

A Theory of Composition for Differential Obliviousness

Mingxun Zhou¹ Elaine Shi¹ T-H. Hubert Chan² Shir Maimon^{3*}

¹Carnegie Mellon University

²The University of Hong Kong

³Cornell University

Abstract

Differential obliviousness (DO) is a privacy notion which guarantees that the access patterns of a program satisfies differential privacy. Differential obliviousness was studied in a sequence of recent works as a relaxation of full obliviousness. Earlier works showed that DO not only allows us to circumvent the logarithmic-overhead barrier of fully oblivious algorithms, in many cases, it also allows us to achieve polynomial speedup over full obliviousness, since it avoids “padding to the worst-case” behavior of fully oblivious algorithms.

Despite the promises of differential obliviousness (DO), a significant barrier that hinders its broad application is the lack of composability. In particular, when we apply one DO algorithm to the output of another DO algorithm, the composed algorithm may no longer be DO (with reasonable parameters). Specifically, the outputs of the first DO algorithm on two neighboring inputs may no longer be neighboring, and thus we cannot directly benefit from the DO guarantee of the second algorithm.

In this work, we are the first to explore a theory of composition for differentially oblivious algorithms. We propose a refinement of the DO notion called (ϵ, δ) -neighbor-preserving-DO, or (ϵ, δ) -NPDO for short, and we prove that our new notion indeed provides nice compositional guarantees. In this way, the algorithm designer can easily track the privacy loss when composing multiple DO algorithms.

We give several example applications to showcase the power and expressiveness of our new NPDO notion. One of these examples is a result of independent interest: we use the compositional framework to prove an optimal privacy amplification theorem for the differentially oblivious shuffle model. In other words, we show that for a class of distributed differentially private mechanisms in the shuffle-model, one can replace the perfectly secure shuffler with a DO shuffler, and nonetheless enjoy almost the same privacy amplification enabled by a shuffler.

*Randomized order. This paper subsumes part of the results in an unpublished manuscript [ZS22] written by a subset of the authors. The work is partially done when Mingxun Zhou was a research intern at NTT Research.

1 Introduction

It is well-known that access patterns to even encrypted or secret-shared data can leak sensitive information [IKK12, CGPR15, NKW15, GLMP19, KKNO16, GMN⁺16]. Initiated by Goldreich and Ostrovsky [GO96, Gol87], *oblivious* algorithms is a line of work that aims to provably obfuscate a program’s access patterns without incurring too much slowdown. In particular, *obliviousness* (also referred to as *full obliviousness*) requires that a program’s access patterns be indistinguishable for any two inputs. It is well-known that oblivious algorithms have broad applications, including in multi-party computation [GKK⁺12, LWN⁺15], secure processors [RYF⁺13, MLS⁺13, LHH⁺15, MPC⁺18, TGS22], secure outsourcing [WST12, SS13], databases [EZ19, AK14, CBC⁺18], blockchains [CZJ⁺17], and so on. In the past decade, our community have significantly improved the efficiency of oblivious algorithms [SCSL11, SvS⁺13, WCS15], leading to large-scale real-world adoption such as Signal’s private contact discovery [Con]. However, as we discuss below, in some applications, the overhead of full obliviousness may still be unacceptable.

Differential Obliviousness (DO), defined by Chan, Chung, Maggs, and Shi [CCMS19], is a relaxed notion of access pattern privacy. DO requires that the program’s access patterns satisfy only differential privacy (DP) [DMNS06], as opposed to a simulation-based notion like in full obliviousness [GO96, Gol87, SCSL11]. Recent works [CCMS19, BKK⁺21, CZSC21, BNZ19, GKLX22] explored DO and showed how DO can allow us to circumvent fundamental performance barriers pertaining to full obliviousness:

- Chan et al. [CCMS19] showed a fundamental separation in terms of efficiency between DO and full obliviousness. Specifically, for a class of common tasks such as compaction, merging, and range query data structures, while full obliviousness is inherently subject to at least $\Omega(\log N)$ multiplicative overhead [LSX19, JLN19, AFKL19, FHLS19] (in comparison with the insecure baseline), using DO allows us to reduce the overhead to only $O(\log \log N)$ where N denotes the data size.
- Not only does DO allow us to overcome the logarithmic barrier for fully oblivious algorithms, another important aspect that is sometime overlooked is that DO allows us to overcome the “worst-case barrier” of fully oblivious algorithms [CZSC21], which leads to *polynomial* speedup over full obliviousness in many applications. Specifically, to achieve full obliviousness, we must pad the running time and output length to the worst case over all possible inputs (of some fixed length), whereas DO algorithms may reveal the *noisy* running time or output length. In many real-world scenarios such as database joins [CZSC21], the common case enjoys much shorter runtime and output length than the worst case. For exactly this reason, there is an entire line of work that focuses on designing algorithms optimized for the common rather than the worst case [Rou20]. In such cases, prior works showed that DO can achieve polynomial speedup over any fully oblivious algorithm [CCMS19, CZSC21]!

Basic DO does NOT lend to composition. Given the promises of DO, we would like to apply DO to more applications. Unfortunately, the status quo of DO hinders its broad applicability due to the lack of *compositional* guarantees. When designing algorithms, it is customary to compose several algorithmic building blocks together. In such cases, it would be nice to say that the composed algorithm also satisfies DO with reasonable parameters as long as the underlying algorithmic building blocks also satisfy DO. Similarly, in some applications, we may need to apply a DO algorithm to the outcome of another (e.g., the SQL database application below). In such cases, we also want to be able to track the privacy loss over time. While the original full obliviousness notion indeed allows such composition, unfortunately, the standard DO notion [CCMS19] does not!

As an explicit example of composition, imagine that we want to build a differentially oblivious database supporting SQL queries. Consider the following natural SQL query where we want to select entries from a table which in itself is the result of a previous **Select** operation¹:

```
Select (id, position) from
  (Select (id, dept, position) from Employees where salary > 200K)
where dept = "CS"
```

To support this query in a differentially oblivious manner, the most natural idea is to use the DO stable compaction algorithm of Chan et al. [CCMS19] to realize each **Select** operator. In stable compaction, we obtain an input array where each element is either a *real* element or a *filler*, and we want to output an array containing all the *real* elements of the input and preserving the order they appear in the input. Unfortunately, this approach completely fails since Chan et al. [CCMS19]’s DO compaction algorithm does NOT compose.

To understand why, we will introduce some basic notation. Let $M : \mathcal{X} \rightarrow \mathcal{Y}$ denote an algorithm, which takes in an input $x \in \mathcal{X}$, and produces an output $y \in \mathcal{Y}$. Consider some neighboring notion $\sim_{\mathcal{X}}$ defined over the input domain $\mathcal{X}$. For example, let $x, x' \in \mathcal{X}$ be two input arrays/tables where each entry corresponds to an individual user. One example is Hamming-distance neighboring: we say that $x \sim_{\mathcal{X}} x'$ iff the Hamming distance of x and x' is at most 1 — this is also the neighboring notion adopted by the DO compaction algorithm of Chan et al. [CCMS19]. The standard DO notion requires the following.

Definition 1.1 (Basic differential obliviousness [CCMS19]). We say that an algorithm M satisfies (ϵ, δ) -DO w.r.t. some symmetric relation $\sim_{\mathcal{X}}$ iff for any $x, x' \in \mathcal{X}$ such that $x \sim_{\mathcal{X}} x'$, for any subset S ,

$$\Pr[\text{View}^M(x) \in S] \leq e^{\epsilon} \cdot \Pr[\text{View}^M(x') \in S] + \delta, \quad (1)$$

where $\text{View}^M(x)$ is a random variable denoting the the memory access patterns observed when running the algorithm M over the input x .

Now, imagine that we have two DO mechanisms $M_1 : \mathcal{X}_1 \rightarrow \mathcal{X}_2$ and $M_2 : \mathcal{X}_2 \rightarrow \mathcal{Y}$ (e.g., think of M_1 and M_2 as Chan et al.’s DO compaction algorithm). We want to apply M_2 to the output of M_1 , and hope that the composed mechanism $M_2 \circ M_1(\cdot)$ satisfies DO. By the DO definition, we know that M_2 offers indistinguishability for two *neighboring* inputs from $\mathcal{X}_2$. Now, consider two neighboring inputs $x \sim_{\mathcal{X}_1} x'$ from $\mathcal{X}_1$, and consider running the mechanism M_1 over x and x' , respectively. Unfortunately, the basic DO notion (of M_1) does *not* guarantee that the outputs $M_1(x)$ and $M_1(x')$ are also neighboring. Therefore, we may not be able to benefit from the DO property of M_2 !

We stress that this is not just a deficiency of the basic DO definition. Natural designs of DO algorithms often do not guarantee that the outputs obtained from two neighboring inputs must be neighboring too. For example, consider the stable compaction algorithm of Chan et al. [CCMS19]. Given two input arrays $x = (1, 2, \perp, 3, 4)$ and $x' = (\perp, 2, \perp, 3, 4)$ with Hamming distance 1 where $\perp$ denotes a filler, the compacted outputs will be $(1, 2, 3, 4)$ and $(2, 3, 4)$, respectively. The outputs have Hamming distance more than 1. While the outputs have large Hamming distance, the edit distance is only one — unfortunately, Chan et al.’s compaction algorithm provides privacy only for Hamming-distance neighboring and the guarantees do not generalize to edit-distance neighboring.

¹Here we write the two **Select** statements in a single query for convenience, In practice, it could be that the first **Select** query is interactively issued and its result stored as a temporary table, and then the second **Select** query is interactively issued.

DP composition theorems do not work for DO. Since DO is essentially DP applied to the memory access patterns, a natural question is: *can we simply use DP composition theorems to reason about the composition DO mechanisms?* The answer is *no* because DP composition and composition of DO mechanisms are of different nature. In DP composition, we have multiple mechanisms $M_1, \dots, M_k$ where M_i satisfies (ϵ_i, δ_i) -DP. The basic DP composition theorem says that the composed mechanism $M(x) := (M_1(x), \dots, M_k(x))$ satisfies $(\sum_{i=1}^k \epsilon_i, \sum_{i=1}^k \delta_i)$ -DP. Here, all these mechanisms are applied to the *same input* x . In DO composition, we want to apply M_2 to the *output* of M_1 instead. More generally, if there are k DO mechanisms $M_1, \dots, M_k$, we want to know whether the composed mechanism $M_k \circ M_{k-1} \circ \dots \circ M_1(x) = M_k(M_{k-1}(\dots M_1(x)))$ is also DO.

Given the status quo, we ask the following natural question:

Can we have suitable and useful refinements of differential obliviousness (DO) that lend to composition?

1.1 Main Contribution: A Theory of Composition for Differential Obliviousness

We are the first to initiate a formal exploration of the composability of differential obliviousness. In this sense, we make an important conceptual contribution: by laying the groundwork for the composition of DO algorithms. we hope that our work can allow DO to have wider applicability.

A new, composable DO notion. Our first contribution is to introduce a new, composable DO notion called *Neighbor-Preserving Differential Obliviousness (NPDO)* that can be viewed as a strengthening of the basic DO by Chan et al. [CCMS19]. Our NPDO notion is composition friendly in the following senses:

- C1. If M_1 satisfies (ϵ_1, δ_1) -NPDO, and M_2 satisfies (ϵ_2, δ_2) -DO (the basic version), then the composed mechanism $M_2 \circ M_1$ satisfies $(\epsilon_1 + \epsilon_2, \delta_1 + \delta_2)$ -DO.
- C2. If M_1 satisfies (ϵ_1, δ_1) -NPDO, and M_2 satisfies (ϵ_2, δ_2) -NPDO, then the composed mechanism $M_2 \circ M_1$ satisfies $(\epsilon_1 + \epsilon_2, \delta_1 + \delta_2)$ -NPDO.

In the above, the first property allows us to apply any basic-DO algorithm M_2 to the output of an NPDO algorithm M_1 , and the composed algorithm $M_2 \circ M_1$ would satisfy basic DO. The second property allows us to perform composition repeatedly. In particular, if both M_1 and M_2 are NPDO, then the composed algorithm $M_2 \circ M_1$ also satisfies NPDO, i.e., it can be further composed with other DO or NPDO algorithms.

Finding the right notion turned out to be non-trivial. We want to capture the intuition that “the algorithm should produce neighboring outputs for neighboring inputs”. However, it is not obvious how to formally capture this idea of “neighbor-preserving” especially when the outputs of the DO algorithm may be randomized. Indeed, naïve ways to define “neighbor-preserving” turned out to be too stringent and preclude many natural and interesting algorithms (see Section 3.1). We instead suggest a more general version that allows us to capture a probabilistic notion of neighbor-preserving. More specifically, our NPDO notion requires that when one applies the algorithm M on two neighboring inputs x and x' , the *joint distribution* of the adversary’s view and the output must be distributionally close in some technical sense, where *closeness is parametrized by some output neighboring relation*. The formal definition is presented below:

Definition 1.2 ((ϵ, δ) -NPDO). We say that an algorithm $M : \mathcal{X} \rightarrow \mathcal{Y}$ with view space $\mathcal{V}$ satisfies (ϵ, δ) -NPDO w.r.t. input relation $\sim_{\mathcal{X}}$ and output relation $\sim_{\mathcal{Y}}$, if for any $x, x' \in \mathcal{X}$ such that

$x \sim_{\mathcal{X}} x'$, for any subset $S \subseteq \mathcal{V} \times \mathcal{Y}$,

$$\Pr[\text{Exec}^M(x) \in S] \leq e^\epsilon \cdot \Pr[\text{Exec}^M(x') \in \mathcal{N}(S)] + \delta.$$

In the above, $\text{Exec}^M(x)$ samples a random execution of M on the input x , and returns the view (i.e., access patterns) as well as the algorithm's output. Further, the notation $\mathcal{N}(S)$, i.e., the neighboring set of S , is defined as follows:

$$\mathcal{N}(S) = \{(v, y) | \exists (v', y') \in S \text{ s.t. } y \sim_{\mathcal{Y}} y'\}$$

Expressiveness of our notion. We give various natural examples to demonstrate the expressiveness and power of our notion. We believe that our NPDO notion is indeed the right notion, given the simplicity in form and its broad applicability. Besides the motivating SQL database example mentioned earlier in this section, other notable examples include the design of a differentially oblivious subsampling algorithm, a stable compaction algorithm that is DO w.r.t. edit distance, and finally, proving an optimal privacy amplification theorem in the differentially oblivious shuffle model. Since the last application is of independent interest even as a standalone result, we will discuss the context and the implications of this result separately in Section 1.2.

Proof of composition theorem. Our second contribution is to prove the composition theorem:

Theorem 1.3 (Composition theorem). *The aforementioned compositional properties C1 and C2 hold, as long as the algorithm M_1 's view space and output space are finite or countably infinite.*

The proof of the composition theorem is rather non-trivial. A key step in the proof is to show the following equivalence (see Lemma 4.1). An algorithm $M : \mathcal{X} \rightarrow \mathcal{Y}$ (with at most countably infinite view space $\mathcal{V}$ and output space $\mathcal{Y}$) satisfies (ϵ, δ) -NPDO w.r.t. $\sim_{\mathcal{X}}$ and $\sim_{\mathcal{Y}}$, *if and only if* for any neighboring inputs $x \sim_{\mathcal{X}} x'$, there exists an (ϵ, δ) -matching between the probability spaces of the random variables $\text{Exec}^M(x) \in \mathcal{V} \times \mathcal{Y}$ and $\text{Exec}^M(x') \in \mathcal{V} \times \mathcal{Y}$. In an (ϵ, δ) -matching, imagine that we have a (possibly countably infinitely large) bipartite graph where one side has the sources, and the other side has the destinations. Both sources and destinations come from the space $\mathcal{V} \times \mathcal{Y}$. If there is an edge of weight w between some source and some destination, we may imagine that the source wants to send w amount of commodity to the destination. Now, each source $(v, y) \in \mathcal{V} \times \mathcal{Y}$ produces an amount of commodity equal to $\Pr[\text{Exec}^M(x) = (v, y)]$, and each destination (v', y') can receive at most $e^\epsilon \cdot \Pr[\text{Exec}^M(x') = (v', y')]$ amount of commodity. Furthermore, a source (v, y) can be matched with a destination (v', y') only if they are neighboring, i.e., $v = v'$ and $y \sim_{\mathcal{Y}} y'$. We want to find a matching such that all but δ amount of commodity is delivered to the destinations. To prove this key equivalence lemma, we are inspired by techniques used to prove the Hall's marriage theorem [Hal35, HJ48]. Once we prove the key equivalence lemma, we then rely on it to prove the composition theorem.

