



Non-malleable Commitments Against Quantum Attacks

Nir Bitansky^{1(✉)}, Huijia Lin², and Omri Shmueli¹

¹ Tel Aviv University, Tel Aviv, Israel
{nirbitan, omrishmueli}@tau.ac.il

² Washington University, Seattle, WA, USA
rachel@cs.washington.edu

Abstract. We construct, under standard hardness assumptions, the first non-malleable commitments secure against quantum attacks. Our commitments are statistically binding and satisfy the standard notion of *non-malleability with respect to commitment*. We obtain a $\log^*(\lambda)$ -round classical protocol, assuming the existence of post-quantum one-way functions.

Previously, non-malleable commitments with quantum security were only known against a restricted class of adversaries known as *synchronizing adversaries*. At the heart of our results is a new general technique that allows to modularly obtain non-malleable commitments from any extractable commitment protocol, obviously of the underlying extraction strategy (black-box or non-black-box) or round complexity. The transformation may also be of interest in the classical setting.

1 Introduction

Commitments are one of the most basic cryptographic primitives. They enable a sender to commit to a string to be *opened* at a later stage. As long as the commitment is not opened, it is *hiding*—efficient receivers learn nothing about the committed value. Furthermore, the commitment is *statistically binding*—with overwhelming probability, the commitment can be opened to a single, information-theoretically determined value in the commitment phase. While these basic security guarantees go a long way in terms of applications, they do not always suffice. In particular, they do not prevent a *man-in-the-middle* adversary from receiving a commitment to a given value v from one party and trying to send to another party a commitment to a related value, say $v-1$ (without knowing the committed value v at all).

Such attacks are called “mauling attacks” and in some settings could be devastating. For instance, consider the scenario where a city opens a bidding process for the construction of a new city hall. Companies are instructed to commit to their proposed bid using a commitment scheme, and these commitments are opened at the end of the bidding period. If the scheme is “malleable”, company A may manage to underbid company B , by covertly mauling B ’s commitment to create their

own commitment to a lower bid. More generally, ensuring independence of private values is vital in many applications of commitments, such as coin tossing, federated learning, and collaborative computation over private data.

In their seminal work, Dolev, Dwork and Naor introduced the concept of *non-malleable commitments* to protect against mauling attacks [DDN03]. They guarantee that the value $\tilde{v}$ a man-in-the-middle adversary commits to is computationally independent of the value v in the commitment it receives (unless the man-in-the-middle simply “copies”, by relaying messages between the honest sender and receiver it interacts with, in which case $\tilde{v} = v$). From its onset, the study of non-malleable cryptography has put stress on achieving solutions without any reliance on trusted parties or any form of trusted setup, and solutions that hold when honest parties may not even be aware of the existence of a man-in-the-middle, and the way it manipulates the messages they send over time. The latter is particularly important in applications where the man-in-the-middle acts “in the dark”. For instance, in the aforementioned example, company A may not be aware of the competing company B .

Since their conception, non-malleable commitments have indeed proved to be a useful and versatile building block for ensuring independence of values. They have been used in coin-tossing protocols, secure multiparty computation protocols, non-malleable proof systems (zero-knowledge, witness indistinguishability, multi-prover interactive proofs), and more. Techniques developed for non-malleable commitments are also useful for building non-malleable codes, non-malleable extractors (and two source extractors), and non-malleable time-lock puzzles. The work of [DDN03] constructed the first non-malleable commitments against classical adversaries based on one-way functions. Since then, a plethora of constructions have been proposed achieving different, sometimes optimal, tradeoffs between round-complexity, efficiency, and underlying assumptions (c.f. [Bar02, PR05a, PPV08, LPV09, PW10, Wee10, Goy11, GLOV12, COSV16, GPR16a, GKS16, Khu17, KS17, LPS17, BL18, KK19, GR19, GKLW20]).

Non-Malleability Against Quantum Adversaries. In contrast to the comprehensive understanding of non-malleability in the classical setting, our understanding of non-malleability against quantum adversaries is very much lacking. The threat of quantum attacks has prompted the development of post-quantum cryptography, and yet despite its important role in cryptography, post-quantum non-malleability has yet to catch up. In this work, we construct, under standard assumptions, the first non-malleable commitments with post-quantum security, namely, the hiding and non-malleability properties hold even against efficient quantum adversaries (and binding continues to be information theoretic).

Prior to our work, post-quantum non-malleable commitments were not known under any assumption. Partial progress was made by Agrawal, Bartusek, Goyal, Khurana, and Malavolta [ABG+20] who, assuming super-polynomial quantum hardness of Learning With Errors, construct post-quantum non-malleable commitments against a restricted class of adversaries known as *synchronizing adversaries*. A synchronizing adversary is limited as follows: When acting as a man-in-the-middle between a sender and a receiver, it is bound to synchronize its interactions with the honest parties; namely, when it receives

the i -th message from the sender, it immediately sends the i -th message to the receiver and vice versa. Such synchronicity may often not exist for example due to network’s asynchronicity, lack of synchronized clocks, or concurrent executions where parties are unaware of the existence of other executions. Enforcing synchronizing behaviour in general requires a trusted setup (like a broadcast channel) and coordination among parties to enforce message ordering.

The gold standard of non-malleability (since its introduction in [DDN03]) requires handling general, *non-synchronizing* adversaries, who can arbitrarily schedule messages in the two interactions (without awareness of the sender and receiver). In this work, for the first time, we achieve this gold standard non-malleability in the post-quantum setting. As we shall explain later on, the challenge stems from the fact that classical techniques previously used to obtain non-malleability against non-synchronizing adversaries (e.g., as robust extraction [LP09], simulation extractability [PR05a, PR05b] and so on) do not generally apply in the quantum setting. This is due to basic quantum phenomena such as *unclonability* [WZ82] and *state disturbance* [FP96].

Our Results in More Detail. We construct statistically binding non-malleable commitments against quantum non-synchronizing adversaries, assuming post-quantum one-way functions. Our main result is a modular construction of post-quantum non-malleable commitments from post-quantum extractable commitments. The latter is a statistically binding commitment protocol that is extractable in the following sense: There exists an efficient quantum extractor-simulator, which given the code of any quantum sender, can simulate the arbitrary output of the sender up to, while extracting the committed value. The construction, in fact, only requires ε -extractability, meaning that the extractor-simulator obtains an additional *simulation accuracy parameter* $1^{1/\varepsilon}$, and the simulation only guarantees ε -indistinguishability

Theorem 1 (Informal). *Assuming k -round post-quantum ε -extractable commitments, there exist $k^{O(1)} \cdot \log^* \lambda$ -round post-quantum non-malleable commitments, where λ is the security parameter.*

By default, when we say “post-quantum” we mean protocols that can be executed by classical parties, but which are secure against quantum adversaries. In particular, starting from a post-quantum classical ε -extractable commitment, we obtain a post-quantum classical non-malleable commitment. Constant-round ε -extractable commitments were constructed by Chia et al. [CCLY21] based on post-quantum one-way functions. Hence, we get the following corollary.

Corollary 1. *Assuming there exist post-quantum one-way functions, there exist $O(\log^* \lambda)$ -round post-quantum non-malleable commitments.*

2 Technical Overview

We now give an overview of the main ideas behind our construction. Following the convention in the non-malleability literature, we refer to the interaction

between Sen and $\mathcal{A}$ as the left interaction/commitment, and that between Rec and $\mathcal{A}$ the right interaction/commitment. Similarly, we refer to v , tg (and $\tilde{v}$, $\tilde{\text{tg}}$) as the left (and right) committed values or tag.

2.1 Understanding the Challenges

Before presenting our base commitments, we explain the main challenges that arise in the quantum setting. First, we recall a basic approach toward proving non-malleability in the classical setting *via extraction*. Here the basic idea is to provide a reduction that given a MIM adversary $\mathcal{A}$, can efficiently extract the value $\tilde{v}$ that $\mathcal{A}$ commits to on the right. Accordingly, if the MIM $\mathcal{A}$ manages to maul the commitment to v on the left and commit to a related value $\tilde{v}$ on the right, the reduction will gain information about v , and be able to break the *hiding* of the commitment.

The Difficulty in MIM Extraction. Extractable commitments allow for efficient extraction from adversarial senders in *the stand-alone setting*. Such extraction is traditionally done by either means of rewinding, or more generally using the sender’s code. In the MIM setting, where $\mathcal{A}$ acts as a sender on the right, while acting as a receiver on the left, extraction from $\mathcal{A}$ is much more challenging. The problem is that the interaction of $\mathcal{A}$ with the receiver Rec on the right may occur concurrently to its interaction with the sender Sen on the left. This means that a reduction attempting to rewind $\mathcal{A}$ to extract the right committed value, may effectively also need to rewind the sender Sen on the left. (This may happen for example if, when the reduction rewinds $\mathcal{A}$ and sends $\mathcal{A}$ a new message, $\mathcal{A}$ also sends a new message in the left commitment and expects a reply from Sen before proceeding in the right commitment.) In such a case, extraction does not generally work—the “actual” sender of the right commitment is essentially the MIM $\mathcal{A}$ *combined with the sender Sen on the left*. However, the reduction does not possess the code of Sen , specifically, it does not possess its randomness. The challenge is to perform such extraction without access to the secret randomness of the sender on the left, and thus without compromising the hiding of the left commitment.

Indeed, classical non-malleable commitments tend to require more than *plain extractable commitments*. A long array of works (c.f., [DDN03, PR05b, PR05a, LP09, PW10, LP11, Goy11]) design various *safe extraction techniques*, which guarantee extraction on the right without compromising hiding of the left committed value. These safe-extraction techniques rely on properties of specific protocols and extraction strategies, rather than general (stand-alone) extractable commitments. For instance, the protocols of [DDN03, LP09, LP11, Goy11, GPR16a] rely on three-message witness-indistinguishable protocols satisfying an extraction guarantee known as *special soundness*, whereas the protocols in [PR05b, PR05a] rely on the specific structure of Barak’s non-black-box zero knowledge protocol.

The Quantum Barrier. The (safe) extraction techniques used to obtain non-malleability in the classical setting fail in the quantum setting. For once, rewinding does not generally work. We cannot record the adversary’s quantum state between rewinding attempts due to the no-cloning theorem [WZ82]. Also, we

cannot simply measure between rewindings, as this disturbs that the adversary's state [FP96]. In this case, even if we do extract, we may not be able to faithfully simulate the adversary's output state in the protocol¹. Similarly, non-black-box techniques do not generally apply. For instance, it is unclear how to apply Barak's non-black-box simulation technique [Bar02], due to the lack of *universal arguments* [BG08] for quantum computations (this is just to mention one difficulty in using Barak's strategy in the quantum setting).

The difficulty of applying classical proof techniques in the setting of quantum adversaries is indeed a well known phenomena, and in some settings, quantum proof techniques have been successfully developed to circumvent this difficulty. Perhaps the most famous example of this is in the context of zero-knowledge simulation. Here Watrous [Wat09] shows that in certain settings quantum rewinding is possible and used it to obtain zero-knowledge protocols. Several other rewinding techniques enable extraction, but disturb the adversary's state in the process [Unr12, CCY20, CMSZ21]. Alternatively, several recent works [AP19, BS20, ABG+20] obtain constant round zero-knowledge via non-black-box quantum techniques, using quantum FHE (and assuming LWE). While post-quantum extractable commitments do exist, they do not satisfy the specific properties that the classical safe-extraction techniques require.

Given the above state of affairs, in this work, we aim to construct post-quantum non-malleable commitments modularly based on *any* post-quantum extractable (or ε -extractable) commitment. The equivalence between extractability and non-malleability is interesting on its own from a theoretical perspective. It turns out that doing so is challenging, and requires designing completely new safe-extraction techniques that work with general quantum extractable commitments, which we explain next.

For the sake of simplicity, and toward highlighting the main new ideas in this work, we ignore the difference between fully-extractable and ε -extractable commitments through the rest of this overview. We note that the transition from full extractable commitments to ε -extractable ones is quite direct and is based on the common knowledge that ε -simulation is sufficient when aiming to achieve indistinguishability-based definitions. Indeed, the definition of non-malleability is an indistinguishability-based definition, and accordingly showing ε -indistinguishability for any inverse polynomial ε is sufficient. In this case, the simulators invoked in the reduction are all still polynomial-time.

