



Fully Malicious Authenticated PIR

Marian Dietz^(✉)  and Stefano Tessaro^{}

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

Abstract. Authenticated PIR enables a server to initially commit to a database of N items, for which a client can later privately obtain individual items with complexity sublinear in N , with the added guarantee that the retrieved item is consistent with the committed database. A crucial requirement is *privacy with abort*, i.e., the server should not learn anything about a query *even* if it learns whether the client aborts.

This problem was recently considered by Colombo et al. (USENIX '23), who proposed solutions secure under the assumption that the database is committed to *honestly*. Here, we close this gap for their DDH-based scheme, and present a solution that tolerates fully malicious servers that provide potentially malformed commitments. Our scheme has communication and client computational complexity $\mathcal{O}_\lambda(\sqrt{N})$, does not require any additional assumptions, and does not introduce heavy machinery (e.g., generic succinct proofs). We do so by introducing *validation queries*, which, from the server's perspective, are computationally indistinguishable from regular PIR queries. Provided that the server succeeds in correctly answering κ such validation queries, the client is convinced with probability $1 - \frac{1}{2^\kappa}$ that the server is unable to break privacy with abort.

1 Introduction

Private Information Retrieval (PIR) [16] allows a client to retrieve items from a database of N items held by a server, without revealing which items are being retrieved. Crucially, both the communication and the client's runtime are sublinear in N . In this paper, we are interested in the more challenging setting of *single server* PIR [35], for which constructions are known from a variety of assumptions (cf. e.g. [1, 3, 11, 18, 21, 22, 24, 30, 35, 38, 42]).

MALICIOUSLY SECURE PIR. The literature has by now surfaced countless applications of PIR, such as anonymous web search [29], anonymous messaging [4, 15], private media delivery [28], certificate transparency logs [30, 39], secure analytics [27], secure breach alerts [2] and more. In most use cases, a PIR service is queried by *multiple* clients, but PIR schemes usually assume a semi-honest server and do not ensure integrity. Hence, clients are not guaranteed consistent answers. This paper considers a setting, referred to as *authenticated PIR*, where the server succinctly *commits* to a database $\mathbf{x}$, and then can only provide answers

consistent with this commitment. At first, it appears easy to solve the problem: The server publishes a succinct vector commitment d to $\mathbf{x}$ (e.g., the root of a Merkle Tree), and then runs a PIR scheme on a modified database $\mathbf{x}'$ where $\mathbf{x}'[i]$ consists of $\mathbf{x}[i]$ and a valid proof π_i of inclusion with respect to d .

While this solution indeed guarantees integrity, Kushilevitz and Ostrovsky [35] already observed that this approach is prone to so-called *selective-failure attacks*, where the server intentionally includes an invalid proof π_i for a position $i \in [N]$, and learning about a failure reveals that the client’s query is for this particular index. Such attacks are realistic, as failures are hard to hide (e.g., because an application will behave differently in an observable way), and they have been considered in a number of settings, such as encryption [8, 40] and two-party computation [12, 31, 33, 37].

AUTHENTICATED PIR. Early proposals to make single-server PIR authentic [41, 44] have not considered selective-failures, and only provide authenticity guarantees that a query is consistent with *some* valid database, but not a specific one the server a-priori committed to. In contrast, multi-server verifiable PIR [43] considers a similar functionality, however without covering selective-failures.

Only very recently, Colombo et al. [17] initiated the study of selective-failure attacks in the context of PIR. They propose in particular single-server PIR schemes based on the LWE and DDH assumptions, both with $O_\lambda(\sqrt{N})$ complexity, which, while less practical than non-authenticated counterparts, offer acceptable performance, integrity, and privacy under selective-failure attacks. However, their model *explicitly disallows* maliciously generated digests.

Still, it is natural and prudent to assume that a data owner willing to deviate maliciously from a PIR protocol would also attempt to do so with help of a malicious digest. Say, the PIR server is maintained by a streaming service, where the server obviously has the power to choose the digest that represents a commitment to the database. Then, it is not possible for a client to obtain the digest with “out-of-band” means as described in [17] without involvement of the malicious server, hence protection against selective abort attacks would not apply here. Or, consider a different database that is used for password breach checking [2]. Again, if the database owner has the power of choosing the digest maliciously, they may infer private data about users’ passwords.

OUR CONTRIBUTIONS. In this paper, we prove that under the DDH assumption, a variant of the scheme from [17] is in fact secure against a fully malicious server generating malicious digests, achieving both integrity and privacy against selective-failure attacks. To achieve the latter property, our solution requires the client to ask λ additional *validation queries* which are indistinguishable from regular PIR queries from the perspective of the server—privacy under selective-failure attacks holds provided these validation queries are successful. We note that our protocol comes with *no* overhead compared to that of [17], except for the λ initial validation queries, after which an *unbounded* number of database queries can be made. This means that our scheme retains communication complexity $O_\lambda(\sqrt{N})$ for PIR queries, and the digest can be made constant size.

Furthermore, by reducing the number of validation queries from λ to any smaller integer κ , the client will still be convinced with probability $1 - \frac{1}{2^\kappa}$ that the server is unable to break privacy. For example, in a scenario where the database is simultaneously accessed by sufficiently many users seeing the same digest, it may be in fact enough to choose $\kappa = 1$, assuming we can rely on a mechanism that allows users to voice complaints should they fail the validation phase. In particular, validation would fail for roughly half of the users, and this is a strong incentive for the server to behave honestly.

We give an overview of our scheme and our analysis in Sect. 1.2. The core of our approach is a new technique which turns an adversary that succeeds in the validation phase into an extraction algorithm able to answer *any* query to the PIR server. This shows that the information of whether a query would abort can be *simulated*, thus reducing selective-failure privacy to regular privacy.

In Sect. 1.1, we first discuss other related and concurrent works, along with less efficient generic baseline solutions using existing cryptographic techniques. This discussion will highlight a unique feature of our work, namely that unlike all alternative solutions, we do not rely on adding additional proof machinery to a PIR scheme to verify validity of a digest. Rather, we show that the much simpler machinery from the scheme of [30] is already sufficient, *as is*, for verifiability.

1.1 Related Works and Generic Baselines

We discuss here some more related work, including concurrent work on authenticated PIR and less efficient baseline solutions to the problem.

CONCURRENT WORK. We first point out that concurrently to this work, de Castro and Lee [14] propose an alternative authenticated PIR scheme that provides security against a maliciously generated digest. It extends SimplePIR [30] by adding proofs to each answer to convince the client that computation has been done correctly and is consistent with a digest. Further, as in our scheme, an additional validation phase ensures validity of the digest. However, there are also several aspects in which VeriSimplePIR differs from our work:

- VeriSimplePIR utilizes lattices assumptions that are more expressive than DDH, and in order avoid several additional rounds of interaction it also requires assuming the Random Oracle Model. Our protocol is based on plain DDH, and therefore gives a theoretically stronger result.
- The two solutions achieve different communication/computation tradeoffs: VeriSimplePIR can be expected to be scheme with better computational efficiency, as it only works with lattices. On the other hand, it also comes at the cost of either large per-query communication cost or large client storage cost. For example, the evaluation in [14] shows that for a database of size $N = 4\text{GiB}$, the preprocessing phase requires almost 2GiB of communication, out of which more than 600MiB need to be permanently stored on the client’s device. Our scheme (instantiated with 256-bit elliptic curves and $\kappa = 80$) would have a validation phase with about 260MiB of communication, and requires no client storage except for the digest of size 6MiB .

Thus, our scheme could even be used on mobile devices, provided that the server has sufficient resources to offset its larger computational cost through parallelization.

- There are several unique features of our protocol: our validation phase is *entirely oblivious* to the server, because each validation query looks identical to a regular PIR query. In fact, we can take an existing server that implements the scheme from [17], and let a client validate that the server will be unable to break privacy with abort. Furthermore, each client will only need to run validation *once* in its lifetime. In contrast, VeriSimplePIR requires restarting the expensive preprocessing phase whenever cheating was detected during a single query, or when the client erases its storage.

GENERIC BASELINES. We also discuss a few alternative solutions based on generic (and heavier) techniques added on top of the authenticated PIR from [17]. For instance, the server could include in each query response a succinct proof of knowledge of an opening of the digest d to which this response is consistent. This could be done using a SNARK (for which constructions under standard falsifiable assumptions hit well-known barriers [25]), or a succinct two-move proof (which are, as of now, not known to be easier to achieve than SNARKs).

Alternatively, leveraging the fact that the schemes from [17] are already secure for honestly generated digests, the validation phase may be used for a succinct interactive proof of knowledge of an opening of the digest d to a *binary* database. (In essence, the same extractor as the one we use to prove our scheme’s security could be used to achieve the crucial notion of answer extractability here.) One could use Kilian’s four-round protocol [32], which has also been shown to be a proof of knowledge [5]. Roughly, the first protocol round, which consists of a description of a collision-resistant hash function, would be included in the public parameters. Then, Kilian’s second round, which is a succinct Merkle commitment to a suitable PCP, would be merged with the digest, whereas the remaining two rounds would constitute the actual validation stage verifying the (committed) PCP by opening a random subset of Merkle paths. However, this solution would have to rely on non-black box techniques or stronger assumptions. Therefore, we cannot expect them to be concretely more efficient than our PIR protocol in terms of communication complexity or computation time.

We also note that for the specific case of the DDH-based scheme from [17] we use in this paper, the inner-product argument from Bulletproofs [10] would provide a fully-algebraic pairing-free solution. This is specific to our setting, in which the digest is supposed to have the form of a commitment $d = \prod_{i \in [N]} \mathbf{h}[i]^{\mathbf{x}[i]}$ given the database $\mathbf{x}$. The techniques would allow us to prove that all elements in the commitment are binary, but this incurs linear computation cost $O(N)$ for the verifier, thus violating an important requirement of a PIR scheme. It also requires either a random oracle (to make the proof non-interactive), or $O(\log N)$ rounds. In contrast, our protocol has a communication complexity of $O(\kappa\sqrt{N})$ group elements per validation phase, only needs a single round-trip, and does not require a random oracle.

Finally, we note that none of the generic baselines we presented here preserves the property of an oblivious validation phase that does not require modification of the server implementation.

OTHER RELATED WORK. Ben David et al. [7] recently studied a different and unrelated notion of verifiable PIR which allows the server to prove properties about the database with sublinear complexity. A number of works [6, 19, 23, 26, 34] studied robustness of PIR against malicious servers, i.e., the goal is to guarantee answers even when some of the servers fail/deviate, but they rely on the availability of some honest server. There have also been impressive advances on single-server PIR with sublinear server complexity [9, 13, 36], but this is not considered in the context of this work.

1.2 Technical Overview

THE CNCWF DDH-BASED SCHEME. Our starting point is the DDH-based scheme from [17] (henceforth referred to as CNCWF), which only provides security against selective-failure attacks when the digest was selected *honestly*. The scheme relies on a cyclic group $\mathbb{G}$ of order q in which the DDH assumption holds. We only discuss here the unbalanced version of the scheme, for which queries consist of N group elements, whereas responses are a single group element. (One can turn this into a scheme with $O(\sqrt{N})$ size queries and responses via standard re-balancing techniques which will preserve all claimed properties.)

In CNCWF, the server initially commits to an N -bit database $\mathbf{x} \in \{0, 1\}^N$ via a digest $d = \prod_{i=1}^N \mathbf{h}[i]^{\mathbf{x}[i]}$, where $\mathbf{h}$ is a vector of N independent generators of $\mathbb{G}$, included in the public parameters of the scheme, along with an extra generator g . To retrieve item $i \in [N]$, the client samples $r \leftarrow \mathbb{Z}_q, \alpha \leftarrow \mathbb{Z}_q^*$, and computes a *query vector* $\tilde{\mathbf{e}} \in \mathbb{G}^N$, where

$$\tilde{\mathbf{e}}[j] = \begin{cases} \mathbf{h}[j]^r \cdot g^\alpha & \text{if } j = i, \\ \mathbf{h}[j]^r & \text{if } j \in [N] \setminus \{i\} \end{cases}.$$

We also write this succinctly as $\tilde{\mathbf{e}} = \mathbf{h}^r \circ (g^{\mathbf{e}_i})^\alpha$, where $\mathbf{e}_i$ is the i -th unit vector, $\circ$ denotes component-wise vector multiplication, and $g^{\mathbf{v}} = (g^{\mathbf{v}[1]}, \dots, g^{\mathbf{v}[N]})$. The server's response is $\tilde{e} = \prod_{j=1}^N \tilde{\mathbf{e}}[j]^{\mathbf{x}[j]}$, and the client computes $e = (d^{-r} \tilde{e})^{\alpha^{-1}}$. It returns 1 if $e = g$, 0 if $e = 1$, and otherwise aborts if $e \notin \{1, g\}$.

INTEGRITY, BUT NO PRIVACY WITH ABORT. Under the DDH assumption, we will show that this scheme already prevents an attacker from providing a malicious digest d along with *two* different valid answers e'_0, e'_1 to the same query. This is a stronger notion of integrity than what is shown in [17].

However, there is a simple selective-failure attack, in which a malicious server can commit to an (invalid) database with $x_1 = 2$ and $x_2 = 1$. An honest client would abort on a query for $i = 1$, as $g^2 \notin \{1, g\}$, but succeed for a query with $i = 2$. Learning whether the operation aborts or not hence tells the adversary whether the query was for $i = 1$ or $i = 2$.

DIGEST VALIDATION. In an attempt to mitigate the above attack, we introduce an initial *digest validation* phase, which consists of a single round-trip of interaction between the client and the server. If successfully completed, we expect the scheme to be private even under selective-failure attacks, a notion we henceforth refer to as “privacy with abort.”

Our validation approach leverages the fact that a client can actually learn $g^{\mathbf{x} \cdot \mathbf{s}}$ (where $\cdot$ denotes inner product) for *any vector* $\mathbf{s}$ via the query $\tilde{\mathbf{s}} = (\mathbf{h})^r \circ (g^{\mathbf{s}})^\alpha$, with no modification to the server. Further, under the DDH assumption, the server cannot distinguish such a query from a regular query, which corresponds to the case $\mathbf{s} = \mathbf{e}_i$. Concretely, the client asks λ queries $\tilde{\mathbf{s}}_1, \dots, \tilde{\mathbf{s}}_\lambda$ associated with $\mathbf{s}_1, \dots, \mathbf{s}_\lambda \leftarrow \{0, 1\}^N$, and accepts if *all of them* return a value in the range $\{g^0, g^1, \dots, g^N\}$.

Intuitively, this ensures that the server answers using a database with entries $\mathbf{x}[i]$ which are small enough (say, in $\{-N, \dots, N\}$), hence preventing the above attack if we additionally modify processing a query answer to output 1 whenever $e \in \{g^{-N}, \dots, g^N\} \setminus \{1\}$, instead of $e = g$ as in the original CNCWF scheme. Of course, we need to prove that any other attack is prevented, and furthermore, we want to do so while relying solely on the DDH assumption.

ANSWER EXTRACTABILITY. The stepping stone to proving privacy with abort, upon passing validation, is proving a notion we call *answer extractability*. To define it, we consider any two-stage adversary $(\mathcal{D}, \mathcal{V})$, where $\mathcal{D}$ initially outputs a digest d , along with a state $\mathbf{st}_{\mathcal{D}}$. Then, $\mathcal{V}$, on input validation queries $\mathbf{s}_1, \dots, \mathbf{s}_\lambda$, as well as $\mathbf{st}_{\mathcal{D}}$, returns answers $\tilde{s}_1, \dots, \tilde{s}_\lambda$. Validation passes if these answers reconstruct to values $s_1, \dots, s_\lambda \in \{g^0, g^1, \dots, g^N\}$.

Answer extractability asks for an extractor $\mathcal{E}_{\text{ans}}$ which, on input $\mathbf{st}_{\mathcal{D}}$ and d , is able to answer each honest PIR query $\tilde{\mathbf{e}} = (\mathbf{h})^r \circ (g^{\mathbf{e}_i})^\alpha$ in a non-aborting way, i.e., $\mathcal{E}_{\text{ans}}(\tilde{\mathbf{e}}, \mathbf{st}_{\mathcal{D}}, d)$ returns with overwhelming probability $\tilde{e} \in \mathbb{G}$ such that

$$(d^{-r} \tilde{e})^{1/\alpha} \in \{g^{-N}, \dots, g^N\}.$$

This is too strong of a requirement, so we will only require $\mathcal{E}_{\text{ans}}$ to succeed whenever the state $\mathbf{st}_{\mathcal{D}}$ is such that $\mathcal{V}$ is able to pass the validation stage with non-negligible probability ν given that state. We also do allow the running time of $\mathcal{E}_{\text{ans}}$ to depend polynomially on $1/\nu$.

When combined with standard PIR privacy and integrity, we will show that answer extractability implies privacy with abort. Intuitively, the idea is to use the extractor to simulate the abort information in the standard privacy experiment. We omit details here, and refer to the body of the paper.

THE EXTRACTOR CONSTRUCTION. The core of the security proof is our construction of an extractor $\mathcal{E}_{\text{ans}}$ for answer extractability, along with its analysis. Given an honest query $\tilde{\mathbf{e}} = \mathbf{h}^r \circ (g^{\mathbf{e}_i})^\alpha$, along with d and $\mathbf{st}_{\mathcal{D}}$, a natural strategy for $\mathcal{E}_{\text{ans}}$ to answer this query is to repeatedly invoke $\mathcal{V}$ as

$$(\tilde{s}_1, \dots, \tilde{s}_\lambda) \leftarrow \mathcal{V}(\mathbf{st}_{\mathcal{D}}, (\tilde{\mathbf{e}}, \tilde{\mathbf{s}}_2, \dots, \tilde{\mathbf{s}}_\lambda)).$$

where $\tilde{\mathbf{s}}_2, \dots, \tilde{\mathbf{s}}_\lambda$ are “fake” validation queries. The hope is that the fact that $\mathcal{V}$ passes validation with non-negligible probability ν also implies that, with

overwhelming probability, $\tilde{s}_1$ is a good answer for $\tilde{\mathbf{e}}$ after $\text{poly}(1/\nu)$ attempts. There are however two obvious challenges:

- (1) $\mathcal{V}$ is only guaranteed to provide accurate answers for validation queries, but it is not clear why this would imply that it provides correct answers for a regular PIR query.
- (2) We need to be able to check when $\tilde{s}_1$ is correct without knowing the randomness r, α .

