>
28. Goyal, R., Goyal, V.: Overcoming cryptographic impossibility results using blockchains. In: Kalai, Y., Reyzin, L. (eds.) TCC 2017, Part I. LNCS, vol. 10677, pp. 529–561. Springer, Cham (2017). [https://doi.org/10.1007/978-3-319-70500-2\\_18](https://doi.org/10.1007/978-3-319-70500-2_18)
29. Goyal, R., Vusirikala, S., Waters, B.: New constructions of hinting PRGs, OWFs with encryption, and more. In: Micciancio, D., Ristenpart, T. (eds.) CRYPTO 2020, Part I. LNCS, vol. 12170, pp. 527–558. Springer, Cham (2020). [https://doi.org/10.1007/978-3-030-56784-2\\_18](https://doi.org/10.1007/978-3-030-56784-2_18)
30. Goyal, V., Ishai, Y., Sahai, A., Venkatesan, R., Wadia, A.: Founding cryptography on tamper-proof hardware tokens. In: Micciancio, D. (ed.) TCC 2010. LNCS, vol. 5978, pp. 308–326. Springer, Heidelberg (2010). [https://doi.org/10.1007/978-3-642-11799-2\\_19](https://doi.org/10.1007/978-3-642-11799-2_19)
31. Harnik, D., Kilian, J., Naor, M., Reingold, O., Rosen, A.: On robust combiners for oblivious transfer and other primitives. In: Cramer, R. (ed.) EUROCRYPT 2005. LNCS, vol. 3494, pp. 96–113. Springer, Heidelberg (2005). [https://doi.org/10.1007/11426639\\_6](https://doi.org/10.1007/11426639_6)
32. Hazay, C., Lindell, Y.: Constructions of truly practical secure protocols using standardsmartcards. In: Ning, P., Syverson, P.F., Jha, S. (eds.) ACM CCS 2008, pp. 491–500. ACM Press, October 2008
33. Intel. Overview on signing and whitelisting for intel software guard extension (SGX) enclaves. <https://www.intel.com/content/dam/develop/external/us/en/documents/overview-signing-whitelisting-intel-sgx-enclaves-737361.pdf>
34. Juels, A., Sudan, M.: A fuzzy vault scheme. *Des. Codes Cryptogr.* **38**(2), 237–257 (2006)
35. Justesen, J.: Class of constructive asymptotically good algebraic codes. *IEEE Trans. Inf. Theory* **18**(5), 652–656 (1972)
36. Kaptchuk, G., Green, M., Miers, I.: Giving state to the stateless: augmenting trustworthy computation with ledgers. In: NDSS 2019 ( 2019)
37. Kensinger, T.: Google and Android have your back by protecting your backups, September 2018. <https://security.googleblog.com/2018/10/google-and-android-have-your-back-by.html>
38. Krassovsky, S., Cadden, G., et al.: Security of End-To-End Encrypted Backups (2021). [https://scontent.whatsapp.net/v/t39.8562-34/241394876\\_546674233234181\\_8907137889500301879\\_n.pdf/WhatsApp\\_Security\\_Encrypted\\_Backups\\_Whitepaper.pdf?ccb=1-5&\\_nc\\_sid=2fbf2a&\\_nc\\_ohc=4K040x7GheAAX\\_-4c-.\\_&\\_nc\\_ht=scontent.whatsapp.net&oh=01\\_AVxDv1cRlVlvg0Fv89URSU\\_XOQUupw70bDPw6o2w0LEWg&oe=6211F5FC](https://scontent.whatsapp.net/v/t39.8562-34/241394876_546674233234181_8907137889500301879_n.pdf/WhatsApp_Security_Encrypted_Backups_Whitepaper.pdf?ccb=1-5&_nc_sid=2fbf2a&_nc_ohc=4K040x7GheAAX_-4c-._&_nc_ht=scontent.whatsapp.net&oh=01_AVxDv1cRlVlvg0Fv89URSU_XOQUupw70bDPw6o2w0LEWg&oe=6211F5FC)

39. Krstić, I.: Behind the scenes with iOS security (2016). <https://www.blackhat.com/docs/us-16/materials/us-16-Krstic.pdf>
40. Lund, J.: December 2019 <https://signal.org/blog/secure-value-recovery/>. Accessed 2 May 2022
41. MacWilliams, F.J., Sloane, N.J.A.: The Theory of Error-Correcting Codes. North-Holland Pub. Co. (1977)
42. McKeen, F., et al.: Intel® software guard extensions (intel® sgx) support for dynamic memory management inside an enclave. In: HASP 2016. ACM (2016)
43. Meier, R., Przydatek, B., Wullschlegel, J.: Robuster combiners for oblivious transfer. In: Vadhan, S.P. (ed.) TCC 2007. LNCS, vol. 4392, pp. 404–418. Springer, Heidelberg (2007). [https://doi.org/10.1007/978-3-540-70936-7\\_22](https://doi.org/10.1007/978-3-540-70936-7_22)
44. Micali, S.: Computationally sound proofs. SIAM J. Comput. **30**(4), 1253–1298 (2000)
45. Murdock, K., Oswald, D.F., Garcia, F.D., Van Bulck, J., Gruss, D., Piessens, F.: Plundervolt: software-based fault injection attacks against intel SGX. In: S&P 2020. IEEE (2020)
46. Pinto, S., Santos, N.: Demystifying arm trustzone: a comprehensive survey. ACM Comput. Surv. (CSUR) **51**(6) (2019)
47. Rosulek, M., Roy, L.: Three halves make a whole? Beating the half-gates lower bound for garbled circuits. In: Malkin, T., Peikert, C. (eds.) CRYPTO 2021, Part I. LNCS, vol. 12825, pp. 94–124. Springer, Cham (2021). [https://doi.org/10.1007/978-3-030-84242-0\\_5](https://doi.org/10.1007/978-3-030-84242-0_5)
48. Savagaonkar, U., Porter, N., Taha, N., Serebrin, B., Mueller, N.: Titan in depth: Security in plaintext (2017). <https://cloud.google.com/blog/products/identity-security/titan-in-depth-security-in-plaintext>
49. Twilio (2022). <https://www.twilio.com/sms/pricing/us>
50. Xu, Q.: ARM-software/tf-issues (2017). <https://github.com/ARM-software/tf-issues/issues/534>
51. Yao, A.C.-C.: How to generate and exchange secrets (extended abstract). In: 27th FOCS, pages 162–167. IEEE Computer Society Press, October 1986
52. yubico. YubiHSM 2. <https://www.yubico.com/product/yubihsm-2/>