The Synchronizing Setting. As observed in [ABG+20], if restricted to synchronizing adversaries, such a modular construction exists using ideas from early works [CR87, DDN03]: When committing under a tag $\mathbf{tg} \in [\tau]$ for $\tau \leq \lambda$, in every round $i \neq \mathbf{tg}$ send an empty message, and in round $\mathbf{tg}$, send an extractable commitment to the value v . Indeed, in the synchronizing setting, a commitment on the left under tag $\mathbf{tg}$ would never interleave with the commitment on the right under tag

¹ Recall that non-malleability requires that the joint distribution of the output state of the adversary and the committed value are indistinguishable regardless of the committed value on the left. Hence the reduction needs to extract the committed value without disturbing the state of the quantum adversary.

$\tilde{\mathbf{t}}\mathbf{g} \neq \mathbf{t}\mathbf{g}$. Thus, safe-extraction opportunities come for free, circumventing the real challenge in achieving non-malleability. It is not hard to see, however, that in the non-synchronizing setting, this approach would completely fail as the adversary can always align the extractable commitment on the right with that on the left. The work of [ABG+20] further constructed constant-round non-malleable commitments for a super-constant number of tags, based on mildly super-polynomial security of quantum FHE and LWE. The non-malleability of the new protocol, however, still relies on the synchronization of the left and right commitments.

2.2 Leveraging Extractable Com in Non-synchronizing Setting

We design a base protocol for a constant number of tags that, using *any* (post-quantum) extractable commitment scheme. The protocol guarantees extraction on the right while preserving hiding on the left, even against a quantum non-synchronizing MIM adversary. In this overview, we explain our base commitments in three steps:

- First, we introduce our basic idea in the simplified *one-sided* non-malleability setting where the MIM is restricted to choose a smaller tag on the right than the tag on the left, $\tilde{\mathbf{t}}\mathbf{g} < \mathbf{t}\mathbf{g}$.
- Then, we extend the basic idea to the general setting where the MIM may also choose a right tag that is larger $\tilde{\mathbf{t}}\mathbf{g} > \mathbf{t}\mathbf{g}$. We illustrate the main ideas here under the simplifying assumption of a certain honest behavior of the adversary.
- Finally, we show how to remove the simplifying assumption on the adversary.

Step 1: One-sided Non-malleability Let us first consider a MIM adversary that given a commitment on the left under tag $\mathbf{t}\mathbf{g}$, produces a commitment on the right under a smaller tag $\tilde{\mathbf{t}}\mathbf{g} < \mathbf{t}\mathbf{g}$. In our commitment, the sender first secret shares the value v to be committed into shares $u_1, \dots, u_n$. It then sequentially sends extractable commitments to each of the shares $u_1, \dots, u_n$ – we refer to the entire batch of these sequential extractable commitments as a *block-commitment to v* . The binding and hiding of this protocol follow directly from those of the underlying extractable commitment. We focus on non-malleability.

To achieve non-malleability, the number of shares n is chosen as a function of the tag $\mathbf{t}\mathbf{g}$. The goal is to guarantee that in every execution where the tag $\tilde{\mathbf{t}}\mathbf{g}$ on the right is smaller than the tag $\mathbf{t}\mathbf{g}$ on the left, there will exist, on the *left*, a commitment to one of the shares u_i that is *free* in the sense that it does not interleave with the interaction on the right; namely, during the commitment to u_i on the left, no message is sent in the right execution (see Fig. 1). Before explaining how freeness is achieved, let us explain how we use it to establish non-malleability.

Extracting While Preserving Hiding and First-Message Binding. To argue non-malleability, we show that we can efficiently extract *all* shares $\tilde{u}_1, \dots, \tilde{u}_{\tilde{n}}$ on the right, while preserving the hiding of the free share u_i on the left, and by the security of secret sharing, also the hiding of the committed value v .

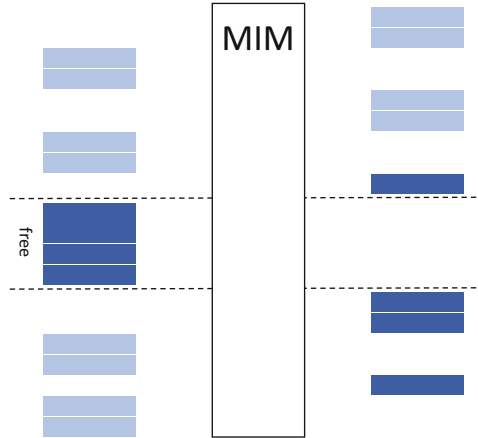


Fig. 1. Freeness Example. Each share commitment has 4 messages and there are $n = 3$ shares on the left, and $\tilde{n} = 2$ shares on the right. The second commitment on the left is free. Note that it splits the second commitment on the right.

Freeness guarantees that almost all commitments on the right do not interleave with the commitment to u_i on the left, more precisely, a single commitment on the right could be “split” by the commitment to u_i on the left (as in Fig. 1), which prevents extraction of that right split commitment. To deal with this, we rely on extractable commitments that are *first-message binding*; namely their first sender message fixes the value of the commitment. This gives rise to a simple extraction strategy: for any commitment on the right, where the first sender’s message is sent before the free commitment (on the left), we can extract the corresponding share *non-uniformly*; for the commitments where the first sender’s message occurs afterwards, we use the efficient extractor. Accordingly, we get a non-uniform reduction to the hiding of the free extractable commitment on the left.

We observe that any extractable commitment can be made first-message binding without any additional assumptions, and while increasing round complexity by at most a constant factor. For simplicity we describe how to achieve this assuming also non-interactive commitments.² We append to the original extractable commitment a first message where the sender sends a non-interactive commitment to the committed value and add at the end a zero-knowledge argument that this commitment is consistent with the commitment in the original extractable commitment. Extractability follows from the extractability of the original scheme and soundness of the argument, whereas hiding follows from that of the original scheme and the zero knowledge property. We note that (post-quantum) zero-knowledge arguments follow from (post-quantum) extractable commitments with a constant round complexity overhead (see e.g. [BS20]), and the same holds for ε -zero-knowledge and ε -extractable commitments, respectively.

² In the body, we observe that Naor commitments [Nao91], which can be obtained from (post-quantum) one-way functions, and thus also from any commitment, are in fact sufficient.

Guaranteeing Freeness. To achieve the required freeness property, it suffices to guarantee that whenever $\tilde{\mathbf{tg}} < \mathbf{tg}$, the number of shares $n(\mathbf{tg})$ (and hence the number of extractable commitments) on the left is larger than the total number of messages on the right, which is $k \cdot n(\tilde{\mathbf{tg}})$, where k is the number of messages in each extractable commitment. Accordingly, we choose $n(\mathbf{tg}) = (k + 1)^{\mathbf{tg}}$.

Step 2: Dealing with General Adversaries. The above commitment does not prevent mauling of commitments under tag $\mathbf{tg}$ to commitments under tags $\tilde{\mathbf{tg}} > \mathbf{tg}$. To deal with general adversaries, we invoke the above idea again in reverse order. That is, the sender now secret shares the value v twice independently: once to n shares $u_1, \dots, u_n$, and again to $\bar{n}$ shares $\bar{u}_1, \dots, \bar{u}_{\bar{n}}$. It then sequentially sends extractable commitments to the shares $u_1, \dots, u_n, \bar{u}_1, \dots, \bar{u}_{\bar{n}}$, that is, sending two sequential block-commitments to v . To understand the basic idea, we assume for simplicity, in this step, that the MIM attacker always commits to shares of the same value $\tilde{v}$ in the two block-commitments on the right (in Step 3, we will remove this assumption using zero-knowledge arguments).

Our goal now is to set the number of shares $n(\cdot), \bar{n}(\cdot)$, based on the tags, to guarantee that there exists a block-commitment on the right with respect to which there exist two extractable commitments to shares u_i and $\bar{u}_{\bar{i}}$ on the left (one from each left block-commitment) that are *free*. This means we can extract every share from that right block-commitment, while keeping the shares u_i and $\bar{u}_{\bar{i}}$, and hence the left committed value, hidden. We say that the corresponding block-commitment on the right is *ideally scheduled* (see Fig. 2).

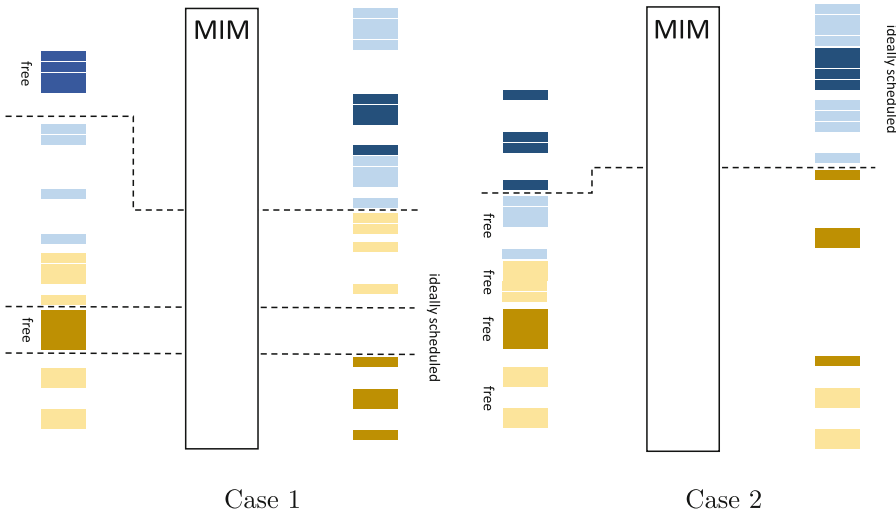


Fig. 2. Examples of an ideally scheduled block of shares (on the right). The first block of share commitments is colored in (light/dark) blue and the second in (light/dark) yellow. We mark the commitments on the left that are free with respect to the ideally scheduled block. (Color figure online)

Once we establish the existence of an ideally scheduled block, we can prove non-malleability using a non-uniform reduction to the hiding of the extractable commitments to u_i and $\bar{u}_i$ similar to the one we used in the first step. Since we are only able to extract from one of the two block-commitments on the right, it is important that both commit to the same value $\tilde{v}$, and thus our reduction would work, regardless of which one of the two it is able to extract from. Before we explain how to enforce this using ZK in Step 3, we explain how the existence of an ideally scheduled block is established.

Guaranteeing an Ideally Scheduled Block-Commitment. We prove that by setting the parameters $n, \bar{n}$ appropriately, an ideally scheduled block of shares always exists. For this purpose we generalize the combinatorial argument from before. Concretely, we set $n, \bar{n}$ to guarantee that:

1. *Either*, the number of shares $n = n(\mathbf{tg})$ in the first left block-commitment is larger than the total number of messages $k \cdot n(\tilde{\mathbf{tg}})$ in the first right block-commitment,
2. *Or*, the number of shares $\bar{n} = \bar{n}(\mathbf{tg})$ in the second left block-commitment is larger than the total number of messages $k \cdot \bar{n}(\tilde{\mathbf{tg}})$ in the second right block-commitment.

In addition, we require that $n, \bar{n}$ are both at least 2. These conditions can be satisfied for example by setting $n = (k+1)^{\mathbf{tg}}, \bar{n} = (k+1)^{\tau-\mathbf{tg}} + 1$, where τ is the total number of tags (namely, $\mathbf{tg} \in [\tau]$).

To see why the above is sufficient, let us assume for instance that Condition 2 of the two above conditions holds (at this point, both are treated symmetrically). We consider two cases:

- **Case 1 (depicted in Fig. 2a):** the commitment to share u_1 (i.e., the first share of the first block-commitment) on the left ends before the second block-commitment starts on the right. In this case, the commitment to u_1 on the left is free with respect to the second block-commitment on the right. Furthermore, since Condition 2 holds, (by the argument in Step 1,) there also exists a commitment to a share $\bar{u}_i$ (in the second block-commitment) on the left that is also free with respect to the second block-commitment on the right. Accordingly, the second block-commitment on the right is ideally scheduled.
- **Case 2 (depicted in Fig. 2b):** the commitment to share u_1 on the left ends after the second block of share commitments starts on the right. In this case, the commitments to shares $u_2, \dots, u_n, \bar{u}_1, \dots, \bar{u}_{\bar{n}}$ on the left are all free with respect to the first block-commitment on the right, and thus it is ideally scheduled. (We use the fact that $n \geq 2$, to deduce that a free share u_2 indeed exists.)

Step 3: Use ZK to Ensure Consistency of Right Block-Commitments.

Recall that in the last step, we made the simplifying assumption that the MIM adversary always commits to the same value $\tilde{v}$ in the two right block-commitments. The expected approach to removing this assumption, would be to

require that the sender gives a (post-quantum) zero-knowledge argument that such consistency indeed holds.