To overcome (2), we will show that the following strategy works: in each iteration we invoke $\mathcal{V}$ *twice*, the second call as

$$(\tilde{s}'_1, \dots, \tilde{s}'_\lambda) \leftarrow \mathcal{V}(\text{st}_{\mathcal{D}}, (\tilde{\mathbf{e}}', \tilde{\mathbf{s}}'_2, \dots, \tilde{\mathbf{s}}'_\lambda)) , \quad (1)$$

where $\tilde{\mathbf{e}}'$ is obtained by randomizing $\tilde{\mathbf{e}}$, i.e., $\tilde{\mathbf{e}}' = \mathbf{h}^{r'} \circ \tilde{\mathbf{e}}^{\alpha'}$, and $\tilde{\mathbf{s}}'_2, \dots, \tilde{\mathbf{s}}'_\lambda$ are fresh validation queries. Then $\mathcal{E}_{\text{ans}}$ outputs $\tilde{s}_1$ whenever for two such calls we have $(d^{-r'} \tilde{s}'_1)^{1/\alpha'} = \tilde{s}_1$. If this never happens, it outputs $\perp$.

To achieve sufficient independence for this argument to go through, the extractor needs to in fact also re-randomize $\tilde{\mathbf{e}}$ itself for each iteration, and then remove this randomization from answers $\tilde{s}_1$ and $\tilde{s}'_1$. This has the added benefit of simplifying notation, as a statistical argument implies the fact that $\mathcal{E}_{\text{ans}}(d, \text{st}_{\mathcal{D}}, \tilde{\mathbf{e}})$ returning a value $\tilde{e}$ such that $(d^{-r} \tilde{e})^{1/\alpha} \in \{g^{-N}, \dots, g^N\}$ is *equivalent* to

$$\mathcal{E}_{\text{ans}}(d, \text{st}_{\mathcal{D}}, g^{\mathbf{e}_i}) \in \{g^{-N}, \dots, g^N\} ,$$

where the input $g^{\mathbf{e}_i}$ is not randomized.

CORRECTNESS. To prove that $\mathcal{E}_{\text{ans}}$ works as desired, hence overcoming (1), a crucial property we establish is that it behaves *homomorphically*—with overwhelming probability, for any $\mathbf{p}_1, \dots, \mathbf{p}_m \in \mathbb{G}^N$ and $t_1, \dots, t_m \in \mathbb{Z}_q$, if we let $\mathbf{p}_{m+1} = \prod_{i=1}^m \mathbf{p}_i^{t_i}$ and compute $p_i \leftarrow \mathcal{E}_{\text{ans}}(d, \text{st}_{\mathcal{D}}, \mathbf{p}_i)$ for $i \in [m+1]$, then

$$p_1, \dots, p_{m+1} \neq \perp \implies p_{m+1} = \prod_{i=1}^m p_i^{t_i} . \quad (2)$$

Then, the proof considers an extended experiment where we initially run $\mathcal{D}$ to sample $(d, \text{st}_{\mathcal{D}})$, and subsequently:

- Compute $e_i \leftarrow \mathcal{E}_{\text{ans}}(d, \text{st}_{\mathcal{D}}, g^{\mathbf{e}_i})$ for all $i \in [N]$.
- Pick $\mathbf{s}_{j,k} \leftarrow_{\$} \{0, 1\}^N$ for $k \in [\lambda]$ and sufficiently many j 's, then run $s_{j,k} \leftarrow \mathcal{E}_{\text{ans}}(d, \text{st}_{\mathcal{D}}, g^{\mathbf{s}_{j,k}})$ for each j and $k \in [\lambda]$.

In this experiment, we prove that, with overwhelming probability, the event that $(d, \text{st}_{\mathcal{D}})$ is good implies the following two events:

- (1) $e_i \neq \perp$ for all $i \in [N]$
- (2) There exists j^* such that $s_{j^*,k} \in \{1, g, \dots, g^N\}$ for all $k \in [\lambda]$

Both properties are non-obvious to prove, but they follow from the fact that $\mathcal{V}$ does well on a good $(d, \text{st}_{\mathcal{D}})$. Then, Eq. 2 yields

$$s_{j^*,k} = \prod_{i=1}^N e_i^{s_{j^*,k}[i]} \in \{1, g, \dots, g^N\} \text{ for all } k \in [\lambda].$$

A simple statistical lemma shows that if $e_1, \dots, e_N$ were *not* all in $\{g^{-N}, \dots, g^N\}$, then due to $s_{j,k}$ being uniformly sampled from $\{0, 1\}^N$, the above was unlikely to have happened. Hence, all e_i 's are in $\{g^{-N}, \dots, g^N\}$, which implies that the extractor can answer any honest query.

1.3 Open Problems

The most obvious open question is to develop a lattice analogue of our results, e.g., on top of the simple LWE-based scheme from [17]. The main challenge here seems to be the step from *answer extractability* to *privacy with abort*. To show the latter, we need to simulate the abort-bit, which we do in the case of DDH by first extracting a “good” answer that correctly answers the query without aborting. However, for lattices, even when it is possible to extract a database $\mathbf{x}$, any honest answer w.r.t. $\mathbf{x}$ would itself already be *noisy*, and therefore comparing it with the actual answer given by the adversary does not necessarily allow for a correctly computed abort-bit. Despite lattices, another interesting question is to add updatability to the scheme.

2 Preliminaries

BASIC NOTATION AND COMPUTATIONAL MODEL. We let λ denote the security parameter, and use $\text{poly}(\lambda)$ and $\text{negl}(\lambda)$ as placeholders for a generic polynomial and negligible function in λ , respectively. All algorithms in this paper are randomized unless otherwise stated. We assume a standard model of computation, and refer to our adversaries as “ppt” to indicate they run in polynomial time. For convenience, we assume these algorithms are non-uniform (i.e., they can also use polynomial-time advice), although it is not hard to extend our treatment to the uniform setting. Throughout this paper, we use $\leftarrow$ to denote the random sampling of an output of a randomized algorithm, $\leftarrow_{\$}$ to denote the uniform sampling of an element from a set, and $:=$ to denote assignment of a value to a variable. We say that two families of distributions $X = \{X_\lambda\}_{\lambda \in \mathbb{N}}$ and $Y = \{Y_\lambda\}_{\lambda \in \mathbb{N}}$ are computationally indistinguishable, denoted $X \approx_c Y$, if the size of X_λ, Y_λ is polynomial in λ , and, for all ppt $\mathcal{A}$, we have $\Pr[\mathcal{A}(X_\lambda) = 1] - \Pr[\mathcal{A}(Y_\lambda) = 1] = \text{negl}(\lambda)$.

GROUP-THEORETIC PRELIMINARIES. We will work with cyclic groups throughout this paper. We will typically use a *group parameter generator* GG to sample parameters. Such an algorithm, on input 1^λ , outputs $(\mathbb{G}, g, q)$ consisting of the description of a group $\mathbb{G}$ with generator g and order $2^\lambda \leq q < 2^{\lambda+1}$. The group

also allows for computing group operations in time $\text{poly}(\lambda)$, and in particular its elements are also represented as polynomially-sized strings.

DECISIONAL DIFFIE-HELLMAN. Our results rely on the Decisional Diffie-Hellman (DDH) assumption, which we repeat below for completeness.

Definition 1 (DDH). *The Decisional Diffie-Hellman assumption (DDH) holds for a group generator GG if, for $\text{pp} = (\mathbb{G}, q, g) \leftarrow \text{GG}(1^\lambda)$,*

$$\{(\text{pp}, g^a, g^b, g^{ab}) \mid a, b \leftarrow_{\$} \mathbb{Z}_q\} \approx_c \{(\text{pp}, g^a, g^b, g^c) \mid a, b, c \leftarrow_{\$} \mathbb{Z}_q\}.$$

We will also repeatedly need the following commonly used lemma.

Lemma 1. *If the DDH assumption holds, then for all polynomials $N(\lambda) = \text{poly}(\lambda)$, and $\text{pp}' = (\mathbb{G}, q, g) \leftarrow \text{GG}(1^\lambda)$, we have*

$$\{(\text{pp}', \mathbf{h}, \mathbf{h}^r) \mid \mathbf{h} \leftarrow_{\$} \mathbb{G}^N, r \leftarrow_{\$} \mathbb{Z}_q\} \approx_c \{(\text{pp}', \mathbf{h}, \bar{\mathbf{h}}) \mid \mathbf{h} \leftarrow_{\$} \mathbb{G}^N, \bar{\mathbf{h}} \leftarrow_{\$} \mathbb{G}^N\}.$$

3 Authenticated PIR: Definitions and Basic Properties

3.1 Basic Definitions

PIR SYNTAX. In this section, we define our notion of authenticated PIR. In contrast to prior work, we allow for a short validation phase. Here, we do not impose any restrictions on how this phase is run, although our concrete scheme (Sect. 4) will guarantee that server is unable to tell validation queries from regular ones, which is a fundamental feature of such a scheme.

Definition 2 (Authenticated PIR). *An authenticated PIR scheme (APIR) APIR consists of the following eight algorithms:*

- *The setup algorithm, on input the security parameter and the database length, both in unary, outputs the public parameters $\text{pp} \leftarrow \text{Setup}(1^\lambda, 1^N)$.*
- *The digest algorithm, on input pp and database $\mathbf{x} \in \{0, 1\}^N$, outputs a digest $d \leftarrow \text{Digest}(\text{pp}, \mathbf{x})$.*
- *The client’s validation query algorithm only takes as input pp , and outputs $(\text{st}, v) \leftarrow \text{VQuery}(\text{pp})$, where st is a state and v is the validation query.*
- *The server’s validation response algorithm takes as input pp , the database $\mathbf{x}$, and a validation query v , and returns a response $a_v \leftarrow \text{VAnswer}(\text{pp}, \mathbf{x}, v)$.*
- *The client’s validation check algorithm takes as inputs pp , the state st , digest d , and a response a_v , and returns a decision bit $\text{VCheck}(\text{pp}, \text{st}, d, a_v) \in \{0, 1\}$.*
- *The client’s query algorithm takes as input pp , along with a query index $i \in [N]$, and outputs a pair $(\text{st}, q) \leftarrow \text{Query}(\text{pp}, i)$ consisting of a state st and a query q .*
- *The server’s query response algorithm takes as input pp , the database $\mathbf{x}$, and the query q , and returns an answer $a \leftarrow \text{Answer}(\text{pp}, \mathbf{x}, q)$.*
- *The client’s reconstruction algorithm takes as inputs pp , the state st , digest d , and an answer a , and returns a value $\text{Rec}(\text{pp}, \text{st}, d, a) \in \{0, 1, \perp\}$.*

Process $\text{Conf}_{\text{APIR}}^{\mathcal{D}}(\lambda, N)$	Game $\text{INTEGRITY}_{\text{APIR}}^{\mathcal{D}, \mathcal{A}}(\lambda, N, i)$
1 : $\text{pp} \leftarrow \text{Setup}(1^\lambda, 1^N)$	1 : $c = (\text{pp}, d, \text{st}_{\mathcal{D}}) \leftarrow \text{Conf}^{\mathcal{D}}(\lambda, N)$
2 : $\text{st}_{\mathcal{D}}, d \leftarrow \mathcal{D}(\text{pp})$	2 : $\text{st}, q \leftarrow \text{Query}(\text{pp}, i)$
3 : return $c := (\text{pp}, d, \text{st}_{\mathcal{D}})$	3 : $a, a' \leftarrow \mathcal{A}(c, q)$
	4 : return $\mathbb{I} \left[\begin{array}{l} a \neq a' \wedge \text{Rec}(\text{st}, d, a) \neq \perp \\ \wedge \text{Rec}(\text{st}, d, a') \neq \perp \end{array} \right]$

Fig. 1. Definition of the configuration generation process (left) and of the integrity definition game (right).

For correctness, we require that for all $\lambda, N \in \mathbb{N}$ and $\text{pp} \in \text{Supp}(\text{Setup}(1^\lambda, 1^N))$, databases $\mathbf{x} \in \{0, 1\}^N$, and indices $i \in [N]$,

$$\Pr \left[\text{VCheck}(\text{pp}, \text{st}, d, a_v) = 1 \mid \begin{array}{l} d \leftarrow \text{Digest}(\text{pp}, \mathbf{x}) \\ \text{st}, v \leftarrow \text{VQuery}(\text{pp}) \\ a_v \leftarrow \text{VAnswer}(\text{pp}, \mathbf{x}, v) \end{array} \right] = 1 ,$$

$$\Pr \left[\text{Rec}(\text{pp}, \text{st}, d, a) = \mathbf{x}[i] \mid \begin{array}{l} d \leftarrow \text{Digest}(\text{pp}, \mathbf{x}) \\ \text{st}, q \leftarrow \text{Query}(\text{pp}, i) \\ a \leftarrow \text{Answer}(\text{pp}, \mathbf{x}, q) \end{array} \right] = 1 .$$

Remark 1. We can assume without loss of generality that VCheck and Rec are both deterministic, as their randomness can be chosen ahead of time by VQuery and Query , and included in the state. We will also often omit pp from the algorithms when they are understood from the context.

DIGEST GENERATORS. We introduce a number of security properties that are meant to hold against *fully malicious* adversaries. These are in particular allowed to choose the digest themselves. All of our adversaries will proceed in stages, and we denote by $\mathcal{D}$ the initial stage of the adversary (common to all of our security games) that generates the digest d and outputs a state $\text{st}_{\mathcal{D}}$ that may be used by later stages of the adversary. We refer to such adversaries as *digest generators*. It will be convenient to introduce the following formalism that will be reused across definitions to capture the initial stage of the game.

Definition 3 (Configuration). For an APIR scheme APIR , and adversarial digest generator $\mathcal{D}$, we define a configuration $c = (\text{pp}, d, \text{st}_{\mathcal{D}})$ as a triple consisting of pp , d , and $\text{st}_{\mathcal{D}}$, where $\text{pp} \in \text{Supp}(\text{Setup}(1^\lambda, 1^N))$ and $(\text{st}_{\mathcal{D}}, d) \in \text{Supp}(\mathcal{D}(\text{pp}))$. We define $\text{Conf}_{\text{APIR}}^{\mathcal{D}}(\lambda, N)$ as the random process generating a configuration c as on the left-hand side of Fig. 1.

INTEGRITY. We target a strong definition of integrity which prevents the adversary from coming up with two different non-aborting answers for the same query, even if they reconstruct to the same value.

Definition 4 (Integrity). *The APIR scheme APIR fulfills integrity if for all ppt adversaries $\mathcal{A}^* = (\mathcal{D}, \mathcal{A})$, database sizes $N = N(\lambda) \leq \text{poly}(\lambda)$, and indices $i = i(\lambda) \in [N(\lambda)]$,*

$$\Pr \left[\text{INTEGRITY}_{\text{APIR}}^{\mathcal{A}^*}(\lambda, N, i) = 1 \right] \leq \text{negl}(\lambda) ,$$

where $\text{INTEGRITY}_{\text{APIR}}^{\mathcal{A}^*}(\lambda, N, i)$ is defined on the right-hand side of Fig. 1.

Remark 2. Another natural definition of integrity would send two *independently generated* queries for the same index i to the adversary $\mathcal{A}$, who then wins if they are able to produce answers that successfully reconstruct to different values. This notion follows from Definition 4 for any APIR scheme that allows to *re-randomize* a query q into another query q' , and to *recover* an answer a' for query q' into an answer a for query q . This is indeed the case for our protocol in Sect. 4.

PRIVACY. We consider two notions of privacy. The former is the standard definition of APIR privacy in a setting where the adversary can generate malicious digests, whereas the latter allows the attacker to additionally learn whether a response to a query aborts. The two notions are defined by the games described in Fig. 2.

Definition 5 (Standard Privacy). *We say the APIR scheme APIR fulfills privacy if for all ppt adversaries $\mathcal{A}^* = (\mathcal{D}, \mathcal{A})$, database sizes $N = N(\lambda) \leq \text{poly}(\lambda)$, and indices $i_0 = i_0(\lambda), i_1 = i_1(\lambda) \in [N(\lambda)]$,*

$$|\Pr[\text{PRIV}_{\text{APIR}}^{\mathcal{A}^*}(\lambda, N, i_0) = 1] - \Pr[\text{PRIV}_{\text{APIR}}^{\mathcal{A}^*}(\lambda, N, i_1) = 1]| \leq \text{negl}(\lambda) ,$$

where $\text{PRIV}_{\text{APIR}}^{\mathcal{A}^*}(\lambda, N, i)$ is defined on the left of Fig. 2.

Definition 6 (Privacy with abort). *The APIR scheme fulfills privacy with abort if for all ppt adversaries $\mathcal{A}^* = (\mathcal{D}, \mathcal{A}_1, \mathcal{A}_2, \mathcal{A}_3)$, database sizes $N = N(\lambda) \leq \text{poly}(\lambda)$, and indices $i_0 = i_0(\lambda), i_1 = i_1(\lambda) \in [N(\lambda)]$,*

$$|\Pr[\text{PRIV}/\mathcal{A}_{\text{APIR}}^{\mathcal{A}^*}(\lambda, N, i_0) = 1] - \Pr[\text{PRIV}/\mathcal{A}_{\text{APIR}}^{\mathcal{A}^*}(\lambda, N, i_1) = 1]| \leq \text{negl}(\lambda) ,$$

where $\text{PRIV}/\mathcal{A}_{\text{APIR}}^{\mathcal{A}^*}(\lambda, N, i)$ is defined on the right of Fig. 2.

One can also consider an alternative definition where validation and the PIR query happen concurrently, and the client learns if either of them aborts. It is not hard to show that this notion is implied by that in Definition 6. (We discuss this in the full version of this paper [20]).

UNIFORM DEFINITIONS. Our definitions quantify over functions of the security parameter, e.g., $N(\lambda)$, $i_0(\lambda)$, $i_1(\lambda)$. Without further restrictions, this leads to non-uniform proofs of security. One could alternatively give uniform counterparts of all definitions by allowing adversarial choice of these values. Our results would easily extend, although with significant notational clutter we seek to avoid here.