In the main body, we primarily focus on proving the composition theorem for *statistical* notions of DO. In Appendix A, we further extend our composition theorem to support suitable, *computational* notions of differential obliviousness as well.

Finally, in our composition theorem, we assume that the view and output spaces of M_1 are at most countably infinitely large. This assumption is reasonable given that we primarily focus on the standard word-RAM model of execution. It is indeed an interesting open question whether we can remove this restriction and prove the composition theorem for uncountably large view and output spaces — this is useful if we consider RAM machines that can handle real arithmetic. In Appendix C, we discuss the additional technicalities that one might encounter if we wish to remove the countable restriction.

1.2 Additional Result: Optimal Privacy Amplification in the DO-Shuffle Model

As an application of our composition framework, we use it to prove an optimal privacy amplification theorem in the differentially oblivious shuffle (DO-shuffle) model. Since this result can be of independent interest on its own, we explain the motivation and context below.

Background: privacy amplification in the shuffle model. To understand the DO-shuffle model, let us first review some background on the so-called shuffle model. Imagine that a set of clients each hold some private data, and an *untrusted* server wants to perform some analytics over the union of the clients’ data, while preserving each individual client’s privacy. Specifically, we want to guarantee that for two neighboring input configurations of the clients denoted $\mathbf{x}$ and $\mathbf{x}'$ respectively, the distributions of the server’s view are “close”.

The shuffle model, first proposed by Bittau et al. [BEM⁺17] in an empirical work, has become a popular model for implementing distributed differentially private mechanisms. The model assumes the existence of a trusted shuffler that takes the union of all clients’ messages, randomly permutes them, and presents the shuffled result to the server. The server then performs some computation and outputs the analytics result. The trusted shuffler guarantees that the server can only see the union of all messages, without knowing the source of an individual message. Numerous earlier works [BBGN19, CSU⁺19, Che21, G GK⁺21, GKMP21] have shown that the shuffle model often enables differentially private mechanisms whose utility approximates the best known algorithms in the *central model* (where the server is trusted and we only need privacy on the outcome of the analytics). Moreover, several works have shown that the trusted shuffler can be efficiently implemented either using trusted hardware [BEM⁺17] or using cryptographic protocols [Cha81, Abe99, BG12, Cha88, CF10, APY20, CBM15, SW21, OS97, GIKM00, ZZZR05, EB21, RSG98, DS18, CL05]. This makes the shuffle model a compelling approach not just in theory, but also in practical applications such as federated learning [GDD⁺21].

A particular useful type of theorem in the shuffle model is called a privacy amplification theorem, which we explain below. Henceforth, let $\mathcal{R}(x_i)$ be some differentially private mechanism each client i applies to randomize its own private input x_i (often called a *locally differentially private* (LDP) randomizer). Roughly speaking, a privacy amplification theorem makes a statement of the following nature where $\mathcal{S}(\cdot)$ denotes the shuffler that outputs a random permutation of the inputs: if each client’s LDP mechanism $\mathcal{R}$ consumes ϵ_0 privacy budget, then shuffler’s outcome $\mathcal{S}(\mathcal{R}(x_1), \dots, \mathcal{R}(x_n))$ satisfies (ϵ, δ) -DP for $\epsilon = \epsilon(\epsilon_0, \delta) \ll \epsilon_0$, i.e., privacy is amplified for the overall shuffle-model mechanism. A line of work [CSU⁺19, EFM⁺19, BBGN19] focused on proving privacy amplification theorems for the shuffle model, culminating in the recent work by Feldman et al. [FMT21], who proved a privacy amplification theorem for any LDP mechanism with optimal parameters.

Connection between the shuffle model and our DO composition framework. We realize that the shuffle model can be expressed with our DO composition framework. Consider a composed mechanism $\mathbf{S} \circ \mathbf{M}_1$. $\mathbf{M}_1 : \mathcal{X}^n \rightarrow \mathcal{Y}^n$ is a local randomization mechanism that takes n clients’ inputs $(x_1, \dots, x_n)$ and outputs the message sequence $(y_1, \dots, y_n)$ where $y_i = \mathcal{R}(x_i)$. $\mathbf{S} : \mathcal{Y}^n \rightarrow \mathcal{Y}^n$ is a shuffling mechanism that takes a message sequence $(y_1, \dots, y_n)$ and outputs a random permutation of the sequence. All computation in $\mathbf{M}_1$ are done by the clients locally, so we define $\text{View}^{\mathbf{M}_1} := \emptyset$. We define the view in $\mathbf{S}$ as exactly its output: the random permutation. Then, the view of the server in the shuffle model is exactly the same as the view of the adversary in $\mathbf{S} \circ \mathbf{M}_1$. Thus, (ϵ, δ) -shuffle-DP guarantee can be expressed by $\mathbf{S} \circ \mathbf{M}_1$ being (ϵ, δ) -DO w.r.t the input neighboring notion $\sim_{\mathcal{X}}$ where $\mathbf{x} \sim_{\mathcal{X}} \mathbf{x}'$ iff the Hamming distance is at most 1.

Can we replace the shuffler with a DO-shuffler? A couple very recent works [GKLX22, BHMS21, ALU18] have suggested a relaxed shuffler model called the *differentially oblivious shuffle* model (or *DO-shuffle model* for short). Unlike the traditional shuffle model which provides full anonymity on the clients' messages, the DO-shuffle model permutes the clients' messages but possibly allowing some differentially private leakage. More concretely, a DO-shuffle protocol guarantees that for two *neighboring* input vectors $\mathbf{x}_H$ and $\mathbf{x}'_H$ corresponding to the set of honest parties, the adversary's views in the protocol execution are computationally or statistically close. The recent works by Gordon et al. [GKLX22] and Bünz et al. [BHMS21] both show that the relaxed DO-shuffle can be asymptotically more efficient to cryptographically realize than a fully anonymous shuffle. It would therefore be desirable to use a DO-shuffler as a drop-in replacement of the perfectly secure shuffle. This raises a couple very natural questions:

- *If we were to replace the shuffler in shuffle-model differentially private mechanisms with a DO-shuffler, can we still get comparable privacy-utility tradeoff?*
- *More specifically, can we prove an optimal privacy amplification theorem for the DO-shuffle model, matching the parameters of Feldman et al. [FMT21]?*

The pioneering work of Gordon et al. [GKLX22] was the first to explore how to use a DO shuffler to design distributed differentially private mechanisms. Gordon et al. [GKLX22] showed two novel results. First, they prove an optimal privacy amplification theorem for the randomized response mechanism in the DO-shuffle model, with parameters that tightly match the shuffle-model counterpart. Next, they generalize their first result, and prove a privacy amplification theorem for any local differentially private (LDP) mechanism — however, this more general result is *non-optimal*, since they rely on the non-optimal shuffle-model amplification theorem from Balle et al. [BBGN19].

Our results. We prove a privacy amplification theorem for *any LDP mechanism* that achieves *optimal* parameters, tightly matching Feldman et al. [FMT21]'s privacy amplification parameters for the shuffle model. This result improves work of Gordon et al. [GKLX22] in the following senses: 1) we asymptotically improve their privacy amplification theorem for any general LDP mechanism; and 2) their privacy amplification theorem for the specific randomized response mechanism can be viewed as a special case of our general theorem. More interestingly, we can prove our result fully under our DO composition framework. The curx of the proof is to show that the local randomization mechanism M_1 is (ϵ, δ) -NPDO w.r.t the output neighboring notion being exactly the DO-shuffler's input neighboring notion. Then, when M_1 composes with an (ϵ_1, δ_1) -DO shuffler, the composed mechanism will be $(\epsilon + \epsilon_1, \delta + \delta_1)$ -DO.

Below, we give a more formal statement of our result. Let Φ denote a DO-shuffling protocol. Given an LDP-randomizer $\mathcal{R}(\cdot)$, we use the notation $\Pi(x_1, \dots, x_n) := \Phi(\mathcal{R}(x_1), \dots, \mathcal{R}(x_n))$ to denote the composed protocol where each of the n parties first applies the local randomizer $\mathcal{R}(\cdot)$ to its own private data, and then invokes an instance of the DO-shuffling protocol Φ on the outcome $\mathcal{R}(x_i)$.

Theorem 1.4 (Optimal privacy amplification for any LDP mechanism in the DO-shuffle model). *Suppose $\epsilon_0 \leq \log\left(\frac{n}{16\log(2/\delta)}\right)$. Given n copies of an ϵ_0 -LDP randomizer $\mathcal{R}$ and an (ϵ_1, δ_1) -DO shuffler Φ resilient to t corrupted parties, the composed protocol $\Pi(x_1, \dots, x_n) := \Phi(\mathcal{R}(x_1), \dots, \mathcal{R}(x_n))$ is $(\epsilon + \epsilon_1, \delta + \delta_1)$ -DO against up to t corrupted parties where*

$$\epsilon = O\left(\frac{(1 - e^{\epsilon_0})e^{\epsilon_0/2}\sqrt{\log(1/\delta)}}{\sqrt{n - t}}\right).$$

Furthermore, if the DO-shuffler satisfies computational (or statistical, resp.) DO, then the composed protocol satisfies computation (or statistical, resp.) DO.

Further, if the underlying DO-shuffle protocol satisfies semi-honest security [GKLX22, ALU18], then the composed protocol is also secure in a semi-honest corruption model. Similarly, if the underlying DO-shuffle satisfies malicious security (e.g., [ALU18, BHMS21]), then the composed protocol is also secure in a malicious model.

2 Model and Preliminaries

2.1 Model of Computation

We consider a standard Random Access Machine (RAM) model of computation. We use the standard word-RAM model where the word size is logarithmic in the space. We assume that addition, multiplication, and boolean operations on words can be done in constant time. We also assume that sampling from truncated geometric distributions can be done in constant time². We assume that the adversary can observe the memory access patterns of the algorithm, including which locations are read or written and in which time steps. The adversary cannot see the contents of the memory tape themselves, which also means that the adversary cannot see the contents of the input and output.

Format of input and output tape. We explain the format of the input and output tape — the modeling technicalities are without loss of generality, and matter if we want to mask the true input and output lengths.

In the most general model, *the algorithm may or may not be able to observe the input and output length, depending on the algorithm*. More specifically, we may assume that the input is written on an input tape — the input tape itself has *unbounded* length and the *actual length of the input is written on some dedicated location*, e.g., address 0, of the input tape. The algorithm can read address 0 to learn the actual input length. During the execution, the algorithm *may read a random number of extraneous locations* on the input tape, such that the adversary may not be able to observe the exact input length. We assume that every extraneous location on the input tape stores a filler symbol $\perp$.

The algorithm must write the output on an output tape. Again, the algorithm, may *write a random number of extraneous locations* on the output tape. For example, if the actual output length is m , the algorithm may actually write $m' > m$ locations on the output tape where m' is a random variable, to mask the true output length. To indicate the actual output length, the algorithm can write the actual output length m on some dedicated location of the output tape. By doing so, the adversary may not be able to observe the exact output length.

2.2 Preliminaries

Mathematical tools. We introduce some basic mathematical tools.

Definition 2.1 (Symmetric geometric distribution). Let $\alpha > 1$. The symmetric geometric distribution $\text{Geom}(\alpha)$ takes integer values such that the probability mass function at k is $\frac{\alpha-1}{\alpha+1} \cdot \alpha^{-|k|}$.

²Based on Appendix B of [CZSC21], we can obviously sample a truncated geometric variable in expected time of $O(\frac{1}{\epsilon} \log \frac{1}{\delta})$. Further, we can sample M truncated geometric variables in time $O((M/\epsilon) \log(1/\delta))$ with probability $1 - \text{negl}(M)$. Our algorithms described in Section 5 only need to sample $\frac{\epsilon L}{(\log^2 L \log(1/\delta))} + O(1)$ truncated geometrics. Therefore, all our runtime bounds (i.e., Theorem 5.2 and Theorem 5.6) hold in expectation and with high probability without assuming sampling in constant time.

In designing DO algorithms, we often pad the true output length with random fillers such that the adversary observes a randomized output length. Below, we define a shifted and truncated geometric distribution which is often used to sample the number of fillers used for padding. In particular, this distribution always gives non-negative and bounded random variables.

Definition 2.2 (Shifted and truncated geometric distribution). Let $\epsilon > 0$ and $\delta \in (0, 1)$ and $\Delta \geq 1$. Let k_0 be the smallest positive integer such that $\Pr[|\text{Geom}(e^{\frac{\epsilon}{\Delta}})| \geq k_0] \leq \delta$, where $k_0 = \frac{\Delta}{\epsilon} \ln \frac{2}{\delta} + O(1)$. The shifted and truncated geometric distribution $\mathcal{G}(\epsilon, \delta, \Delta)$ has support $[0, 2(k_0 + \Delta - 1)]$, and is defined as:

$$\min\{\max\{0, k_0 + \Delta - 1 + \text{Geom}(e^\epsilon)\}, 2(k_0 + \Delta - 1)\}$$

For the special case $\Delta = 1$, we write $\mathcal{G}(\epsilon, \delta) := \mathcal{G}(\epsilon, \delta, 1)$.

Common distance notions. We will also use a couple common distance notions in our examples, including Hamming distance and edit distance.

Definition 2.3 (Hamming distance neighboring $\sim_H$). We say that two arrays x, x' are neighboring by the Hamming distance iff 1) they have the same length; and 2) they differ in at most one position.

Definition 2.4 (Edit distance neighboring $\sim_E$). We say that two arrays x, x' are neighboring by the edit distance iff x' can be obtained from x through either one insertion, one deletion, or one substitution. Note that x and x' need not have the same length.

Notations for randomized execution. Given randomized mechanisms $M_1 : \mathcal{X} \rightarrow \mathcal{Y}$ and $M_2 : \mathcal{Y} \times \mathcal{Z}$, the composed mechanism $M_2 \circ M_1 : \mathcal{X} \rightarrow \mathcal{Z}$ works as follows: for input $x \in \mathcal{X}$, we first apply $M_1(x)$ to produce an intermediate $y \in \mathcal{Y}$, and then we apply $M_2(y)$.

Henceforth, given an algorithm $M : \mathcal{X} \rightarrow \mathcal{Y}$, and an input $x \in \mathcal{X}$, we often use the following random variables:

- The random variable $\text{View}^M(x) : \mathcal{X} \rightarrow \mathcal{V}$ denotes the memory access patterns (also called the **view**) observed by the adversary when M receives the input x , where $\mathcal{V}$ is the view space for M .
- The notation $\text{Exec}^M(x) : \mathcal{X} \rightarrow \mathcal{V} \times \mathcal{Y}$ is a random variable that outputs the view and the output over a random execution of $M(x)$.

3 A Composition Framework for DO

In this section, we explore what kind of DO notions are composition-friendly. As a warmup, we first suggest a simple notion called strongly neighbor-preserving (or strongly NP for short), and show that any DO algorithm that is strongly NP lends to composition. The strong NP notion, however, is too stringent. We then propose a more general notion called (ϵ, δ) -neighbor-preserving differential obliviousness or (ϵ, δ) -NPDO for short, which captures a probabilistically approximate notion of neighbor-preserving. We then present our main composition theorem which states that any algorithm that satisfies NPDO lends to composition. Along the way, we give several simple motivating examples to demonstrate the usefulness our compositional framework.

3.1 Strongly Neighbor-Preserving

3.1.1 Definition and Composition Theorem

Earlier, in Section 1, we argued why basic DO algorithms do not lend to composition, because neighboring inputs may lead to very dissimilar outputs. One (somewhat imprecise) intuition is the following: if a DO mechanism is additionally *neighbor-preserving*, i.e., neighboring inputs lead to neighboring outputs, then it should lend to composition.

We first define strongly neighbor-preserving that running the algorithm over two neighboring inputs produces neighboring outputs with probability 1.

Definition 3.1 (Strongly neighbor-preserving). We say that a randomized algorithm $M : \mathcal{X} \rightarrow \mathcal{Y}$ is strongly neighbor-preserving w.r.t. $\sim_{\mathcal{X}}$ and $\sim_{\mathcal{Y}}$, iff for any two inputs $x, x' \in \mathcal{X}$ such that $x \sim_{\mathcal{X}} x'$,

$$\Pr[y \leftarrow M(x), y' \leftarrow M(x') : y \sim_{\mathcal{Y}} y'] = 1.$$

We can prove that if an algorithm satisfies both DO and strongly neighbor-preserving, then it is composable, formally stated below.