While the soundness of the argument guarantees the required consistency on the right, the addition of a zero knowledge proof brings about new challenges in the reduction of non-malleability to hiding on the left, due to non-synchronizing adversaries. Indeed, in the proof of non-malleability, before using the hiding of the extractable commitments on the left, we must use the zero knowledge property on the left to argue that the proof does not compromise the hidden shares. The problem is that the zero-knowledge argument on the left might interleave with our ideally scheduled block-commitment on the right, and thus with our extraction procedure. For instance, if the extractor wants to rewind the MIM, it might have to rewind the zero knowledge prover on the left, which is not possible. More generally, there could be a circular dependency: The zero-knowledge simulation needs to be applied to the verifier's code which depends on the extractor's code; however, extraction needs to be applied to the sender's code which depends on the simulator's code.

To circumvent this difficulty, we would like to guarantee that an ideally scheduled block-commitment would also be free of the zero knowledge messages on the left, namely, during its execution, no zero knowledge messages should be sent in the left execution (see Fig. 3). Indeed, if this is the case, then we can apply the zero knowledge simulator to the verifier that when needed runs the extractor on the right *in its head*. Note that since the right block-commitment is free from zero knowledge messages on the left, the code of the extractor, and induced verifier, is independent of the simulator's code, breaking the circularity.

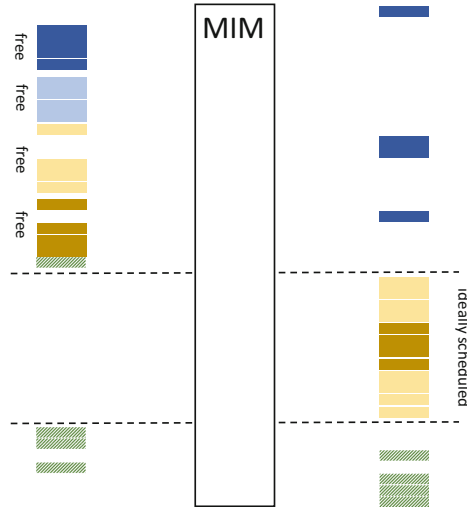


Fig. 3. The zero knowledge argument on the left is colored in green. The ideally scheduled block of shares on the right is required to be free of any zero knowledge messages (as well as satisfy the same conditions as before). (Color figure online)

Guaranteeing (the Stronger Form of) Ideal Scheduling. To achieve the stronger form of ideal scheduling, we augment the protocol yet again. Specifically, we repeat sequentially for $\ell + 1$ times the second block-commitment to shares $\bar{u}_1, \dots, \bar{u}_{\bar{n}}$, where ℓ is the number of rounds in the zero knowledge protocol. We now require that there is a block-commitment I among the $\ell + 2$ right block-commitments (one of $u_1, \dots, u_n$, and $\ell + 1$ of $\bar{u}_1, \dots, \bar{u}_{\bar{n}}$) that is ideally scheduled in the following stronger sense:

1. There exist shares u_i and $\bar{u}_{\bar{i}}$ such that *all* commitments to these shares (one to u_i and $\ell + 1$ ones to $\bar{u}_{\bar{i}}$) on the left, are free of the I 'th right block-commitment.
2. The I 'th right block-commitment is free of the zero knowledge argument on the left.

We provide a more involved combinatorial argument (and choice of parameters $n, \bar{n}$) showing that an ideally scheduled right block-commitment I always exists. Concretely, we set $n, \bar{n}$ to guarantee that:

1. *Either*, the number of shares $n = n(\text{tg})$ in the first left block-commitment as well as the number of shares $\bar{n} = \bar{n}(\text{tg})$ in each of the left block-commitments $2, \dots, \ell + 2$ are *both* larger than the total number of messages $k \cdot n(\tilde{\text{tg}})$ in the first right block-commitment.
2. *Or*, the number of shares $\bar{n} = \bar{n}(\text{tg})$ in each of the left block-commitments $2, \dots, \ell + 2$ is larger than the total number of messages $k \cdot \bar{n}(\tilde{\text{tg}})$ in each of the right block-commitments $2, \dots, \ell + 2$.

Again, we also require that $n, \bar{n}$ are both at least 2. The above conditions can be satisfied for example by setting $n = (k + 1)^{\text{tg}}, \bar{n} = (k + 1)^{2\tau - \text{tg}} + 1$, where τ is the total number of tags. The above two conditions can no longer be treated symmetrically as before. We explain separately, how each one of them implies the existence of an ideally scheduled block on the right (in the stronger sense defined above).

- **Case 1 (applies for either one of the two conditions):** the first block-commitment on the right ends after the knowledge argument on the left had started. In this case, block commitments $2, \dots, \ell + 2$ on the right do not interleave with any of the block commitments on the left. Thus, we only need to establish that one of them does not interleave with the zero knowledge argument on the left. This follows from the fact that there are $\ell + 1$ of them, but only ℓ messages in the zero knowledge argument.
- **Case 2:** Condition 1 holds, but Case 1 above does not hold. First, since Case 1 does not hold, the first right block commitment does not interleave the zero knowledge argument on the left (which only starts after this block commitments ends). Accordingly, it is left to establish that there exist share commitments u_i in left block commitment 1 and $\bar{u}_{\bar{i}}$ in each of the left block commitments $2, \dots, \ell + 2$ that are free with respect to the first right block commitment. This is where we use Condition 1—since the number of messages in this right block is strictly smaller than the number of shares $n, \bar{n}$ in each left block, the required free share commitments are guaranteed to exist.

- **Case 3 (applies for either one of the two conditions):** the commitment to share u_1 on the left ends after the second block of share commitments starts on the right. In this case, the commitments to shares $u_2, \dots, u_n, \bar{u}_1, \dots, \bar{u}_{\bar{n}}$, as well as the zero knowledge argument on the left are all free with respect to the first block-commitment on the right, and thus it is ideally scheduled. (This case is similar to the simplified case depicted in Fig. 2a.)
- **Case 4:** Condition 2 holds, but Case 3 above does not hold. First, since Case 3 does not hold, all the right block commitments $2, \dots, \ell + 2$ do not interleave with the commitment to share u_1 in the first left block commitment. Furthermore, one of these right blocks $\text{blk} \in \{2, \dots, \ell + 2\}$ does not interleave with the zero knowledge argument on the left (which consists of ℓ messages). To deduce that blk is ideally schedule, it is left to show that there is a free share $\bar{u}_{\bar{i}}$ in each of the left blocks $2, \dots, \ell + 2$. Here we invoke Condition 2—the number of messages in blk is strictly smaller than the number of shares $\bar{n}$ in each of the left blocks $2, \dots, \ell + 2$, the required free share commitments are again guaranteed to exist.

2.3 Tag Amplification

We now briefly overview the tag amplification process, which takes a non-malleable commitment $\langle \text{Sen}, \text{Rec} \rangle$ for $t \in [3, O(\log \lambda)]$ bit tags and transforms it into $\langle \text{Sen}, \bar{\text{Rec}} \rangle$ for $T = 2^{t-1}$ bit tags. The amplification procedure is an adaptation of existing procedures from the literature mostly similar to [KS17, ABG+20] which in turn is based on that of [DDN03]; however, unlike the first of the two, it relies on polynomial hardness assumptions, and avoids complexity leveraging, and unlike the second, it works against non-synchronizing adversaries and not only synchronizing ones.

The basic way that previous amplification schemes work is as follows: to commit to a value v , under a tag $\mathbf{tg} \in \{0, 1\}^T$ for $T = 2^{t-1}$, we consider $t - 1$ tags of the form $\mathbf{tg}_i = (i, \mathbf{tg}[i]) \in \{0, 1\}^t$ corresponding to the base scheme (here $\mathbf{tg}[i]$ is the i -th bit of $\mathbf{tg}$). The committer then sends $t - 1$ commitments to the value v in parallel under each one of the tags $\mathbf{tg}_i$, using the base protocol $\langle \text{Sen}, \text{Rec} \rangle$. Finally, a proof that all $t - 1$ commitments are consistent is added.

The basic idea behind the transformation is that if all the commitments are consistent, then in order to maul a commitment to value v under tag $\mathbf{tg}$ to a commitment to a related value $\bar{v}$ under tag $\mathbf{tg}' \neq \mathbf{tg}$, the MIM must create a commitment to $\bar{v}$ using the base protocol under tag $\mathbf{tg}'_i = (i, \mathbf{tg}'[i])$ for every $i \in [t - 1]$, by potentially mauling from some of the left commitments to v under tags $\{\mathbf{tg}_i = (i, \mathbf{tg}[i])\}_{i \in [t-1]}$. However, the fact that $\mathbf{tg} \neq \mathbf{tg}'$ means that they differ on at least one bit, that is, $\mathbf{tg}[j] \neq \mathbf{tg}'[j]$ for some j . Thus, tag $\mathbf{tg}'_j = (j, \mathbf{tg}'[j])$ on the right is different from all the tags $\{\mathbf{tg}_i = (i, \mathbf{tg}[i])\}_{i \in [t-1]}$ on the left. By the non-malleability of the base protocol, the value committed to under tag $\mathbf{tg}'_j = (j, \mathbf{tg}'[j])$ on the right must be independent of the value v committed under tags $\{\mathbf{tg}_i = (i, \mathbf{tg}[i])\}$ on the left. Given additionally that the

values committed to in all base commitments on the right are the same, the non-malleability of $\langle \widehat{\text{Sen}}, \widehat{\text{Rec}} \rangle$ with respect to $\hat{\text{tg}}, \hat{\text{tg}}'$ then follows.

In the setting of synchronizing MIM adversaries the above intuition can be formalized as expected, when the proof of consistency is instantiated with a zero-knowledge argument. In the more general setting of non-synchronizing adversaries, things become more subtle. Specifically, if the zero knowledge argument on the left interleaves with the non-malleable commitments on the right, then it is not clear how to leverage the non-malleability of the base protocol $\langle \text{Sen}, \text{Rec} \rangle$. (More specifically, we need to apply zero-knowledge simulation on the code of the verifier, which however might depend on the honest receiver Rec's code. Then, we can no longer reduce to the non-malleability of the base protocol.)

To overcome this difficulty, we rely on the Feige-Lapidot-Shamir trapdoor paradigm [FLS99]. The first receiver message in our protocol sets up a trapdoor (a solution to a hard problem), and the final proof of consistency is a witness indistinguishable (WI) proof that either: (1) the $t - 1$ commitments are consistent, or (2) the sender “knows” the trapdoor (where formally knowledge is enforced using an extractable commitment). The idea behind the FLS paradigm is that the trapdoor cannot be obtained by a sender running the protocol, and thus the validity of assertion (1) is guaranteed on right. In contrast, we would like to ensure that the reduction of non-malleability to hiding on the left would be able to obtain the trapdoor and use it in order to simulate the WI proof.

We can show that the reduction can indeed do this, but only provided certain scheduling conditions. Specifically, the trapdoor on the left should be set up before the non-malleable commitment on the right occurs. In this case, we can non-uniformly obtain the witness. To deal with the other case, we augment the protocol yet again, adding a plain non-interactive commitment to the committed value v between the trapdoor set up phase and the non-malleable commitment phase. In case the non-malleable commitment on the right starts before the trapdoor set up on the left, then in particular the plain commitment on the right occurs before any commitment was made on the left. In this case, we have a direct reduction from non-malleability to hiding, which non-uniformly obtains the value of the plain commitment on the right (this is akin to our earlier use of “first-message binding”). We refer the reader to Fig. 5 for the amplification scheme and Sect. 5 for the proof.

Robustness. One challenge in the proof above is that even in the case that we can obtain the trapdoor witness on the left, it is not immediate that non-malleability holds when the commitments on the right interleave with the proof. For this, we require that the base non-malleable commitment satisfies an extra property known as *r-robustness* [LP12]. This property essentially says that the committed value on the right can be extracted without rewinding an arbitrary r -message protocol (the WI proof in our case) executed concurrently. This allows to switch the witness used in the WI on the left, and argue that the right committed value stays the same after the switch.

We show that our base protocol (described in Sect. 2.2) is indeed robust for an appropriate choice of parameters. We further show that the tag amplification transformation described here, preserves *r-robustness*.