Game $\text{PRIV}_{\text{APIR}}^{\mathcal{D}, \mathcal{A}}(\lambda, N, i)$	Game $\text{PRIV}/\mathcal{A}_{\text{APIR}}^{\mathcal{D}, \mathcal{A}_1, \mathcal{A}_2, \mathcal{A}_3}(\lambda, N, i)$
1 : $c = (\text{pp}, d, \text{st}_{\mathcal{D}}) \leftarrow \text{Conf}^{\mathcal{D}}(\lambda, N)$	1 : $c = (\text{pp}, d, \text{st}_{\mathcal{D}}) \leftarrow \text{Conf}^{\mathcal{D}}(\lambda, N)$
2 : $\text{st}, q \leftarrow \text{Query}(\text{pp}, i)$	2 : $\text{st}_v, v \leftarrow \text{VQuery}(\text{pp})$
3 : return $\mathcal{A}(c, q)$	3 : $\text{st}_{\mathcal{A}}, a_v \leftarrow \mathcal{A}_1(c, v)$
	4 : if $\text{VCheck}(\text{st}_v, d, a_v) = 0$, return 0
	5 : $\text{st}_q, q \leftarrow \text{Query}(\text{pp}, i)$
	6 : $\text{st}'_{\mathcal{A}}, a_q \leftarrow \mathcal{A}_2(\text{st}_{\mathcal{A}}, q)$
	7 : $\text{abort} := \mathbb{I}[\text{Rec}(\text{st}_q, d, a_q) = \perp]$
	8 : return $\mathcal{A}_3(\text{st}'_{\mathcal{A}}, \text{abort})$

Fig. 2. Authenticated PIR Privacy Notions. The two above experiments define our standard privacy (left) and privacy with abort (right) security notions.

3.2 Answer Extractability

Our approach to proving that an APIR scheme fulfills privacy with abort relies on an intermediate notion which we refer to as *answer extractability*. Informally, this notion means that ability to pass the validation phase implies ability to answer arbitrary PIR queries. The correct definition of this notion is rather tricky, along with showing that answer extractability, when combined with (regular) privacy and integrity, implies privacy with abort. Our definition relies on measuring the adversary’s probability of passing validation, conditioned on a particular configuration c having been produced, which we now define.

Definition 7 (Validation success probability). *Let APIR be an APIR scheme, $(\mathcal{D}, \mathcal{V})$ be an adversary, where $\mathcal{D}$ generates a digest and $\mathcal{V}$ attempts to pass validation queries. Then, for any fixed $\lambda, N \in \mathbb{N}$ and configuration $c = (\text{pp}, d, \text{st}_{\mathcal{D}}) \in \text{Supp}(\text{Conf}_{\text{APIR}}^{\mathcal{D}}(\lambda, N))$, the adversary’s validation success probability is defined as*

$$\nu_{\mathcal{D}, \mathcal{V}}^{\text{APIR}}(c) := \Pr[\text{VCheck}(\text{st}, d, \mathcal{V}(c, v)) = 1 \mid \text{st}, v \leftarrow \text{VQuery}(\text{pp})] .$$

For a pair $(\mathcal{D}, \mathcal{V})$, a corresponding *answer extractor* $\mathcal{E}_{\text{ans}}$ is a ppt algorithm that takes a running time parameter 1^r , the configuration c output by $\text{Conf}_{\text{APIR}}^{\mathcal{D}}(\lambda, N)$, and an arbitrary query q . Its task is to give a non-aborting answer for q . The extractor construction will typically depend on $\mathcal{V}$ (this is the case for our construction in Sect. 4.4), and it is natural that its response time is inversely proportional to $\nu_{\mathcal{D}, \mathcal{V}}^{\text{APIR}}(c)$, i.e., when $\mathcal{V}$ has a low validation success probability for some configuration c , then this will also make it “harder” for $\mathcal{E}_{\text{ans}}$ to extract the answer for any given query. For this reason, the actual definition of answer extractability given next sets the running time parameter 1^r so that the running time of $\mathcal{E}_{\text{ans}}$ may grow with the inverse of $\nu_{\mathcal{D}, \mathcal{V}}^{\text{APIR}}(c)$.

Definition 8 (Answer extractability). *The validation step of an APIR scheme APIR fulfills answer extractability if for all ppt adversaries $(\mathcal{D}, \mathcal{V})$, there*

exists an answer extractor $\mathcal{E}_{\text{ans}}$ s.t. for all database sizes $N(\lambda) \leq \text{poly}(\lambda)$, indices $i(\lambda) \in [N]$, and polynomials $r(\lambda)$:

$$\Pr \left[\nu_{\mathcal{D}, \mathcal{V}}^{\text{APIR}}(c) \geq \frac{1}{r(\lambda)} \mid \begin{array}{l} c \leftarrow \text{Conf}_{\text{APIR}}^{\mathcal{D}}(\lambda, N) \\ \text{st}_q, q \leftarrow \text{Query}(\text{pp}, i) \end{array} \right] \geq 1 - \text{negl}(\lambda). \quad (3)$$

PROVING PRIVACY WITH ABORT. Together with *integrity* and *standard privacy*, our notion of *answer extractability* implies *privacy with abort*: We can transform an adversary $\mathcal{A}^*$ for privacy with abort into an adversary $\mathcal{B}^*$ that breaks standard privacy. The idea is that $\mathcal{B}^*$ makes use of the extractor $\mathcal{E}_{\text{ans}}$ guaranteed by extractability, to simulate the *abort-bit* by itself without knowing whether the answer given by $\mathcal{A}^*$ actually aborts. In particular, $\mathcal{B}^*$ simply asks $\mathcal{E}_{\text{ans}}$ for the “correct” answer to the given query, and compare it with the actual answer to check whether there will be an abort.

Lemma 2. *If an APIR scheme APIR fulfills (1) integrity, (2) standard privacy, and (3) answer extractability, then it also fulfills privacy with abort.*

We give the proof of Lemma 2 in the full version [20].

3.3 Alternative Notion for *Relaxed Validation*

In our definition of privacy with abort, passing the validation phase is *very strict*: if the server succeeds, then the client may assume that for all subsequent queries, the server’s distinguishing advantage will be *negligible*. Below we sketch a way of formalizing a tradeoff, which potentially allows for better efficiency in cases where it is acceptable to have a lower confidence in the server’s honesty.

First, we augment the VQuery algorithm for initiating the validation phase, by letting it take an additional parameter κ . Then, privacy with abort can be relaxed by decoupling the adversary into (1) $\mathcal{V}$ attempting to pass the validation phase, and (2) $\mathcal{A}$ answering queries and later making a decision based on query and abort-bit. Denoting the success probability of $\mathcal{V}$ by $\nu(c)$, and the advantage of $\mathcal{A}$ by $\delta(c)$, privacy with abort now states that for any polynomial $p(\lambda)$ it is highly unlikely that simultaneously

$$\delta(c) > \frac{1}{p(\lambda)} \quad \text{and} \quad \nu(c) > \frac{1}{p(\lambda)} + \frac{1}{2^\kappa}$$

hold (over the random choice of the configuration $c \leftarrow \text{Conf}^{\mathcal{D}}(\lambda, N)$).

Intuitively, this means that with probability $\frac{1}{2^\kappa}$, the validation phase may succeed despite the server’s ability to break privacy in any of its subsequent, unbounded number of queries (i.e., validation has *false negatives*). In other words, κ serves as a tunable *confidence* parameter: higher κ means more confidence in the validation phase (e.g., choosing $\kappa = \lambda$ means overwhelming confidence), while lower κ may allow for a more efficient APIR scheme.

<u>Setup($1^\lambda, 1^N$)</u> 1. return $\text{APIR.Setup}(1^\lambda, 1^{\sqrt{N}})$	<u>Digest($\mathbf{x}$)</u> 1. $d_j \leftarrow \text{APIR.Digest}(\mathbf{x}_j) \quad \forall j \in [\sqrt{N}]$ 2. return $d := (d_1, \dots, d_{\sqrt{N}})$
Validation <u>VQuery()</u> 1. return $\text{st}, v \leftarrow \text{APIR.VQuery}()$ <u>VAnswer($\mathbf{x}, v$)</u> 1. $a_j \leftarrow \text{APIR.VAnswer}(\mathbf{x}_j, v) \quad \forall j \in [\sqrt{N}]$ 2. return $a := (a_1, \dots, a_{\sqrt{N}})$ <u>VCheck($\text{st}, d = \{d_j\}, a = \{a_j\}$)</u> 1. $t_j \leftarrow \text{APIR.VCheck}(\text{st}, d_j, a_j) \quad \forall j \in [\sqrt{N}]$ 2. return $\mathbb{I}[t_1 \wedge \dots \wedge t_{\sqrt{N}}]$	Query <u>Query(i)</u> 1. $\text{st}, q \leftarrow \text{APIR.Query}(i_{\text{row}})$ 2. return $\overline{\text{st}} := (\text{st}, i_{\text{col}}), q$ <u>Answer($\mathbf{x}, q$)</u> 1. $a_j \leftarrow \text{APIR.Answer}(\mathbf{x}_j, q) \quad \forall j \in [\sqrt{N}]$ 2. return $a := (a_1, \dots, a_{\sqrt{N}})$ <u>Rec($(\text{st}, i_{\text{col}}), d = \{d_j\}, a = \{a_j\}$)</u> 1. $t_j \leftarrow \text{APIR.Rec}(\text{st}, d_j, a_j) \quad \forall j \in [\sqrt{N}]$ 2. if $\exists j \in [\sqrt{N}]$ with $t_j = \perp$, return $\perp$ 3. else return $t_{i_{\text{col}}}$

Fig. 3. Rebalancing $\overline{\text{APIR}}$ of an PIR scheme APIR. We assume that the public parameters pp output by **Setup** are available to all algorithms, but omit them for ease of readability.

Similar to Lemma 2, this relaxed notion of privacy with abort can be proven under the assumption of *integrity*, *standard privacy*, and a modified type of *answer extractability*. The latter requires the probability in Eq. (3) to hold with probability $1 - \frac{1}{2^\kappa} - \text{negl}(\lambda)$ only, but also imposes a restriction that requires all non- $\perp$ answers given by $\mathcal{E}_{\text{ans}}$ to reconstruct correctly (except with negligible probability).

Concretely, in our APIR construction (Sect. 4), κ may correspond to the number of individual “challenges” that the client includes in a single validation query. Since each such challenge has a success probability of at most $\frac{1}{2} + \text{negl}(\lambda)$ if the server is cheating, the client’s confidence in the server’s inability to break privacy with abort will be $1 - \frac{1}{2^\kappa} - \text{negl}(\lambda)$. Given this, it is possible to prove the modified answer extractability sketched above, but for the purpose of a clean presentation we will stick to the setting with $\kappa = \lambda$ for the remainder of this paper.

3.4 Rebalancing

We revisit the standard rebalancing trick, originally proposed in [35], in the context of Authenticated PIR. Given a PIR scheme APIR, we construct a *rebalanced* scheme $\overline{\text{APIR}}$. It splits a database of size N into $\sqrt{N}$ chunks of size $\sqrt{N}$

<u>Setup($1^\lambda, 1^N$)</u> 1. $(\mathbb{G}, q, g) = \text{pp}' \leftarrow \text{GG}(1^N)$ 2. $\mathbf{h} \leftarrow_{\$} \mathbb{G}^N$ 3. return $\text{pp} := (\mathbb{G}, q, g, \mathbf{h})$	<u>Digest($\mathbf{x}$)</u> 1. return $d := \prod_{i \in [N]} \mathbf{h}[i]^{\mathbf{x}[i]}$
Validation <u>VQuery()</u> 1. $\mathbf{s}_j \leftarrow_{\$} \{0, 1\}^N \quad \forall j \in [\lambda]$ 2. $\text{st}_j, \tilde{\mathbf{s}}_j \leftarrow \text{randomize}(g^{\mathbf{s}_j}) \quad \forall j \in [\lambda]$ 3. return $\text{st} := (\text{st}_j)_{j \in [\lambda]}, v := (\tilde{\mathbf{s}}_j)_{j \in [\lambda]}$ <u>VAnswer($\mathbf{x}, v = (\tilde{\mathbf{s}}_j)_{j \in [\lambda]}$)</u> 1. $\tilde{s}_j := \prod_{i \in [N]} (\tilde{\mathbf{s}}_j[i])^{\mathbf{x}[i]} \quad \forall j \in [\lambda]$ 2. return $a := (\tilde{s}_j)_{j \in [\lambda]}$ <u>VCheck($\text{st} = (\text{st}_j)_{j \in [\lambda]}, d, a = (\tilde{s}_j)_{j \in [\lambda]}$)</u> 1. $s_j \leftarrow \text{recover}(\text{st}_j, d, \tilde{s}_j) \quad \forall j \in [\lambda]$ 2. return $\mathbb{I}[s_1, \dots, s_\lambda \in \{g^0, \dots, g^N\}]$	Query <u>Query(i)</u> 1. $\text{st}, \tilde{\mathbf{e}} \leftarrow \text{randomize}(g^{\mathbf{e}_i})$ 2. return $\text{st}, q := \tilde{\mathbf{e}}$ <u>Answer($\mathbf{x}, q = \tilde{\mathbf{e}}$)</u> 1. return $a := \tilde{e} = \prod_{i \in [N]} (\tilde{\mathbf{e}}[i])^{\mathbf{x}[i]}$ <u>Rec($\text{st}, d, a = \tilde{e}$)</u> 1. $e \leftarrow \text{recover}(\text{st}, d, \tilde{e})$ 2. if $e \in \{g^{-N}, \dots, g^N\} \setminus \{g^0\}$, return 1 3. else if $e = g^0$, return 0 4. else return $\perp$

Fig. 4. Our PIR scheme. We assume that the public parameters pp output by **Setup** are available to all algorithms, but omit them for ease of readability.

(w.l.o.g. we may assume that N is a perfect square, as the database can always be padded with dummy values).

For database $\mathbf{x} \in \{0, 1\}^N$, we denote by $\mathbf{x}_i$ the i -th chunk (which is a vector in $\{0, 1\}^{\sqrt{N}}$), s.t. $\mathbf{x} = (\mathbf{x}_1^T, \dots, \mathbf{x}_{\sqrt{N}}^T)$. Any index $i \in [N]$ can be split into two values $i_{\text{row}} \in [\sqrt{N}]$ and $i_{\text{col}} \in [\sqrt{N}]$, s.t. $\mathbf{x}[i] = \mathbf{x}_{i_{\text{col}}}[i_{\text{row}}]$ (i_{col} denotes the chunk that i is in, and i_{row} denotes the position within that chunk). The full scheme is described in Fig. 3.

It is easy to see that if APIR fulfills *integrity*, then $\overline{\text{APIR}}$ fulfills integrity as well, and if APIR fulfills *standard privacy*, then $\overline{\text{APIR}}$ fulfills standard privacy. Furthermore, we can prove that if APIR fulfills *answer extractability*, then $\overline{\text{APIR}}$ does so as well.

Lemma 3. *If APIR fulfills answer extractability, then the re-balanced scheme $\overline{\text{APIR}}$ fulfills answer extractability as well.*

We give the proof of Lemma 3 in the full version [20].

Remark 3. As in [17], we can compress the digest by letting $\overline{\text{APIR}}.\text{Digest}$ return a short hash of $(d_1, \dots, d_{\sqrt{N}})$ (any collision-resistant hash function chosen by

$\overline{\text{APIR.Setup}}$ suffices). Then, $\overline{\text{APIR.VAnswer}}$ and $\overline{\text{APIR.Answer}}$ would both output the full digest $(d_1, \dots, d_{\sqrt{N}})$ (in addition to the actual answers), and $\overline{\text{APIR.VCheck}}$ and $\overline{\text{APIR.Rec}}$ test whether it matches the digest hash before doing anything else.

4 The Authenticated PIR Scheme and Its Security

4.1 Description and Security

This section describes our APIR scheme, and states the security theorems we then prove below. A description of the scheme is given in Fig. 4.

MAIN IDEAS BEHIND THE SCHEME. We discuss first the unbalanced version of our scheme, which we refer to as APIR. We obtain a re-balanced version $\overline{\text{APIR}}$ which achieve $O(\sqrt{N})$ complexity following the transformation of Sect. 3.4. The scheme is based on the DDH scheme from Colombo et al. [17]. The server is in fact *identical*: It initially publishes a Pedersen hash $d = \prod_{i \in [N]} \mathbf{h}[i]^{\mathbf{x}[i]}$ of the database as the digest, where $\mathbf{h} \in \mathbb{G}^N$ is a vector of independent random generators for $\mathbb{G}$. The server then answers any query vector $\mathbf{p} \in \mathbb{G}^N$ by returning the product $\prod_{i \in [N]} \mathbf{p}[i]^{\mathbf{x}[i]}$.

ITEM RETRIEVAL. To retrieve an item $i \in [N]$ privately, the client sends a randomized version of the vector $\mathbf{e} = g^{\mathbf{e}_i}$, where $\mathbf{e}_i$ is the i -th unit vector. The easiest way to hide the query, under the DDH assumption, is to multiply $\mathbf{e}$ component-wise with a blinding vector $\mathbf{h}^r$. Then, the client would receive the response $\tilde{e} = d^r \cdot g^{\mathbf{x}[i]}$, and $e = g^{\mathbf{x}[i]}$ is recovered by dividing off d^r . To achieve integrity, however, we will need to additionally randomize $\mathbf{e}$ by exponentiating it with a non-zero $\alpha \leftarrow \mathbb{Z}_q^*$, where q is the group order. Then, if the server answers honestly with a value $\tilde{e}$, we get the final answer as $e = (d^{-r} \tilde{e})^{1/\alpha}$. However, in contrast to [17], to accommodate the validation process we describe next, we will need to allow e to take values in $\{g^{-N}, \dots, g^N\}$. We think of any value different than $1 = g^0$ as corresponding to the output 1, whereas $1 = g^0$ is mapped to 0.