Theorem 3.2 (Strongly neighbor-preserving + DO gives composition). *Suppose that $M_1 : \mathcal{X} \rightarrow \mathcal{Y}$ is (ϵ_1, δ_1) -DO w.r.t. $\sim_{\mathcal{X}}$ and strongly neighbor-preserving w.r.t. $\sim_{\mathcal{X}}$ and $\sim_{\mathcal{Y}}$, and moreover, suppose that $M_2 : \mathcal{Y} \rightarrow \mathcal{Z}$ is (ϵ_2, δ_2) -DO w.r.t. $\sim_{\mathcal{Y}}$, then $M_2 \circ M_1$ satisfies $(\epsilon_1 + \epsilon_2, \delta_1 + \delta_2)$ -DO w.r.t. $\sim_{\mathcal{X}}$.*

Furthermore, if M_2 is additionally strongly neighbor-preserving w.r.t. $\sim_{\mathcal{Y}}$ and $\sim_{\mathcal{Z}}$, then $M_2 \circ M_1$ is also strongly neighbor-preserving w.r.t. $\sim_{\mathcal{X}}$ and $\sim_{\mathcal{Z}}$.

Proof. Later in Lemma 3.7 of Section 3.4, we will prove that (ϵ_1, δ_1) -DO plus strongly neighbor-preserving is a special case of our more general notion (ϵ_1, δ_1) -NPDO. In this sense, this composition theorem can be viewed as a special case of our main composition theorem for NPDO (Theorem 3.6). $\square$

3.1.2 Composition Examples

Example 1. Recall that in Section 1, we gave a natural SQL database example that required applying one compaction algorithm on the output of another compaction algorithm. We pointed out that two sequential instances of Chan et al.'s DO compaction algorithm [CCMS19] do not give (tight) composable guarantees. In Example 1, we will see that if we replace the second instance with a modification of the compaction algorithm such that it is DO w.r.t. edit distance (as opposed to Hamming distance), the two instances would compose nicely.

Specifically, let M_1 be Chan et al.'s DO compaction algorithm [CCMS19]. Recall that the algorithm receives an input array where each element is either a *real* element or a *filler*, and outputs an array containing all the real elements in the input and preserving the order they appear in the input. M_1 is (ϵ_1, δ_1) -DO w.r.t. $\sim_H$ (i.e., Hamming distance). Now, suppose we can construct another compaction algorithm denoted M_2 that is (ϵ_2, δ_2) -DO w.r.t. to $\sim_E$ (i.e., *edit* distance). How to construct such an M_2 while preserving efficiency turns out to be non-trivial, and we defer the construction to Section 5 — interestingly, designing M_2 itself demonstrates the usefulness of our composition framework, too.

Observe that given a fixed input array x , the output of $M_1(x)$ must be an ordered list of real elements contained in x plus an appropriate number of fillers, and the total length of the

output³ is the same as the input x . Thus, for any neighboring inputs $x \sim_H x'$, it must be that $M_1(x) \sim_E M_2(x')$. Therefore, we conclude that M_1 is strongly neighbor-preserving w.r.t. the input relation $\sim_H$ and the output relation $\sim_E$. Applying Theorem 3.2, we conclude that the composed mechanism $M_2 \circ M_1$ satisfies $(\epsilon_1 + \epsilon_2, \delta_1 + \delta_2)$ -DO.

Example 2. Let M_1 be an algorithm that merges two sorted input arrays (x_0, x_1) , where each element in the input array has a payload besides the sort-key. Suppose that M_1 satisfies (ϵ_1, δ_1) differential obliviousness w.r.t. $\overset{2}{\sim}_E$, i.e., two inputs (x_0, x_1) and (x'_0, x'_1) are considered neighboring iff for $b \in \{0, 1\}$, $|x_b| = |x'_b|$, and x_b and x'_b have edit distance at most 2 (i.e., $x_b \overset{2}{\sim}_E x'_b$). Such an DO merge algorithm was proposed by Chan et al. [CCMS19], and moreover, their algorithm always outputs an array whose length is the sum of the input arrays. Notice that for neighboring inputs, M_1 always produces outputs that have edit distance at most 4.

Let M_2 be a stable tight compaction algorithm that selects elements from the input array whose payload string satisfies a certain predicate (e.g., entries corresponding to students in the computer science department). Suppose that M_2 satisfies (ϵ_2, δ_2) -DO w.r.t. $\overset{4}{\sim}_E$, i.e., where neighboring inputs are those with edit distance at most 4 — such an M_2 is described in Section 5.

By Theorem 3.2, we conclude that the composed mechanism $M_2 \circ M_1$ satisfies $(\epsilon_1 + \epsilon_2, \delta_1 + \delta_2)$ -DO w.r.t. $\overset{2}{\sim}_E$.

Remark 3.3 (Capturing k -neighboring relations). Recall that our strongly neighbor-preserving definition (i.e., Definition 3.1) is parametrized with the input and output relations. Example 2 is used to illustrate the case when these input/output relations are parametrized with a k -neighboring notion (rather than 1-neighboring) — this shows the generality of the approach. For example, later in Section 5, we will construct an efficient stable compaction algorithm that is (ϵ, δ) -DO w.r.t. to $\overset{1}{\sim}_E$ neighboring. Applying the standard group privacy theorem of differential privacy [DR14], we can get a compaction algorithm that is $(4\epsilon, 4e^{4\epsilon}\delta)$ -DO w.r.t. to $\overset{4}{\sim}_E$ neighboring.

3.1.3 Limitations

The strong neighbor-preserving requirement (i.e., Definition 3.1) is natural and directly captures our intuition that if a DO mechanism maps neighboring inputs to neighboring outputs, then it is composable. The strongly neighbor-preserving requirement is often suitable when the output computed by the algorithm is deterministic (i.e., uniquely determined by the inputs), even though the algorithm itself may be randomized, like Examples 1 and 2.

However, the strongly neighbor-preserving requirement may be too stringent especially when the output of the algorithm may be randomized. For example, consider the following DO subsampling algorithm.

Example 3. We consider the task of subsampling, which is widely used in private data analytics [BBG18, WBK21]: given an input array x , we want to sample each entry with probability p , and generate a new array with only the sampled elements. Consider a subsampling algorithm where n denotes the length of the input array x :

1. Call $M_1(x) := \text{InPlaceSample}(x)$ which is defined as follows: Scan the input array x . For each real element encountered, append it to the output tape with probability p and append a filler element otherwise. For each filler element encountered, just append a filler to the output tape.

³Even though the algorithm M_1 itself is randomized, the output of M_1 is deterministic and unique given the input.

2. Apply M_2 , a compaction algorithm that is (ϵ', δ') -DO w.r.t. $\sim_H$ to the output of the above step.

We want to prove that the above algorithm satisfies DO w.r.t. $\sim_H$ through composition — intuitively, this should be true. In particular, the first subroutine $M_1 := \text{InPlaceSample}$ has deterministic access patterns. We explicitly denote $M_1(\cdot; \rho)$ to fix the random tape ρ consumed by M_1 . For any fixed random tape ρ , and any neighboring inputs $x \sim_H x'$, $M_1(x; \rho)$ and $M_1(x'; \rho)$ output two arrays with Hamming distance 1. Therefore, intuitively, as long as the compaction algorithm in the second step is (ϵ', δ') -DO w.r.t. Hamming distance, the entire subsampling algorithm should be (ϵ', δ') -DO as well. Unfortunately, we cannot directly use strong neighbor-preserving to prove this composition here, since a random execution of $M_1(x)$ and a random execution of $M_1(x')$ are not guaranteed to always output Hamming-distance-neighboring outputs — it depends on which subset of elements are selected.

This motivates us to relax the strongly neighbor-preserving to make it more general, such that our compositional framework can be more expressive. However, before we do so, we introduce another more general example, Example 4, which is a variation of Example 3. Specifically, in Example 3, although the output of $M_1 := \text{InPlaceSample}$ is randomized, the view of M_1 is deterministic. In Example 4, both the view and the output of the first algorithm are randomized.

Example 4. The main difference between Examples 3 and 4 is that Example 4 aims to have a subsampling algorithm that is DO w.r.t. *edit distance*, whereas Example 3 aims to be DO w.r.t. *Hamming distance*. To achieve this, in Example 4, we need to mask the true length of the input and output by reading/writing a random number of extraneous locations on the input tape. Further, the compaction algorithm we call must now be DO w.r.t. edit distance too. The detailed algorithm is described below. The key differences are highlighted by underlining.

1. Call $M_1(x) := \text{InPlaceSample}_{\epsilon, \delta}(x)$ which is defined as follows:
 - Sample $r \xleftarrow{\$} \mathcal{G}(\epsilon, \delta, \Delta = 1)$, let $n' = n + r$ be the noisy input length.
 - Scan n' locations on the input tape. For each real element encountered, append it to the output tape with probability p and append a filler element otherwise. For each filler encountered, append a filler to the output tape.
 - The output array is defined to be the first n elements of the output tape.
Write down its length n at a fixed dedicated location on the output tape.
2. Apply M_2 , a compaction algorithm that is (ϵ', δ') -DO w.r.t. $\sim_E$ to the output of the above step, i.e., the compaction algorithm treats the output tape of M_1 as its own input tape.

We later prove that Examples 3 and 4 satisfy DO with our new framework.

3.2 (ϵ, δ) -Neighbor-Preserving Differential Obliviousness (NPDO)

Recognizing the limitations of strongly neighbor-preserving (Definition 3.1), we would like to make the compositional framework more general. In particular, the above Examples 3 and 4 can serve as simple motivating examples.

Given a mechanism M whose view space is $\mathcal{V}$ and output space is $\mathcal{Y}$, given some symmetric relation $\sim_Y$ over the output space, and given a set $S \subseteq \mathcal{V} \times \mathcal{Y}$, we define the following notation for denoting **neighbor sets**:

$$\mathcal{N}(S) := \{(v, y') | \exists (v, y) \in S \text{ s.t. } y' \sim_Y y\}$$

Definition 3.4 ((ϵ, δ) -NPDO). Given a mechanism $M : \mathcal{X} \rightarrow \mathcal{Y}$ with view space $\mathcal{V}$, we say that it satisfies (ϵ, δ) -neighbor-preserving differential obliviousness, or (ϵ, δ) -NPDO for short, w.r.t. symmetric relations $\sim_{\mathcal{X}}$ and $\sim_{\mathcal{Y}}$, respectively, iff for all $x \sim_{\mathcal{X}} x'$, for every $S \subseteq \mathcal{V} \times \mathcal{Y}$,

$$\Pr[\text{Exec}^M(x) \in S] \leq e^\epsilon \cdot \Pr[\text{Exec}^M(x') \in \mathcal{N}(S)] + \delta. \quad (2)$$

Our NPDO definition looks similar in form as the standard differential privacy notion, with some important observations: 1) the notion is defined over the Cartesian product of the view and the output of the mechanism, which is important for composition to hold; 2) on the right-hand-side of Equation (2), we consider the probability of $M(x')$ landing in the *neighboring* set $\mathcal{N}(S)$ on a neighboring input $x' \sim_{\mathcal{X}} x$ — this is important for capturing a probabilistic notion of neighbor-preserving.

It is not hard to see that if an algorithm satisfies (ϵ, δ) -NPDO, it must satisfy (ϵ, δ) -DO, as stated in the following theorem.

Theorem 3.5. *Suppose that $M : \mathcal{X} \rightarrow \mathcal{Y}$ satisfies (ϵ, δ) -NPDO w.r.t. $\mathcal{X}$ and $\mathcal{Y}$. Then, M satisfies (ϵ, δ) -DO w.r.t. $\mathcal{X}$.*

Proof. Let $\mathcal{V}$ be the sample space of View^M . Consider all those $S_v \subseteq \mathcal{V}$. We know $\mathcal{N}(S_v \times \mathcal{Y}) = S_v \times \mathcal{Y}$. Therefore,

$$\begin{aligned} \Pr[\text{View}^M(x) \in S_v] &= \Pr[\text{Exec}^M(x) \in S_v \times \mathcal{Y}] \\ &\leq e^\epsilon \Pr[\text{Exec}^M(x') \in \mathcal{N}(S_v \times \mathcal{Y})] + \delta = e^\epsilon \Pr[\text{View}^M(x') \in S_v] + \delta. \end{aligned}$$

□

3.3 Main Composition Theorem

One main technical contribution of our paper is to prove a composition theorem for our NPDO notion, as stated below.

Theorem 3.6 (Main composition theorem). *Suppose that an algorithm $M_1 : \mathcal{X} \rightarrow \mathcal{Y}$ satisfies (ϵ_1, δ_1) -NPDO w.r.t. $\sim_{\mathcal{X}}$ and $\sim_{\mathcal{Y}}$. Further, suppose that the algorithm M_1 's view space $\mathcal{V}$ and the output space $\mathcal{Y}$ are finite or countably infinite. Then, the following composition statements hold:*

1. *Suppose that $M_2 : \mathcal{Y} \rightarrow \mathcal{Z}$ satisfies (ϵ_2, δ_2) -DO w.r.t. $\sim_{\mathcal{Y}}$. Then, the composed mechanism $M_2 \circ M_1 : \mathcal{X} \rightarrow \mathcal{Z}$ satisfies $(\epsilon_1 + \epsilon_2, \delta_1 + \delta_2)$ -DO.*
2. *Suppose that $M_2 : \mathcal{Y} \rightarrow \mathcal{Z}$ satisfies (ϵ_2, δ_2) -NPDO w.r.t. $\sim_{\mathcal{Y}}$ and $\sim_{\mathcal{Z}}$. Then, the composed mechanism $M_2 \circ M_1 : \mathcal{X} \rightarrow \mathcal{Z}$ satisfies $(\epsilon_1 + \epsilon_2, \delta_1 + \delta_2)$ -NPDO.*

The proof of Theorem 3.6 is presented in Section 4. We can use Theorem 3.6 to prove that the algorithms in the earlier Examples 1 to 4 satisfy DO. Before doing so, let us first introduce some helpful tools for proving an algorithm NPDO.

3.4 Helpful Tools for Proving NPDO

To use our main composition theorem, we need to prove that some algorithm satisfies NPDO. The following lemmas provide helpful tools for this purpose.

Strongly NP + DO $\implies$ NPDO. First, it is not hard to see that if an algorithm satisfies the earlier strongly neighbor-preserving notion (Definition 3.1) as well as DO, then it also satisfies NPDO as stated below:

Lemma 3.7 (Strongly NP and DO imply NPDO). *Suppose that an algorithm $M : \mathcal{X} \rightarrow \mathcal{Y}$ is strongly neighbor-preserving w.r.t. $\sim_{\mathcal{X}}$ and $\sim_{\mathcal{Y}}$, as well as (ϵ, δ) -DO w.r.t. $\sim_{\mathcal{X}}$. Then, M satisfies (ϵ, δ) -NPDO.*

Proof. Fix a pair of neighboring input x, x' . By the strongly neighbor-preserving definition, for any $y, y' \in \mathcal{Y}$ such that $\Pr[M(x) = y] > 0$ and $\Pr[M(x') = y'] > 0$, it must be that $y \sim_{\mathcal{Y}} y'$. For any subset $S \subseteq \mathcal{V} \times \mathcal{Y}$, define the partial set $V(S) := \{v \mid \exists (v, y) \in S\}$. Then, we have

$$\begin{aligned} \Pr[\text{Exec}^M(x) \in S] &\leq \Pr[\text{View}^M(x) \in V(S)] \\ &\leq e^\epsilon \Pr[\text{View}^M(x') \in V(S)] + \delta \leq e^\epsilon \Pr[\text{Exec}^M(x') \in \mathcal{N}(S)] + \delta. \end{aligned}$$

□

(ϵ, δ) -NP. Next, we define another notion that captures the idea of “probabilistically approximate neighbor-preserving” called (ϵ, δ) -neighbor-preserving, or (ϵ, δ) -NP for short. We show that if an algorithm satisfies (ϵ, δ) -NP as well as (ϵ', δ') -DO, then it also satisfies $(\epsilon + \epsilon', \delta + \delta')$ -NPDO.

Definition 3.8 ((ϵ, δ) -NP). Given a mechanism $M : \mathcal{X} \rightarrow \mathcal{Y}$ whose view space is $\mathcal{V}$, we say that it satisfies (ϵ, δ) -neighbor-preserving, or (ϵ, δ) -NP for short, w.r.t. $\sim_{\mathcal{X}}$ and $\sim_{\mathcal{Y}}$, iff for all $x \sim_{\mathcal{X}} x'$, for every view $v^* \in \mathcal{V}$ that happens with non-zero probability in $\text{Exec}^M(x)$ as well as $\text{Exec}^M(x')$, for every $Y \subseteq \mathcal{Y}$,

$$\begin{aligned} \Pr[(v, y) \leftarrow \text{Exec}^M(x) : y \in Y \mid v = v^*] \\ \leq e^\epsilon \cdot \Pr[(v', y') \leftarrow \text{Exec}^M(x') : y' \in \mathcal{N}(Y) \mid v' = v^*] + \delta \end{aligned} \quad (3)$$

where $\mathcal{N}(Y)$ contains all y' such that $y' \sim_{\mathcal{Y}} y$ for all $y \in Y$.