Two-Sided Extraction via Watrous’ Rewinding Lemma. One challenge in our analysis of both the base protocol and the tag amplification procedure is that the adversary’s scheduling of messages is *adaptive*. In particular, even though the protocol’s design guarantees that executions always contain certain *extraction opportunities*, we do not know ahead of time when they will occur. This is not a problem in the classical setting, where one can typically run first the so called main thread to identify the extraction opportunities and then rewind back to extract. However, such rewinding in the quantum setting might disturb the adversary’s state.

The analysis of our base scheme circumvents this difficulty by showing a reduction to adversaries that commit ahead of time to the timing of the so called extraction opportunities. This reduction strongly relies on the fact that the definition of non-malleability is an indistinguishability-based definition. In contrast, r -robustness is a simulation based definition—it requires a simulator that given the code of the MIM adversary can extract on the right, while interacting with an r -message protocol on the left. Let us briefly explain the difficulty in this setting.

To achieve r -robustness, we make sure there are more than r extraction opportunities on the right. Consider a simplified scenario where the MIM gives $r + 1$ extractable commitments, and we want to extract from the “free” extractable commitment that does not interleave with any of the r left messages—we refer to this as *non-interleaving extraction*. The difficulty is that the simulator does not know which extractable commitment would be “free”. If the simulator starts an extractable commitment without applying the extractor, it might miss the sole extraction opportunity. On the other hand, if it always applies the extractor, extraction may halt when the adversary expects a message on the left, and the simulator should give up extraction but still faithfully simulate the left and right interactions from here. To resolve this conundrum, we need the extractor of an extractable commitment protocol to be able to interchangeably simulate two types of interactions, ones that will eventually constitute an extraction opportunity and ones that will turn out not to be extractable due to the adversary’s scheduling.

Toward this, we prove a *two-sided simulation lemma* for extractable commitments. This lemma shows that we can always enhance the extractor so that in case the sender in the commitment prematurely aborts, not only can we simulate the sender’s state at that point, *but also the state of the receiver* (in case of abort, extraction is not required); otherwise, the extractor simulates the sender’s state and extracts the committed value as usual (without simulating the state of the receiver). Using this two-sided extractor we can deal with cases where a commitment on the right turns out not to be extractable due to scheduled messages on the left by viewing this event as a premature abort, and then using the simulated state of the receiver to faithfully continue the interaction (without extracting).

The proof of the lemma is inspired by [BS20] and uses the fact that up to the point of abort a real execution and an execution simulated by the extractor are indistinguishable. Our two-sided extractor first tosses a random coin to decide whether to simulate with extraction or to honestly simulate the receiver anticipating an abort; if the guess failed, it tries again (the expected number of trials

is negligibly close to two). While this works smoothly in the classical setting, in the quantum setting it should be done with care, as rewinding without state disturbance is typically a problem. In this specific setting, however, we meet the conditions of Watrous' *quantum rewinding lemma* [Wat09]—our extractor is guaranteed to succeed with probability close to $1/2$, obviously of the quantum internal state of the adversarial sender.

3 Preliminaries

We rely on standard notions of classical Turing machines and Boolean circuits:

- A PPT algorithm is a probabilistic polynomial-time Turing machine.
- For a PPT algorithm M , we denote by $M(x; r)$ the output of M on input x and random coins r . For such an algorithm and any input x , we write $m \in M(x)$ to denote the fact that m is in the support of $M(x; \cdot)$.

We follow standard notions from quantum computation.

- A QPT algorithm is a quantum polynomial-time Turing machine.
- An interactive algorithm M , in a two-party setting, has input divided into two registers and output divided into two registers. For the input, one register I_m is for an input message from the other party, and a second register I_a is an auxiliary input that acts as an inner state of the party. For the output, one register O_m is for a message to be sent to the other party, and another register O_a is again for auxiliary output that acts again as an inner state. For a quantum interactive algorithm M , both input and output registers are quantum.

The Adversarial Model. Throughout, efficient adversaries are modeled as quantum circuits with non-uniform quantum advice (i.e. quantum auxiliary input). Formally, a *polynomial-size adversary* $\mathcal{A} = \{\mathcal{A}_\lambda, \rho_\lambda\}_{\lambda \in \mathbb{N}}$, consists of a polynomial-size non-uniform sequence of quantum circuits $\{\mathcal{A}_\lambda\}_{\lambda \in \mathbb{N}}$, and a sequence of polynomial-size mixed quantum states $\{\rho_\lambda\}_{\lambda \in \mathbb{N}}$.

For an interactive quantum adversary in a classical protocol, it can be assumed without loss of generality that its output message register is always measured in the computational basis at the end of computation. This assumption is indeed without the loss of generality, because whenever a quantum state is sent through a classical channel then qubits decohere and are effectively measured in the computational basis.

3.1 Indistinguishability in the Quantum Setting

- Let $f : \mathbb{N} \rightarrow [0, 1]$ be a function.
 - f is negligible if for every constant $c \in \mathbb{N}$ there exists $N \in \mathbb{N}$ such that for all $n > N$, $f(n) < n^{-c}$.
 - f is noticeable if there exists $c \in \mathbb{N}$, $N \in \mathbb{N}$ such that for every $n \geq N$, $f(n) \geq n^{-c}$.

- f is overwhelming if it is of the form $1 - \mu(n)$, for a negligible function μ .
- We may consider random variables over bit strings or over quantum states. This will be clear from the context.
- For two random variables X and Y supported on quantum states, quantum distinguisher circuit D with, quantum auxiliary input ρ , and $\mu \in [0, 1]$, we write $X \approx_{D, \rho, \mu} Y$ if

$$|\Pr[D(X; \rho) = 1] - \Pr[D(Y; \rho) = 1]| \leq \mu.$$

- Two ensembles of random variables $\mathcal{X} = \{X_i\}_{\lambda \in \mathbb{N}, i \in I_\lambda}$, $\mathcal{Y} = \{Y_i\}_{\lambda \in \mathbb{N}, i \in I_\lambda}$ over the same set of indices $I = \cup_{\lambda \in \mathbb{N}} I_\lambda$ are said to be *computationally indistinguishable*, denoted by $\mathcal{X} \approx_c \mathcal{Y}$, if for every polynomial-size quantum distinguisher $D = \{D_\lambda, \rho_\lambda\}_{\lambda \in \mathbb{N}}$ there exists a negligible function $\mu(\cdot)$ such that for all $\lambda \in \mathbb{N}, i \in I_\lambda$,

$$X_i \approx_{D_\lambda, \rho_\lambda, \mu(\lambda)} Y_i.$$

For a (non-negligible) function $\varepsilon(\lambda) \in [0, 1]$, the ensembles $\mathcal{X}, \mathcal{Y}$ are ε -indistinguishable if the above requirement is replaced with

$$X_i \approx_{D_\lambda, \rho_\lambda, \varepsilon(\lambda) + \mu(\lambda)} Y_i.$$

- The trace distance between two distributions X, Y supported over quantum states, denoted $\text{TD}(X, Y)$, is a generalization of statistical distance to the quantum setting and represents the maximal distinguishing advantage between two distributions supported over quantum states, by unbounded quantum algorithms. We thus say that ensembles $\mathcal{X} = \{X_i\}_{\lambda \in \mathbb{N}, i \in I_\lambda}$, $\mathcal{Y} = \{Y_i\}_{\lambda \in \mathbb{N}, i \in I_\lambda}$, supported over quantum states, are statistically indistinguishable (and write $\mathcal{X} \approx_s \mathcal{Y}$), if there exists a negligible function $\mu(\cdot)$ such that for all $\lambda \in \mathbb{N}, i \in I_\lambda$,

$$\text{TD}(X_i, Y_i) \leq \mu(\lambda).$$

Standard Tools. Due to the lack of space, some of the basic definitions such as Witness Indistinguishability, Zero Knowledge, and Commitments, are omitted and can be found in the full version of the paper.

3.2 Non-malleable Commitments

Standard commitment schemes are defined in the full version of the paper. Let $\langle \text{Sen}, \text{Rec} \rangle$ be a commitment scheme. In an interaction between a malicious sender Sen^* and honest receiver Rec , we say that Sen^* is *non-aborting* if the Rec accepts (i.e., outputs 1) at the end of the commitment stage. Let $\text{open}_{\langle \text{Sen}, \text{Rec} \rangle}(c, v, d)$ be the function for verifying decommitments of $\langle \text{Sen}, \text{Rec} \rangle$. Define the following value function:

$$\text{val}(c) = \begin{cases} v & \text{if } \exists \text{ unique } v \text{ s.t. } \exists d, \text{open}_{\langle \text{Sen}, \text{Rec} \rangle}(c, v, d) = 1 \\ \perp & \text{otherwise} \end{cases}$$

A commitment c is *valid* if $\text{val}(c) \neq \perp$, and otherwise *invalid*.

Tag-Based Commitment Scheme. Following [DDN03, PR05b], we consider *tag-based commitment schemes* where, in addition to the security parameter, the sender and the receiver also receive a “tag”—a.k.a. the identity— tg as common input.

We recall the definition of non-malleability from [LPV08], adapted to quantum polynomial-size man-in-the-middle adversaries.

Let $\langle \text{Sen}, \text{Rec} \rangle$ be a tag-based commitment scheme, and let $\lambda \in \mathbb{N}$ be a security parameter. Consider a man-in-the-middle (MIM) adversary $\mathcal{A}$ that participates in one left and one right interactions simultaneously. In the left interactions the MIM adversary $\mathcal{A}$, on auxiliary quantum state ρ , interacts with Sen , receiving commitments to value v , using a tag $\text{tg} \in [T]$ of its choice. In the right interactions $\mathcal{A}$ interacts with Rec attempting to commit to a related value $\tilde{v}$, again using a tag $\tilde{\text{tg}}$ of length t of its choice. If the right commitment is invalid, or $\tilde{\text{tg}} = \text{tg}$, set $\tilde{v}_i = \perp$ —i.e., choosing the same tags in the left and right interactions is considered invalid. Let $\text{mim}_{\langle \text{Sen}, \text{Rec} \rangle}(\mathcal{A}, \rho, v)$ denote a random variable that describes the value $\tilde{v}$ along with the quantum output of $\mathcal{A}(\rho)$ at the end of the interaction where Sen commits to v on the left.

Definition 1. A commitment scheme $\langle \text{Sen}, \text{Rec} \rangle$ is said to be *non-malleable* if for every quantum polynomial-size man-in-the-middle adversary $A = \{A_\lambda, \rho_\lambda\}_{\lambda \in \mathbb{N}}$ and a polynomial $\ell : \mathbb{N} \rightarrow \mathbb{N}$,

$$\{\text{mim}_{\langle \text{Sen}, \text{Rec} \rangle}(A_\lambda, \rho_\lambda, v)\}_{\lambda, v, v'} \approx_c \{\text{mim}_{\langle \text{Sen}, \text{Rec} \rangle}(A_\lambda, \rho_\lambda, v')\}_{\lambda, v, v'} \quad ,$$

where $\lambda \in \mathbb{N}$ is the security parameter and $v, v' \in \{0, 1\}^{\ell(\lambda)}$ are two committed values by the honest sender.

Generally, the distributions in the MIM experiment include a quantum algorithm with a quantum auxiliary state. A standard strengthening of indistinguishability definitions for distributions of the above-mentioned type is to let the distinguisher prepare an entangled register, which is entangled with the register that contains the auxiliary state of the quantum algorithm in the distribution. In our specific case of MIM distributions the stronger definition (defined below) is equivalent as we prove next.

Definition 2 (Stronger Definition of Non-malleability). A commitment scheme $\langle \text{Sen}, \text{Rec} \rangle$ is said to be *non-malleable (with respect to entanglement)* if for every quantum polynomial-size man-in-the-middle adversary $A = \{A_\lambda\}_{\lambda \in \mathbb{N}}$ that can obtain a quantum auxiliary state, a polynomial-size quantum state $\sigma = \{\sigma_\lambda\}_{\lambda \in \mathbb{N}}$ of size at least what A obtains, and a polynomial $\ell : \mathbb{N} \rightarrow \mathbb{N}$,

$$\{\text{mim}_{\langle \text{Sen}, \text{Rec} \rangle}(A_\lambda, \sigma_{1, \lambda}, v), \sigma_{2, \lambda}\}_{\lambda, v, v'} \approx_c \{\text{mim}_{\langle \text{Sen}, \text{Rec} \rangle}(A_\lambda, \sigma_{1, \lambda}, v'), \sigma_{2, \lambda}\}_{\lambda, v, v'} \quad ,$$

where $\lambda \in \mathbb{N}$ is the security parameter, $v, v' \in \{0, 1\}^{\ell(\lambda)}$ are two committed values by the honest sender and σ_1 is the first register of the state σ such that it is in the size of the auxiliary state for A and σ_2 is the rest of the state.