In summary, we define a general query randomization mechanism (and the associated recovery operation), used for both validation and normal queries in Fig. 4, as

$\text{randomize}(\mathbf{p})$	$\text{recover}(\text{st} = (r, \alpha), d, \tilde{p})$
1 : $r \leftarrow \mathbb{Z}_q$	1 : return $p := (d^{-r} \cdot \tilde{p})^{1/\alpha}$
2 : $\alpha \leftarrow \mathbb{Z}_q^*$	
3 : return $\text{st} := (r, \alpha), \tilde{\mathbf{p}} := \mathbf{h}^r \circ \mathbf{p}^\alpha$	

GENERALIZED QUERIES. It will be convenient to think of queries in a more general sense, where the client can ask for any *selection vector* $\mathbf{p} \in \mathbb{Z}_q^N$, turning it into a corresponding vector of group elements $\mathbf{p} = g^{\mathbf{p}}$, which is then randomized

as $\tilde{\mathbf{p}}$. We usually denote the answer as $\tilde{p}$, and the de-randomized answer as p . This defines some notational conventions we follow throughout our proofs.

DIGEST VALIDATION. Our digest validation procedure consists of λ parallel generalized queries where the selection vectors are independent random binary vectors. The client accepts if all results of these queries are in the set $\{g^0, g^1, \dots, g^N\}$. Note that if the server produces the digest and answers honestly, the check passes. It also passes on a “slightly malformed” database, in which the sum of any subset of database components is between 0 and N (say one database entry is N , all others are 0.) This is the main reason why we do relax the allowable range of values for e for PIR queries.

SECURITY OF APIR. Below, we prove the following three theorems in Sects. 4.2, 4.3, and 4.4, respectively.

Theorem 1. *Assuming that the DDH assumption holds, APIR fulfills integrity.*

Theorem 2. *Assuming that the DDH assumption holds, APIR fulfills standard privacy.*

Theorem 3. *Assuming that the DDH assumption holds, APIR fulfills answer extractability.*

Combining these theorems with Lemmas 2 and 3 immediately yields the following corollary.

Corollary 1. *Assuming that the DDH assumption holds, both APIR and its rebalancing $\overline{\text{APIR}}$ fulfill integrity, standard privacy, and privacy with abort.*

EFFICIENCY. After rebalancing our scheme as described in Sect. 3.4, a single query and its answer both have size proportional to that of $\sqrt{N}$ group elements. Each validation consists of λ individual queries, each with cost identical to that of a regular query (thus, v consists of $\lambda\sqrt{N}$ group elements in total). The relaxed security notion described in Sect. 3.3 is achieved by performing only κ instead of λ individual queries per validation phase.

Note that if we apply the rebalancing transformation naïvely, $\overline{\text{APIR}}.\text{Rec}$ would run $\text{APIR}.\text{Rec}$ for $\sqrt{N}$ times, each of them re-computing the same $2\sqrt{N}+1$ powers of g during the call to `recover`. This would result in computational complexity $O(N)$. Obviously, a client should compute these powers only *once*, store them in a suitable data structure (e.g., a hash table), and perform a single constant-time look up. While we did not perform own benchmark, it should be clear that the same performance profile as in [17] would emerge here.

Remark 4. In order to remove clutter, when clear from context (e.g., configuration $c = (\text{pp}, d, \text{st}_{\mathcal{D}})$ is fixed) we will typically omit the parameter d passed to `recover`($\text{st}, d, \tilde{p}$) when proving security of APIR in the following sections.

In some places, an answer $\tilde{p} \in \mathbb{G} \cup \{\perp\}$ may be equal to $\perp$. Any operations involving $\tilde{p} = \perp$ will result in $\perp$, e.g. `recover`($\text{st}, d, \perp$) = $\perp$.

4.2 Proof of Theorem 1 (Integrity)

Lemma 4 establishes a slightly stronger version of integrity that will be useful within proofs of other theorems below. Specifically:

1. We let $\mathcal{A}$ choose the query vector $\mathbf{p} \in \mathbb{G}^N$ that will be randomized. For integrity, it suffices to consider $\mathbf{p} = g^{\mathbf{e}_i}$.
2. We let $\mathcal{A}$ choose a *small* set S , along with its answers $\tilde{p}, \tilde{p}'$. The adversary wins whenever their two answers $\tilde{p}$ and $\tilde{p}'$ are distinct, and *the ratio between the two recovered answers* is in S .

We show that the probability of winning is small, and we refer to this as “Evasive ratio” lemma. Below, we show integrity easily follows.

Lemma 4 (Evasive Ratio). *Assuming that the DDH assumption holds, then for any ppt adversaries $\mathcal{D}, \mathcal{A}_1, \mathcal{A}_2$ (where $\mathcal{A}_1(c)$ outputs state $\text{st}_{\mathcal{A}}$ and a query vector $\mathbf{p} \in \mathbb{G}^N$, and $\mathcal{A}_2(\text{st}_{\mathcal{A}}, \tilde{\mathbf{p}})$ outputs two answers $\tilde{p}, \tilde{p}' \in \mathbb{G} \cup \{\perp\}$ and a set $S \subseteq \mathbb{G}$ of size $|S| \leq \text{poly}(\lambda, N)$) and database sizes $N(\lambda) \leq \text{poly}(\lambda)$:*

$$\Pr \left[\begin{array}{l} \tilde{p} \neq \tilde{p}' \wedge \\ \frac{\text{recover}(\text{st}, \tilde{p}')}{\text{recover}(\text{st}, \tilde{p})} \in S \end{array} \middle| \begin{array}{l} c = (\text{pp}, d, \text{st}_{\mathcal{D}}) \leftarrow \text{Conf}^{\mathcal{D}}(\lambda, N) \\ \text{st}_{\mathcal{A}}, \mathbf{p} \leftarrow \mathcal{A}_1(c) \\ \text{st}, \tilde{\mathbf{p}} \leftarrow \text{randomize}(\mathbf{p}) \\ \tilde{p}, \tilde{p}', S \leftarrow \mathcal{A}_2(\text{st}_{\mathcal{A}}, \tilde{\mathbf{p}}) \end{array} \right] \leq \text{negl}(\lambda) .$$

Proof. By unrolling the calls to `randomize` and `recover` in the probability above, we need to show

$$\Pr[\text{Hyb}_0(\lambda, N) = 1] \leq \text{negl}(\lambda) ,$$

where Hyb_0 is defined as follows.

$$\begin{array}{l} \text{Hyb}_0(\lambda, N) \\ \hline 1 : c = (\text{pp}, d, \text{st}_{\mathcal{D}}) \leftarrow \text{Conf}^{\mathcal{D}}(\lambda, N) \\ 2 : \text{st}_{\mathcal{A}}, \mathbf{p} \leftarrow \mathcal{A}_1(c) \\ 3 : (r, \alpha) \leftarrow \mathbb{S} \mathbb{Z}_q \times \mathbb{Z}_q^* \\ 4 : \tilde{\mathbf{p}} := \mathbf{h}^r \circ \mathbf{p}^\alpha \\ 5 : \tilde{p}, \tilde{p}', S \leftarrow \mathcal{A}_2(\text{st}_{\mathcal{A}}, \tilde{\mathbf{p}}) \\ 6 : \text{return } \mathbb{I} \left[\tilde{p} \neq \tilde{p}' \wedge \left(\frac{\tilde{p}'}{\tilde{p}} \right)^{1/\alpha} \in S \right] \end{array}$$

We do so by a hybrid argument.

Hyb₁ Note that in Hyb_0 , the random $r \leftarrow \mathbb{S} \mathbb{Z}_q$ is *only* used for computing $\mathbf{h}^r$. By DDH (Lemma 1), this means that $\mathbf{h}^r$ looks identical to a uniformly random vector in $\mathbb{G}^N$. Therefore, in Hyb_1 , we replace the lines

$$\begin{array}{ll} (r, \alpha) \leftarrow \mathbb{S} \mathbb{Z}_q \times \mathbb{Z}_q^* & \alpha \leftarrow \mathbb{S} \mathbb{Z}_q^* \\ \tilde{\mathbf{p}} := \mathbf{h}^r \circ \mathbf{p}^\alpha & \text{by } \begin{array}{l} \bar{\mathbf{h}} \leftarrow \mathbb{S} \mathbb{G}^N \\ \tilde{\mathbf{p}} := \bar{\mathbf{h}} \circ \mathbf{p}^\alpha \end{array} \end{array} .$$

We use the following ppt adversary $\mathcal{B}$ distinguishing between the two distributions in Lemma 1:

- $\mathcal{B}(\text{pp}', \bar{\mathbf{h}})$: Run $(\text{st}_{\mathcal{D}}, d) \leftarrow \mathcal{D}(\text{pp})$ (with $\text{pp} := (\text{pp}', \bar{\mathbf{h}})$) to create the configuration $c := (\text{pp}, d, \text{st}_{\mathcal{D}})$, and run $\text{st}_{\mathcal{A}}, \mathbf{p} \leftarrow \mathcal{A}_1(c)$. Then, sample $\alpha \leftarrow \mathbb{Z}_q^*$ and compute $\tilde{\mathbf{p}} := \bar{\mathbf{h}} \circ \mathbf{p}^\alpha$. Run $\tilde{p}, \tilde{p}', S \leftarrow \mathcal{A}_2(\text{st}_{\mathcal{A}}, \tilde{\mathbf{p}})$, and return $\mathbb{I}[(\frac{\tilde{p}'}{\tilde{p}})^{1/\alpha} \in S \setminus \{1\}]$.

Note that running $\mathcal{B}$ on left-hand side distribution in Lemma 1 is identical to Hyb_0 (i.e., $\bar{\mathbf{h}} := \mathbf{h}^r$ for $r \leftarrow \mathbb{Z}_q$), and that running $\mathcal{B}$ on the right-hand side distribution is identical to Hyb_1 (i.e., $\bar{\mathbf{h}} \leftarrow \mathbb{G}^N$). Thus, by Lemma 1, we get

$$|\Pr[\text{Hyb}_1(\lambda, N) = 1] - \Pr[\text{Hyb}_0(\lambda, N) = 1]| \leq \text{negl}(\lambda) .$$

Hyb_2 Note that $\tilde{\mathbf{p}} := \bar{\mathbf{h}} \circ \mathbf{p}^\alpha$ for random $\bar{\mathbf{h}} \leftarrow \mathbb{G}^N$ is identically distributed as a uniformly random vector in $\mathbb{G}^N$. Thus, in Hyb_2 , we replace the lines

$$\begin{array}{l} \bar{\mathbf{h}} \leftarrow \mathbb{G}^N \\ \tilde{\mathbf{p}} := \bar{\mathbf{h}} \circ \mathbf{p}^\alpha \end{array} \quad \text{by} \quad \tilde{\mathbf{p}} \leftarrow \mathbb{G}^N .$$

We have $\text{Hyb}_2 \equiv \text{Hyb}_1$.

Note that in Hyb_2 , $\alpha \leftarrow \mathbb{Z}_q^*$ is only used in the very last line, which returns $\mathbb{I}[\tilde{p} \neq \tilde{p}' \wedge (\frac{\tilde{p}'}{\tilde{p}})^{1/\alpha} \in S]$. Furthermore, if $\tilde{p} \neq \tilde{p}'$ (i.e., $\frac{\tilde{p}'}{\tilde{p}} \neq 1$), then $(\frac{\tilde{p}'}{\tilde{p}})^{1/\alpha}$ will be $\perp$ or a uniformly random element in $\mathbb{G} \setminus \{1\}$. Thus, the probability of returning 1 is at most $\frac{|S|}{q-1}$. Because of $|S| \leq \text{poly}(\lambda, N)$, we get

$$\Pr[\text{Hyb}_2(\lambda, N) = 1] \leq \text{negl}(\lambda) ,$$

which concludes the proof by our hybrid argument. $\square$

Now we prove integrity: let $\mathcal{D}, \mathcal{A}$ be an adversary for the integrity game, $N(\lambda) \leq \text{poly}(\lambda)$ the database size, and $i(\lambda) \in [N]$ an index. We need to show

$$\Pr \left[\begin{array}{l} \tilde{p} \neq \tilde{p}' \wedge \\ \text{Rec}(\text{st}, d, \tilde{p}) \neq \perp \wedge \\ \text{Rec}(\text{st}, d, \tilde{p}') \neq \perp \end{array} \middle| \begin{array}{l} c = (\text{pp}, d, \text{st}_{\mathcal{D}}) \leftarrow \text{Conf}^{\mathcal{D}}(\lambda, N) \\ \text{st}, \tilde{\mathbf{p}} \leftarrow \text{randomize}(g^{e_i}) \\ \tilde{p}, \tilde{p}' \leftarrow \mathcal{A}(c, \tilde{\mathbf{p}}) \end{array} \right] \leq \text{negl}(\lambda) .$$

Note that the event $\text{Rec}(\text{st}, d, \tilde{p}) \neq \perp$ holds iff $\text{recover}(\text{st}, \tilde{p}) \in \{g^{-N}, \dots, g^N\}$ (and analogously for $\tilde{p}'$). Therefore,

$$\begin{array}{l} \tilde{p} \neq \tilde{p}' \wedge \\ \text{Rec}(\text{st}, d, \tilde{p}) \neq \perp \wedge \\ \text{Rec}(\text{st}, d, \tilde{p}') \neq \perp \end{array} \quad \Rightarrow \quad \frac{\text{recover}(\text{st}, \tilde{p}')}{\text{recover}(\text{st}, \tilde{p})} \in \{g^{-2N}, \dots, g^{2N}\} ,$$

and it suffices to prove that the probability of the event on the right-hand side is negligible. This follows immediately by applying Lemma 4 to the following adversary $\mathcal{D}, \mathcal{B}_1, \mathcal{B}_2$.

- $\mathcal{B}_1(c)$ outputs $\text{st}_{\mathcal{B}} := c, \mathbf{p} := g^{e_i}$.
- $\mathcal{B}_2(\text{st}_{\mathcal{B}}, \tilde{\mathbf{p}})$ runs $\tilde{p}, \tilde{p}' \leftarrow \mathcal{A}(c, \tilde{\mathbf{p}})$, chooses $S := \{g^{-2N}, \dots, g^{2N}\}$, and outputs $\tilde{p}, \tilde{p}', S$.

4.3 Proof of Theorem 2 (Standard Privacy)

We prove a stronger notion of privacy, called *chosen-vector indistinguishability*, which will be useful elsewhere, and whose security game is as follows:

Game $\text{CVA}^{\mathcal{A}^*}(\lambda, N, b)$	
1 : $c = (\mathbf{pp}, d, \mathbf{st}_{\mathcal{D}}) \leftarrow \text{Conf}^{\mathcal{D}}(\lambda, N)$	
2 : $\mathbf{st}_{\mathcal{A}}, \mathbf{v}_0, \mathbf{v}_1 \leftarrow \mathcal{A}_1(c); r \leftarrow \$_{\mathbb{Z}_q}$	
3 : return $\mathcal{A}_2(\mathbf{st}_{\mathcal{A}}, \mathbf{h}^r \circ \mathbf{v}_b)$	

Lemma 5 (Chosen Vector Indistinguishability). *Assuming that the DDH assumption holds, then for any ppt adversaries $\mathcal{A}^* = (\mathcal{D}, \mathcal{A}_1, \mathcal{A}_2)$ and database sizes $N(\lambda) \leq \text{poly}(\lambda)$,*

$$|\Pr[\text{CVA}^{\mathcal{A}^*}(\lambda, N, 0) = 1] - \Pr[\text{CVA}^{\mathcal{A}^*}(\lambda, N, 1) = 1]| \leq \text{negl}(\lambda).$$

Remark 5. In order to simplify reductions involving $\text{CVA}^{\mathcal{A}^*}$, we defined this experiment to sample a configuration $c = (\mathbf{pp}, d, \mathbf{st}_{\mathcal{D}}) \leftarrow \text{Conf}^{\mathcal{D}}(\lambda, N)$. Lemma 5 would still hold if $\text{CVA}^{\mathcal{A}^*}$ does not run the digest generator $\mathbf{st}_{\mathcal{D}}, d \leftarrow \mathcal{D}(\mathbf{pp})$, and instead only provides $\mathbf{pp} \leftarrow \text{Setup}(1^\lambda, 1^N)$ as input to $\mathcal{A}_1$.

Proof. We show that both $\text{CVA}^{\mathcal{A}^*}(\lambda, N, 0)$ and $\text{CVA}^{\mathcal{A}^*}(\lambda, N, 1)$ are close to Hyb defined below.

Hyb(λ, N)	
1 : $c = (\mathbf{pp}, d, \mathbf{st}_{\mathcal{D}}) \leftarrow \text{Conf}^{\mathcal{D}}(\lambda, N)$	
2 : $\mathbf{st}_{\mathcal{A}}, \mathbf{v}_0, \mathbf{v}_1 \leftarrow \mathcal{A}_1(c)$	
3 : $\tilde{\mathbf{v}} \leftarrow \$_{\mathbb{G}^N}$	
4 : return $\mathcal{A}_2(c, \tilde{\mathbf{v}})$	

By a hybrid argument, it suffices to show

$$|\Pr[\text{CVA}^{\mathcal{A}^*}(\lambda, N, b) = 1] - \Pr[\text{Hyb}(\lambda, N) = 1]| \leq \text{negl}(\lambda), \quad (4)$$

for both $b = 0, 1$.

To show that Eq. 4 holds, we construct an adversary $\mathcal{B}$ that distinguishes between the two distributions stated in Lemma 1 in the following way:

- $\mathcal{B}(\mathbf{pp}', \bar{\mathbf{h}})$: Run $(\mathbf{st}_{\mathcal{D}}, d) \leftarrow \mathcal{D}(\mathbf{pp})$ (with $\mathbf{pp} := (\mathbf{pp}', \bar{\mathbf{h}})$) to create the configuration $c := (\mathbf{pp}, d, \mathbf{st}_{\mathcal{D}})$. Then, compute $\mathbf{st}_{\mathcal{A}}, \mathbf{v}_0, \mathbf{v}_1 \leftarrow \mathcal{A}_1(c)$. Run and output $\mathcal{A}_2(\mathbf{st}_{\mathcal{A}}, \bar{\mathbf{h}} \circ \mathbf{v}_b)$.