Intuitively, (ϵ, δ) -NP requires that *conditioned on any view*, the algorithm, on neighboring inputs, must output probabilistically approximately close outputs.

Lemma 3.9 ((ϵ, δ) -NP and DO imply NPDO). *Suppose that an algorithm $M : \mathcal{X} \rightarrow \mathcal{Y}$ is (ϵ_1, δ_1) -DO and (ϵ_2, δ_2) -neighbor-preserving w.r.t. $\sim_{\mathcal{X}}$ and $\sim_{\mathcal{Y}}$. Then, M satisfies $(\epsilon_1 + \epsilon_2, \delta_1 + \delta_2)$ -NPDO w.r.t. $\sim_{\mathcal{X}}$ and $\sim_{\mathcal{Y}}$.*

Proof. We slightly abuse the notation of $\mathcal{N}(\cdot)$ to denote the “neighbor set” for $\mathcal{Y}$ and also the $\mathcal{V} \times \mathcal{Y}$, as defined in Definition 3.8 and Definition 3.4. Fix a pair of neighboring input x, x' . For any $S \subseteq \mathcal{V} \times \mathcal{Y}$, denote the partial set $S_v = \{y \mid \exists (v, y) \in S\}$. Let $\mu(v) = \min(\Pr[\text{View}^M(x) = v], e^{\epsilon_1} \Pr[\text{View}^M(x') = v])$. It’s trivial that $\sum_{v \in \mathcal{V}} \mu(v) \leq 1$. Also, since M is (ϵ_1, δ_1) -DO, we have

$\sum_{v \in \mathcal{V}} \mu(v) \geq 1 - \delta_1$. Then,

$$\begin{aligned}
& \Pr[\text{Exec}^M(x) \in S] \\
&= \sum_{v^* \in \mathcal{V}} \Pr[\text{View}^M(x) = v^*] \Pr[(v, y) \leftarrow \text{Exec}^M(x) : y \in S_{v^*} \mid v = v^*] \\
&\leq \delta_1 + \sum_{v^* \in \mathcal{V}} \mu(v^*) \Pr[(v, y) \leftarrow \text{Exec}^M(x) : y \in S_{v^*} \mid v = v^*] \\
&\leq \delta_1 + \sum_{v^* \in \mathcal{V}} \mu(v^*) (e^\epsilon \Pr[(v', y') \leftarrow \text{Exec}^M(x') : y' \in \mathcal{N}(S_{v^*}) \mid v' = v^*] + \delta_2) \\
&\leq \delta_1 + \delta_2 + \sum_{v^* \in \mathcal{V}} \mu(v^*) (e^\epsilon \Pr[(v', y') \leftarrow \text{Exec}^M(x') : y' \in \mathcal{N}(S_{v^*}) \mid v' = v^*]) \\
&\leq \delta_1 + \delta_2 + \sum_{v^* \in \mathcal{V}} e^{\epsilon_1} \Pr[\text{View}^M(x') = v^*] e^{\epsilon_2} \Pr[(v', y') \leftarrow \text{Exec}^M(x') : y' \in \mathcal{N}(S_{v^*}) \mid v' = v^*] \\
&= \delta_1 + \delta_2 + e^{\epsilon_1 + \epsilon_2} \Pr[\text{Exec}^M(x') \in \mathcal{N}(S)].
\end{aligned}$$

□

3.5 Our Composition Theorem in Action

Using the simple motivating examples introduced so far, we can see our composition theorems in action.

Examples 1 and 2. As mentioned earlier, the first algorithm M_1 in either Example 1 or Example 2 satisfies strongly neighbor-preserving as well as (ϵ_1, δ_1) -DO. Therefore, they can be viewed as a special case of (ϵ, δ) -NPDO. Since M_2 in Example 1 or 2 satisfies (ϵ_2, δ_2) -DO, by our main composition theorem, we immediately reach the conclusion that the composed algorithm $M_2 \circ M_1$ satisfies $(\epsilon_1 + \epsilon_2, \delta_1 + \delta_2)$ -DO.

Example 3. We can use Theorem 3.6 to prove that the subsampling algorithm of Example 3 satisfies (ϵ', δ') -DO w.r.t. $\sim_H$. To accomplish this, it suffices to show that the first algorithm, $M_1 := \text{InPlaceSample}$, satisfies $(0, 0)$ -NPDO w.r.t. $\sim_H$ and $\sim_H$. Observe that in Example 3, two inputs are neighboring if their Hamming distance is at most 1, which implies that neighboring inputs must have the same length. Also, M_1 always generates a deterministic view that depends only on the length of the input. Therefore, to prove that $M_1(x)$ satisfies $(0, 0)$ -NPDO, it suffices to show that for any pair of neighboring inputs $x \sim_H x'$, for any subset of outputs $Y \subseteq \mathcal{Y}$ where $\mathcal{Y}$ is the output space of M_1 ,

$$\Pr[M_1(x) \in Y] \leq \Pr[M_1(x') \in \mathcal{N}(Y)], \quad (4)$$

where $\mathcal{N}(Y)$ denotes the set of all output arrays that are neighboring to some array in Y . Observe also that for any possible output y of $M_1(x)$, let ρ be the random coins used for subsampling that led to the result y , then, if the same random coins ρ is encountered in an execution of $M_1(x')$ on some neighboring $x' \sim_H x$, the outcome must be neighboring to y . Therefore, Equation (4) holds.

Example 4. Similarly, we can use Theorem 3.6 to prove that the subsampling algorithm of Example 4 satisfies $(\epsilon + \epsilon', \delta + \delta')$ -DO w.r.t. $\sim_E$. By Theorem 3.6, it suffices to show that the $M_1 := \text{InPlaceSample}_{\epsilon, \delta}$ algorithm in Example 4 satisfies (ϵ, δ) -NPDO w.r.t. $\sim_E$ being both of the

input and output neighboring notion. Recall that M_1 pads the input array with a random number of elements, such that the noisy length is n' . Then, it simply scans through the n' elements and either writes down the element if it is a real element and has been sampled, or writes down $\perp$. To show that M_1 satisfies (ϵ, δ) -NPDO, we will prove that M_1 satisfies $(0, 0)$ -NP and (ϵ, δ) -DO, respectively, and then the conclusion follows from Lemma 3.9. It is easy to prove that M_1 satisfies (ϵ, δ) -DO. To see this, observe that the view depends only on the noisy input length where the noise is sampled according to a truncated geometric distribution.

Therefore, we focus on showing that M_1 satisfies $(0, 0)$ -NP. Observe that in M_1 , the random coins that determine the view and those that determine the output are independent. Therefore, it suffices to show that for any $x \sim_E x'$, for any $Y \subseteq \mathcal{Y}$ where $\mathcal{Y}$ is the output space of M_1 ,

$$\Pr[M_1(x) \in Y] \leq \Pr[M_1(x') \in \mathcal{N}(Y)].$$

Since $x \sim_E x'$, there can be at most one element in x that is not in x' (e.g., the element that is added or modified in x), and vice versa. Henceforth, we use $\text{Common}(x, x')$ to denote the list of common elements that appear both in x and x' . Let $G(Y)$ be the event that there exists some $y \in Y$, such that the elements in $\text{Common}(x, x')$ receive the same sampling decision as in y . We also say that $G(Y)$ represents the event that $\text{Common}(x, x')$ receive coins compatible with Y . Therefore, we have that

$$\Pr[M_1(x) \in Y] \leq \Pr[M_1(x) : G(Y)] \leq \Pr[M_1(x') \in \mathcal{N}(Y)].$$

In the above, the second inequality holds since conditioned on $\text{Common}(x, x')$ receiving coins compatible with Y in a random execution of $M_1(x')$, the outcome must be neighboring to some element in Y with probability 1.

Additional applications. Later in Section 5, we use our composition framework to design a differentially oblivious stable compaction algorithm w.r.t. the edit distance — this building block was needed in Examples 1, 2, 4. Last but not the least, in Section 6, we use our composition framework to prove an optimal privacy amplification theorem for the DO-shuffle model.

4 Proof of Main Composition Theorem

In this section, we shall prove our main composition theorem, that is, Theorem 3.6. A key stepping stone is the following equivalence lemma.

Lemma 4.1 (Equivalence of (ϵ, δ) -NPDO and existence of an (ϵ, δ) -matching). *Assume the axiom of choice. Given a finite or countable infinite sample space Ω and a symmetric relation $\sim$ on Ω , consider two random variables $A, B \in \Omega$. The following statements are equivalent:*

1. *For every $S \subseteq \Omega$, $\Pr[A \in S] \leq e^\epsilon \cdot \Pr[B \in \mathcal{N}(S)] + \delta$, where the neighbor set $\mathcal{N}(S)$ is defined as $\mathcal{N}(S) := \{b \in \Omega \mid \exists a \in S, a \sim b\}$.*
2. *There exists an (ϵ, δ) -matching $w : \Omega \times \Omega \rightarrow [0, 1]$ satisfying the following conditions:*
 - (a) *For all $a, b \in \Omega$, $w(a, b) > 0$ only if $a \sim b$;*
 - (b) *For all $a \in \Omega$, $\sum_{b \in \Omega, b \sim a} w(a, b) \leq \Pr[A = a]$;*
 - (c) *For all $b \in \Omega$, $\sum_{a \in \Omega, a \sim b} w(a, b) \leq e^\epsilon \cdot \Pr[B = b]$;*
 - (d) $\sum_{a, b \in \Omega} w(a, b) \geq 1 - \delta$.

Graph interpretation. Lemma 4.1 has a similar flavor as the Hall's theorem for bipartite graphs. The Hall's theorem says that if for each subset S of one component of a bipartite graph, the size of its neighbor set satisfies $|\mathcal{N}(S)| \geq |S|$, then we can find a perfect matching in the graph. The proof of Lemma 4.1 is also inspired by the proof of the Hall's theorem.

We think of a bipartite graph where vertices on the left and right both come from the set Ω , and $w(a, b)$ defines the weight on edge (a, b) . Imagine that each vertex $a \in \Omega$ on the left is a factory that produces $\Pr[A = a]$ amount of produce, and each vertex $b \in \Omega$ on the right is a warehouse that can store up to $e^\epsilon \cdot \Pr[B = b]$ amount of produce. Condition (a) says that a factory is only allowed to route its produce to neighboring warehouses. The function w effectively defines a fractional flow such that almost all, i.e., $1 - \delta$ amount of produce is routed to some warehouse, and moreover, none of the warehouses exceed their capacity. For this reason, we also call w an (ϵ, δ) -**matching**. The full proof of Lemma 4.1 is deferred to Appendix B.2. Below, we prove our main composition theorem assuming that Lemma 4.1 holds.

Proof of Theorem 3.6. We directly prove the more general case when M_2 is (ϵ_2, δ_2) -NPDO. When M_2 is only (ϵ_2, δ_2) -DO, we can prove $M_2 \circ M_1$ is $(\epsilon_1 + \epsilon_2, \delta_1 + \delta_2)$ -DO with nearly the same argument.

Fix any neighboring input x, x' . By Lemma 4.1, there exists an (ϵ_1, δ_1) -matching $w : (\mathcal{V}_1 \times \mathcal{Y}) \times (\mathcal{V}_1 \times \mathcal{Y}) \rightarrow [0, 1]$ w.r.t the natural neighbor notion $\sim$ in the product space $\mathcal{V}_1 \times \mathcal{Y}$: $(v_1, y) \sim (v'_1, y')$ when $v_1 = v'_1$ and $y \sim_{\mathcal{Y}} y'$. We want to prove that, for any subset $S \subseteq \mathcal{V}_1 \times \mathcal{V}_2 \times \mathcal{Z}$,

$$\Pr[\text{Exec}^{M_2 \circ M_1}(x) \in S] \leq e^{\epsilon_1 + \epsilon_2} \Pr[\text{Exec}^{M_2 \circ M_1}(x') \in \mathcal{N}(S)] + \delta_1 + \delta_2.$$

Define the partial set $S_{v_1} := \{(v_2, z) | \exists (v_1, v_2, z) \in S\}$ for any $v_1 \in \mathcal{V}_1$. Then,

$$\begin{aligned} & \Pr[\text{Exec}^{M_2 \circ M_1}(x) \in S] \\ &= \sum_{(v_1, y) \in \mathcal{V}_1 \times \mathcal{Y}} \Pr[\text{Exec}^{M_1}(x) = (v_1, y)] \cdot \Pr[\text{Exec}^{M_2}(y) \in S_{v_1}] \\ & \quad (\text{Use condition (a), (b) and (d) of the matching}) \\ &\leq \sum_{(v_1, y) \in \mathcal{V}_1 \times \mathcal{Y}, y' \sim_{\mathcal{Y}} y} w((v_1, y), (v_1, y')) \cdot \Pr[\text{Exec}^{M_2}(y) \in S_{v_1}] + \delta_1 \\ &\leq \sum_{(v_1, y) \in \mathcal{V}_1 \times \mathcal{Y}, y' \sim_{\mathcal{Y}} y} w((v_1, y), (v_1, y')) \cdot \left(e^{\epsilon_2} \Pr[\text{Exec}^{M_2}(y') \in \mathcal{N}(S_{v_1})] + \delta_2 \right) + \delta_1 \\ & \quad (\text{Use condition (b) of the matching}) \\ &\leq \sum_{(v_1, y) \in \mathcal{V}_1 \times \mathcal{Y}, y' \sim_{\mathcal{Y}} y} w((v_1, y), (v_1, y')) \cdot \left(e^{\epsilon_2} \Pr[\text{Exec}^{M_2}(y') \in \mathcal{N}(S_{v_1})] \right) + \delta_2 + \delta_1 \\ & \quad (\text{Use condition (c) of the matching}) \\ &\leq \sum_{(v_1, y') \in \mathcal{V}_1 \times \mathcal{Y}} e^{\epsilon_1} \Pr[\text{Exec}^{M_1}(x') = (v_1, y')] \cdot \left(e^{\epsilon_2} \Pr[\text{Exec}^{M_2}(y') \in \mathcal{N}(S_{v_1})] \right) + \delta_2 + \delta_1 \\ &= e^{\epsilon_1 + \epsilon_2} \Pr[\text{Exec}^{M_2 \circ M_1}(x') \in \mathcal{N}(S)] + \delta_2 + \delta_1. \end{aligned}$$

□

5 Application: DO Compaction w.r.t. Edit Distance

Earlier in our Examples 1, 2, and 4, we assumed a stable compaction algorithm that is differentially oblivious w.r.t. the *edit* distance. Chan et al. [CCMS19] showed how to construct a

stable compaction algorithm that is (ϵ, δ) -DO w.r.t. the *Hamming* distance [CCMS19], taking $O(n(\log \log n + \log \log \frac{1}{\delta}))$ time to compact an array of size n (assuming that ϵ is a constant). However, we are not aware of any straightforward way to modify their algorithm to work for edit distance. Another naïve approach is to use oblivious sorting directly but this would incur $\Theta(n \log n)$ runtime which is asymptotically worse. In this section, we fill in this missing piece that is needed by Examples 1, 2, and 4. We will describe a stable compaction algorithm that works for edit distance and it preserves the runtime of Chan et al. [CCMS19]. Intriguingly, the design of our new compaction algorithm turns out to be a great example that demonstrates the power of our compositional framework.

5.1 Additional Preliminaries

Stable compaction. Recall that in stable compaction, we are given an input array which is written on an input tape. Some elements in the input array are *real* elements, and others are *fillers*. We want to output an array that contains only the real elements, and they must appear in the same order as the input array. We assume that the input array is written on the input tape, and its true length is written on some designated location on the input tape. The algorithm should write the output array to an output tape, and the true length of the output array should be written to some dedicated location on the output array.

Stable Oblivious Sorting. Suppose we are given an input array I containing a list of m elements with a key attached to each element. Earlier works [CGLS18, LSX19] showed how to obliviously sort the array according to the keys in $O(m \log m)$ runtime while maintaining the stable property: the elements will be ordered by their relative order in the original array when their keys are the same.

Differentially private prefix sum. Given an input array I containing a list of m integers, we want to its prefix sums. We say that two inputs I , and I' are neighboring iff 1) they have the same length and 2) they differ in at most one position j , and $|I[j] - I'[j]| \leq 1$. Earlier works [DNPR10, CSS10] showed how to construct a prefix sum mechanism that satisfies (ϵ, δ) -differential privacy, and moreover, the mechanism satisfies the following properties: 1) The access patterns (i.e., view) of the algorithm depend only on the input length; 2) The additive error is upper bounded by $O(\frac{1}{\epsilon}(\log |I|)^{1.5} \log \frac{1}{\delta})$ with probability 1.