Claim. Any commitment scheme $\langle \text{Sen}, \text{Rec} \rangle$ satisfying security Definition 1 also satisfy security Definition 2.

Proof. Assume $\langle \text{Sen}, \text{Rec} \rangle$ is secure with respect to Definition 1 and assume toward contradiction that it is not secure with respect to Definition 2. Let $A = \{A_\lambda\}_{\lambda \in \mathbb{N}}$ a MIM adversary and let $D = \{D_\lambda, \sigma_\lambda\}$ a distinguisher that distinguishes between,

$$\{\text{mim}_{\langle \text{Sen}, \text{Rec} \rangle}(A_\lambda, \sigma_{1,\lambda}, v), \sigma_{2,\lambda}\}_{\lambda, v, v'} \quad , \quad \{\text{mim}_{\langle \text{Sen}, \text{Rec} \rangle}(A_\lambda, \sigma_{1,\lambda}, v'), \sigma_{2,\lambda}\}_{\lambda, v, v'} \quad ,$$

for some v, v' . Consider A' a new MIM adversary: A' has quantum auxiliary state σ . The MIM execution of A' is to run A with auxiliary state σ_1 , and keep the rest of σ , which we denote by σ_2 , untouched on the side. D can thus distinguish between the distributions

$$\{\text{mim}_{\langle \text{Sen}, \text{Rec} \rangle}(A'_\lambda, \sigma_\lambda, v)\}_{\lambda, v, v'} \quad , \quad \{\text{mim}_{\langle \text{Sen}, \text{Rec} \rangle}(A'_\lambda, \sigma_\lambda, v')\}_{\lambda, v, v'} \quad ,$$

in contradiction to the security of $\langle \text{Sen}, \text{Rec} \rangle$ with respect to Definition 1.

3.3 Committed Value Oracle

Let $\langle \text{Sen}, \text{Rec} \rangle$ be a (possibly tag-based) commitment scheme. A sequential committed-value oracle $\mathcal{O}^\infty[\langle \text{Sen}, \text{Rec} \rangle]$ of $\langle \text{Sen}, \text{Rec} \rangle$ acts as follows in interaction with a sender Sen^* : it interacts with Sen^* in many *sequential* sessions; in each session,

- it participates with Sen^* in the commit phase of $\langle \text{Sen}, \text{Rec} \rangle$ as the honest receiver Rec (using a tag chosen adaptively by Sen^*), obtaining a commitment c , and
- if Sen^* is *non-aborting* in the commit phase and sends request *break*, it returns $\text{val}(c)$.

The single-session oracle $\mathcal{O}^1[\langle \text{Sen}, \text{Rec} \rangle]$ is similar to $\mathcal{O}^\infty$, except that it interacts with the adversary in a single session.

Throughout, when the commitment scheme is clear from the context, we write $\mathcal{O}^\infty, \mathcal{O}^1$ for simplicity.

3.4 Extractable Commitments

We define the standard notion of post-quantum extractable commitments (and ε -extractable) along with several enhancements of this notion. These enhancements of extractable commitments are for both the ε -extractable and (fully) extractable versions.

Definition 3. Let $\langle \text{ExCom.Sen}, \text{ExCom.Rec} \rangle$ be a (possibly tag-based) commitment scheme and $\mathcal{O}^1$ its (single-session) committed value oracle. We say that $\langle \text{ExCom.Sen}, \text{ExCom.Rec} \rangle$ is ε -extractable if there exists a QPT simulator Sim^1 , such that, for every quantum polynomial-size sender $\text{Sen}^* = \{\text{Sen}^*_\lambda, \rho_\lambda\}_{\lambda \in \mathbb{N}}$ and function $\varepsilon(\lambda) \in [0, 1]$,

– For every quantum polynomial-time distinguisher $D^* = \{D_\lambda^*, \rho_\lambda\}_{\lambda \in \mathbb{N}}$,

$$\left\{ \text{OUT}_{\text{Sen}_\lambda^*} \left(\text{Sen}_\lambda^{*\mathcal{O}^1}(\rho_\lambda) \right) \right\}_{\lambda \in \mathbb{N}} \approx_\varepsilon \left\{ \text{Sim}^1(\text{Sen}_\lambda^*, \rho_\lambda, 1^{1/\varepsilon}) \right\}_{\lambda \in \mathbb{N}} .$$

We say the scheme is (fully) extractable if there is a QPT simulator Sim^1 , such that, for every quantum polynomial-size sender $\text{Sen}^* = \{\text{Sen}_\lambda^*, \rho_\lambda\}_{\lambda \in \mathbb{N}}$,

$$\left\{ \text{OUT}_{\text{Sen}_\lambda^*} \left(\text{Sen}_\lambda^{*\mathcal{O}^1}(\rho_\lambda) \right) \right\}_{\lambda \in \mathbb{N}} \approx_c \left\{ \text{Sim}^1(\text{Sen}_\lambda^*, \rho_\lambda) \right\}_{\lambda \in \mathbb{N}} .$$

Sequential Extraction. We analogously define sequential extractability.

Definition 4. Let $\langle \text{ExCom.Sen}, \text{ExCom.Rec} \rangle$ be a (possibly tag-based) commitment scheme and $\mathcal{O}^\infty$ its sequential committed value oracle. We say that $\langle \text{ExCom.Sen}, \text{ExCom.Rec} \rangle$ is sequentially extractable if there exists a QPT simulator Sim^∞ , such that, for every quantum polynomial-size sender $\text{Sen}^* = \{\text{Sen}_\lambda^*, \rho_\lambda\}_{\lambda \in \mathbb{N}}$,

$$\left\{ \text{OUT}_{\text{Sen}_\lambda^*} \left(\text{Sen}_\lambda^{*\mathcal{O}^\infty}(\rho_\lambda) \right) \right\}_{\lambda \in \mathbb{N}} \approx_c \left\{ \text{Sim}^\infty(\text{Sen}_\lambda^*, \rho_\lambda) \right\}_{\lambda \in \mathbb{N}} .$$

Sequential ε -extractability is defined analogously when considering ε -indistinguishability instead of (plain) computational indistinguishability.

Constructions of post-quantum extractable commitments with have been known for a while either in polynomially many rounds assuming post-quantum oblivious transfer [HSS15, LN11] or in constant rounds assuming Learning with Errors in quantum fully homomorphic encryption [BS20]. More recently Chia et al. [CCLY21] constructed post-quantum ε -extractable commitments with in constant rounds, assuming the existence of post-quantum one-way functions. (Lombardi, Ma, and Spooner [LMS21] also construct such commitments, but relying super-polynomial hardness of the one-way functions.)

These constructions address the single-session oracle. However, a standard proof shows that sequential extraction follows.

Lemma 1. Any extractable commitment is sequentially extractable. This applies also for ε -extractability.

r -Robustness. The work of [LP12] introduced the notion of r -robustness w.r.t. committed value oracle, following similar notions of r -robustness introduced in [CLP16, LP09]. We here recall their definition, adapted to working with quantum polynomial-size adversaries. Let $\langle \text{Sen}, \text{Rec} \rangle$ be a (possibly tag-based) commitment scheme. Consider a man-in-the-middle adversary that participates in an arbitrary left interaction with a limited number r of rounds, while having access to the committed value oracle $\mathcal{O}^\infty[\langle \text{Sen}, \text{Rec} \rangle]$. Roughly speaking, $\langle \text{Sen}, \text{Rec} \rangle$ is r -robust if the output of $\mathcal{A}$ in any r -round interaction, with access to the oracle $\mathcal{O}^\infty[\langle \text{Sen}, \text{Rec} \rangle]$, can be simulated without the oracle. In other words, having access to the oracle does not help the adversary in breaking the security in any r -round protocol much.

Definition 5 (r -robust extraction). Let $\langle \text{Sen}, \text{Rec} \rangle$ be a (possibly tag-based) commitment scheme. We say that $\langle \text{Sen}, \text{Rec} \rangle$ is r -robust w.r.t. the committed-value oracle, if there exists a QPT simulator Sim_r , such that, for every QPT adversary $\mathcal{A} = \{\mathcal{A}_\lambda, \rho_\lambda\}_{\lambda \in \mathbb{N}}$, the following holds:

– *Simulation:* For every PPT r -round machine B ,

$$\left\{ \text{OUT}_{\mathcal{A}_\lambda} \langle B(z, 1^\lambda), A_\lambda^{\mathcal{O}^\infty[\langle \text{Sen}, \text{Rec} \rangle]}(\rho_\lambda) \rangle \right\}_{\lambda \in \mathbb{N}, z \in \{0,1\}^*} \\ \approx_c \left\{ \text{OUT}_{\text{Sim}} \langle B(z, 1^\lambda), \text{Sim}_r(A_\lambda, \rho_\lambda) \rangle \right\}_{\lambda \in \mathbb{N}, z \in \{0,1\}^*} .$$

(ε, r) -robustness is defined analogously when considering ε -indistinguishability instead of (plain) computational indistinguishability.

First-Message Binding. We define an additional property of extractable commitments which will come in handy later in the construction of post-quantum non-malleable commitments. The property, which we call first-message binding, asserts that the first message of the sender determines the committed value. Additionally, if the first message in the extractable commitment protocol is a receiver message, then the extractor simulates it honestly, in particular, independently of the malicious sender's circuit.

Definition 6. Let $\langle \text{ExCom.Sen}, \text{ExCom.Rec} \rangle$ be an extractable commitment scheme. We say that the scheme has first-message binding if:

1. With overwhelming probability over the choice of the honest receiver randomness, the first sender message in the protocol fixes the committed value.
2. If the first message in the protocol is a receiver message, in a simulated session, the extractor ExCom.Ext samples this message by invoking the honest receiver (independently of the malicious sender circuit).

We observe that every extractable commitment can easily be turned into an extractable commitment with first-message binding. A proof sketch is provided in supplemental material.

Lemma 2. Let $\langle \text{ExCom.Sen}, \text{ExCom.Rec} \rangle$ be an extractable commitment scheme. Then there exists an extractable commitment scheme $\langle \text{Sen}, \text{Rec} \rangle$ with first-message binding. Furthermore, the sequential extractor Sim^∞ for the scheme also satisfies Property 2 in the above definition. The same also holds for ε -extractability.

3.5 Two-Sided Extraction

In this section, we state a *two-sided extraction lemma* for any extractable commitment. We then use it to prove a *non-interleaving extraction lemma*, which we later rely on.

Two-Sided Extractor. We define the following variant $\mathcal{O}_\perp^1$ of the committed value oracle $\mathcal{O}^1$. Recall that $\mathcal{O}^1$ participates in a session of the commit phase of $\langle \text{ExCom.Sen}, \text{ExCom.Rec} \rangle$ with Sen^* , acting as the honest receiver ExCom.Rec . If Sen^* is non-aborting in the commit phase and requests **break**, $\mathcal{O}^1$ returns the value $\text{val}(c)$ committed in the produced commitment c .

$\mathcal{O}_\perp^1$ does the same, except that in the case that Sen^* aborts, it sends back the internal state of the honest receiver ExCom.Rec in that session. That is,

$$\mathcal{O}_\perp^1 \text{ returns } \begin{cases} \text{internal state of ExCom.Rec} & \text{if Sen}^* \text{ aborts} \\ \text{val}(c) & \text{if Sen}^* \text{ is non-aborting in } c \text{ and requests break} \\ \text{nothing} & \text{otherwise} \end{cases}$$

In the full version of this work, we prove that every extractable commitment satisfies such *two-sided extractability*:

Claim. Let $\langle \text{ExCom.Sen}, \text{ExCom.Rec} \rangle$ be an extractable commitment scheme and $\mathcal{O}_\perp^1$ its enhanced committed value oracle. There exists a QPT simulator $\text{Sim}_\perp^1$, such that, for every quantum polynomial-size sender $\text{Sen}^* = \{\text{Sen}_\lambda^*, \rho_\lambda\}_{\lambda \in \mathbb{N}}$, the following two ensembles are computationally indistinguishable,

$$\left\{ \text{OUT}_{\text{Sen}_\lambda^*} \left(\text{Sen}_\lambda^{*\mathcal{O}_\perp^1}(\rho_\lambda) \right) \right\}_{\lambda \in \mathbb{N}} \approx_c \left\{ \text{Sim}_\perp^1(\text{Sen}_\lambda^*, \rho_\lambda) \right\}_{\lambda \in \mathbb{N}}.$$

The same also holds for ε -extractability.