Note that running $\mathcal{B}$ on the left-hand side distribution of Lemma 1 is identical to $\text{CVA}^{\mathcal{A}^*}(\lambda, N, b)$. Furthermore, running $\mathcal{B}$ on the right-hand side distribution of Lemma 1 is identical to $\text{Hyb}(\lambda, N)$. Equation 4 follows. $\square$

Procedure $\mathcal{E}_{\text{ans}}(1^r, c = (\text{pp}, d, \text{st}_{\mathcal{D}}), \mathbf{p})$	Procedure $\bar{\mathcal{V}}(c, \tilde{\mathbf{p}})$
1 : repeat $\lambda \cdot r^3$ times	1 : $\mathbf{s}_j \leftarrow_{\$} \{0, 1\}^N \quad \forall j \in [\lambda] \setminus \{1\}$
2 : $\text{st}, \tilde{\mathbf{p}} \leftarrow \text{randomize}(\mathbf{p})$	2 : $_, \tilde{\mathbf{s}}_j \leftarrow \text{randomize}(g^{\mathbf{s}_j}) \quad \forall j \in [\lambda] \setminus \{1\}$
3 : $\tilde{p} \leftarrow \bar{\mathcal{V}}(c, \tilde{\mathbf{p}})$	3 : $(\tilde{p}, _, \dots, _) \leftarrow \mathcal{V}(c, (\tilde{\mathbf{p}}, \tilde{\mathbf{s}}_2, \dots, \tilde{\mathbf{s}}_\lambda))$
4 : $\text{st}', \tilde{\mathbf{p}}' \leftarrow \text{randomize}(\tilde{\mathbf{p}})$	4 : return $\tilde{p}$
5 : $\tilde{p}' \leftarrow \text{recover}(\text{st}', \bar{\mathcal{V}}(c, \tilde{\mathbf{p}}'))$	
6 : if $\tilde{p} = \tilde{p}'$, return $\text{recover}(\text{st}, \tilde{p})$	
7 : return $\perp$	

Fig. 5. Description of the extractor $\mathcal{E}_{\text{ans}}$ (left-hand side) used in the proof of Theorem 3. The right-hand side describes a sub-procedure used within $\mathcal{E}_{\text{ans}}$.

Now we prove standard privacy. Let $\mathcal{D}, \mathcal{A}$ be an adversary for the privacy game, $N(\lambda) \leq \text{poly}(\lambda)$ the database size, and $i_0(\lambda), i_1(\lambda) \in [N]$ indices.

We plug in the following adversary $\mathcal{B}^* := (\mathcal{D}, \mathcal{B}_1, \mathcal{B}_2)$ into Lemma 5:

- $\mathcal{B}_1(c)$: sample $\alpha \leftarrow_{\$} \mathbb{Z}_q^*$ and output $\text{st}_{\mathcal{A}} := c$, $\mathbf{v}_0 := (g^{e_{i_0}})^\alpha$, $\mathbf{v}_1 := (g^{e_{i_1}})^\alpha$.
- $\mathcal{B}_2(c, \hat{\mathbf{v}})$: run and output $\mathcal{A}(c, \hat{\mathbf{v}})$.

Privacy follows from Lemma 5 because of $\text{CVA}^{\mathcal{B}^*}(\lambda, N, b) \equiv \text{PRIV}^{\mathcal{D}, \mathcal{A}}(\lambda, N, i_b)$.

4.4 Proof of Theorem 3 (Answer Extractability)

Let $(\mathcal{D}, \mathcal{V})$ be an adversary against answer extractability for our PIR scheme, where $\mathcal{D}$ is the digest generator, and $\mathcal{V}$ answers validation queries. Our answer extractor $\mathcal{E}_{\text{ans}}(1^r, c, \mathbf{p})$ is described in Fig. 5. The core idea is to repeatedly invoke $\mathcal{V}$ to get answers for $\mathbf{p}$. Because $\mathcal{V}$ expects a list $v = (\tilde{\mathbf{s}}_1, \dots, \tilde{\mathbf{s}}_\lambda)$ of λ validation queries, we use a wrapper procedure $\bar{\mathcal{V}}(c, \tilde{\mathbf{p}})$ which embeds $\tilde{\mathbf{p}}$ as the first component of a vector consisting of otherwise honestly generated $\lambda - 1$ validation queries.

Throughout λr^3 iterations, the extractor $\mathcal{E}_{\text{ans}}$ queries first $\bar{\mathcal{V}}(c, \tilde{\mathbf{p}})$ on a fresh randomization $\tilde{\mathbf{p}}$ of $\mathbf{p}$, and obtains answer $\tilde{p}$. Then, $\mathcal{E}_{\text{ans}}$ generates a fresh randomization $\tilde{\mathbf{p}}'$ of $\tilde{\mathbf{p}}$, and runs $\bar{\mathcal{V}}(c, \tilde{\mathbf{p}}')$. Intuitively, we use the fact that the recovered answer $\tilde{p}'$ matches $\tilde{p}$ as an indication that $\bar{\mathcal{V}}$ has answered correctly. Therefore, in this case, $\mathcal{E}_{\text{ans}}$ returns $\text{recover}(\text{st}, \tilde{p})$. However, if no iteration is successful, then $\mathcal{E}_{\text{ans}}$ returns $\perp$.

ROADMAP. We now give an informal description of why $\mathcal{E}_{\text{ans}}$ indeed works. We refer to several lemmas that we will state formally afterwards, followed by the full proof of answer extractability using these lemmas.

First, Lemma 9 establishes what we refer to as the *answer guarantee* property, i.e., the fact that $\mathcal{E}_{\text{ans}}$ never answers with $\perp$ on any (even adversarially chosen) query, as long as $\nu_{\mathcal{D}, \mathcal{V}}(c)$ is large enough. The proof will use crucially the indistinguishability of any two randomized queries.

The bulk of the proof considers the scenario in which we query $\mathcal{E}_{\text{ans}}$ on g^{e_i} for all $i \in [N]$, obtaining answers

$$e_i \leftarrow \mathcal{E}_{\text{ans}}(1^{r(\lambda)}, c, g^{e_i}) \quad \forall i \in [N].$$

The queries g^{e_i} are not randomized, which is going to simplify our analysis—Lemma 6 (*randomization independence*) shows that this is equivalent to randomizing g^{e_i} first, and later recovering $\mathcal{E}_{\text{ans}}$'s answer. Therefore, to infer answer extractability, we simply need to show that $e_i \in \{g^{-N}, \dots, g^N\}$ for each $i \in [N]$.

Also, assume that we already know that $\mathcal{E}_{\text{ans}}$ answers *validations* successfully, i.e., for λ random vectors $\mathbf{s}_1, \dots, \mathbf{s}_\lambda \in \{0, 1\}^N$ the answers $s'_k \leftarrow \mathcal{E}_{\text{ans}}(1^{r(\lambda)}, c, g^{\mathbf{s}_k})$ will fulfill $s'_1, \dots, s'_\lambda \in \{g^0, \dots, g^N\}$. To draw a connection to the individual answers e_i , we need a *homomorphism* property (Lemma 8), which states that

$$s'_k = \prod_{i \in [N]} e_i^{s_k[i]} \quad \forall k \in [\lambda].$$

Because vectors $\mathbf{s}_k$ are random, a simple argument (Lemma 10, *random subset testing*) shows that, due to the small range of all s'_k , with probability at least $1 - \frac{1}{2^\lambda}$, all the individual answers are also in a small range: $e_i \in \{g^{-N}, \dots, g^N\}$.

It remains to ensure that there is a good chance of $\mathcal{E}_{\text{ans}}$ answering all validation queries $\mathbf{s}_1, \dots, \mathbf{s}_\lambda$ correctly. Indeed, whenever $\bar{\mathcal{V}}(c, \tilde{\mathbf{p}})$ manages to reply to the randomization $\tilde{\mathbf{p}}$ of a query $g^{\mathbf{s}_k}$ with a value $s_k \in \{g^0, \dots, g^N\}$, then w.h.p., $\mathcal{E}_{\text{ans}}$ will output the exact same value $s'_k = s_k$ (Lemma 7, *$\mathcal{E}_{\text{ans}}$ agrees with small answers*). Therefore, whenever $\nu_{\mathcal{D}, \mathcal{V}}(c)$ is large enough, we can be sure that (after testing sufficiently many different validations), $\mathcal{E}_{\text{ans}}$'s responses fulfill $s'_1, \dots, s'_\lambda \in \{g^0, \dots, g^N\}$, thereby implying $e_1, \dots, e_N \in \{g^{-N}, \dots, g^N\}$.

BASIC EXTRACTOR PROPERTIES. We state four essential properties of our $\mathcal{E}_{\text{ans}}$ construction that suffice to prove answer extractability. We defer the proofs of the first three properties to Sects. 4.5, 4.6, and 4.7, and the lengthy proof of Lemma 9 can be found in the full version of this paper [20].

Lemma 6 (Randomization Independence). *For all $\lambda, N \in \mathbb{N}$ and configurations $c = (\text{pp}, d, \text{st}_{\mathcal{D}}) \in \text{Supp}(\text{Conf}^{\mathcal{D}}(\lambda, N))$, queries $\mathbf{p} \in \mathbb{G}^N$, randomizations $(\text{st}, \tilde{\mathbf{p}}) \in \text{Supp}(\text{randomize}(\mathbf{p}))$ of $\mathbf{p}$, and runtime parameters $r \in \mathbb{N}$, the following two distributions are exactly the same:*

$$\mathcal{E}_{\text{ans}}(1^r, c, \mathbf{p}) \quad \text{and} \quad \text{recover}(\text{st}, \mathcal{E}_{\text{ans}}(1^r, c, \tilde{\mathbf{p}}))$$

Lemma 7 (Agrees with Small Answers). *Assuming the DDH assumption holds, then for all ppt adversaries $\mathcal{A}_1, \mathcal{A}_2$ (where $\mathcal{A}_1(c)$ outputs a state $\text{st}_{\mathcal{A}}$ and a query vector $\mathbf{p} \in \mathbb{G}^N$, and $\mathcal{A}_2(\text{st}_{\mathcal{A}}, \tilde{\mathbf{p}}^*)$ outputs an answer $p^* \in \mathbb{G} \cup \{\perp\}$), database sizes $N(\lambda) \leq \text{poly}(\lambda)$, and polynomials $r(\lambda)$:*

$$\Pr \left[\begin{array}{l} p^* \in \{g^0, \dots, g^N\} \\ \Rightarrow p \in \{p^*, \perp\} \end{array} \middle| \begin{array}{l} c = (\text{pp}, d, \text{st}_{\mathcal{D}}) \leftarrow \text{Conf}^{\mathcal{D}}(\lambda, N) \\ \text{st}_{\mathcal{A}}, \mathbf{p} \leftarrow \mathcal{A}_1(c) \\ \text{st}^*, \tilde{\mathbf{p}}^* \leftarrow \text{randomize}(\mathbf{p}) \\ p^* \leftarrow \text{recover}(\text{st}^*, \mathcal{A}_2(\text{st}_{\mathcal{A}}, \tilde{\mathbf{p}}^*)) \\ p \leftarrow \mathcal{E}_{\text{ans}}(1^{r(\lambda)}, c, \mathbf{p}) \end{array} \right] \geq 1 - \text{negl}(\lambda).$$

Lemma 8 (Homomorphism). *Assuming the DDH assumption holds, then for all ppt adversaries $\mathcal{A}$ (where $\mathcal{A}(c)$ outputs m query vectors $\mathbf{p}_1, \dots, \mathbf{p}_m \in \mathbb{G}^N$ and m exponents $t_1, \dots, t_m \in \mathbb{Z}_q$), and $N(\lambda), r(\lambda) \leq \text{poly}(\lambda)$,*

$$\Pr \left[\begin{array}{l} p, p_1, \dots, p_m \neq \perp \\ \Rightarrow p = \prod_{i \in [m]} p_i^{t_i} \end{array} \middle| \begin{array}{l} c = (\mathbf{pp}, d, \text{st}_{\mathcal{D}}) \leftarrow \text{Conf}^{\mathcal{D}}(\lambda, N) \\ \mathbf{p}_1, \dots, \mathbf{p}_m, t_1, \dots, t_m \leftarrow \mathcal{A}(c) \\ p \leftarrow \mathcal{E}_{\text{ans}}(1^{r(\lambda)}, c, \prod_{i \in [m]} \mathbf{p}_i^{t_i}) \\ p_i \leftarrow \mathcal{E}_{\text{ans}}(1^{r(\lambda)}, c, \mathbf{p}_i) \quad \forall i \in [m] \end{array} \right] \geq 1 - \text{negl}(\lambda) .$$

Lemma 9 (Answer Guarantee). *Assuming the DDH assumption holds, then for all ppt adversaries $\mathcal{A}$ (where $\mathcal{A}(c)$ outputs a query vector $\mathbf{p} \in \mathbb{G}^N$), database sizes $N(\lambda) \leq \text{poly}(\lambda)$, and polynomials $r(\lambda)$,*

$$\Pr \left[\begin{array}{l} \nu_{\mathcal{D}, \nu}(c) \geq \frac{1}{r(\lambda)} \\ \Rightarrow \mathcal{E}_{\text{ans}}(1^{r(\lambda)}, c, \mathbf{p}) \neq \perp \end{array} \middle| \begin{array}{l} c = (\mathbf{pp}, d, \text{st}_{\mathcal{D}}) \leftarrow \text{Conf}^{\mathcal{D}}(\lambda, N) \\ \mathbf{p} \leftarrow \mathcal{A}(c) \end{array} \right] \geq 1 - \text{negl}(\lambda) .$$

RANDOM SUBSET TESTING. We will also make use of the following lemma, which is proved in Sect. 4.8.

Lemma 10 (Random Subset Testing). *Let $N \in \mathbb{N}$, let $\mathbb{G}$ be a cyclic group with generator $g \in \mathbb{G}$, and let $e_1, \dots, e_N \in \mathbb{G}$ be arbitrary group elements. Further, assume that for at least one $i^* \in [N]$, we have $e_{i^*} \notin \{g^{-N}, \dots, g^N\}$. Then, for $\mathbf{s}_1, \dots, \mathbf{s}_\lambda \leftarrow_{\$} \{0, 1\}^N$,*

$$\Pr \left[\forall k \in [\lambda] : \prod_{i \in [N]} e_i^{\mathbf{s}_k[i]} \in \{g^0, \dots, g^N\} \right] \leq \frac{1}{2^\lambda} .$$

THE ACTUAL PROOF. We are now ready to combine all of the above elements and prove answer extractability (Eq. 3), i.e., for all database sizes $N(\lambda) \leq \text{poly}(\lambda)$, indices $i(\lambda) \in [N(\lambda)]$, polynomials $r(\lambda) \leq \text{poly}(\lambda)$,

$$\Pr \left[\begin{array}{l} \nu_{\mathcal{D}, \nu}(c) \geq \frac{1}{r(\lambda)} \Rightarrow \\ \text{Rec}(\text{st}, d, \tilde{e}) \neq \perp \end{array} \middle| \begin{array}{l} c = (\mathbf{pp}, d, \text{st}_{\mathcal{D}}) \leftarrow \text{Conf}^{\mathcal{D}}(\lambda, N) \\ \text{st}, \tilde{\mathbf{e}} \leftarrow \text{Query}(i) \\ \tilde{e} \leftarrow \mathcal{E}_{\text{ans}}(1^{r(\lambda)}, c, \tilde{\mathbf{e}}) \end{array} \right] \geq 1 - \text{negl}(\lambda) .$$

Using Lemma 6, this simplifies to

$$\Pr \left[\begin{array}{l} \nu_{\mathcal{D}, \nu}(c) \geq \frac{1}{r(\lambda)} \Rightarrow \\ e \in \{g^{-N}, \dots, g^N\} \end{array} \middle| \begin{array}{l} c = (\mathbf{pp}, d, \text{st}_{\mathcal{D}}) \leftarrow \text{Conf}^{\mathcal{D}}(\lambda, N) \\ e \leftarrow \mathcal{E}_{\text{ans}}(1^{r(\lambda)}, c, g^{\mathbf{e}_i}) \end{array} \right] \geq 1 - \text{negl}(\lambda) .$$

In fact, we will prove something even stronger, which is that for all database sizes $N(\lambda) \leq \text{poly}(\lambda)$ and polynomials $r(\lambda)$, the answers of $\mathcal{E}_{\text{ans}}$ on the whole database (not just a single index i) is in $\{g^{-N}, \dots, g^N\}$:

$$\Pr \left[\begin{array}{l} \nu_{\mathcal{D}, \nu}(c) \geq \frac{1}{r(\lambda)} \Rightarrow \\ \forall i \in [N] : e_i \in \{g^{-N}, \dots, g^N\} \end{array} \middle| \begin{array}{l} c = (\mathbf{pp}, d, \text{st}_{\mathcal{D}}) \leftarrow \text{Conf}^{\mathcal{D}}(\lambda, N) \\ e_i \leftarrow \mathcal{E}_{\text{ans}}(1^{r(\lambda)}, c, g^{\mathbf{e}_i}) \quad \forall i \in [N] \end{array} \right] \geq 1 - \text{negl}(\lambda) . \quad (5)$$

To prove this, it is helpful to consider the following augmented experiment, which additionally introduces several validation queries, and both $\mathcal{V}$'s and $\mathcal{E}_{\text{ans}}$'s answers to those:

1. First, sample configuration $c = (\text{pp}, d, \text{st}_{\mathcal{D}}) \leftarrow \text{Conf}^{\mathcal{D}}(\lambda, N)$.
2. Second, query $\mathcal{E}_{\text{ans}}$ on the whole database:

$$e_i \leftarrow \mathcal{E}_{\text{ans}}(1^{r(\lambda)}, c, g^{e_i}) \quad \forall i \in [N]$$

3. Third, sample $\lambda \cdot r(\lambda)$ random validations, and run $\mathcal{V}$ on them:

$$\begin{aligned} \mathbf{s}_{j,k} &\leftarrow \$ \{0, 1\}^N & \forall j \in [\lambda \cdot r(\lambda)], k \in [\lambda] \\ \text{st}_{j,k}, \tilde{\mathbf{s}}_{j,k} &\leftarrow \text{randomize}(g^{\mathbf{s}_{j,k}}) & \forall j \in [\lambda \cdot r(\lambda)], k \in [\lambda] \\ \tilde{s}_{j,1}, \dots, \tilde{s}_{j,\lambda} &\leftarrow \mathcal{V}(c, (\tilde{\mathbf{s}}_{j,1}, \dots, \tilde{\mathbf{s}}_{j,\lambda})) & \forall j \in [\lambda \cdot r(\lambda)] \\ s_{j,k} &\leftarrow \text{recover}(\text{st}_{j,k}, \tilde{s}_{j,k}) & \forall j \in [\lambda \cdot r(\lambda)], k \in [\lambda] \end{aligned}$$

4. And finally, we run $\mathcal{E}_{\text{ans}}$ on the same validation queries:

$$s'_{j,k} \leftarrow \mathcal{E}_{\text{ans}}(1^{r(\lambda)}, c, g^{\mathbf{s}_{j,k}}) \quad \forall j \in [\lambda \cdot r(\lambda)], k \in [\lambda]$$

We are now going to apply Lemmas 7, 8, and 9. In each of the following probabilities, we refer to the distribution of all c , e_i , $\mathbf{s}_{j,k}$, $s_{j,k}$, and $s'_{j,k}$ as described above.