5.2 Roadmap and Intuition

Our algorithm **Compact** is the composition of the following two algorithms, i.e., $\text{Compact}(\cdot) = \text{CompactBin} \circ \text{RandBin}(\cdot)$. Suppose we can prove that **RandBin** is (ϵ_1, δ_1) -NPDO and prove that **CompactBin** is (ϵ_2, δ_2) -NPDO, we have **Compact** is $(\epsilon_1 + \epsilon_2, \delta_1 + \delta_2)$ -NPDO due to our main composition theorem 3.6.

1. **RandBin:** Given an input array I containing real elements and fillers, and whose true length is stored in a dedicated location on the input tape, **RandBin** outputs a list of B bins denoted $(\text{Bin}_i^{(Z)} : i \in [B])$, each of capacity Z . Each bin contains a random number of real elements and the rest are fillers. Furthermore, the ordered list of all real elements in all bins is the same as the ordered list of real elements in the input. The algorithm should output the parameters B and Z to some dedicated location on the output tape.
2. **CompactBin:** Given a list of B bins denoted $(\text{Bin}_i^{(Z)} : i \in [B])$ each of capacity Z , where the parameters B and Z are stored in some dedicated location on the input tape, the **CompactBin**

algorithm outputs a compacted array containing only the real elements in the input bins, and preserving the same order they appear in the input bins. The algorithm outputs the true output length to some dedicated location on the output tape.

In short, **RandBin** is a pre-processing step that takes the input array and converts it into bin format, and **CompactBin** takes the bin representation, and performs the actual compaction. The informal intuition is as follows. From earlier work [CCMS19], we know how to construct an efficient DO stable compaction algorithm for Hamming distance. However, in our case, we have two inputs I and I' that have *edit* distance 1. The difficulty with edit distance is when I is obtained by inserting an extra element into I' at position j , the two inputs I and I' will differ in *every* position after j . Our idea is to leverage **RandBin** to “probabilistically localize” this difference caused by a single insertion operation. In particular, if some bin representation occurs for input I with some probability p , we want that under the neighboring input I' , with probability close to p , we should encounter a similar bin representation where the difference is localized to only one or two bins. If we can accomplish this, then hopefully we can adapt ideas that worked for Hamming distance [CCMS19] to compact the resulting bin representation.

Below, we will define an appropriate neighboring notion $\sim_B$ on the bin representation. We want to show that the **RandBin** pre-processing step satisfies NPDO w.r.t. the input relation $\sim_E$ and output relation $\sim_B$ for the bin representation. Further, we want to show that **CompactBin** satisfies NPDO w.r.t. $\sim_B$ and $\sim_H$. Then, the composed algorithm should be NPDO by our composition theorem.

Neighboring relation for bin representation $\sim_B$. Specifically, the neighboring relation $\sim_B$ is defined as below. Two lists of bins $(\text{Bin}_i^{(Z)} : i \in [B])$ and $(\text{Bin}_i'^{(Z')} : i \in [B'])$ are said to be neighboring, iff the following all hold:

- they have compatible dimensions, i.e., $B = B'$ and $Z = Z'$;
- After removing all fillers and concatenating the real elements in the list of bins, the resulting outcomes have edit distance at most one;
- There are at most *two* bins that have different bin loads (defined to be the number of real elements in the bin), further, for both of them, the difference in load is at most one.

5.3 RandBin Algorithm

We now describe the **RandBin** algorithm, which preprocesses the input array into a bin representation.

RandBin $^{\epsilon, \delta}(I)$: // Let $\epsilon_1 = \epsilon_2 = \epsilon_3 = \frac{\epsilon}{3}$, and $\delta_1 = \delta_2 = \delta_3 = \frac{\delta}{3}$.

- Sample $G \xleftarrow{\$} \mathcal{G}(\epsilon_1, \delta_1)$. Let $L = |I| + G$ be the noisy input length. Let $s = \frac{15}{\epsilon} \log^2 L \log \frac{1}{\delta}$ be an upper bound on the support of $\mathcal{G}(\epsilon_2, \delta_2)$ and also the additive error of **PrefixSum** $^{\epsilon_3, \delta_3}$ on at most L integers. Let the maximum bin load be $Z = 2s$ and $B = \lceil \frac{2L}{Z} \rceil + 1$ be the number of bins.
- For $i = 1$ to B , let $\rho_i \xleftarrow{\$} s + \mathcal{G}(\epsilon_2, \delta_2) \in [\frac{Z}{2} \dots Z]$. Let $\boldsymbol{\rho} := (\rho_1, \rho_2, \dots, \rho_B)$.
- Let $\text{cnt} := \text{PrefixSum}^{\epsilon_3, \delta_3}(\boldsymbol{\rho}) \in \mathbb{Z}^B$. Let **Buf** := $\emptyset$ be a working buffer.
- For $i = 1$ to B :

- fetch the unvisited elements in the input array up to index $\text{cnt}[i] + s$ and add them to Buf ; mark them as visited.
- if the current length of Buf is less than Z , append enough fillers such that its length is at least Z ;
- perform stable oblivious sorting on Buf such that the first Z positions contains only the real elements coming from the first $\sum_{j \leq i} \rho_j$ positions in the input and fillers; all remaining elements are moved to the end of Buf .
- pop the first Z elements of Buf to Bin_i .
- perform stable oblivious sorting on Buf to move all the fillers to the end, if necessary, truncate Buf such that its length is at most $2s$.
- Output the bin representation $(\text{Bin}_i^{(Z)} : i \in [B])$, and store the parameters B and Z in some dedicated location on the output tape.

^aWe may assume that any location in the input array beyond the original length $|I|$ is occupied by a filler.

Roughly speaking, the **RandBin** algorithm generates a list of random counts $\boldsymbol{\rho} := (\rho_1, \dots, \rho_B)$. Then, all real elements contained in the first ρ_1 positions of the input are moved into Bin_1 , all real elements contained in the next ρ_2 positions of the input are moved into Bin_2 , and so on. To guarantee differential obliviousness, the algorithm cannot directly reveal the vector $\boldsymbol{\rho}$ — instead, it reveals only the noisy prefix sum of $\boldsymbol{\rho}$. Specifically, we apply an (ϵ_3, δ_3) -differentially private prefix sum algorithm to the vector y , i.e., $\text{cnt} := \text{PrefixSum}^{\epsilon_3, \delta_3}(Y)$. In other words, $\text{cnt}[i]$ stores a noisy version of $\sum_{j \leq i} \rho_j$, and it is guaranteed that the estimation error is at most s . Now, in each step i of the algorithm, we want to populate Bin_i . To do so, we simply fetch the next batch of elements in the input array up to position $\text{cnt}[i]$ into a poly-logarithmically sized working buffer Buf . Buf also contains previously fetched elements that have not been placed into any bin yet. We can now obviously sort Buf to create the next Bin_i . At the end of each step i , it is guaranteed that there are at most $2s$ real elements remaining in Buf . Therefore, we can obviously sort Buf and compact its length to $2s$. This makes sure that Buf is always poly-logarithmic in size. Finally, to make the algorithm secure, we also need to mask the true input length, and we can accomplish this by adding a truncated geometric random noise to the true length, and revealing only the noisy length. Note that the number of bins B is a random variable that depends only on the noisy input length.

Theorem 5.1. *The **RandBin** algorithm always outputs the correct bin representation: For all $i \in [B]$, all real elements from $I \left[\sum_{j < i} \rho_j + 1 \right]$ to $I \left[\sum_{j \leq i} \rho_j \right]$ are moved to Bin_i . Also, all real elements in the input array are moved to the bins.*

Theorem 5.2. *Assuming $|I| \geq \frac{c}{\epsilon} \log \frac{1}{\delta}$ for any fix constant c , **RandBin** has a worst-case runtime of $O \left(|I| \left(\log \log |I| + \log \frac{1}{\epsilon} + \log \log \frac{1}{\delta} \right) \right)$.*

Theorem 5.3. ***RandBin** is (ϵ, δ) -NPDO w.r.t. the input neighboring notion $\sim_E$ and the output neighboring notion $\sim_B$.*

The above theorems' proofs are deferred to Appendix [B.3](#).

5.4 CompactBin Algorithm

We now describe the **CompactBin** algorithm which takes in a bin representation, outputs a compacted array, and writes the true length of the output to some dedicated location on the output tape.

$\text{CompactBin}^{\epsilon, \delta} \left((\text{Bin}_i^{(Z)} : i \in [B]) \right)$:

// Let $\epsilon_1 = \frac{\epsilon}{2}$, $\delta_1 = \frac{\delta}{2(1+\epsilon_1)}$.

- Let $s = \frac{4}{\epsilon_1} \log^2 B \log \frac{2}{\delta_1}$, be an upper bound of the additive error of (ϵ_1, δ_1) -differentially private prefix sums on at most B integers.
- Let $R := (R_i : i \in [B])$, where R_i is the number of real elements in Bin_i . Call $\text{cnt} := \text{PrefixSum}^{\epsilon_1, \delta_1}(R)$.
- Let Buf and the output array be initially empty. For $i = 1$ to B :
 - Read the i -th bin and append it to the end of Buf .
 - Perform stable oblivious sorting on Buf such that all real elements are moved to the front.
 - Let L be the current length of the output array. Remove an appropriate number of elements from the beginning of Buf and append them to the output array, such that the output array has length exactly $\max(\text{cnt}[i] - s, L)$.
 - Truncate Buf if necessary such that its length is at most $2s$.
- Append Buf to the end of the output array. Write the true output length $\sum_{i \in [B]} R_i$ to some dedicated location on the output tape.

To gain some intuition, basically in each step i , the **CompactBin** algorithm reads the next bin i , and tries to copy the real elements in bin i to the end of the output array. To achieve differential obliviousness, the algorithm cannot reveal the true number of real elements inside each bin. Therefore, it calls a differentially private prefix sum mechanism to compute an array $\text{cnt}[1 : B]$ where $\text{cnt}[i]$ is an estimate of the number of real elements contained in the first i bins. The prefix sum algorithm guarantees that the estimation error is upper bounded by s . Therefore, at the end of the i -th step, the algorithm should have written exactly $\text{cnt}[i] - s$ number of real elements to the output array. To accomplish this, the algorithm makes use of a temporary working buffer Buf that is used to store the real elements that have been fetched from the input bins but have not been appended to the output array. It guarantees that at the end of each step, there are at most $2s$ real elements leftover in Buf .

Theorem 5.4. *With probability 1, the output of **CompactBin** includes all the real elements from the B input bins with their order preserved and the filler elements in the output array only appear after the last real element.*

Theorem 5.5. $\text{CompactBin}^{\epsilon, \delta}$ is (ϵ, δ) -NPDO w.r.t. the input neighboring relation $\sim_B$ and output neighboring relation $\sim_E$.

Theorem 5.6. **CompactBin** has a worst-case runtime of $O(B(Z + s) \log(Z + s))$.

The above theorems' proofs are deferred to Appendix [B.3](#).

From the **RandBin** algorithm, $BZ = O(|I|)$, $Z = \Theta(s)$, and $s = O(\frac{1}{\epsilon} \log^2 |I| \log \frac{1}{\delta})$. By Theorem [5.2](#) and Theorem [5.6](#), the following corollary holds:

Corollary 5.7. *Assuming $|I| \geq \frac{c}{\epsilon} \log \frac{1}{\delta}$ for any fix constant c , then the full compaction algorithm $\text{Compact}^{\epsilon, \delta} := \text{CompactBin}^{\epsilon/2, \delta/2} \circ \text{RandBin}^{\epsilon/2, \delta/2}$ has a worst-case runtime of $O(|I| (\log \log |I| + \log \frac{1}{\epsilon} + \log \log \frac{1}{\delta}))$.*

6 Application: Optimal Privacy Amplification in the Differentially Oblivious Shuffle Model

In this section, we use our composition framework to prove a privacy amplification theorem for the differentially oblivious shuffle (DO-shuffle) model.

Background. Consider a distributed setting with n clients and an untrusted server who wants to learn some statistics of the clients' private data. To achieve this, the differential privacy literature proposed two models for achieving this. In the *local model*, each client adds some noise to its own data by running an ϵ_0 -locally differentially private (LDP) mechanism, and sends the noisy result to the server. The server then computes the desired statistics using each client's noisy input. In the local model, the server's view satisfies ϵ_0 -DP.

Some more recent works [CSU⁺19, Che21, G GK⁺21, GKMP21] considered a new model called the *shuffle model*. In this model, we assume that a trusted shuffler can shuffle all of the clients' messages, and the server only sees the permuted messages (without learning the permutation). Interestingly, earlier works [CSU⁺19, Che21, G GK⁺21, GKMP21] showed that the shuffle model can *amplify* privacy. In particular, suppose each client still runs an ϵ_0 -locally differentially private (LDP) mechanism before sending the noisy outcomes to the shuffler, then the server's view would satisfy (ϵ, δ) -DP where ϵ can be much smaller than ϵ_0 . Notably, the recent work of Feldman et al. [FMT21] proved optimal parameters for privacy amplification in a perfectly secure shuffle model, that is, it can achieve (ϵ, δ) -DP with any $\delta > 0$ and $\epsilon = O\left((1 - e^{-\epsilon_0})e^{\epsilon_0/2}\sqrt{\frac{\log(1/\delta)}{n}}\right)$. In this section, our

goal is to show that the *perfectly secure* shuffle in privacy amplification can be replaced with a much weaker, (ϵ, δ) -*differentially oblivious* shuffle, without degrading the amplification guarantees (except for extra ϵ and δ additive factors that arise from the differentially oblivious shuffler itself).

To benefit from the shuffle model, we need to realize the shuffle either using trusted hardware or through a cryptographic protocol. Some recent works [GKLX22, BHMS21, ALU18] showed that it may be asymptotically more efficient to realize a relaxed shuffler that satisfies differential obliviousness than a perfectly secure shuffler. Therefore, a natural question is whether we can also enjoy the same degree of privacy amplification with a differentially oblivious shuffler rather than a perfectly secure shuffler. We explore this question in the remainder of this section.

6.1 Definitions

Suppose that the server and the clients jointly execute a protocol to realize a shuffler. The syntax of a shuffle protocol is defined below.

Definition 6.1 (Syntax of a shuffle protocol). A protocol between a server and n clients each with some input from $\mathcal{X}$ is said to be a shuffle protocol, iff under an *honest execution*, the server outputs a random permutation of the clients' inputs.

Before defining security, we need to define the adversary's capabilities and the view of the adversary. We assume that an adversary $\mathcal{A}$ may control up to t clients as well as the server, we define the random variable $\text{View}^{\mathcal{A}}(\mathbf{x}_H)$ to mean the view of the adversary during an execution where the honest clients' inputs are $\mathbf{x}_H \in \mathcal{X}^{n-t}$. The view of the adversary $\mathcal{A}$ should include whatever the adversary can observe during the execution. Specifically, the view include the server's output, all messages sent and received by the corrupted clients and the server. Further, the view may include any additional information the adversary can observe. For example, if the adversary can observe honest-to-honest communication (e.g., a network adversary), then, the view should also

include the honest-to-honest communication. For a protocol secure in the *semi-honest* model, we assume that the corrupt players will honestly follow the protocol. For the protocol secure in the *malicious* model, we assume that the corrupt players can send arbitrary messages and the adversary $\mathcal{A}$ controls the messages sent by corrupt players.

Remark 6.2. Different DO-shuffle protocols may provide security guarantees under differing adversarial power. For example, of Gordon et al. [GKLX22] assumes a semi-honest adversary cannot observe honest-to-honest communication, whereas Bünz et al. [BHMS21] assumes a malicious adversary who can observe the entire network communication. Our privacy amplification theorem does not care about the exact modeling choice made by the underlying DO-shuffle protocol, and the composed DO-shuffle-model mechanism essentially inherits the same assumptions as the underlying DO-shuffle.

Next, we define the notion of differential obliviousness for a shuffle protocol. We first need to define what neighboring means.

Neighboring by swapping. Given some set $\mathcal{D}$ and two vectors $\mathbf{y}, \mathbf{y}' \in \mathcal{D}^m$, we say that $\mathbf{y} \sim_S \mathbf{y}'$, iff either $\mathbf{y} = \mathbf{y}'$, or $\mathbf{y}'$ can be obtained from $\mathbf{y}$ by swapping the values of two coordinates.