ε -Extractability vs Full Extractability. To simplify notation, the technical sections in this extended abstract are based on fully extractability (and corresponding indistinguishability) rather than full extractability. As mentioned in the introduction, the transition between the two is quite direct. In more detail, our final goal is to achieve an indistinguishability-based definition of non-malleability. The proof toward that is based on a fixed polynomial number $h(\lambda) = \text{poly}(\lambda)$ of hybrid distributions that depends only on the security parameter. Thus when relying on indistinguishability between a simulated execution and a real execution, the corresponding indistinguishability between hybrids is only ε indistinguishability. Accordingly, for any polynomial $p(\lambda)$, to overall obtain $1/p(\lambda)$ indistinguishability, we can set $\varepsilon = 1/(h(\lambda) \cdot p(\lambda))$. All corresponding simulators still run in polynomial time, and hence all intermediate reductions still hold.

4 Post-quantum Non-malleable Commitment for Few Tags

In this section, we present our construction of a classical post-quantum non-malleable commitment protocol with at most a logarithmic number of tags τ . It makes use of A quantumly-extractable classical commitment scheme (ExCom.Sen, ExCom.Rec) with first-message binding, and a post-quantum classical zero-knowledge argument (P, V). We describe the protocol in Fig. 4.

Using post-quantum ε -extractable commitments with k rounds one can obtain post-quantum ε -zero-knowledge arguments with $k + O(1)$ rounds [Ros04, BS20]. It follows that the number of rounds in Protocol 4 is $k^{O(\tau)}$. Statistical binding of the commitment scheme follows readily from the statistical binding of the extractable commitment scheme. Hiding of any commitment scheme follows directly from non-malleability, so it remains for us to show that our commitment protocol is non-malleable. Later, we also show that our commitment scheme satisfies r -robustness, a property of the commitment protocol which we use in our tag amplification scheme in Sect. 5.

Proposition 1. *The protocol in Fig. 4 is non malleable.*

4.1 Ideally-Scheduled Block Commitments

Before turning to prove Proposition 1, we state and prove a combinatorial claim regarding the structure of executions. We first fix relevant terminology for addressing different parts of the protocol.

Block Commitments. For $m, N \in \mathbb{N}$, a block commitment of length N and sub-block length m for a string $s = s_1, s_2, \dots, s_N \in \{0, 1\}^{m \times N}$ (such that $\forall i \in s_i \in \{0, 1\}^m$) consists of N sequential extractable commitment to each of the strings $s_1, \dots, s_N$ in their respective order. In particular, note that in Phase 1 of our Protocol 4, the sender gives one block commitment of length n with sub-block length $|v|$ to $\mathbf{u} = (u_1, \dots, u_n)$ and $\ell + 1$ block commitments to $\bar{\mathbf{u}} = (\bar{u}_1, \dots, \bar{u}_{\bar{n}})$, each of length $\bar{n}$ and sub-block length $|v|$.

Ideally Scheduled Block Commitments. Consider a two-sided MIM execution of Protocol 4; that is, the MIM adversary $\mathcal{A}$ interacts with Sen on the left and Rec on the right.

We call an execution of a block commitment on the left *free* on index i with respect to a given block commitment on the right, if interaction during the i -th extractable commitment in that block commitment does not interleave with the interaction during the given right block commitment. We call an execution of a block commitment on the right *free* if it does not interleave with the interaction during Phase 2 of the protocol on the left.

An execution I of a block commitment on the right is *ideally scheduled* if all of the above hold:

- It is free (with respect to the second phase on the left).

Protocol 6

Parameters: λ is the security parameter. r is the robustness parameter. k is the total number of messages in the extractable commitment protocol. $\tau \leq O(\log_k(\lambda))$ is the number of tags. ℓ is the maximum between (1) the robustness parameter r , and (2) the total number of messages in the zero knowledge argument system.

Common input: Security parameter $\lambda \in \mathbb{N}$, robustness parameter $r \leq \text{poly}(\lambda)$, an identification tag $\text{tg} \in [\tau]$ for the sender.

Sender private input: A value $v \in \{0, 1\}^*$ to commit to.

Phase 1: Commitments to Secret Shares of Value:

- Let $n := (k + 1)^{\text{tg}}$ and $\bar{n} := (k + 1)^{2\tau - \text{tg}}$.
- **Sen** secret-shares the value v twice, first into n shares and second into $\bar{n}$ shares: $\mathbf{u} = (u_1, \dots, u_n)$ and $\bar{\mathbf{u}} = (\bar{u}_1, \dots, \bar{u}_{\bar{n}})$, respectively.
- **Sen** provides extractable commitments to the two sequences of shares:
 1. An extractable commitment to u_i , for every $i \in [n]$, sequentially, one after the other.
 2. An extractable commitment to $\bar{u}_i$, for every $i \in [\bar{n}]$, sequentially, one after the other. This sequential commitment to $\bar{\mathbf{u}}$ is repeated $\ell + 1$ times, sequentially.

Phase 2: Zero-knowledge Argument of Consistency: The protocol ends with **Sen** giving a ZK argument that its generated transcript is consistent; namely, there exists private input and randomness for the honest sender inducing the transcript.

Decommitment. If the interaction ends in an accepting proof, the decommitment information includes the shares $u_1, \dots, u_n$ along with the decommitment information for each of their corresponding extractable commitments. The decommitment verification algorithm checks that the shares yield the value v and then runs the decommitment verification algorithm of the extractable commitment on each of the shares and its decommitment information. If the ZK argument is not accepting, or the sender prematurely aborts, the verification algorithm rejects, regardless of the decommitment information given.

Fig. 4. A τ -tag post-quantum non-malleable commitment (**Sen**, **Rec**).

- There is some index $i \in [n]$ such that the block commitment to $\mathbf{u}$ on the left is free on index i with respect to I .
- There is some index $j \in [\bar{n}]$ such that *all* $\ell + 1$ block commitments to $\bar{\mathbf{u}}$ on the left are free on *the same* index j with respect to I .

In case, the MIM adversary aborts, we assume w.l.o.g it keeps sending messages $\perp$ according to some schedule, so that the above notion is always defined. The proof of the following claim is provided in the full version of this work.

Claim. In every MIM execution of Protocol 4 with tag $\mathbf{tg}$ on the left and tag $\tilde{\mathbf{tg}}$ on the right, if $\mathbf{tg} \neq \tilde{\mathbf{tg}}$, there is an ideally scheduled execution of a block commitment on the right.

4.2 Adversaries with Predetermined Ideal Schedule

Before proving Proposition 1, we prove a lemma that basically says that we can restrict attention to MIM adversaries that always announce ahead of time the structure of the ideal schedule. This lemma will later simplify our proof of Proposition 1.

In what follows, let N be a bound on the size of $n := (k + 1)^{\mathbf{tg}}, \bar{n} := (k + 1)^{2\tau - \mathbf{tg}}$, for every possible $\mathbf{tg}$. We consider *configurations* of the form

$$C = (i, c, \bar{c}, w) \in [\ell + 2] \times [N] \times [N] \times \{\mathbf{IP}_2, \mathbf{P}_2\mathbf{I}\} \quad .$$

We say that a given MIM execution is consistent with such a configuration C if:

- The i -th block commitment on the right is the first ideally scheduled block.
- The commitment to u_c (in the first block) on the left is free with respect to the ideally scheduled block i .
- The commitment to $\bar{u}_{\bar{c}}$ in every one of the blocks $2, \dots, \ell + 2$ on the left is free with respect to the ideally scheduled block i .
- If the first ideally scheduled block i on the right ends before Phase 2 on the left begins, $w = \mathbf{IP}_2$. Otherwise (Phase 2 on the left begins before the first ideally scheduled block i has ended), $w = \mathbf{P}_2\mathbf{I}$. Note that in case $w = \mathbf{P}_2\mathbf{I}$, due to the fact that block i on the right is ideally scheduled and in particular is continuous with respect to Phase 2 on the left, we can also say that block i on the right *begins* after Phase 2 on the left has started (rather than say that it only ends after the beginning of Phase 2 on the left).

Note that the number of possible configurations is bounded by $\Delta := (\ell + 2) \times N \times N \times 2 = \text{poly}(\lambda)$.

Definition 7 (MIM with predetermined ideal schedule). A MIM QPT adversary $\mathcal{A} = \{\mathcal{A}_\lambda, \rho_\lambda\}_\lambda$ has a predetermined ideal schedule $C = \{C_\lambda\}_\lambda$, if any execution in which $\mathcal{A}_\lambda$ participates is consistent with configuration C_λ .

Lemma 3. If the protocol in Fig. 4 is secure against MIM QPT adversaries with predetermined ideal schedule, then it is also secure against arbitrary MIM QPT adversaries.

Proof. Given an arbitrary MIM QPT $\mathcal{A}$ and QPT distinguisher D that break non-malleability for some values v, v' with advantage δ , we construct an MIM QPT adversary with predetermined schedule, which breaks the scheme with probability δ/Δ .

Consider an adversary $\mathcal{A}'$ that first samples uniformly at random a configuration $C \leftarrow [\ell + 2] \times [N] \times [N] \times \{\text{IP}_2, \text{P}_2\text{I}\}$. It then emulates $\mathcal{A}$, and if at any point the execution is about to become inconsistent with C , $\mathcal{A}'$ stops emulating $\mathcal{A}$, completes the execution consistently with C , and eventually outputs $\perp$. If the emulation of $\mathcal{A}$ is completed consistently with C , $\mathcal{A}'$ outputs the same as $\mathcal{A}$.

Then, since every execution has an ideally scheduled block (Claim 4.1), $\mathcal{A}'$ breaks non-malleability with probability exactly δ/Δ (with respect to the same distinguisher D and v, v'). Finally, by an averaging argument, we fix the choice of $\mathcal{A}'$ for a configuration to be the configuration C that maximizes D 's distinguishing advantage. We obtain a corresponding MIM with predetermined ideal schedule with the same advantage δ/Δ .

4.3 Proof of Proposition 1

We prove the Proposition by a hybrid argument, specifically, we show that the MIM experiment output distribution for any value v on the left is indistinguishable from an experiment independent of v . Following Lemma 3, we restrict attention to a MIM adversary with a predetermined ideal schedule $C = (i, c, \bar{c}, w)$.

$\mathcal{H}_0$: The original MIM experiment output. This includes the output of the MIM adversary in the experiment and the committed value on the right.

$\mathcal{H}_1$: Inefficient extraction from ideally-scheduled block. In this hybrid, instead of the committed value $\tilde{v}$ on the right, we consider the value $\tilde{v}_1$ reconstructed from the shares of the ideally scheduled block i on the right. If the value of any of the commitments to these shares is $\perp$, we set $\tilde{v}_1 = \perp$. $\mathcal{H}_0$ and $\mathcal{H}_1$ are statistically indistinguishable following from the soundness of the ZK argument that $\mathcal{A}$ gives to the receiver on the right in Phase 2.

$\mathcal{H}_2$: Alternative description via oracle extraction. In this hybrid we consider an augmented adversary $\mathcal{A}_2^{\mathcal{O}^\infty}$, which is given access to the sequential committed-value oracle $\mathcal{O}^\infty = \mathcal{O}^\infty[\text{ExCom.Sen}, \text{ExCom.Rec}]$ and acts as follows:

- $\mathcal{A}_2$ emulates $\mathcal{A}$. On the left, $\mathcal{A}_2$ relays all messages between $\mathcal{A}$ and the sender. On the right,
 - During the ideally scheduled block i , $\mathcal{A}_2$ interacts with its oracle $\mathcal{O}^\infty$, in every extractable commitment. Recall that $\mathcal{O}^\infty$ acts as the honest receiver, and answers **break** requests with the corresponding committed value. $\mathcal{A}_2$ submits such a **break** request after each of the commitments and stores the received share value.
 - In any other (than i) block in Phase 1, $\mathcal{A}_2$ internally emulates the receiver on the right.
 - In Phase 2, $\mathcal{A}_2$ internally emulates the zero knowledge verifier on the right.

- Eventually, $\mathcal{A}_2$ outputs the output of $\mathcal{A}$ as well as the value $\tilde{v}_1$ reconstructed from the ideal block shares obtained from the oracle $\mathcal{O}^\infty$.