- Note that Step 3 performs $\lambda \cdot r(\lambda)$ independently sampled validations identically to $\text{VQuery}()$. Thus, under the assumption that configuration c has a high success probability ($\nu_{\mathcal{D}, \mathcal{V}}(c) \geq \frac{1}{r(\lambda)}$), there will be at least one validation that succeeds: The probability that for no $j^* \in [\lambda \cdot r(\lambda)]$, the j^* -th validation succeeds, is at most

$$\left(1 - \frac{1}{r(\lambda)}\right)^{\lambda \cdot r(\lambda)} \leq e^{-\frac{1}{r(\lambda)} \cdot \lambda \cdot r(\lambda)} = e^{-\lambda} \leq \text{negl}(\lambda).$$

A successful validation (as defined by $\text{VCheck}(\{\text{st}_{j^*,k}\}, d, \{\tilde{s}_{j^*,k}\})$) means that $s_{j^*,1}, \dots, s_{j^*,\lambda} \in \{g^0, \dots, g^N\}$, and therefore:

$$\Pr \left[\nu_{\mathcal{D}, \mathcal{V}}(c) \geq \frac{1}{r(\lambda)} \Rightarrow \exists j^* \in [\lambda \cdot r(\lambda)] \text{ with } s_{j^*,1}, \dots, s_{j^*,\lambda} \in \{g^0, \dots, g^N\} \right] \geq 1 - \text{negl}(\lambda) \quad (6)$$

- For all $j \in [\lambda \cdot r(\lambda)]$ and $k \in [\lambda]$, we apply Lemma 7 ($\mathcal{E}_{\text{ans}}$ agrees with small answers) to the following adversary $\mathcal{A}$:

- $\mathcal{A}_1(c)$ samples $\mathbf{s}_{j,k} \leftarrow \$ \{0, 1\}^N$ and returns $\text{st}_{\mathcal{A}} := (c, \mathbf{s}_{j,k})$ and $g^{\mathbf{s}_{j,k}}$.
- $\mathcal{A}_2(\text{st}_{\mathcal{A}}, \tilde{\mathbf{s}}_{j,k})$, for all $k' \in [\lambda] \setminus \{k\}$, samples $\mathbf{s}_{j,k'} \leftarrow \$ \{0, 1\}^N$ and randomizes $\text{st}_{j,k'}, \tilde{\mathbf{s}}_{j,k'} \leftarrow \text{randomize}(g^{\mathbf{s}_{j,k'}})$, then runs $\tilde{s}_{j,1}, \dots, \tilde{s}_{j,\lambda} \leftarrow \mathcal{V}(c, (\tilde{\mathbf{s}}_{j,1}, \dots, \tilde{\mathbf{s}}_{j,\lambda}))$ and outputs $\tilde{s}_{j,k}$.

Then, by Lemma 7, for $s_{j,k}$ produced by $\mathcal{V}$ that is within the set $\{g^0, \dots, g^N\}$, we can be certain that $\mathcal{E}_{\text{ans}}$'s answer $s'_{j,k}$ is equal to $s_{j,k}$ or $\perp$:

$$\Pr \left[s_{j,k} \in \{g^0, \dots, g^N\} \Rightarrow s'_{j,k} \in \{s_{j,k}, \perp\} \right] \geq 1 - \text{negl}(\lambda) \quad \forall j \in [\lambda \cdot r(\lambda)], k \in [\lambda] \quad (7)$$

- For all $j \in [\lambda \cdot r(\lambda)]$ and $k \in [\lambda]$, we apply Lemma 8 (*homomorphism*) to the following adversary $\mathcal{A}$:

- $\mathcal{A}(c)$ samples $\mathbf{s}_{j,k} \leftarrow_{\$} \{0,1\}^N$, and outputs the N unit queries $g^{\mathbf{e}_1}, \dots, g^{\mathbf{e}_N}$, and exponents $\mathbf{s}_{j,k}[1], \dots, \mathbf{s}_{j,k}[N]$.

Then, by Lemma 8, we can be certain that $\mathcal{E}_{\text{ans}}$'s answer $s'_{j,k}$ to validation query $g^{\mathbf{s}_{j,k}}$ is equal to the product $\prod_{i \in [N]} e_i^{\mathbf{s}_{j,k}[i]}$ of answers e_i (as long as none of these answers is $\perp$):

$$\Pr \left[\begin{array}{l} s'_{j,k}, e_1, \dots, e_N \neq \perp \Rightarrow \\ s'_{j,k} = \prod_{i \in [N]} e_i^{\mathbf{s}_{j,k}[i]} \end{array} \right] \geq 1 - \text{negl}(\lambda) \quad \forall j \in [\lambda \cdot r(\lambda)], k \in [\lambda] \quad (8)$$

- For all $i \in [N]$, we apply Lemma 9 (*guaranteed to answer*) to the following adversary $\mathcal{A}$:

- $\mathcal{A}(c)$ returns $g^{\mathbf{e}_i}$.

Then, by Lemma 9, whenever c has high validation success probability (i.e., $\nu_{\mathcal{D}, \nu}(c) \geq \frac{1}{r(\lambda)}$), then w.h.p. $\mathcal{E}_{\text{ans}}$ answered all queries without ever returning $\perp$:

$$\Pr \left[\nu_{\mathcal{D}, \nu}(c) \geq \frac{1}{r(\lambda)} \Rightarrow e_i \neq \perp \right] \geq 1 - \text{negl}(\lambda) \quad \forall i \in [N] \quad (9)$$

Similarly, Lemma 9 also shows that none of the answers $s'_{j,k}$ of $\mathcal{E}_{\text{ans}}$ is $\perp$:

$$\Pr \left[\nu_{\mathcal{D}, \nu}(c) \geq \frac{1}{r(\lambda)} \Rightarrow s'_{j,k} \neq \perp \right] \geq 1 - \text{negl}(\lambda) \quad \forall j \in [\lambda \cdot r(\lambda)], k \in [\lambda] \quad (10)$$

We can now take the union bound of the probabilities in Eqs. 6 through 10. If all of the given events hold, and furthermore $\nu_{\mathcal{D}, \nu}(c) \geq \frac{1}{r(\lambda)}$ for the sampled configuration c , then:

- By Eq. 6, there is a $j^* \in [\lambda \cdot r(\lambda)]$, s.t. $\mathcal{V}$'s answers to the j^* -th validation are good: $s_{j^*,1}, \dots, s_{j^*,\lambda} \in \{g^0, \dots, g^N\}$.
- By Eq. 7, this in turn means that $\mathcal{E}_{\text{ans}}$ answers to the same queries are good or $\perp$: $s'_{j^*,1}, \dots, s'_{j^*,\lambda} \in \{\perp, g^0, \dots, g^N\}$.
- By Eq. 10, we know that none of $s'_{j^*,k}$ is $\perp$, and therefore we can infer: $s'_{j^*,1}, \dots, s'_{j^*,\lambda} \in \{g^0, \dots, g^N\}$.
- By Eqs. 9 and 8, we can infer that not only $s'_{j^*,k}$'s themselves, but also the answers to $g^{\mathbf{s}_{j^*,k}}$ computed as combinations of e_i 's are all good. That is, for all $k \in [\lambda]$: $\prod_{i \in [N]} e_i^{\mathbf{s}_{j^*,k}[i]} \in \{g^0, \dots, g^N\}$.

In summary, this yields:

$$\Pr \left[\begin{array}{l} \nu_{\mathcal{D}, \nu}(c) \geq \frac{1}{r(\lambda)} \Rightarrow \\ \exists j^* \text{ s.t. } \forall k \in [\lambda] : \prod_{i \in [N]} e_i^{\mathbf{s}_{j^*,k}[i]} \in \{g^0, \dots, g^N\} \end{array} \right] \geq 1 - \text{negl}(\lambda) \quad (11)$$

Now we use Lemma 10, which says that for any set of fixed e_i 's, if at least one of them is not within $\{g^{-N}, \dots, g^N\}$, then w.h.p. at least one of the products $\prod_{i \in [N]} e_i^{\mathbf{s}_{j,k}^{[i]}}$ will be outside of $\{g^0, \dots, g^N\}$. More precisely:

$$\Pr \left[\begin{array}{l} \forall i \in [N] : e_i \in \{g^{-N}, \dots, g^N\} \vee \\ \exists k \in [\lambda] : \prod_{i \in [N]} e_i^{\mathbf{s}_{j,k}^{[i]}} \notin \{g^0, \dots, g^N\} \end{array} \right] \geq 1 - \text{negl}(\lambda) \quad \forall j \in [\lambda \cdot r(\lambda)] \quad (12)$$

Taking the union bound of Eqs. 11 and 12, we get that w.h.p., whenever $\nu_{\mathcal{D}, \nu}(c) \geq \frac{1}{r(\lambda)}$, then $e_i \in \{g^{-N}, \dots, g^N\}$ for all $i \in [N]$ (Eq. 5). This concludes the proof of Theorem 3.

Remark 6. By applying the statistical Lemma 10 to κ instead of λ , we would achieve the *relaxed* version of privacy with abort (see Sect. 3.3), because $\mathcal{E}_{\text{ans}}$ would return a non-aborting answer with probability $1 - \frac{1}{2^\kappa} - \text{negl}(\lambda)$ only.

4.5 Proof of Lemma 6 (Randomization Independence)

The *randomization independence* property of $\mathcal{E}_{\text{ans}}$ follows from the simple statistical observation that randomizing a query vector multiple times is identical to randomizing it a single time. Any new randomization “absorbs” the previous randomization, as captured by the following lemma.

Lemma 11 (Double Randomization). *For all $\lambda, N \in \mathbb{N}$, configurations $c = (\text{pp}, d, \text{st}_{\mathcal{D}}) \in \text{Supp}(\text{Conf}^{\mathcal{D}}(\lambda, N))$, queries $\mathbf{p} \in \mathbb{G}^N$ and its randomizations $(\text{st}^*, \mathbf{p}^*) \in \text{Supp}(\text{randomize}(\mathbf{p}))$, and unbounded adversaries $\mathcal{A}$ (with output in $\mathbb{G} \cup \{\perp\}$), the following two distributions are identical:*

$$\begin{aligned} & \{\text{recover}(\text{st}, \mathcal{A}(\tilde{\mathbf{p}})) \mid (\text{st}, \tilde{\mathbf{p}}) \leftarrow \text{randomize}(\mathbf{p})\} \\ & \equiv \{\text{recover}(\text{st}^*, \text{recover}(\text{st}, \mathcal{A}(\tilde{\mathbf{p}}^*))) \mid (\text{st}, \tilde{\mathbf{p}}^*) \leftarrow \text{randomize}(\mathbf{p}^*)\} \end{aligned}$$

Proof. Let $\text{st}^* = (r^*, \alpha^*)$ and $\mathbf{p}^* = \mathbf{h}^{r^*} \circ \mathbf{p}^{\alpha^*}$. By unrolling *randomize* and *recover*, the two distributions are equal to

$$\begin{aligned} & \left\{ (d^{-r} \cdot \mathcal{A}(\mathbf{h}^r \circ \mathbf{p}^\alpha))^{1/\alpha} \mid (r, \alpha) \leftarrow_{\$} \mathbb{Z}_q \times \mathbb{Z}_q^* \right\} \\ \text{and} \quad & \left\{ \left(d^{-r-r^*\alpha} \cdot \mathcal{A}(\mathbf{h}^{r+r^*\alpha} \circ \mathbf{p}^{\alpha^*\alpha}) \right)^{1/(\alpha^*\alpha)} \mid (r, \alpha) \leftarrow_{\$} \mathbb{Z}_q \times \mathbb{Z}_q^* \right\}. \end{aligned}$$

When substituting r by $r + r^*\alpha$, and α by $\alpha^*\alpha$ in the first distribution, the result is exactly the second one. This substitution does not have any impact on distribution, because for any fixed $r^* \in \mathbb{Z}_q$ and $\alpha^* \in \mathbb{Z}_q^*$, the following identity holds.

$$\{(r, \alpha) \mid (r, \alpha) \leftarrow_{\$} \mathbb{Z}_q \times \mathbb{Z}_q^*\} \equiv \left\{ \underbrace{(r + r^*\alpha)}_{\text{over } \mathbb{Z}_q}, \underbrace{(\alpha^*\alpha)}_{\text{over } \mathbb{Z}_q^*} \mid (r, \alpha) \leftarrow_{\$} \mathbb{Z}_q \times \mathbb{Z}_q^* \right\}$$

This concludes the proof. $\square$

Now we can prove *randomization independence* of $\mathcal{E}_{\text{ans}}$, meaning that for any query $\mathbf{p} \in \mathbb{G}^N$, and its randomization $(\text{st}^*, \mathbf{p}^*) \in \text{Supp}(\text{randomize}(\mathbf{p}))$, applying $\mathcal{E}_{\text{ans}}$ to $\mathbf{p}$ yields the same outcome as applying $\mathcal{E}_{\text{ans}}$ to $\mathbf{p}^*$, and then recovering:

$$\mathcal{E}_{\text{ans}}(1^r, c, \mathbf{p}) \equiv \text{recover}(\text{st}^*, \mathcal{E}_{\text{ans}}(1^r, c, \mathbf{p}^*))$$

To see that this holds, consider the k -th iteration (where $k \in [\lambda \cdot r^3]$) of $\mathcal{E}_{\text{ans}}$, and apply Lemma 11 with the following adversary.

- $\mathcal{A}(\tilde{\mathbf{p}})$ simulates lines 3–5 of $\mathcal{E}_{\text{ans}}$ by running $\tilde{p} \leftarrow \bar{\mathcal{V}}(c, \tilde{\mathbf{p}})$, $\text{st}', \tilde{\mathbf{p}}' \leftarrow \text{randomize}(\tilde{\mathbf{p}})$, and $\tilde{p}' \leftarrow \text{recover}(\text{st}', \bar{\mathcal{V}}(c, \tilde{\mathbf{p}}'))$. Then, if $\tilde{p} = \tilde{p}'$, output $\tilde{p}$. Otherwise, output $\perp$.

The k -th iteration of $\mathcal{E}_{\text{ans}}$ can be rewritten as

```

st,  $\tilde{\mathbf{p}} \leftarrow \text{randomize}(\mathbf{p})$ 
 $p \leftarrow \text{recover}(\text{st}, \mathcal{A}(\tilde{\mathbf{p}}))$ 
if  $p \neq \perp$ , return  $p$ 

```

By Lemma 11, the behavior stays the same when replacing $\mathbf{p}$ by $\mathbf{p}^*$ and returning $\text{recover}(\text{st}^*, p)$ instead of p (since $p \neq \perp$ is equivalent to $\text{recover}(\text{st}^*, p) \neq \perp$). The result is exactly $\text{recover}(\text{st}^*, \mathcal{E}_{\text{ans}}(1^r, c, \mathbf{p}^*))$, thereby concluding the proof of $\mathcal{E}_{\text{ans}}$'s *randomization independence*.

Remark 7. Note that in both *randomization independence* of $\mathcal{E}_{\text{ans}}$, and *double randomization* for arbitrary adversary $\mathcal{A}$, the first randomization $(\text{st}^*, \mathbf{p}^*) \in \text{Supp}(\text{randomize}(\mathbf{p}))$ does not need to be sampled by actually calling $\text{randomize}(\mathbf{p})$. There simply need to exist $r^* \in \mathbb{Z}_q$ and $\alpha^* \in \mathbb{Z}_q^*$ with $\text{st}^* = (r^*, \alpha^*)$ and $\mathbf{p}^* = \mathbf{h}^{r^*} \circ \mathbf{p}^{\alpha^*}$.

4.6 Proof of Lemma 7 (Agrees with Small Answers)

In this section, we prove Lemma 7, i.e., that if an adversary $\mathcal{A}$ is capable of answering a query $\mathbf{p}$ with $p^* \in \{g^0, \dots, g^N\}$, then the value p returned by $\mathcal{E}_{\text{ans}}$ on $\mathbf{p}$ is equal to p^* (if it is not $\perp$).

To do so, we open up the definition of $\mathcal{E}_{\text{ans}}$: let $\tilde{p}$ and $\tilde{p}'$ be the two answers to $\tilde{\mathbf{p}}$ as defined in Fig. 5. Furthermore, let $p := \text{recover}(\text{st}, \tilde{p})$ and $p' := \text{recover}(\text{st}, \tilde{p}')$ be the two recovered answers. Note that $\mathcal{E}_{\text{ans}}$ returns p whenever $\tilde{p} = \tilde{p}'$ (or equivalently, $p = p'$, due to the bijectivity of $\text{recover}(\text{st}, \cdot)$).

Therefore, by a union bound over all $\lambda \cdot r(\lambda)$ iterations of $\mathcal{E}_{\text{ans}}$, we just need to prove the following (which states that it is unlikely that $\mathcal{E}_{\text{ans}}$ terminates in the current iteration if the answer does not match p^*):

$$\Pr \left[\begin{array}{l} p^* \in \{g^0, \dots, g^N\} \\ \wedge p = p' \\ \wedge p \neq p^* \end{array} \middle| \begin{array}{l} c = (\text{pp}, d, \text{st}_{\mathcal{D}}) \leftarrow \text{Conf}^{\mathcal{D}}(\lambda, N) \\ \text{st}_{\mathcal{A}}, \mathbf{p} \leftarrow \mathcal{A}_1(c) \\ (p, p', p^*) \leftarrow \text{Answers}(\text{st}_{\mathcal{A}}, \mathbf{p}) \end{array} \right] \leq \text{negl}(\lambda). \quad (13)$$

To remove clutter in our calculations, the three answers p, p', p^* are generated by the following procedure $\text{Answers}(\lambda, N)$, which computes p and p' exactly as an iteration of $\mathcal{E}_{\text{ans}}$ would, and p^* as the answer computed by $\mathcal{A}_2$.