Definition 6.3 (Differential obliviousness of a shuffle protocol). A shuffle protocol is said to satisfy *statistical* (ϵ, δ) -differential obliviousness in the presence of $t \leq n$ corruptions, iff the following holds: for any adversary $\mathcal{A}$ controlling the server and at most t clients, for any two honest input configurations $\mathbf{y}_H, \mathbf{y}'_H \in \mathcal{Y}^{n-t}$ such that $\mathbf{y}_H \sim_S \mathbf{y}'_H$, for any subset $S \subseteq \mathcal{V}$ where $\mathcal{V}$ denotes the view space, it holds that

$$\Pr [\text{View}^{\mathcal{A}}(\mathbf{y}_H) \in S] \leq e^\epsilon \cdot \Pr [\text{View}^{\mathcal{A}}(\mathbf{y}'_H) \in S] + \delta$$

If we set $\epsilon = 0$ and $\delta = 0$, the above notion becomes equivalent to the security of a perfectly secure shuffle.

We next define a computational variant of the DO notion, since some known DO shuffler instantiations enjoy computational security [GKLX22, BHMS21, ALU18].

Definition 6.4 (Computational DO for a shuffle protocol). A shuffle protocol Φ is said to satisfy computational (ϵ, δ) -differential obliviousness in the presence of $t \leq n$ corruptions, iff for any non-uniform probabilistic polynomial-time (p.p.t.) adversary $\mathcal{A}$ controlling the server and at most t clients, for any two neighboring honest input configurations $\mathbf{y}_H \sim_S \mathbf{y}'_H$, it holds that

$$\Pr [\text{Expt}^{\mathcal{A}}(1^\lambda, \mathbf{y}_H) = 1] \leq e^\epsilon \cdot \Pr [\text{Expt}^{\mathcal{A}}(1^\lambda, \mathbf{y}'_H) = 1] + \delta$$

where $\text{Expt}^{\mathcal{A}}(1^\lambda, \mathbf{y})$ is the randomized experiment where we execute the protocol using security parameter λ and interacting with the adversary $\mathcal{A}$, and at the end we output whatever $\mathcal{A}$ outputs.

Next, we define the notion of a locally differentially private (LDP) mechanism. The main privacy amplification theorem we want to prove in this section asks the following question: if each client computes its message by applying an ϵ_0 -LDP mechanism to its private input, and then a DO shuffler shuffles all clients' messages before revealing the shuffled result to the server, can we prove that the server's view satisfies (ϵ, δ) -DP where ϵ is much smaller than ϵ_0 ?

Definition 6.5 (ϵ_0 -LDP mechanism). $\mathcal{R} : \mathcal{X} \rightarrow \mathcal{Y}$ is an ϵ_0 -LDP mechanism if for any $x, x' \in \mathcal{X}$ and any subset $S \subseteq \mathcal{Y}$, $\Pr[\mathcal{R}(x) \in S] \leq e^{\epsilon_0} \Pr[\mathcal{R}(x') \in S]$.

6.2 Privacy Amplification in the DO-Shuffle Model

We formally restate Theorem 1.4 when the DO-shuffler Φ satisfies statistical DO as the following theorem. In Appendix A.1, we will extend our composition framework to support the case when the DO-shuffler satisfies computational differential obliviousness (Definition 6.4).

Theorem 6.6. *Let $\mathcal{R} : \mathcal{X} \rightarrow \mathcal{Y}$ be an ϵ_0 -LDP mechanism to be run by each client over its private input. Suppose $\mathcal{A}$ is an adversary controlling the server and at most t clients. Let Φ be a statistical (ϵ_1, δ_1) -DO shuffling protocol. Define the random experiment $\text{Expt}^{\mathcal{A}}(x_1, \dots, x_n)$ as the following where each $x_i \in \mathcal{X}$ denotes the private input of client $i \in [n]$:*

1. *Each honest client i treats the output of $\mathcal{R}(x_i)$ as the input in the next step;*
2. *Execute the DO-shuffling protocol Φ with the presence of the adversary $\mathcal{A}$ and let $\mathcal{A}$ also observe the outcome of the shuffling.*
3. *Output whatever $\mathcal{A}$ outputs.*

Let $\mathbf{x} = (x_1, \dots, x_n) \in \mathcal{X}^n$ and $\mathbf{x}' = (x'_1, \dots, x'_n) \in \mathcal{X}^n$ be any two neighboring input configurations that differ in at most one client's input. For any $\delta > 0$ that $\epsilon_0 \leq \log\left(\frac{n-t}{16 \log(2/\delta)}\right)$, we have that

$$\Pr[\text{Expt}^{\mathcal{A}}(\mathbf{x}) = 1] \leq e^{\epsilon + \epsilon_1} \cdot \Pr[\text{Expt}^{\mathcal{A}}(\mathbf{x}') = 1] + \delta + \delta_1$$

for $\epsilon = O\left((1 - e^{-\epsilon_0})e^{\epsilon_0/2} \sqrt{\frac{\log(1/\delta)}{n-t}}\right)$.

Since we want to use our composition framework to prove optimal privacy amplification in the DO-shuffle model, we can define the first and second mechanism M_1 and M_2 as follows:

- The first mechanism $M_1 : \mathcal{X}^n \rightarrow \mathcal{Y}^n$ is where the n clients each apply the ϵ_0 -LDP mechanism $\mathcal{R} : \mathcal{X} \rightarrow \mathcal{Y}$ to their private data, respectively. The mechanism generates no view observable by the adversary, and moreover, its output is the concatenation of all clients' outputs.
- The second mechanism $M_2 : \mathcal{Y}^n \rightarrow \mathcal{Y}^n$ is the DO-shuffler itself. Here, the view of the adversary is its view in the DO-shuffle protocol, and the output is the shuffled outcome. In the main body, we shall first assume that M_2 satisfies *statistical* differential obliviousness (Definition 6.3).

It is easy to see that the adversary has the same view in the composed mechanism $M_2 \circ M_1$ and in the random experiment described in Theorem 6.6. So we only need to prove that $M_2 \circ M_1$ is $(\epsilon + \epsilon_1, \delta + \delta_1)$ -DO. The crux is to show that M_1 satisfies (ϵ, δ) -NPDO when at most t clients are corrupted, as more formally stated in the following lemma:

Lemma 6.7. *Suppose $\epsilon_0 \leq \log\left(\frac{n-t}{16 \log(2/\delta)}\right)$. The above mechanism M_1 satisfies (ϵ, δ) -NPDO w.r.t. the input relation $\sim_H$ (i.e., two vectors are neighboring if they have the same length and differ in at most one position) and the output relation $\sim_S$ (i.e., neighboring by swapping).*

We can directly apply our composition theorem (Theorem 3.6) and Lemma 6.7 to get the desired result.

6.3 Proof Sketch of Lemma 6.7

The full proof is deferred to Appendix B.4.

We want to prove that the mechanism M_1 as defined above satisfies (ϵ, δ) -NPDO. Instead of directly proving this statement, we first define a related mechanism M'_1 which is otherwise the same as M_1 , except that we augment the adversary's view with some extra information. Specifically, recall that the mechanism M_1 does not generate any view observable by the adversary. Now, in M'_1 , we assume the adversary can observe some additional auxiliary information which we shall define shortly. We shall prove that M'_1 satisfies (ϵ, δ) -NPDO by proving that M'_1 satisfies (ϵ, δ) -DO and $(0, 0)$ -NP. This immediately implies that the original M_1 satisfies (ϵ, δ) -NPDO as well.

Equivalent view of M_1 . Without loss of generality, we may assume that client 1 is the client whose input differs in the two neighboring input configurations $\mathbf{x} = (x_1, x_2, \dots, x_n)$ and $\mathbf{x}' = (x'_1, x_2, \dots, x_n)$. Feldman et al. [FMT21] showed that M_1 can be equivalently viewed as the following randomized experiment. Recall that in M_1 , every client i runs a ϵ_0 -DP local randomizer $\mathcal{R}$ on its private data x_i . An equivalent way to view $\mathcal{R}(x_i)$ where $i \neq 1$ is the following:

$$\mathcal{R}(x_i) = \begin{cases} \mathcal{R}(x_1) & w.p. \frac{1}{2e^{\epsilon_0}}, \\ \mathcal{R}(x'_1) & w.p. \frac{1}{2e^{\epsilon_0}}, \\ \mathcal{Q}_{\{x_1, x'_1\}}(x_i), & w.p. 1 - \frac{1}{e^{\epsilon_0}}. \end{cases}$$

With this equivalent view, we can imagine that every client besides client 1 first flips a coin to decide whether it wants to be a clone of $\mathcal{R}(x_1)$ or $\mathcal{R}(x'_1)$. If so, it chooses to be a clone of $\mathcal{R}(x_1)$ or $\mathcal{R}(x'_1)$ with equal probability. If not, it samples from a suitable leftover distribution $\mathcal{Q}_{\{x_1, x'_1\}}(x_i)$. The client 1 is always a clone of itself, i.e., of $\mathcal{R}(x_1)$ in the first world, and of $\mathcal{R}(x'_1)$ in the second world.

Definition of M'_1 . In M_1 , the adversary's view is empty. We define an augmented mechanism that is almost identical to M_1 except that we give the adversary some extra information, i.e., we augment the adversary's view. In particular, we will give the adversary the extra information (U, V, T) , where

- V contains the indices of the clones, i.e., the clients (including 1) who clone either $\mathcal{R}(x_1)$ or $\mathcal{R}(x'_1)$,
- U is the multi-set of the messages (i.e., outcomes of the local randomizer) sent by the clones in V ,
- T contains the messages sent by all the non-clones (i.e., clients not in V) along with their indices.

It suffices to prove that M'_1 satisfies (ϵ, δ) -NPDO. To do so, we will prove that M'_1 satisfies (ϵ, δ) -DO and $(0, 0)$ -NP. It turns out that the fact that M'_1 is (ϵ, δ) -DO is already implied by the key lemma of Feldman et al. [FMT21], as stated in Lemma B.1 of Appendix B.4. The intuition is that if we can secretly permute all clients' messages, then client 1 is well-hidden among the clones — in particular, observe that even when we directly tell the adversary all the messages sent by non-clones, client 1 is still well-hidden.

Therefore, the crux of our proof is to show that M'_1 satisfies $(0, 0)$ -NP.

Proving that M'_1 satisfies $(0, 0)$ -NP. This is the most technical part of our proof. To prove that M'_1 satisfies $(0, 0)$ -NP, we will be using Lemma 4.1. Specifically, we will consider the probability space conditioned on the adversary’s view (U, V, T) , and we want to find an $(0, 0)$ -matching in a bipartite graph where the vertices on both sides represent a point in the output space of M'_1 . One can think of the vertices on the left as producers, and the amount of commodity produced by each producer is exactly the *conditional* probability of this output. The vertices on the right are consumers, and each consumer’s maximum capacity is also the *conditional* probability of this output. Now, each producer can only route to neighboring consumers, and we want to make sure that all but $1 - \delta$ amount of the commodity is routed to some consumer. In our proof, it is equivalent to just ignore the term T , and imagine that the output space is merely all permutations of the messages in U . This part of the proof is somewhat more involved, and we defer the actual proof to Appendix B.4.

Acknowledgments

This work is in part supported by a grant from ONR, a gift from Cisco, NSF awards under grant numbers CIF-1705007, 2128519 and 2044679, and a Packard Fellowship. The work is also supported by DARPA SIEVE research program. T-H. Hubert Chan was partially supported by the Hong Kong RGC under the grants 17201220, 17202121 and 17203122.

References

- [Abe99] Masayuki Abe. Mix-networks on permutation networks. In Kwok-Yan Lam, Eiji Okamoto, and Chaoping Xing, editors, *ASIACRYPT’99*, volume 1716 of *LNCS*, pages 258–273. Springer, Heidelberg, November 1999.
- [ABG⁺11] Ron Aharoni, Eli Berger, Agelos Georgakopoulos, Amitai Perlstein, and Philipp Sprüssel. The max-flow min-cut theorem for countable networks. *Journal of Combinatorial Theory, Series B*, 101(1):1–17, 2011.
- [AFKL19] Peyman Afshani, Casper Benjamin Freksen, Lior Kamma, and Kasper Green Larsen. Lower bounds for multiplication via network coding. In Christel Baier, Ioannis Chatzigiannakis, Paola Flocchini, and Stefano Leonardi, editors, *ICALP 2019*, volume 132 of *LIPIcs*, pages 10:1–10:12. Schloss Dagstuhl, July 2019.
- [AK14] Arvind Arasu and Raghav Kaushik. Oblivious query processing. In *Proc. 17th International Conference on Database Theory (ICDT), Athens, Greece, March 24-28, 2014*, pages 26–37. OpenProceedings.org, 2014.
- [ALU18] Megumi Ando, Anna Lysyanskaya, and Eli Upfal. Practical and provably secure onion routing. In Ioannis Chatzigiannakis, Christos Kaklamanis, Dániel Marx, and Donald Sannella, editors, *ICALP 2018*, volume 107 of *LIPIcs*, pages 144:1–144:14. Schloss Dagstuhl, July 2018.
- [APY20] Ittai Abraham, Benny Pinkas, and Avishay Yanai. Blinder - scalable, robust anonymous committed broadcast. In Jay Ligatti, Xinming Ou, Jonathan Katz, and Giovanni Vigna, editors, *ACM CCS 2020*, pages 1233–1252. ACM Press, November 2020.
- [BBG18] Borja Balle, Gilles Barthe, and Marco Gaboardi. Privacy amplification by subsampling: Tight analyses via couplings and divergences. *NeurIPS*, 2018.

- [BBGN19] Borja Balle, James Bell, Adrià Gascón, and Kobbi Nissim. The privacy blanket of the shuffle model. In Alexandra Boldyreva and Daniele Micciancio, editors, *CRYPTO 2019, Part II*, volume 11693 of *LNCS*, pages 638–667. Springer, Heidelberg, August 2019.
- [BEM⁺17] Andrea Bittau, Úlfar Erlingsson, Petros Maniatis, Ilya Mironov, Ananth Raghunathan, David Lie, Mitch Rudominer, Ushasree Kode, Julien Tinnes, and Bernhard Seefeld. Prochlo: Strong privacy for analytics in the crowd. In *SOSP*, page 441–459, New York, NY, USA, 2017. Association for Computing Machinery.
- [BG12] Stephanie Bayer and Jens Groth. Efficient zero-knowledge argument for correctness of a shuffle. In David Pointcheval and Thomas Johansson, editors, *EUROCRYPT 2012*, volume 7237 of *LNCS*, pages 263–280. Springer, Heidelberg, April 2012.
- [BHMS21] Benedikt Bünz, Yuncong Hu, Shin’ichiro Matsuo, and Elaine Shi. Non-interactive differentially anonymous router. Cryptology ePrint Archive, Report 2021/1242, 2021. <https://eprint.iacr.org/2021/1242>.
- [BKK⁺21] Dmytro Bogatov, Georgios Kellaris, George Kollios, Kobbi Nissim, and Adam O’Neill. epsolute: Efficiently querying databases while providing differential privacy. In Giovanni Vigna and Elaine Shi, editors, *ACM CCS 2021*, pages 2262–2276. ACM Press, November 2021.
- [BNZ19] Amos Beimel, Kobbi Nissim, and Mohammad Zaheri. Exploring differential obliviousness. In *Approx/Random*, volume 145 of *LIPICs*, pages 65:1–65:20, 2019.
- [CBC⁺18] Natacha Crooks, Matthew Burke, Ethan Cecchetti, Sitar Harel, Rachit Agarwal, and Lorenzo Alvisi. Obladi: Oblivious serializable transactions in the cloud. In *13th USENIX Symposium on Operating Systems Design and Implementation, OSDI 2018, Carlsbad, CA, USA, October 8-10, 2018*, pages 727–743. USENIX Association, 2018.
- [CBM15] Henry Corrigan-Gibbs, Dan Boneh, and David Mazières. Riposte: An anonymous messaging system handling millions of users. In *2015 IEEE Symposium on Security and Privacy*, pages 321–338. IEEE Computer Society Press, May 2015.
- [CCMS19] T.-H. Hubert Chan, Kai-Min Chung, Bruce M. Maggs, and Elaine Shi. Foundations of differentially oblivious algorithms. In Timothy M. Chan, editor, *30th SODA*, pages 2448–2467. ACM-SIAM, January 2019.
- [CF10] Henry Corrigan-Gibbs and Bryan Ford. Dissent: accountable anonymous group messaging. In Ehab Al-Shaer, Angelos D. Keromytis, and Vitaly Shmatikov, editors, *ACM CCS 2010*, pages 340–350. ACM Press, October 2010.
- [CGLS18] T.-H. Hubert Chan, Yue Guo, Wei-Kai Lin, and Elaine Shi. Cache-oblivious and data-oblivious sorting and applications. In Artur Czumaj, editor, *29th SODA*, pages 2201–2220. ACM-SIAM, January 2018.
- [CGPR15] David Cash, Paul Grubbs, Jason Perry, and Thomas Ristenpart. Leakage-abuse attacks against searchable encryption. In Indrajit Ray, Ninghui Li, and Christopher Kruegel, editors, *ACM CCS 2015*, pages 668–679. ACM Press, October 2015.
- [Cha81] David L. Chaum. Untraceable electronic mail, return addresses, and digital pseudonyms. *Communications of the ACM*, 24(2):84–90, February 1981.