The output of this hybrid is the output of $\mathcal{A}_2$. It follows directly from the construction of $\mathcal{A}_2^{\mathcal{O}^\infty}$ and the definition of $\mathcal{O}^\infty$ that $\mathcal{H}_1 \equiv \mathcal{H}_2$.

$\mathcal{H}_3$: Efficient extraction on the right when $w = \text{P}_2\text{I}$. This hybrid, differs from the previous hybrid only if $w = \text{P}_2\text{I}$; namely, Phase 2 on the left begins before the ideally scheduled block commitment i on the right had started. In such executions, for the ideally scheduled block commitment i , we perform sequential extraction to obtain the corresponding shares.

In more detail, let ψ be the (quantum) state of $\mathcal{A}_2$ when it initiates the ideally scheduled block i on the right, and let $\bar{\mathcal{A}}_2^{\mathcal{O}^\infty}$ be the adversary that starting from ψ , emulates $\mathcal{A}_2^{\mathcal{O}^\infty}$ during block i and outputs its state at the end (Note that since block i is ideally scheduled and also starts after Phase 2 on the left, it follows that $\bar{\mathcal{A}}_2$ does not perform *any* interaction on the left during the right block i).

In $\mathcal{H}_3$, we consider another augmented adversary $\mathcal{A}_3$ that acts like $\mathcal{A}_2$, only that instead of executing $\bar{\mathcal{A}}_2^{\mathcal{O}^\infty}$ during block i , it invokes the sequentially-extracting simulator $\text{Sim}^\infty(\bar{\mathcal{A}}_2, \psi)$, given by Lemma 1, which eliminates the use of the commitment oracle $\mathcal{O}^\infty$. Computational indistinguishability of $\mathcal{H}_2$ and $\mathcal{H}_3$ follows directly from the sequential extraction guarantee (Lemma 1).

$\mathcal{H}_4$: Simulating the ZK argument on the left. In this hybrid, the ZK argument on the left is generated by the zero knowledge simulator.

Specifically, let ψ be the state of $\mathcal{A}_3$ when the zero-knowledge argument is initiated on the left. We consider the zero-knowledge verifier $\mathbf{V}^*$ that starting from ψ emulates $\mathcal{A}_3$ in the rest of the interaction while forwarding its messages on the left to the zero-knowledge prover, and eventually outputs the same. In particular, if $w = \text{P}_2\text{I}$ then the code of $\mathbf{V}^*$ includes the code of the simulator Sim^∞ , which is applied to $(\bar{\mathcal{A}}_2, \psi)$ as part of the execution of $\mathcal{A}_3$. Note that in both cases $w = \text{IP}_2$ and $w = \text{P}_2\text{I}$, once Phase 2 on the left starts, $\mathcal{A}_3$ no longer makes oracle calls to $\mathcal{O}^\infty$, so the code of $\mathbf{V}^*$ is fully specified and executes in polynomial time.

In $\mathcal{H}_4$, we consider an augmented adversary $\mathcal{A}_4$ that acts as $\mathcal{A}_3$, only that when Phase 2 starts on the left, instead of executing $\mathbf{V}^*$ and interacting on the left with the zero knowledge prover, $\mathcal{A}_4$ runs the zero knowledge simulator $\text{Sim}(\mathbf{V}^*, \psi)$, and outputs the same.

$\mathcal{H}_3 \approx_c \mathcal{H}_4$. This is because by construction, the output of $\mathbf{V}^*$ is identically distributed to the output of $\mathcal{H}_3$. Computational indistinguishability of $\mathcal{H}_3$ and $\mathcal{H}_4$ now follows from the zero knowledge simulation guarantee (we note that any use of the inefficient oracle $\mathcal{O}^\infty$, in case $w = \text{IP}_2$, occurs before Phase 2 on the left, and can thus be non-uniformly fixed).

$\mathcal{H}_5$ and $\mathcal{H}_6$: Interchangeably, changing left committed values and efficient extraction threshold. As a preliminary high-level explanation to the next step, at this point in our hybrid distributions, we consider the $1 + (\ell + 1)$ block commitments given to the MIM adversary on the left, and in each block, we'll switch a commitment for a secret share (of v), to a commitment for a string of zeros. For this, we will need to use the computational hiding property of the extractable commitments. The point, however, is to be able to use the hiding of the extractable commitments while still being able to *efficiently* extract the value $\tilde{v}_1$ from the right interaction with the MIM adversary³.

Formally, we next define two sequences of hybrids $\mathcal{H}_{5,j}$ and $\mathcal{H}_{6,j}$ (for $j \in [\ell + 3]$) that augment one another interchangeably:

$$\mathcal{H}_4 = \mathcal{H}_{5,\ell+3} \rightarrow \mathcal{H}_{6,\ell+2} \rightarrow \mathcal{H}_{5,\ell+2} \rightarrow \cdots \rightarrow \mathcal{H}_{5,2} \rightarrow \mathcal{H}_{6,1} \rightarrow \mathcal{H}_{5,1} .$$

In what follows, recall that $\mathcal{A}_4$ in $\mathcal{H}_4$ is following a predetermined ideal schedule $C = (i, c, \bar{c}, w)$.

$\mathcal{H}_{5,j}$, for $j = \ell + 3, \dots, 1$: Swapping one more free commitment to zeros. In this hybrid, we simulate the most bottom free commitment on the left. Formally:

- $\mathcal{H}_{5,\ell+3}$ is defined as $\mathcal{H}_4$.
- For $j \leq \ell + 2$, $\mathcal{H}_{5,j}$ is defined exactly as $\mathcal{H}_{6,j}$, except that the left extractable commitment c_j (to share u_c or $\bar{u}_{\bar{c}}$) in the left block j is replaced with a commitment to $0^{|v|}$.

$\mathcal{H}_{6,j}$, for $j = \ell + 2, \dots, 1$: Raising the threshold for efficient extraction. Recall $\mathcal{A}_4$ in $\mathcal{H}_4$ interacts with the sender in Phase 1 on the left and in case $w = \text{IP}_2$, namely, the ideally scheduled block on the right ends before Phase 2 on the left begins, $\mathcal{A}_4$ interacts with the sequential commitment oracle $\mathcal{O}^\infty$ on the right during block i . For a left block index $j \in [\ell + 2]$, we denote by c_j the corresponding free extractable commitment; namely, $c_j = c$ if $j = 1$, and $c_j = \bar{c}$ if $j \geq 2$.

Informally, in hybrid $\mathcal{H}_{6,j}$, we move to simulating the oracle $\mathcal{O}^\infty$ in any right extractable commitment that starts after the free left commitment c_j . Formally, $\mathcal{H}_{6,j}$ is different from $\mathcal{H}_{5,j+1}$ only if $w = \text{IP}_2$. In this case, we consider an augmented adversary $\mathcal{A}_{6,j}$ defined as follows for $j \in [\ell + 2]$:

- $\mathcal{A}_{6,j}$ acts as $\mathcal{A}_{6,j+1}^{\mathcal{O}^\infty}$ until the first right extractable commitment (in the ideally scheduled right block i) in which the first sender message is sent after the free left commitment c_j .
- $\mathcal{A}_{6,j}$ simulates the remaining calls to $\mathcal{O}^\infty$ as follows:
 - Let ψ be the state of $\mathcal{A}_{6,j+1}^{\mathcal{O}^\infty}$ at the abovementioned point, just before the right extractable commitment begins.

³ Recall that currently, if $w = \text{IP}_2$, we extract $\tilde{v}_1$ inefficiently using the sequential committed-value oracle $\mathcal{O}^\infty = \mathcal{O}^\infty[\text{ExCom.Sen}, \text{ExCom.Rec}]$. If $w = \text{P}_2\text{I}$ we don't have this problem, as the ideally scheduled right block commitment i starts after the beginning of Phase 2 on the left.

- Let $\bar{\mathcal{A}}_{6,j+1}^{\mathcal{O}^\infty}$ be the adversary that starting from ψ emulates $\mathcal{A}_{6,j+1}^{\mathcal{O}^\infty}$ in the following right extractable commitments, up to those that are already simulated, while internally emulating the sender in any left commitment.
 - $\mathcal{A}_{6,j}$ invokes the sequentially-extracting simulator $\text{Sim}^\infty(\bar{\mathcal{A}}_{6,j+1}, \psi)$ to remove the use $\mathcal{O}^\infty$.
- $\mathcal{A}_{6,j}$ then completes the execution as $\mathcal{A}_{6,j+1}$ and outputs the same.

In the full version of this work, we prove the following claim, which concludes Proposition 1,

Claim. 1) The output of $\mathcal{H}_{5,1}$ is independent of the committed value v . 2) $\forall j \in [\ell + 2] : \mathcal{H}_{5,j+1} \approx_c \mathcal{H}_{6,j}$. 3) $\forall j \in [\ell + 2] : \mathcal{H}_{6,j} \approx_c \mathcal{H}_{5,j}$.

Due to space limits, we prove that our protocol is robust in the full version of this work.

5 Tag Amplification

In this section, we present a tag amplification transformation that converts a non-malleable commitment scheme $\langle \text{Sen}, \text{Rec} \rangle$ for $t \in [3, O(\log(\lambda))]$ bit tags into a non-malleable commitment scheme $\langle \widehat{\text{Sen}}, \widehat{\text{Rec}} \rangle$ for $T = 2^{t-1}$ bit tags. The transformation can be applied iteratively to amplify the number of tags from constant to exponential in the security parameter λ ,

The transformation uses the following ingredients: 1) A post-quantum secure one-way function f . 2) Naor's 2-message statistically binding commitment [Nao91] instantiated with a post-quantum secure pseudo-random generator, which in turn can be based on post-quantum one-way functions. The receiver of Naor's protocol is public coin and sends a random string a as the first message, the sender then responds with $c = \text{Com}_a(m; d)$ depending on a ; the decommitment is simply sender's private random coins. The receiver can reuse a across many commitments sent to it, and we can effectively use the second message of Naor's commitments as a non-interactive commitment. 3) A post-quantum secure ε -extractable commitment scheme ECom. Let k_1 be the number of rounds in this commitment scheme. 4) A post-quantum secure WI protocol which can be based on any post-quantum one-way functions. Let k_2 be the number of rounds of WI. 5) A non-malleable commitment scheme $\langle \text{Sen}, \text{Rec} \rangle$ for $t \geq 3$ bit tags that is also r -robust for $r = k_1 + k_2$. Let n be the length of messages the scheme can commit to. The transformed non-malleable commitment $\langle \widehat{\text{Sen}}, \widehat{\text{Rec}} \rangle$ for $T = 2^{t-1}$ tags is presented in Fig. 5.

In the full version of this work, we show that $\langle \widehat{\text{Sen}}, \widehat{\text{Rec}} \rangle$ is statistically binding, r -robust and post-quantum non-malleable as well as the detailed analysis of the complexity growth and security loss.

Protocol 7

Common Input: Security parameter $\lambda \in \mathbb{N}$ and a tag $\hat{\mathbf{t}}\mathbf{g} \in \{0,1\}^T$ for the sender, where $T = 2^{t-1}$.

Sen's private input: A message $m \in \{0,1\}^n$ to commit to.

1. **Trapdoor Setup:** $\widehat{\text{Rec}}$ sends two random images $y_1 = f(u_1)$ and $y_2 = f(u_2)$ of the one-way function f , where $u_1 \leftarrow \{0,1\}^\lambda$, $u_2 \leftarrow \{0,1\}^\lambda$. $\widehat{\text{Rec}}$ proves using WI that either y_1 or y_2 is in the image of f for λ -bit inputs. We refer to u_1 and u_2 as the trapdoors.
2. **Initial Commitment:** $\widehat{\text{Rec}}$ sends the first message a of Naor's commitment. $\widehat{\text{Sen}}$ commits to m using Naor's commitment $c = \text{Com}_a(m; d)$ w.r.t. receiver's message a , using random coins d .
3. **(Sen, Rec) commitments:** For every bit $\hat{\mathbf{t}}\mathbf{g}_i$ in the $T = 2^{t-1}$ bit tag $\hat{\mathbf{t}}\mathbf{g}$, define tag $\mathbf{t}\mathbf{g}_i = (i, \hat{\mathbf{t}}\mathbf{g}_i)$, which has exactly t bits. For every $i \in [T]$, $\widehat{\text{Sen}}$ commits to m using $\langle \text{Sen}, \text{Rec} \rangle$ and tag $\mathbf{t}\mathbf{g}_i$; let c_i denote the produced commitment and d_i the decommitment. All commitments are sent in parallel.
4. **Proof of Consistency:** $\widehat{\text{Sen}}$ first commits to 0^λ using the extractable commitment scheme ECom . Let c_e denote the produced commitment. $\widehat{\text{Sen}}$ proves using WI that either $c, c_1, \dots, c_T$ are all valid commitments to m , or c_e commits to a preimage of y_1 or y_2 . Formally, it proves that the statement $X = (a, c, c_1, \dots, c_T, c_e, y_1, y_2)$ belongs to the language $\mathcal{L}$ defined by the following witness relation:

$$\begin{aligned} \mathcal{R}_{\mathcal{L}}(X, W) = (m, d, d_1, \dots, d_T, d_e, u) = 1 \text{ iff} \\ \text{Either } c = \text{Com}_a(m; d) \wedge \forall i \in [T], \text{open}_{\langle \text{Sen}, \text{Rec} \rangle}(c_i, m, d_i) = 1, \\ \text{Or } \text{open}_{\text{ECom}}(c_e, u, d_e) = 1 \text{ and } (y_1 = f(u) \text{ or } y_2 = f(u)) \end{aligned}$$
5. **Receiver's Decision:** $\widehat{\text{Rec}}$ accepts the commitment iff the proof of consistency is accepting.
6. **Decommitment:** $\widehat{\text{Sen}}$ outputs decommitment d . The decommitment is accepted if $c = \text{Com}_a(m; d)$.