Procedure Answers($\text{st}_{\mathcal{A}}, \mathbf{p}$)

- 1 : $p^* \leftarrow \text{recover}(\text{st}^*, \mathcal{A}_2(\text{st}_{\mathcal{A}}, \tilde{\mathbf{p}}^*)), \text{ where } \text{st}^*, \tilde{\mathbf{p}}^* \leftarrow \text{randomize}(\mathbf{p})$
- 2 : $p \leftarrow \text{recover}(\text{st}, \bar{\mathcal{V}}(c, \tilde{\mathbf{p}})), \text{ where } \text{st}, \tilde{\mathbf{p}} \leftarrow \text{randomize}(\mathbf{p})$
- 3 : $p' \leftarrow \text{recover}(\text{st}, \text{recover}(\text{st}', \bar{\mathcal{V}}(c, \tilde{\mathbf{p}}'))), \text{ where } \text{st}', \tilde{\mathbf{p}}' \leftarrow \text{randomize}(\tilde{\mathbf{p}})$
- 4 : **return** (p, p', p^*)

To provide some intuition for the remainder of this proof, note that our generalized notion of integrity (*evasive ratio*, Lemma 4) shows that for any two *different* ways of answering a query, it is impossible to predict their ratio. However, if the probability above were not negligible, then there is (1) the option of answering with p^* , given by $\mathcal{A}_2$, and (2) the option of answering with p' , given by $\bar{\mathcal{V}}$. Furthermore, it would be *easy* to predict the ratio $\frac{p'}{p^*}$: we just need to query $\bar{\mathcal{V}}$ again to get another answer p . Then, with non-negligible probability, the ratio would be in the small set $\{\frac{p}{g^0}, \dots, \frac{p}{g^N}\}$.

To formalize this idea, note that the probability in Eq. 13 is upper bounded by

$$\Pr \left[\begin{array}{l} p' \neq p^* \\ \wedge \frac{p'}{p^*} \in \left\{ \frac{p}{g^0}, \dots, \frac{p}{g^N} \right\} \end{array} \middle| \begin{array}{l} c = (\text{pp}, d, \text{st}_{\mathcal{D}}) \leftarrow \text{Conf}^{\mathcal{D}}(\lambda, N) \\ \text{st}_{\mathcal{A}}, \mathbf{p} \leftarrow \mathcal{A}_1(c) \\ (p, p', p^*) \leftarrow \text{Answers}(\text{st}_{\mathcal{A}}, \mathbf{p}) \end{array} \right]. \quad (14)$$

We face one technical difficulty preventing us from applying Lemma 4: we would need an adversary that computes two answers $\tilde{p}'$ and $\tilde{p}^*$ from the *same* randomization $\text{st}^*, \tilde{\mathbf{p}}^* \leftarrow \text{randomize}(\mathbf{p})$ of query $\mathbf{p}$ (without knowing st^*). While $\tilde{p}^* = \mathcal{A}_2(\text{st}_{\mathcal{A}}, \tilde{\mathbf{p}}^*)$ is an answer to $\tilde{\mathbf{p}}^*$, the value $\tilde{p}' \leftarrow \text{recover}(\text{st}', \bar{\mathcal{V}}(c, \tilde{\mathbf{p}}'))$ is an answer to $\tilde{\mathbf{p}}$, a separate randomization of $\mathbf{p}$. We fix this issue using *double randomization* (Lemma 11) twice with adversary $\bar{\mathcal{V}}(c, \cdot)$: we may replace line 3 of Answers($\text{st}_{\mathcal{A}}, \mathbf{p}$) by

$$p' \leftarrow \text{recover}(\text{st}^*, \text{recover}(\text{st}', \bar{\mathcal{V}}(c, \tilde{\mathbf{p}}'))), \text{ where } \text{st}', \tilde{\mathbf{p}}' \leftarrow \text{randomize}(\tilde{\mathbf{p}}^*),$$

to obtain the procedure Answers'($\text{st}_{\mathcal{A}}, \mathbf{p}$) which behaves identical to Answers($\text{st}_{\mathcal{A}}, \mathbf{p}$). Now, we use Lemma 4 (*evasive ratio*) with the following adversary $\mathcal{B}_1, \mathcal{B}_2$.

- $\mathcal{B}_1(c)$: run $\text{st}_{\mathcal{A}}, \mathbf{p} \leftarrow \mathcal{A}_1(c)$, and output $\text{st}_{\mathcal{B}} := (\text{st}_{\mathcal{A}}, c)$ and $\mathbf{p}$.
- $\mathcal{B}_2(\text{st}_{\mathcal{B}}, \tilde{\mathbf{p}}^*)$: Compute the first answer as $\tilde{p}^* \leftarrow \mathcal{A}_2(\text{st}_{\mathcal{A}}, \tilde{\mathbf{p}}^*)$, and the second answer by running $\text{st}', \tilde{\mathbf{p}}' \leftarrow \text{randomize}(\tilde{\mathbf{p}}^*)$ and $\tilde{p}' \leftarrow \text{recover}(\text{st}', \bar{\mathcal{V}}(c, \tilde{\mathbf{p}}'))$. Choose the set $S := \left\{ \frac{p}{g^0}, \dots, \frac{p}{g^N} \right\}$ by running $p \leftarrow \text{recover}(\text{st}, \bar{\mathcal{V}}(c, \tilde{\mathbf{p}}))$ with $\text{st}, \tilde{\mathbf{p}} \leftarrow \text{randomize}(\tilde{\mathbf{p}})$. Output $\tilde{p}^*, \tilde{p}', S$.

Lemma 4, applied to adversary $\mathcal{B}_1, \mathcal{B}_2$ shows that

$$\Pr \left[\frac{\tilde{p}' \neq \tilde{p}^* \wedge \frac{\text{recover}(\text{st}^*, \tilde{p}')}{\text{recover}(\text{st}^*, \tilde{p}^*)} \in \left\{ \frac{p}{g^0}, \dots, \frac{p}{g^N} \right\}}{c = (\text{pp}, d, \text{st}_{\mathcal{D}}) \leftarrow \text{Conf}^{\mathcal{D}}(\lambda, N) \begin{array}{l} \text{st}_{\mathcal{A}}, \mathbf{p} \leftarrow \mathcal{A}_1(c) \\ \text{st}^*, \tilde{\mathbf{p}}^* \leftarrow \text{randomize}(\mathbf{p}) \\ \tilde{p}^* \leftarrow \mathcal{A}_2(\text{st}_{\mathcal{A}}, \tilde{\mathbf{p}}^*) \\ \text{st}', \tilde{\mathbf{p}}' \leftarrow \text{randomize}(\tilde{\mathbf{p}}^*) \\ \tilde{p}' \leftarrow \text{recover}(\text{st}', \bar{\mathcal{V}}(c, \tilde{\mathbf{p}}')) \\ \text{st}, \tilde{\mathbf{p}} \leftarrow \text{randomize}(\mathbf{p}) \\ p \leftarrow \text{recover}(\text{st}, \bar{\mathcal{V}}(c, \tilde{\mathbf{p}})) \end{array}} \right] \leq \text{negl}(\lambda). \quad (15)$$

Note that for $p' := \text{recover}(\text{st}^*, \tilde{p}')$ and $p^* := \text{recover}(\text{st}^*, \tilde{p}^*)$, the equivalence

$$\tilde{p}' \neq \tilde{p}^* \Leftrightarrow p' \neq p^*$$

holds due to the bijectivity of $\text{recover}(\text{st}^*, \cdot)$. Therefore, we can observe that the probability in Eq. 15 is identical to the one in Eq. 14 (after replacing $\text{Answers}(\text{st}_{\mathcal{A}}, \mathbf{p})$ by the equivalent $\text{Answers}'(\text{st}_{\mathcal{A}}, \mathbf{p})$). This concludes the proof of Lemma 7. $\square$

4.7 Proof of Lemma 8 (Homomorphism)

In this section, we prove the *homomorphism* property of $\mathcal{E}_{\text{ans}}$. Let $\mathcal{A}$ be a ppt adversary, and let $N(\lambda)$ and $r(\lambda)$ be polynomials. W.l.o.g., we may assume that $t_{i^*} \neq 0$ holds for each t_{i^*} chosen by $\mathcal{A}$ (otherwise, the adversary would only *increase* its success probability by *not* returning $\mathbf{p}_{i^*}$ and t_{i^*} in the first place, because neither $p = \mathcal{E}_{\text{ans}}(1^{r(\lambda)}, c, \prod_{i \in [m]} \mathbf{p}_i^{t_i})$ nor $\prod_{i \in [m]} p_i^{t_i}$ would be affected by doing so.)

Consider the following hybrid experiment $\text{Hyb}_{i^*}(\lambda, N)$ (for any $i^* \in \{0, \dots, m\}$) that corresponds to the statement of Lemma 8, except that the first i^* queries have been changed from $\mathbf{p}_i$ to $\mathbf{1}$ (where $\mathbf{1}$ is the vector consisting of N times the neutral group element 1).

$$\begin{array}{l} \text{Hyb}_{i^*}(\lambda, N) \\ \hline 1 : c = (\text{pp}, d, \text{st}_{\mathcal{D}}) \leftarrow \text{Conf}^{\mathcal{D}}(\lambda, N) \\ 2 : \mathbf{p}_1, \dots, \mathbf{p}_m, t_1, \dots, t_m \leftarrow \mathcal{A}(c) \\ 3 : p \leftarrow \mathcal{E}_{\text{ans}}(1^{r(\lambda)}, c, \prod_{i \in \{i^*+1, \dots, m\}} \mathbf{p}_i^{t_i}) \\ 4 : p_i \leftarrow \mathcal{E}_{\text{ans}}(1^{r(\lambda)}, c, \mathbf{1}) \quad \forall i \in \{1, \dots, i^*\} \\ 5 : p_i \leftarrow \mathcal{E}_{\text{ans}}(1^{r(\lambda)}, c, \mathbf{p}_i) \quad \forall i \in \{i^*+1, \dots, m\} \\ 6 : \text{return } \mathbb{I} \left[p, p_1, \dots, p_m \neq \perp \Rightarrow p = \prod_{i \in [m]} p_i^{t_i} \right] \end{array}$$

Note that the statement of Lemma 8 is identical to

$$\Pr[\text{Hyb}_0(\lambda, N) = 1] \geq 1 - \text{negl}(\lambda),$$

which is what we will show with a hybrid argument. That is, we need

$$|\Pr[\text{Hyb}_{i^*-1}(\lambda, N) = 1] - \Pr[\text{Hyb}_{i^*}(\lambda, N) = 1]| \leq \text{negl}(\lambda) \quad \forall i \in [w], \quad (16)$$

and

$$\Pr[\text{Hyb}_m(\lambda, N) = 1] \geq 1 - \text{negl}(\lambda). \quad (17)$$

We start with proving Eq. 16, by applying Lemma 5 (*chosen vector indistinguishability*) to the following ppt adversary $\mathcal{B}_1, \mathcal{B}_2$:

- $\mathcal{B}_1(c)$: Run $\mathbf{p}_1, \dots, \mathbf{p}_m, t_1, \dots, t_m \leftarrow \mathcal{A}(c)$. Output state $\text{st}_{\mathcal{B}} := (c, \{\mathbf{p}_i\}, \{t_i\})$ and vectors $\mathbf{v}_0 := \mathbf{p}_{i^*}^{t_{i^*}}, \mathbf{v}_1 := \mathbf{1}$.
- $\mathcal{B}_2(\text{st}_{\mathcal{B}}, \tilde{\mathbf{v}})$: Compute the two answers $\tilde{p} \leftarrow \mathcal{E}_{\text{ans}}(1^{r(\lambda)}, c, \tilde{\mathbf{v}} \cdot \prod_{i \in \{i^*+1, \dots, m\}} \mathbf{p}_i^{t_i})$ and $\tilde{p}_{i^*} \leftarrow \mathcal{E}_{\text{ans}}(1^{r(\lambda)}, c, \tilde{\mathbf{v}})$. Further, compute answers $p_i \leftarrow \mathcal{E}_{\text{ans}}(1^{r(\lambda)}, c, \mathbf{1})$ for each $i \in [i^* - 1]$ and $p_i \leftarrow \mathcal{E}_{\text{ans}}(1^{r(\lambda)}, c, \mathbf{p}_i)$ for each $i \in \{i^* + 1, \dots, m\}$. Then output

$$\mathbb{I} \left[\tilde{p}, p_1, \dots, p_{i^*-1}, \tilde{p}_{i^*}, p_{i^*+1}, \dots, p_m \neq \perp \Rightarrow \tilde{p} = \tilde{p}_{i^*} \cdot \prod_{i \in [m] \setminus \{i^*\}} p_i^{t_i} \right].$$

Note that the chosen-vector experiment $\text{CVA}^{\mathcal{D}, \mathcal{B}_1, \mathcal{B}_2}(\lambda, N, 0)$ for bit $b = 0$ (as defined in Lemma 5) is exactly the following:

Game $\text{CVA}^{\mathcal{D}, \mathcal{B}_1, \mathcal{B}_2}(\lambda, N, 0)$

- 1 : $c = (\text{pp}, d, \text{st}_{\mathcal{D}}) \leftarrow \text{Conf}^{\mathcal{D}}(\lambda, N)$
- 2 : $\mathbf{p}_1, \dots, \mathbf{p}_m, t_1, \dots, t_m \leftarrow \mathcal{A}(c)$
- 3 : $r \leftarrow \mathbb{S}_{\mathbb{Z}_q}$
- 4 : $\tilde{p} \leftarrow \mathcal{E}_{\text{ans}}(1^{r(\lambda)}, c, \mathbf{h}^r \circ \prod_{i \in \{i^*, \dots, m\}} \mathbf{p}_i^{t_i})$
- 5 : $p_i \leftarrow \mathcal{E}_{\text{ans}}(1^{r(\lambda)}, c, \mathbf{1}) \quad \forall i \in \{1, \dots, i^* - 1\}$
- 6 : $\tilde{p}_{i^*} \leftarrow \mathcal{E}_{\text{ans}}(1^{r(\lambda)}, c, \mathbf{h}^r \circ \mathbf{p}_{i^*}^{t_{i^*}})$
- 7 : $p_i \leftarrow \mathcal{E}_{\text{ans}}(1^{r(\lambda)}, c, \mathbf{p}_i) \quad \forall i \in \{i^* + 1, \dots, m\}$
- 8 : **return** $\mathbb{I} \left[\tilde{p}, p_1, \dots, p_{i^*-1}, \tilde{p}_{i^*}, p_{i^*+1}, \dots, p_m \neq \perp \Rightarrow \tilde{p} = \tilde{p}_{i^*} \cdot \prod_{i \in [m] \setminus \{i^*\}} p_i^{t_i} \right]$

By *randomization independence* (Lemma 6) (which we can apply due to $t_{i^*} \neq 0$), the following two equivalences hold:

$$\begin{aligned} \tilde{p} = \mathcal{E}_{\text{ans}}(1^{r(\lambda)}, c, \mathbf{h}^r \circ \prod_{i \in \{i^*, \dots, m\}} \mathbf{p}_i^{t_i}) &\equiv d^r \cdot \mathcal{E}_{\text{ans}}(1^{r(\lambda)}, c, \prod_{i \in \{i^*, \dots, m\}} \mathbf{p}_i^{t_i}), \\ \tilde{p}_{i^*} = \mathcal{E}_{\text{ans}}(1^{r(\lambda)}, c, \mathbf{h}^r \circ \mathbf{p}_{i^*}^{t_{i^*}}) &\equiv d^r \cdot (\mathcal{E}_{\text{ans}}(1^{r(\lambda)}, c, \mathbf{p}_{i^*}))^{t_{i^*}}. \end{aligned}$$

Thus, without affecting the distribution, in $\text{CVA}^{\mathcal{D}, \mathcal{B}_1, \mathcal{B}_2}(\lambda, N, 0)$ we can replace line 4 by $p \leftarrow \mathcal{E}_{\text{ans}}(1^{r(\lambda)}, c, \prod_{i \in \{i^*, \dots, m\}} \mathbf{p}_i^{t_i})$, line 6 by $p_{i^*} \leftarrow \mathcal{E}_{\text{ans}}(1^{r(\lambda)}, c, \mathbf{p}_{i^*})$, and the return value by

$$\mathbb{I} \left[d^r \cdot p, p_1, \dots, p_{i^*-1}, d^r \cdot p_{i^*}^{t_{i^*}}, p_{i^*+1}, \dots, p_m \neq \perp \Rightarrow d^r \cdot p = d^r \cdot \prod_{i \in [m]} p_i^{t_i} \right].$$

We note that $d^r \cdot p = \perp \Leftrightarrow p = \perp$ and $d^r \cdot p_{i^*}^{t_{i^*}} \Leftrightarrow p_{i^*} = \perp$, and that the factor d^r cancels out on both sides in the equality test. The result is identical to

$\text{Hyb}_{i^*-1}(\lambda, N)$, i.e., $\text{CVA}^{\mathcal{D}, \mathcal{B}_1, \mathcal{B}_2}(\lambda, N, 0) \equiv \text{Hyb}_{i^*-1}(\lambda, N)$. Analogously, we also get $\text{CVA}^{\mathcal{D}, \mathcal{B}_1, \mathcal{B}_2}(\lambda, N, 1) \equiv \text{Hyb}_{i^*}(\lambda, N)$. Therefore, Eq. 16 follows from Lemma 5.

It remains to prove Eq. 17. Note that in $\text{Hyb}_m(\lambda, N)$, $\mathcal{E}_{\text{ans}}$ is only ever called on query $\mathbf{1}$. If all those calls to $\mathcal{E}_{\text{ans}}$ return $\perp$ or 1, then $\text{Hyb}_m(\lambda, N)$ will output 1. Indeed, we do have such a property, captured by the following lemma. Applying it to all $m+1$ calls to $\mathcal{E}_{\text{ans}}(1^{r(\lambda)}, c, \mathbf{1})$ in $\text{Hyb}_m(\lambda, N)$ concludes the proof of *homomorphism*.