- [Cha88] David Chaum. The dining cryptographers problem: Unconditional sender and recipient untraceability. *Journal of Cryptology*, 1(1):65–75, January 1988.
- [Che21] Albert Cheu. Differential privacy in the shuffle model: A survey of separations. *arXiv preprint arXiv:2107.11839*, 2021.
- [CL05] Jan Camenisch and Anna Lysyanskaya. A formal treatment of onion routing. In Victor Shoup, editor, *CRYPTO 2005*, volume 3621 of *LNCS*, pages 169–187. Springer, Heidelberg, August 2005.
- [Con] Graeme Connell. Technology deep dive: Building a faster oram layer for enclaves. <https://signal.org/blog/building-faster-oram/>.
- [CSS10] T.-H. Hubert Chan, Elaine Shi, and Dawn Song. Private and continual release of statistics. In Samson Abramsky, Cyril Gavoille, Claude Kirchner, Friedhelm Meyer auf der Heide, and Paul G. Spirakis, editors, *ICALP 2010, Part II*, volume 6199 of *LNCS*, pages 405–417. Springer, Heidelberg, July 2010.
- [CSU⁺19] Albert Cheu, Adam D. Smith, Jonathan Ullman, David Zeber, and Maxim Zhilyaev. Distributed differential privacy via shuffling. In Yuval Ishai and Vincent Rijmen, editors, *EUROCRYPT 2019, Part I*, volume 11476 of *LNCS*, pages 375–403. Springer, Heidelberg, May 2019.
- [CZJ⁺17] Ethan Cecchetti, Fan Zhang, Yan Ji, Ahmed E. Kosba, Ari Juels, and Elaine Shi. Solidus: Confidential distributed ledger transactions via PVORM. In Bhavani M. Thuraisingham, David Evans, Tal Malkin, and Dongyan Xu, editors, *ACM CCS 2017*, pages 701–717. ACM Press, October / November 2017.
- [CZSC21] Shumo Chu, Danyang Zhuo, Elaine Shi, and T.-H. Hubert Chan. Differentially oblivious database joins: Overcoming the worst-case curse of fully oblivious algorithms. In Stefano Tessaro, editor, *ITC*, 2021.
- [DMNS06] Cynthia Dwork, Frank McSherry, Kobbi Nissim, and Adam Smith. Calibrating noise to sensitivity in private data analysis. In Shai Halevi and Tal Rabin, editors, *TCC 2006*, volume 3876 of *LNCS*, pages 265–284. Springer, Heidelberg, March 2006.
- [DNPR10] Cynthia Dwork, Moni Naor, Toniann Pitassi, and Guy N. Rothblum. Differential privacy under continual observation. In Leonard J. Schulman, editor, *42nd ACM STOC*, pages 715–724. ACM Press, June 2010.
- [DR14] Cynthia Dwork and Aaron Roth. The algorithmic foundations of differential privacy. *Found. Trends Theor. Comput. Sci.*, 9(3–4):211–407, 2014.
- [DS18] Jean Paul Degabriele and Martijn Stam. Untagging Tor: A formal treatment of onion encryption. In Jesper Buus Nielsen and Vincent Rijmen, editors, *EUROCRYPT 2018, Part III*, volume 10822 of *LNCS*, pages 259–293. Springer, Heidelberg, April / May 2018.
- [EB21] Saba Eskandarian and Dan Boneh. Clarion: Anonymous communication from multiparty shuffling protocols. Cryptology ePrint Archive, Report 2021/1514, 2021. <https://eprint.iacr.org/2021/1514>.

- [EFM⁺19] Úlfar Erlingsson, Vitaly Feldman, Ilya Mironov, Ananth Raghunathan, Kunal Talwar, and Abhradeep Thakurta. Amplification by shuffling: From local to central differential privacy via anonymity. In Timothy M. Chan, editor, *30th SODA*, pages 2468–2479. ACM-SIAM, January 2019.
- [EZ19] Saba Eskandarian and Matei Zaharia. Oblidb: Oblivious query processing for secure databases. *Proc. VLDB Endow.*, 13(2):169–183, 2019.
- [FHLS19] Alireza Farhadi, MohammadTaghi Hajiaghayi, Kasper Green Larsen, and Elaine Shi. Lower bounds for external memory integer sorting via network coding. In Moses Charikar and Edith Cohen, editors, *51st ACM STOC*, pages 997–1008. ACM Press, June 2019.
- [FMT21] Vitaly Feldman, Audra McMillan, and Kunal Talwar. Hiding among the clones: A simple and nearly optimal analysis of privacy amplification by shuffling. In *FOCS*, 2021.
- [GDD⁺21] Antonious M. Girgis, Deepesh Data, Suhas Diggavi, Peter Kairouz, and Ananda Theertha Suresh. Shuffled model of federated learning: Privacy, accuracy and communication trade-offs. *IEEE Journal on Selected Areas in Information Theory*, 2(1):464–478, 2021.
- [GGK⁺21] Badih Ghazi, Noah Golowich, Ravi Kumar, Rasmus Pagh, and Ameya Velingker. On the power of multiple anonymous messages: Frequency estimation and selection in the shuffle model of differential privacy. In Anne Canteaut and François-Xavier Standaert, editors, *EUROCRYPT 2021, Part III*, volume 12698 of *LNCS*, pages 463–488. Springer, Heidelberg, October 2021.
- [GIKM00] Yael Gertner, Yuval Ishai, Eyal Kushilevitz, and Tal Malkin. Protecting data privacy in private information retrieval schemes. *JCSS*, 2000.
- [GKK⁺12] S. Dov Gordon, Jonathan Katz, Vladimir Kolesnikov, Fernando Krell, Tal Malkin, Mariana Raykova, and Yevgeniy Vahlis. Secure two-party computation in sublinear (amortized) time. In Ting Yu, George Danezis, and Virgil D. Gligor, editors, *ACM CCS 2012*, pages 513–524. ACM Press, October 2012.
- [GKLX22] S. Dov Gordon, Jonathan Katz, Mingyu Liang, and Jiayu Xu. Spreading the privacy blanket: - differentially oblivious shuffling for differential privacy. In Giuseppe Ateniese and Daniele Venturi, editors, *ACNS 22*, volume 13269 of *LNCS*, pages 501–520. Springer, Heidelberg, June 2022.
- [GKMP21] Badih Ghazi, Ravi Kumar, Pasin Manurangsi, and Rasmus Pagh. Private counting from anonymous messages: Near-optimal accuracy with vanishing communication overhead, 2021.
- [GLMP19] Paul Grubbs, Marie-Sarah Lacharité, Brice Minaud, and Kenneth G. Paterson. Learning to reconstruct: Statistical learning theory and encrypted database attacks. In *2019 IEEE Symposium on Security and Privacy*, pages 1067–1083. IEEE Computer Society Press, May 2019.
- [GMN⁺16] Paul Grubbs, Richard McPherson, Muhammad Naveed, Thomas Ristenpart, and Vitaly Shmatikov. Breaking web applications built on top of encrypted data. In Edgar R.

- Weippl, Stefan Katzenbeisser, Christopher Kruegel, Andrew C. Myers, and Shai Halevi, editors, *ACM CCS 2016*, pages 1353–1364. ACM Press, October 2016.
- [GO96] Oded Goldreich and Rafail Ostrovsky. Software protection and simulation on oblivious rams. *J. ACM*, 43(3):431–473, may 1996.
 - [Gol87] Oded Goldreich. Towards a theory of software protection and simulation by oblivious RAMs. In Alfred Aho, editor, *19th ACM STOC*, pages 182–194. ACM Press, May 1987.
 - [Hal35] P Hall. On representatives of subsets. *Journal of the London Mathematical Society*, 1(1):26–30, 1935.
 - [HJ48] Marshall Hall Jr. Distinct representatives of subsets. *Bulletin of the American Mathematical Society*, 54(10):922–926, 1948.
 - [IKK12] Mohammad Saiful Islam, Mehmet Kuzu, and Murat Kantarcioglu. Access pattern disclosure on searchable encryption: Ramification, attack and mitigation. In *NDSS 2012*. The Internet Society, February 2012.
 - [JLN19] Riko Jacob, Kasper Green Larsen, and Jesper Buus Nielsen. Lower bounds for oblivious data structures. In Timothy M. Chan, editor, *30th SODA*, pages 2439–2447. ACM-SIAM, January 2019.
 - [KKNO16] Georgios Kellaris, George Kollios, Kobbi Nissim, and Adam O’Neill. Generic attacks on secure outsourced databases. In Edgar R. Weippl, Stefan Katzenbeisser, Christopher Kruegel, Andrew C. Myers, and Shai Halevi, editors, *ACM CCS 2016*, pages 1329–1340. ACM Press, October 2016.
 - [LHH⁺15] Chang Liu, Michael Hicks, Austin Harris, Mohit Tiwari, Martin Maas, and Elaine Shi. Ghost rider: A hardware-software system for memory trace oblivious computation. In *ASPLOS*, 2015.
 - [Loc22] Andreas Lochbihler. A mechanized proof of the max-flow min-cut theorem for countable networks with applications to probability theory. *Journal of Automated Reasoning*, pages 1–26, 2022.
 - [LSX19] Wei-Kai Lin, Elaine Shi, and Tiancheng Xie. Can we overcome the $n \log n$ barrier for oblivious sorting? In Timothy M. Chan, editor, *30th SODA*, pages 2419–2438. ACM-SIAM, January 2019.
 - [LWN⁺15] Chang Liu, Xiao Shaun Wang, Kartik Nayak, Yan Huang, and Elaine Shi. OblivM: A programming framework for secure computation. In *2015 IEEE Symposium on Security and Privacy*, pages 359–376. IEEE Computer Society Press, May 2015.
 - [MLS⁺13] Martin Maas, Eric Love, Emil Stefanov, Mohit Tiwari, Elaine Shi, Krste Asanovic, John Kubiawicz, and Dawn Song. PHANTOM: practical oblivious computation in a secure processor. In Ahmad-Reza Sadeghi, Virgil D. Gligor, and Moti Yung, editors, *ACM CCS 2013*, pages 311–324. ACM Press, November 2013.
 - [MPC⁺18] Pratyush Mishra, Rishabh Poddar, Jerry Chen, Alessandro Chiesa, and Raluca Ada Popa. Obliv: An efficient oblivious search index. In *2018 IEEE Symposium on Security and Privacy*, pages 279–296. IEEE Computer Society Press, May 2018.

- [MPRV09] Ilya Mironov, Omkant Pandey, Omer Reingold, and Salil P. Vadhan. Computational differential privacy. In Shai Halevi, editor, *CRYPTO 2009*, volume 5677 of *LNCS*, pages 126–142. Springer, Heidelberg, August 2009.
- [NKW15] Muhammad Naveed, Seny Kamara, and Charles V. Wright. Inference attacks on property-preserving encrypted databases. In Indrajit Ray, Ninghui Li, and Christopher Kruegel, editors, *ACM CCS 2015*, pages 644–655. ACM Press, October 2015.
- [OS97] Rafail Ostrovsky and Victor Shoup. Private information storage (extended abstract). In *29th ACM STOC*, pages 294–303. ACM Press, May 1997.
- [Rou20] Tim Roughgarden. *Beyond the Worst-Case Analysis of Algorithms*. Cambridge University Press, 2020.
- [RSG98] M.G. Reed, P.F. Syverson, and D.M. Goldschlag. Anonymous connections and onion routing. *IEEE Journal on Selected Areas in Communications*, 16(4):482–494, 1998.
- [RYF⁺13] Ling Ren, Xiangyao Yu, Christopher W. Fletcher, Marten van Dijk, and Srinivas Devadas. Design space exploration and optimization of path oblivious RAM in secure processors. In *ISCA*, pages 571–582, 2013.
- [SCSL11] Elaine Shi, T.-H. Hubert Chan, Emil Stefanov, and Mingfei Li. Oblivious RAM with $O((\log N)^3)$ worst-case cost. In Dong Hoon Lee and Xiaoyun Wang, editors, *ASIACRYPT 2011*, volume 7073 of *LNCS*, pages 197–214. Springer, Heidelberg, December 2011.
- [SS13] Emil Stefanov and Elaine Shi. ObliviStore: High performance oblivious cloud storage. In *2013 IEEE Symposium on Security and Privacy*, pages 253–267. IEEE Computer Society Press, May 2013.
- [SvS⁺13] Emil Stefanov, Marten van Dijk, Elaine Shi, Christopher W. Fletcher, Ling Ren, Xiangyao Yu, and Srinivas Devadas. Path ORAM: an extremely simple oblivious RAM protocol. In Ahmad-Reza Sadeghi, Virgil D. Gligor, and Moti Yung, editors, *ACM CCS 2013*, pages 299–310. ACM Press, November 2013.
- [SW21] Elaine Shi and Ke Wu. Non-interactive anonymous router. In Anne Canteaut and François-Xavier Standaert, editors, *EUROCRYPT 2021, Part III*, volume 12698 of *LNCS*, pages 489–520. Springer, Heidelberg, October 2021.
- [TGS22] Afonso Tinoco, Sixiang Gao, and Elaine Shi. EnigMap: Signal should use oblivious algorithms for private contact discovery. Cryptology ePrint Archive, Report 2022/1083, 2022. <https://eprint.iacr.org/2022/1083>.
- [WBK21] Yu-Xiang Wang, Borja Balle, and Shiva Kasiviswanathan. Subsampled rényi differential privacy and analytical moments accountant. *Journal of Privacy and Confidentiality*, 10(2), Feb. 2021.
- [WCS15] Xiao Wang, T.-H. Hubert Chan, and Elaine Shi. Circuit ORAM: On tightness of the Goldreich-Ostrovsky lower bound. In Indrajit Ray, Ninghui Li, and Christopher Kruegel, editors, *ACM CCS 2015*, pages 850–861. ACM Press, October 2015.

- [WST12] Peter Williams, Radu Sion, and Alin Tomescu. PrivateFS: a parallel oblivious file system. In Ting Yu, George Danezis, and Virgil D. Gligor, editors, *ACM CCS 2012*, pages 977–988. ACM Press, October 2012.
- [ZS22] Mingxun Zhou and Elaine Shi. The power of the differentially oblivious shuffle in distributed privacy mechanisms. Cryptology ePrint Archive, Report 2022/177, 2022. <https://eprint.iacr.org/2022/177>.
- [ZZR05] Li Zhuang, Feng Zhou, Ben Y. Zhao, and Antony Rowstron. Cashmere: Resilient anonymous routing. In *NSDI*, 2005.

A Computational Differential Obliviousness and Composition

In the main body, we focused on *statistical* notions of NPDO and DO. In this section, we discuss how to extend our composition theorem to computational notions of NPDO and DO.

As a first step (Appendix [A.1](#)), we consider the case where the second mechanism satisfies a standard, indistinguishability-based computational DO notion, and the first mechanism still satisfies statistical NPDO. We particularly care about this scenario, since in our DO-shuffle-model application, M_1 is the step where the clients each run their LDP algorithm, and indeed M_1 enjoys statistical security. However, known instantiations [\[ALU18, GKLX22, BHMS21\]](#) of the DO-shuffler, i.e., M_2 , satisfy only computational notions of security.

Next, in Appendix [A.2](#) we will consider the case where both the first and the second mechanisms enjoy computational security. To do so, we will have to first introduce a computational variant of our NPDO notion.

A.1 When M_2 is Computationally Secure

We first define an indistinguishability-based notion of DO, in the same flavor of the indistinguishability-based, computational DP notion of Mironov et al. [\[MPRV09\]](#).