Fig. 5. Post-quantum tag amplification.

Acknowledgements. Nir Bitansky is a Member of the Check Point Institute of Information Security, supported by ISF grants 18/484 and 19/2137, by Len Blavatnik and the Blavatnik Family Foundation, by the Blavatnik Interdisciplinary Cyber Research Center at Tel Aviv University, and by the European Union Horizon 2020 Research and Innovation Program via ERC Project REACT (Grant 756482). Huijia is supported by NSF grant CNS-1936825 (CAREER), CNS-2026774, a Hellman Fellowship, a JP Morgan AI Research Award, a Simons Collaboration grant on the Theory of Algorithmic Fairness. Omri Shmueli is supported by a Clore Fellowship, ISF grants 18/484 and 19/2137, by Len

Blavatnik and the Blavatnik Family Foundation, and by the European Union Horizon 2020 Research and Innovation Program via ERC Project REACT (Grant 756482).

The authors are grateful to Fang Song for valuable discussions.

References

- ABG+20. Agarwal, A., Bartusek, J., Goyal, V., Khurana, D., Malavolta, G.: Post-quantum multi-party computation in constant rounds. *CoRR*, abs/2005.12904 (2020)
- AP19. Ananth, P., La Placa, R.L.: Secure quantum extraction protocols. *IACR Cryptol. ePrint Arch.* **2019**, 1323 (2019)
- Bar02. Barak, B.: Constant-round coin-tossing with a man in the middle or realizing the shared random string model, pp. 345–355 (2002)
- BG08. Barak, B., Goldreich, O.: Universal arguments and their applications. *SIAM J. Comput.* **38**(5), 1661–1694 (2008)
- BL18. Bitansky, N., Lin, H.: One-message zero knowledge and non-malleable commitments. In: Beimel, A., Dziembowski, S. (eds.) *TCC 2018, Part I*. LNCS, vol. 11239, pp. 209–234. Springer, Cham (2018). https://doi.org/10.1007/978-3-030-03807-6_8
- BS20. Bitansky, N., Shmueli, O.: Post-quantum zero knowledge in constant rounds. In: Makarychev, K., Makarychev, Y., Tulsiani, M., Kamath, G., Chuzhoy, J. (eds.) *Proceedings of the 52nd Annual ACM SIGACT Symposium on Theory of Computing, STOC 2020, Chicago, IL, USA, 22–26 June 2020*, pp. 269–279. ACM (2020)
- CCLY21. Chia, N.-H., Chung, K.-M., Liang, X., Yamakawa, T.: Post-quantum simulatable extraction with minimal assumptions: black-box and constant-round. *arXiv preprint* [arXiv:2111.08665](https://arxiv.org/abs/2111.08665) (2021)
- CCY20. Chia, N.-H., Chung, K.-M., Yamakawa, T.: Black-box approach to post-quantum zero-knowledge in constant round. *arXiv preprint* [arXiv:2011.02670](https://arxiv.org/abs/2011.02670) (2020)
- CLP16. Canetti, R., Lin, H., Pass, R.: Adaptive hardness and composable security in the plain model from standard assumptions. *SIAM J. Comput.* **45**(5), 1793–1834 (2016)
- CMSZ21. Chiesa, A., Ma, F., Spooner, N., Zhandry, M.: Post-quantum succinct arguments. *arXiv preprint* [arXiv:2103.08140](https://arxiv.org/abs/2103.08140) (2021)
- COSV16. Ciampi, M., Ostrovsky, R., Siniscalchi, L., Visconti, I.: Concurrent non-malleable commitments (and more) in 3 rounds. In: Robshaw, M., Katz, J. (eds.) *CRYPTO 2016*. LNCS, vol. 9816, pp. 270–299. Springer, Heidelberg (2016). https://doi.org/10.1007/978-3-662-53015-3_10
- CR87. Chor, B., Rabin, M.O.: Achieving independence in logarithmic number of rounds. In: *Proceedings of the Sixth Annual ACM Symposium on Principles of Distributed Computing*, Vancouver, British Columbia, Canada, 10–12 August 1987, pp. 260–268 (1987)
- DDN03. Dolev, D., Dwork, C., Naor, M.: Nonmalleable cryptography. *SIAM Rev.* **45**(4), 727–784 (2003)
- FLS99. Feige, U., Lapidot, D., Shamir, A.: Multiple noninteractive zero knowledge proofs under general assumptions. *SIAM J. Comput.* **29**(1), 1–28 (1999)
- FP96. Fuchs, C.A., Peres, A.: Quantum-state disturbance versus information gain: uncertainty relations for quantum information. *Phys. Rev. A* **53**, 2038–2045 (1996)

- GKLW20. Garg, R., Khurana, D., Lu, G., Waters, B.: Black-box non-interactive non-malleable commitments. *Cryptology ePrint Archive, Report 2020/1197* (2020). <https://eprint.iacr.org/2020/1197>
- GKS16. Goyal, V., Khurana, D., Sahai, A.: Breaking the three round barrier for non-malleable commitments, pp. 21–30 (2016)
- GLOV12. Goyal, V., Lee, C.-K., Ostrovsky, R., Visconti, I.: Constructing non-malleable commitments: a black-box approach, pp. 51–60 (2012)
- Goy11. Goyal, V.: Constant round non-malleable protocols using one way functions, pp. 695–704 (2011)
- GPR16a. Goyal, V., Pandey, O., Richelson, S.: Textbook non-malleable commitments. In: *Proceedings of the 48th Annual ACM SIGACT Symposium on Theory of Computing, STOC 2016, Cambridge, MA, USA, 18–21 June 2016*, pp. 1128–1141 (2016)
- GR19. Goyal, V., Richelson, S.: Non-malleable commitments using Goldreich-Levin list decoding. In: *60th IEEE Annual Symposium on Foundations of Computer Science, FOCS 2019, Baltimore, Maryland, USA, 9–12 November 2019*, pp. 686–699 (2019)
- HSS15. Hallgren, S., Smith, A., Song, F.: Classical cryptographic protocols in a quantum world. *Int. J. Quant. Inf.* **13**(04), 1550028 (2015). Preliminary version appeared in *IACR Crypto 2011*
- Khu17. Khurana, D.: Round optimal concurrent non-malleability from polynomial hardness. In: Kalai, Y., Reyzin, L. (eds.) *TCC 2017, Part II. LNCS*, vol. 10678, pp. 139–171. Springer, Cham (2017). https://doi.org/10.1007/978-3-319-70503-3_5
- KK19. Kalai, Y.T., Khurana, D.: Non-interactive non-malleability from quantum supremacy. In: Boldyreva, A., Micciancio, D. (eds.) *CRYPTO 2019. LNCS*, vol. 11694, pp. 552–582. Springer, Cham (2019). https://doi.org/10.1007/978-3-030-26954-8_18
- KS17. Khurana, D., Sahai, A.: How to achieve non-malleability in one or two rounds. In: *58th IEEE Annual Symposium on Foundations of Computer Science, FOCS 2017, Berkeley, CA, USA, 15–17 October 2017*, pp. 564–575 (2017)
- LMS21. Lombardi, A., Ma, F., Spooner, N.: Post-quantum zero knowledge, revisited (or: How to do quantum rewinding undetectably). *arXiv preprint arXiv:2111.12257* (2021)
- LN11. Lunemann, C., Nielsen, J.B.: Fully simulatable quantum-secure coin-flipping and applications. In: Nitaj, A., Pointcheval, D. (eds.) *AFRICACRYPT 2011. LNCS*, vol. 6737, pp. 21–40. Springer, Heidelberg (2011). https://doi.org/10.1007/978-3-642-21969-6_2
- LP09. Lin, H., Pass, R.: Non-malleability amplification, pp. 189–198 (2009)
- LP11. Lin, H., Pass, R.: Constant-round non-malleable commitments from any one-way function, pp. 705–714 (2011)
- LP12. Lin, H., Pass, R.: Black-box constructions of composable protocols without set-up. In: Safavi-Naini, R., Canetti, R. (eds.) *CRYPTO 2012. LNCS*, vol. 7417, pp. 461–478. Springer, Heidelberg (2012). https://doi.org/10.1007/978-3-642-32009-5_27
- LPS17. Lin, H., Pass, R., Soni, P.: Two-round and non-interactive concurrent non-malleable commitments from time-lock puzzles, pp. 576–587 (2017)

- LPV08. Lin, H., Pass, R., Venkatasubramanian, M.: Concurrent non-malleable commitments from any one-way function. In: Canetti, R. (ed.) TCC 2008. LNCS, vol. 4948, pp. 571–588. Springer, Heidelberg (2008). https://doi.org/10.1007/978-3-540-78524-8_31
- LPV09. Lin, H., Pass, R., Venkatasubramanian, M.: A unified framework for concurrent security: universal composability from stand-alone non-malleability. In: Mitzenmacher, M. (ed.) Proceedings of the 41st Annual ACM Symposium on Theory of Computing, STOC 2009, Bethesda, MD, USA, May 31–June 2 2009, pp. 179–188. ACM (2009)
- Nao91. Naor, M.: Bit commitment using pseudorandomness. *J. Cryptol.* **4**(2), 151–158 (1991). <https://doi.org/10.1007/BF00196774>
- PPV08. Pandey, O., Pass, R., Vaikuntanathan, V.: Adaptive one-way functions and applications. In: Wagner, D. (ed.) CRYPTO 2008. LNCS, vol. 5157, pp. 57–74. Springer, Heidelberg (2008). https://doi.org/10.1007/978-3-540-85174-5_4
- PR05a. Pass, R., Rosen, A.: Concurrent non-malleable commitments, pp. 563–572 (2005)
- PR05b. Pass, R., Rosen, A.: New and improved constructions of non-malleable cryptographic protocols, pp. 533–542 (2005)
- PW10. Pass, R., Wee, H.: Constant-round non-malleable commitments from sub-exponential one-way functions. In: Gilbert, H. (ed.) EUROCRYPT 2010. LNCS, vol. 6110, pp. 638–655. Springer, Heidelberg (2010). https://doi.org/10.1007/978-3-642-13190-5_32
- Ros04. Rosen, A.: A note on constant-round zero-knowledge proofs for NP. In: Naor, M. (ed.) TCC 2004. LNCS, vol. 2951, pp. 191–202. Springer, Heidelberg (2004). https://doi.org/10.1007/978-3-540-24638-1_11
- Unr12. Unruh, D.: Quantum proofs of knowledge. In: Pointcheval, D., Johansson, T. (eds.) EUROCRYPT 2012. LNCS, vol. 7237, pp. 135–152. Springer, Heidelberg (2012). https://doi.org/10.1007/978-3-642-29011-4_10
- Wat09. Watrous, J.: Zero-knowledge against quantum attacks. *SIAM J. Comput.* **39**(1), 25–58 (2009)
- Wee10. Wee, H.: Efficient chosen-ciphertext security via extractable hash proofs. In: Rabin, T. (ed.) CRYPTO 2010. LNCS, vol. 6223, pp. 314–332. Springer, Heidelberg (2010). https://doi.org/10.1007/978-3-642-14623-7_17
- WZ82. Wootters, W.K., Zurek, W.H.: A single quantum cannot be cloned. *Nature* **299**, 802–803 (1982)