Lemma 12. *For any $\lambda \in \mathbb{N}$ and polynomials $N(\lambda) \leq \text{poly}(\lambda)$ and $r(\lambda) \leq \text{poly}(\lambda)$:*

$$\Pr \left[\mathcal{E}_{\text{ans}}(1^{r(\lambda)}, c, \mathbf{1}) \in \{1, \perp\} \mid c \leftarrow \text{Conf}^{\mathcal{D}}(\lambda, N) \right] \geq 1 - \text{negl}(\lambda)$$

Proof. We show that in each of the $\lceil \lambda \cdot (r(\lambda))^3 \rceil$ iterations of $\mathcal{E}_{\text{ans}}(1^{r(\lambda)}, c, \mathbf{1})$, the probability of returning anything $\neq 1$ is negligible. The full lemma then follows by a union bound.

By unrolling **randomize** and **recover**, we can rewrite a single iteration of $\mathcal{E}_{\text{ans}}$ as on the left-hand side in the following equation.

$$\begin{array}{ll} r, r' \leftarrow \$_{\mathbb{Z}_q} \text{ and } \alpha, \alpha' \leftarrow \$_{\mathbb{Z}_q^*} & r, r' \leftarrow \$_{\mathbb{Z}_q} \text{ and } \alpha, \alpha' \leftarrow \$_{\mathbb{Z}_q^*} \\ \tilde{p} \leftarrow \overline{\mathcal{V}}(c, \mathbf{h}^r) & \tilde{p} \leftarrow \overline{\mathcal{V}}(c, \mathbf{h}^r) \\ \tilde{p}' \leftarrow (d^{-r'} \cdot \overline{\mathcal{V}}(c, \mathbf{h}^{r'+r\alpha'}))^{1/\alpha'} & \tilde{p}' \leftarrow \overline{\mathcal{V}}(c, \mathbf{h}^{r'}) \\ \text{if } \tilde{p} = \tilde{p}', \text{return } (d^{-r} \cdot \tilde{p})^{1/\alpha} & \text{if } d^{-r} \cdot \tilde{p} = (d^{-r'} \cdot \tilde{p}')^{1/\alpha'}, \text{return } (d^{-r} \cdot \tilde{p})^{1/\alpha} \end{array} \quad \equiv$$

By replacing r' with $(r' - r\alpha') \bmod q$ (which are identically distributed for $r' \leftarrow \$_{\mathbb{Z}_q}$), we obtain the right-hand side. Note that α' is *only* used for checking the final **if**-condition. For any $x, y \in \mathbb{G}$ with $x \neq 1$, it is easy to see that (due to the group size q being prime)

$$\Pr \left[x = y^{1/\alpha'} \mid \alpha' \leftarrow \$_{\mathbb{Z}_q^*} \right] \leq \frac{1}{q-1} \leq \text{negl}(\lambda).$$

Thus, plugging in $x = d^{-r} \cdot \tilde{p}$ and $y = d^{-r'} \cdot \tilde{p}'$ shows that the probability of $\mathcal{E}_{\text{ans}}$ returning a value $\neq 1$ is negligible. $\square$

4.8 Proof of Lemma 10 (Random Subset Testing)

Note that, since all $\mathbf{s}_k$ are sampled independently of each other, it suffices to show the following:

$$\Pr_{\mathbf{s} \leftarrow \{0,1\}^N} \left[\prod_{i \in [N]} e_i^{\mathbf{s}[i]} \in \{g^0, \dots, g^N\} \right] \leq \frac{1}{2}.$$

Fix $\mathbf{s}[i] \in \{0, 1\}$ for all $i \in [N] \setminus \{i^*\}$. It suffices to show that

$$Pr_{\mathbf{s}[i^*] \leftarrow \mathbf{s}\{0,1\}} \left[e_{i^*}^{\mathbf{s}[i^*]} \cdot \prod_{i \in [N] \setminus \{i^*\}} e_i^{\mathbf{s}[i]} \in \{g^0, \dots, g^N\} \right] \leq \frac{1}{2},$$

i.e., at most one of

$$\prod_{i \in [N] \setminus \{i^*\}} e_i^{\mathbf{s}[i]} \quad \text{and} \quad e_{i^*} \cdot \prod_{i \in [N] \setminus \{i^*\}} e_i^{\mathbf{s}[i]}$$

is in $\{g^0, \dots, g^N\}$. Note that if this is the case for the left term, then the right term will be equal to $e_{i^*} \cdot g^m$ for some $m \in \{0, \dots, N\}$.

Now assume that, additionally, $e_{i^*} \cdot g^m \in \{g^0, \dots, g^N\}$. This implies that $e_{i^*} = g^{m'}$ for some $m' \in \{-N, \dots, N\}$, which contradicts this lemma's assumptions. $\square$

Acknowledgments. We thank Chenzhi Zhu for his involvement in the initial stage of this project. This research was partially supported by NSF grants CNS-2026774, CNS-2154174, a JP Morgan Faculty Award, a CISCO Faculty Award, and a gift from Microsoft.

References

1. Aguilar-Melchor, C., Barrier, J., Fousse, L., Killijian, M.O.: XPIR: private information retrieval for everyone. *PoPETs* **2016**(2), 155–174 (2016). <https://doi.org/10.1515/popets-2016-0010>
2. Ali, A., Lepoint, T., Patel, S., Raykova, M., Schoppmann, P., Seth, K., Yeo, K.: Communication-computation trade-offs in PIR. In: Bailey, M., Greenstadt, R. (eds.) *USENIX Security 2021*, pp. 1811–1828. USENIX Association, August 2021. <https://www.usenix.org/conference/usenixsecurity21/presentation/ali>
3. Angel, S., Chen, H., Laine, K., Setty, S.T.V.: PIR with compressed queries and amortized query processing. In: *2018 IEEE Symposium on Security and Privacy*, pp. 962–979. IEEE Computer Society Press, May 2018. <https://doi.org/10.1109/SP.2018.00062>
4. Angel, S., Setty, S.T.V.: Unobservable communication over fully untrusted infrastructure. In: Keeton, K., Roscoe, T. (eds.) *12th USENIX Symposium on Operating Systems Design and Implementation, OSDI 2016*, Savannah, GA, USA, 2–4 November 2016, pp. 551–569. USENIX Association (2016). <https://www.usenix.org/conference/osdi16/technical-sessions/presentation/angel>
5. Barak, B., Goldreich, O.: Universal arguments and their applications. *SIAM J. Comput.* **38**(5), 1661–1694 (2008). <https://doi.org/10.1137/070709244>
6. Beimel, A., Stahl, Y.: Robust information-theoretic private information retrieval. *J. Cryptol.* **20**(3), 295–321 (2007). <https://doi.org/10.1007/s00145-007-0424-2>
7. Ben-David, S., Kalai, Y.T., Paneth, O.: Verifiable private information retrieval. In: Kiltz, E., Vaikuntanathan, V. (eds.) *TCC 2022, Part III*. LNCS, vol. 13749, pp. 3–32. Springer, Heidelberg (2022). https://doi.org/10.1007/978-3-031-22368-6_1

8. Bleichenbacher, D.: Chosen ciphertext attacks against protocols based on the RSA encryption standard PKCS #1. In: Krawczyk, H. (ed.) CRYPTO 1998. LNCS, vol. 1462, pp. 1–12. Springer, Heidelberg (1998). <https://doi.org/10.1007/BFb0055716>
9. Boyle, E., Ishai, Y., Pass, R., Wootters, M.: Can we access a database both locally and privately? In: Kalai, Y., Reyzin, L. (eds.) TCC 2017, Part II. LNCS, vol. 10678, pp. 662–693. Springer, Cham (2017). https://doi.org/10.1007/978-3-319-70503-3_22
10. Bünz, B., Bootle, J., Boneh, D., Poelstra, A., Wuille, P., Maxwell, G.: Bulletproofs: short proofs for confidential transactions and more. In: 2018 IEEE Symposium on Security and Privacy, pp. 315–334. IEEE Computer Society Press, May 2018. <https://doi.org/10.1109/SP.2018.00020>
11. Cachin, C., Micali, S., Stadler, M.: Computationally private information retrieval with polylogarithmic communication. In: Stern, J. (ed.) EUROCRYPT 1999. LNCS, vol. 1592, pp. 402–414. Springer, Heidelberg (1999). https://doi.org/10.1007/3-540-48910-X_28
12. Camenisch, J., Neven, G., Shelat, A.: Simulatable adaptive oblivious transfer. In: Naor, M. (ed.) EUROCRYPT 2007. LNCS, vol. 4515, pp. 573–590. Springer, Heidelberg (2007). https://doi.org/10.1007/978-3-540-72540-4_33
13. Canetti, R., Holmgren, J., Richelson, S.: Towards doubly efficient private information retrieval. In: Kalai, Y., Reyzin, L. (eds.) TCC 2017, Part II. LNCS, vol. 10678, pp. 694–726. Springer, Cham (2017). https://doi.org/10.1007/978-3-319-70503-3_23
14. de Castro, L., Lee, K.: VeriSimplePIR: verifiability in SimplePIR at no online cost for honest servers. In: USENIX Security 2024. USENIX Association, August 2024. <https://www.usenix.org/conference/usenixsecurity24/presentation/de-castro>
15. Cheng, R., et al.: Talek: private group messaging with hidden access patterns. In: ACSAC 2020: Annual Computer Security Applications Conference, Virtual Event/Austin, TX, USA, 7–11 December 2020, pp. 84–99. ACM (2020). <https://doi.org/10.1145/3427228.3427231>
16. Chor, B., Goldreich, O., Kushilevitz, E., Sudan, M.: Private information retrieval. In: 36th FOCS, pp. 41–50. IEEE Computer Society Press, October 1995. <https://doi.org/10.1109/SFCS.1995.492461>
17. Colombo, S., Nikitin, K., Corrigan-Gibbs, H., Wu, D.J., Ford, B.: Authenticated private information retrieval. In: Calandrino, J.A., Troncoso, C. (eds.) USENIX Security 2023, pp. 3835–3851. USENIX Association, August 2023. <https://www.usenix.org/conference/usenixsecurity23/presentation/colombo>
18. Davidson, A., Pestana, G., Celi, S.: FrodoPIR: simple, scalable, single-server private information retrieval. PoPETs **2023**(1), 365–383 (2023). <https://doi.org/10.56553/popets-2023-0022>
19. Devet, C., Goldberg, I., Heninger, N.: Optimally robust private information retrieval. In: Kohno, T. (ed.) USENIX Security 2012, pp. 269–283. USENIX Association, August 2012. <https://www.usenix.org/conference/usenixsecurity12/technical-sessions/presentation/devet>
20. Dietz, M., Tessaro, S.: Fully malicious authenticated PIR. Cryptology ePrint Archive, Report 2023/1804 (2023). <https://eprint.iacr.org/2023/1804>
21. Dong, C., Chen, L.: A fast single server private information retrieval protocol with low communication cost. In: Kutyłowski, M., Vaidya, J. (eds.) ESORICS 2014, Part I. LNCS, vol. 8712, pp. 380–399. Springer, Cham (2014). https://doi.org/10.1007/978-3-319-11203-9_22

22. Doröz, Y., Sunar, B., Hammouri, G.: Bandwidth efficient PIR from NTRU. In: Böhme, R., Brenner, M., Moore, T., Smith, M. (eds.) FC 2014. LNCS, vol. 8438, pp. 195–207. Springer, Heidelberg (2014). https://doi.org/10.1007/978-3-662-44774-1_16
23. Eriguchi, R., Kurosawa, K., Nuida, K.: Multi-server PIR with full error detection and limited error correction. In: Dachman-Soled, D. (ed.) 3rd Conference on Information-Theoretic Cryptography (ITC 2022). Leibniz International Proceedings in Informatics (LIPIcs), vol. 230, pp. 1:1–1:20. Schloss Dagstuhl – Leibniz-Zentrum für Informatik, Dagstuhl, Germany (2022). <https://doi.org/10.4230/LIPIcs.ITC.2022.1>
24. Gentry, C., Halevi, S.: Compressible FHE with applications to PIR. In: Hofheinz, D., Rosen, A. (eds.) TCC 2019, Part II. LNCS, vol. 11892, pp. 438–464. Springer, Cham (2019). https://doi.org/10.1007/978-3-030-36033-7_17
25. Gentry, C., Wichs, D.: Separating succinct non-interactive arguments from all falsifiable assumptions. In: Fortnow, L., Vadhan, S.P. (eds.) 43rd ACM STOC, pp. 99–108. ACM Press, June 2011. <https://doi.org/10.1145/1993636.1993651>
26. Goldberg, I.: Improving the robustness of private information retrieval. In: 2007 IEEE Symposium on Security and Privacy, pp. 131–148. IEEE Computer Society Press, May 2007. <https://doi.org/10.1109/SP.2007.23>
27. Green, M., Ladd, W., Miers, I.: A protocol for privately reporting ad impressions at scale. In: Weippl, E.R., Katzenbeisser, S., Kruegel, C., Myers, A.C., Halevi, S. (eds.) ACM CCS 2016, pp. 1591–1601. ACM Press, October 2016. <https://doi.org/10.1145/2976749.2978407>
28. Gupta, T., Crooks, N., Mulhern, W., Setty, S.T.V., Alvisi, L., Walfish, M.: Scalable and private media consumption with popcorn. In: Argyraki, K.J., Isaacs, R. (eds.) 13th USENIX Symposium on Networked Systems Design and Implementation, NSDI 2016, Santa Clara, CA, USA, 16–18 March 2016, pp. 91–107. USENIX Association (2016). <https://www.usenix.org/conference/nsdi16/technical-sessions/presentation/gupta-trinabh>
29. Henzinger, A., Dauterman, E., Corrigan-Gibbs, H., Zeldovich, N.: Private web search with Tiptoe. In: Flinn, J., Seltzer, M.I., Druschel, P., Kaufmann, A., Mace, J. (eds.) Proceedings of the 29th Symposium on Operating Systems Principles, SOSP 2023, Koblenz, Germany, 23–26 October 2023, pp. 396–416. ACM (2023). <https://doi.org/10.1145/3600006.3613134>
30. Henzinger, A., Hong, M.M., Corrigan-Gibbs, H., Meiklejohn, S., Vaikuntanathan, V.: One server for the price of two: simple and fast single-server private information retrieval. In: Calandrino, J.A., Troncoso, C. (eds.) 32nd USENIX Security Symposium, USENIX Security 2023, Anaheim, CA, USA, 9–11 August 2023. USENIX Association (2023). <https://www.usenix.org/conference/usenixsecurity23/presentation/henzinger>
31. Huang, Y., Katz, J., Evans, D.: Efficient secure two-party computation using symmetric cut-and-choose. In: Canetti, R., Garay, J.A. (eds.) CRYPTO 2013. LNCS, vol. 8043, pp. 18–35. Springer, Heidelberg (2013). https://doi.org/10.1007/978-3-642-40084-1_2
32. Kilian, J.: On the complexity of bounded-interaction and noninteractive zero-knowledge proofs. In: 35th FOCS, pp. 466–477. IEEE Computer Society Press, November 1994. <https://doi.org/10.1109/SFCS.1994.365744>
33. Kiraz, M., Schoenmakers, B.: A protocol issue for the malicious case of Yao’s garbled circuit construction. In: 27th Symposium on Information Theory in the Benelux, vol. 29, pp. 283–290 (2006)

34. Kurosawa, K.: How to correct errors in multi-server PIR. In: Galbraith, S.D., Moriai, S. (eds.) ASIACRYPT 2019. LNCS, vol. 11922, pp. 564–574. Springer, Cham (2019). https://doi.org/10.1007/978-3-030-34621-8_20
35. Kushilevitz, E., Ostrovsky, R.: Replication is NOT needed: SINGLE database, computationally-private information retrieval. In: 38th FOCS, pp. 364–373. IEEE Computer Society Press, October 1997. <https://doi.org/10.1109/SFCS.1997.646125>
36. Lin, W.K., Mook, E., Wicks, D.: Doubly efficient private information retrieval and fully homomorphic RAM computation from ring LWE. In: 55th ACM STOC, pp. 595–608. ACM Press, June 2023. <https://doi.org/10.1145/3564246.3585175>
37. Lindell, Y., Pinkas, B.: Secure two-party computation via cut-and-choose oblivious transfer. In: Ishai, Y. (ed.) TCC 2011. LNCS, vol. 6597, pp. 329–346. Springer, Heidelberg (2011). https://doi.org/10.1007/978-3-642-19571-6_20
38. Lipmaa, H.: An oblivious transfer protocol with log-squared communication. In: Zhou, J., Lopez, J., Deng, R.H., Bao, F. (eds.) ISC 2005. LNCS, vol. 3650, pp. 314–328. Springer, Heidelberg (2005). https://doi.org/10.1007/11556992_23
39. Ryan, M.D.: Enhanced certificate transparency and end-to-end encrypted mail. In: NDSS 2014. The Internet Society, February 2014
40. Vaudenay, S.: Security flaws induced by CBC padding — applications to SSL, IPSEC, WTLS... In: Knudsen, L.R. (ed.) EUROCRYPT 2002. LNCS, vol. 2332, pp. 534–545. Springer, Heidelberg (2002). https://doi.org/10.1007/3-540-46035-7_35
41. Wang, X., Zhao, L.: Verifiable single-server private information retrieval. In: Naccache, D., et al. (eds.) ICICS 2018. LNCS, vol. 11149, pp. 478–493. Springer, Cham (2018). https://doi.org/10.1007/978-3-030-01950-1_28
42. Yi, X., Kaosar, M.G., Paulet, R., Bertino, E.: Single-database private information retrieval from fully homomorphic encryption. *IEEE Trans. Knowl. Data Eng.* **25**(5), 1125–1134 (2013). <https://doi.org/10.1109/TKDE.2012.90>
43. Zhang, L.F., Safavi-Naini, R.: Verifiable multi-server private information retrieval. In: Boureanu, I., Owesarski, P., Vaudenay, S. (eds.) ACNS 2014. LNCS, vol. 8479, pp. 62–79. Springer, Cham (2014). https://doi.org/10.1007/978-3-319-07536-5_5
44. Zhao, L., Wang, X., Huang, X.: Verifiable single-server private information retrieval from LWE with binary errors. *Inf. Sci.* **546**, 897–923 (2021). <https://doi.org/10.1016/j.ins.2020.08.071>