Definition A.1 (CDO Mechanism). A mechanism Φ is said to satisfy computational (ϵ, δ) -differential obliviousness iff for any probabilistic polynomial-time (p.p.t.) adversary $\mathcal{A}$ for any two neighboring input configurations $x \sim_{\mathcal{X}} x'$, it holds that

$$\Pr \left[\text{Expt}_{\Phi}^{\mathcal{A}}(1^{\lambda}, x) = 1 \right] \leq e^{\epsilon} \cdot \Pr \left[\text{Expt}_{\Phi}^{\mathcal{A}}(1^{\lambda}, x') = 1 \right] + \delta + \text{negl}(\lambda)$$

where $\text{Expt}_{\Phi}^{\mathcal{A}}(1^{\lambda}, x)$ is the randomized experiment where we execute the mechanism Φ using security parameter λ and interacting with the adversary $\mathcal{A}$, and at the end we output whatever $\mathcal{A}$ outputs.

We extend our composition theorem to work even when the second mechanism M_2 satisfies the above CDO notion.

Corollary A.2 (Composition when M_2 satisfies CDO). *Suppose that an algorithm $M_1 : \mathcal{X} \rightarrow \mathcal{Y}$ satisfies (ϵ_1, δ_1) -NPDO w.r.t. $\sim_{\mathcal{X}}$ and $\sim_{\mathcal{Y}}$, and moreover, the view space $\mathcal{V}$ and output space $\mathcal{Y}$ of M_1 is finite or countably infinite. Suppose that $M_2 : \mathcal{Y} \rightarrow \mathcal{Z}$ satisfies (ϵ_2, δ_2) -CDO w.r.t. $\sim_{\mathcal{Y}}$. Then, the composed mechanism $M_2 \circ M_1 : \mathcal{X} \rightarrow \mathcal{Z}$ satisfies $(\epsilon_1 + \epsilon_2, \delta_1 + \delta_2)$ -CDO.*

Proof. Fix any neighboring input x, x' . By Lemma 4.1, there exists an (ϵ_1, δ_1) -matching $w : (\mathcal{V}_1 \times \mathcal{Y}) \times (\mathcal{V}_1 \times \mathcal{Y}) \rightarrow [0, 1]$ w.r.t natural neighbor notion $\sim$ in the product space $\mathcal{V}_1 \times \mathcal{Y}$: $(v_1, y) \sim (v'_1, y')$ when $v_1 = v'_1$ and $y \sim_{\mathcal{Y}} y'$. We want to prove that, for any p.p.t. adversary $\mathcal{A}$,

$$\Pr \left[\text{Expt}_{M_2 \circ M_1}^{\mathcal{A}}(1^\lambda, x) = 1 \right] \leq e^{\epsilon_1 + \epsilon_2} \cdot \Pr \left[\text{Expt}_{M_2 \circ M_1}^{\mathcal{A}}(1^\lambda, x') = 1 \right] + \delta_1 + \delta_2 + \text{negl}(\lambda).$$

For any p.p.t. adversary $\mathcal{A}$, we denote adversary $\mathcal{A}(v_1)$, whose behavior is the same as the adversary $\mathcal{A}$ when it is given the auxiliary information v_1 (which should be polynomial-length).

We have

$$\begin{aligned} & \Pr \left[\text{Expt}_{M_2 \circ M_1}^{\mathcal{A}}(1^\lambda, x) = 1 \right] \\ = & \sum_{(v_1, y) \in \mathcal{V}_1 \times \mathcal{Y}} \Pr[\text{Exec}^{M_1}(x) = (v_1, y)] \cdot \Pr \left[\text{Expt}_{M_2}^{\mathcal{A}(v_1)}(1^\lambda, y) = 1 \right] \\ & \text{(Use condition (b) and (d) of the matching)} \\ \leq & \sum_{(v_1, y) \in \mathcal{V}_1 \times \mathcal{Y}, y' \sim_{\mathcal{Y}} y} w((v_1, y), (v_1, y')) \cdot \Pr \left[\text{Expt}_{M_2}^{\mathcal{A}(v_1)}(1^\lambda, y) = 1 \right] + \delta_1 \\ & \text{(M}_2 \text{ is } (\epsilon_2, \delta_2)\text{-CDO)} \\ \leq & \sum_{(v_1, y) \in \mathcal{V}_1 \times \mathcal{Y}, y' \sim_{\mathcal{Y}} y} w((v_1, y), (v_1, y')) \cdot \left(e^{\epsilon_2} \Pr \left[\text{Expt}_{M_2}^{\mathcal{A}(v_1)}(1^\lambda, y') = 1 \right] + \delta_2 + \text{negl}(\lambda) \right) + \delta_1 \\ & \text{(Use condition (b) of the matching)} \\ \leq & \sum_{(v_1, y) \in \mathcal{V}_1 \times \mathcal{Y}, y' \sim_{\mathcal{Y}} y} w((v_1, y), (v_1, y')) \cdot \left(e^{\epsilon_2} \Pr \left[\text{Expt}_{M_2}^{\mathcal{A}(v_1)}(1^\lambda, y') = 1 \right] \right) + \text{negl}(\lambda) + \delta_2 + \delta_1 \\ & \text{(Use condition (c) of the matching)} \\ \leq & \sum_{(v_1, y') \in \mathcal{V}_1 \times \mathcal{Y}} e^{\epsilon_1} \Pr[\text{Exec}^{M_1}(x') = (v_1, y')] \cdot \left(e^{\epsilon_2} \Pr \left[\text{Expt}_{M_2}^{\mathcal{A}(v_1)}(1^\lambda, y') = 1 \right] \right) + \text{negl}(\lambda) + \delta_2 + \delta_1 \\ = & e^{\epsilon_1 + \epsilon_2} \Pr \left[\text{Expt}_{M_2 \circ M_1}^{\mathcal{A}}(1^\lambda, x') = 1 \right] + \text{negl}(\lambda) + \delta_2 + \delta_1. \end{aligned}$$

□

A.2 When Both M_1 and M_2 Are Computationally Secure

We now consider the case when both M_1 and M_2 are computational secure. To understand this case, we have to first define a computational variant of our NPDO notion. The most straightforward way is to use a simulation-based notion, i.e., an algorithm satisfies (ϵ, δ) -computational-NPDO iff the view (i.e., access patterns) and the output generated by the algorithm can be simulated by those generated by another algorithm which satisfies (ϵ, δ) -statistical-NPDO. We first define the computationally indistinguishability of mechanisms:

Definition A.3 (Computationally indistinguishable mechanisms). We say two mechanisms M and M' are computationally indistinguishable iff for any p.p.t. adversary $\mathcal{A}$,

$$\left| \Pr \left[x \leftarrow \mathcal{A}; (v, y) \leftarrow \text{Exec}^M(1^\lambda, x) : \mathcal{A}(v, y) = 1 \right] - \Pr \left[x \leftarrow \mathcal{A}; (v, y) \leftarrow \text{Exec}^{M'}(1^\lambda, x) : \mathcal{A}(v, y) = 1 \right] \right| \leq \text{negl}(\lambda).$$

Definition A.4 ((ϵ, δ) -computational-NPDO). We say that an algorithm $M(1^\lambda, \cdot) : \mathcal{X} \rightarrow \mathcal{Y}$ with view space $\mathcal{V}$ satisfies (ϵ, δ) -computational-NPDO, iff there exists an (ϵ, δ) -statistical-NPDO algorithm $M'(1^\lambda, \cdot) : \mathcal{X} \rightarrow \mathcal{Y}$ such that M and M' are computationally indistinguishable.

We now extend our composition framework to the case when both M_1 and M_2 satisfy computational notions of security.

Corollary A.5 (Composition when both M_1 and M_2 have computational security). *Suppose that an algorithm $M_1 : \mathcal{X} \rightarrow \mathcal{Y}$ satisfies (ϵ_1, δ_1) -computational-NPDO w.r.t. $\sim_{\mathcal{X}}$ and $\sim_{\mathcal{Y}}$, and moreover, the view space $\mathcal{V}$ and output space $\mathcal{Y}$ of M_1 is finite or countably infinite. The following two statements hold:*

1. *If $M_2 : \mathcal{Y} \rightarrow \mathcal{Z}$ satisfies (ϵ_2, δ_2) -CDO w.r.t. $\sim_{\mathcal{Y}}$, then, the composed mechanism $M_2 \circ M_1 : \mathcal{X} \rightarrow \mathcal{Z}$ satisfies $(\epsilon_1 + \epsilon_2, \delta_1 + \delta_2)$ -CDO.*
2. *If M_2 satisfies (ϵ_2, δ_2) -computational-NPDO w.r.t. $\sim_{\mathcal{Y}}$ and $\sim_{\mathcal{Z}}$, then $M_2 \circ M_1$ satisfies $(\epsilon_1 + \epsilon_2, \delta_1 + \delta_2)$ -computation-NPDO.*

Before we prove this corollary, we first provide the following lemma:

Lemma A.6. *Suppose mechanism M is computationally indistinguishable from mechanism M' and M' is (ϵ, δ) -CDO, then M is also (ϵ, δ) -CDO.*

Proof. Since M is computationally indistinguishable from mechanism M' , we have that for any p.p.t. $\mathcal{A}$ and any input x ,

$$\Pr \left[\text{Expt}_M^{\mathcal{A}}(1^\lambda, x) = 1 \right] \leq \Pr \left[\text{Expt}_{M'}^{\mathcal{A}}(1^\lambda, x) = 1 \right] + \text{negl}(\lambda).$$

Thus, for any neighboring input $x \sim_{\mathcal{X}} x'$

$$\begin{aligned} & \Pr \left[\text{Expt}_M^{\mathcal{A}}(1^\lambda, x) = 1 \right] \\ & \leq \Pr \left[\text{Expt}_{M'}^{\mathcal{A}}(1^\lambda, x) = 1 \right] + \text{negl}(\lambda) \\ & \leq e^\epsilon \left(\Pr \left[\text{Expt}_{M'}^{\mathcal{A}}(1^\lambda, x') = 1 \right] \right) + \delta + \text{negl}(\lambda) \\ & \leq e^\epsilon \left(\Pr \left[\text{Expt}_M^{\mathcal{A}}(1^\lambda, x') = 1 \right] + \text{negl}(\lambda) \right) + \delta + \text{negl}(\lambda) \\ & \leq e^\epsilon \Pr \left[\text{Expt}_M^{\mathcal{A}}(1^\lambda, x') = 1 \right] + \delta + \text{negl}(\lambda). \end{aligned}$$

□

Proof of Corollary A.5. When M_2 is CDO: By the definition of computational-NPDO, we know that there exists an (ϵ_1, δ_1) -statistical NPDO mechanism M'_1 that is computationally indistinguishable from M_1 . We now prove that $M_2 \circ M_1$ is computationally indistinguishable from $M_2 \circ M'_1$ with a standard argument for “post-processing” (where M_2 is the post-process step). Suppose there is a p.p.t. adversary $\mathcal{A}$ that can distinguish $M_2 \circ M_1$ and $M_2 \circ M'_1$ with a non-negligible advantage. Then, we simply build an p.p.t. adversary $\mathcal{B}$ that takes the challenge input x from $\mathcal{A}$, runs the execution of the first mechanism, $\text{Exec}^{M_1}(x)$, takes the output of the first mechanism as y and receives an execution from the environment ($\text{Exec}^{M_2}(1^\lambda, y)$ or $\text{Exec}^{M'_1}(1^\lambda, y)$, depending on the experiment). It then sends the view in $\text{Exec}^{M_1}(x)$ and the whole execution of the second mechanism to $\mathcal{A}$ and outputs what $\mathcal{A}$ outputs. $\mathcal{B}$ correctly simulates the input for $\mathcal{A}$, so $\mathcal{B}$ can distinguish M_1

and M'_1 with non-negligible advantage, which is a contradiction. Thus, we know that $M_2 \circ M_1$ is computationally indistinguishable from $M_2 \circ M'_1$, which is an $(\epsilon_1 + \epsilon_2, \delta_1 + \delta_2)$ -CDO mechanism by Corollary A.2. Finally, by Lemma A.6, $M_2 \circ M_1$ is $(\epsilon_1 + \epsilon_2, \delta_1 + \delta_2)$ -CDO.

When M_2 is computational-NPDO: By the definition of computational-NPDO, we know that there exists an (ϵ_1, δ_1) -statistical NPDO mechanism M'_1 that is computationally indistinguishable from M_1 . Also, there exists an (ϵ_2, δ_2) -statistical NPDO mechanism M'_2 that is computationally indistinguishable from M_2 . We only need to show that $M_2 \circ M_1$ is computationally indistinguishable from $M'_2 \circ M'_1$ and then by our main composition theorem for statistical NPDO mechanism 3.6, we can prove the statement.

Let a hybrid mechanism be $M'_2 \circ M_1$. By a similar argument as the first case in this proof, we have that $M'_2 \circ M_1$ is computationally indistinguishable from $M'_2 \circ M'_1$. We only need to prove that $M'_2 \circ M_1$ is computationally indistinguishable from $M_2 \circ M_1$. Suppose there is a p.p.t. $\mathcal{A}$ that distinguishes $M'_2 \circ M_1$ and $M_2 \circ M_1$. We can build a p.p.t. $\mathcal{B}$ that distinguishes M'_2 and M_2 with non-negligible advantage: $\mathcal{B}$ takes the challenge input x from $\mathcal{A}$, samples an execution from M_1 and then submits the output of M_1 as the challenge input. When $\mathcal{B}$ gets the view and the output from M_2 or M'_2 , it sends the view, the output and also the view of $\text{Exec}^{M_1}(x)$ to $\mathcal{A}$. Then $\mathcal{B}$ outputs what $\mathcal{A}$ outputs. Since $\mathcal{B}$ correctly simulates the input distribution for $\mathcal{A}$, $\mathcal{B}$ can distinguish M'_2 and M_2 with a non-negligible advantage, which is a contradiction. $\square$

B Deferred Proofs

B.1 Proof for basic NPDO properties

Proof of Theorem 3.5: (ϵ, δ) -NPDO implies (ϵ, δ) -DO. Let $\mathcal{V}$ be the sample space of View^M . Consider all those $S_v \subseteq \mathcal{V}$. We know $\mathcal{N}(S_v \times \mathcal{Y}) = S_v \times \mathcal{Y}$. Therefore,

$$\begin{aligned} \Pr[\text{View}^M(x) \in S_v] &= \Pr[\text{Exec}^M(x) \in S_v \times \mathcal{Y}] \\ &\leq e^\epsilon \Pr[\text{Exec}^M(x') \in \mathcal{N}(S_v \times \mathcal{Y})] + \delta = e^\epsilon \Pr[\text{View}^M(x') \in S_v] + \delta. \end{aligned}$$

$\square$

Proof of Lemma 3.7: Strongly NP and DO imply NPDO. Fix a pair of neighboring input x, x' . By the strongly neighbor-preserving definition, for any $y, y' \in \mathcal{Y}$ such that $\Pr[M(x) = y] > 0$ and $\Pr[M(x') = y'] > 0$, it must be that $y \sim_{\mathcal{Y}} y'$. For any subset $S \subseteq \mathcal{V} \times \mathcal{Y}$, define the partial set $V(S) := \{v \mid \exists (v, y) \in S\}$. Then, we have

$$\begin{aligned} \Pr[\text{Exec}^M(x) \in S] &\leq \Pr[\text{View}^M(x) \in V(S)] \\ &\leq e^\epsilon \Pr[\text{View}^M(x') \in V(S)] + \delta \leq e^\epsilon \Pr[\text{Exec}^M(x') \in \mathcal{N}(S)] + \delta. \end{aligned}$$

$\square$

Proof of Lemma 3.9: (ϵ, δ) -NP and DO imply NPDO. We slightly abuse the notation of $\mathcal{N}(\cdot)$ to denote the “neighbor set” for $\mathcal{Y}$ and also the $\mathcal{V} \times \mathcal{Y}$, as defined in Definition 3.8 and Definition 3.4. Fix a pair of neighboring input x, x' . For any $S \subseteq \mathcal{V} \times \mathcal{Y}$, denote the partial set $S_v = \{y \mid \exists (v, y) \in S\}$.

Let $\mu(v) = \min(\Pr[\text{View}^M(x) = v], e^{\epsilon_1} \Pr[\text{View}^M(x') = v])$. It's trivial that $\sum_{v \in \mathcal{V}} \mu(v) \leq 1$. Also, since M is (ϵ_1, δ_1) -DO, we have $\sum_{v \in \mathcal{V}} \mu(v) \geq 1 - \delta_1$. Then,

$$\begin{aligned}
& \Pr[\text{Exec}^M(x) \in S] \\
&= \sum_{v^* \in \mathcal{V}} \Pr[\text{View}^M(x) = v^*] \Pr[(v, y) \leftarrow \text{Exec}^M(x) : y \in S_{v^*} \mid v = v^*] \\
&\leq \delta_1 + \sum_{v^* \in \mathcal{V}} \mu(v^*) \Pr[(v, y) \leftarrow \text{Exec}^M(x) : y \in S_{v^*} \mid v = v^*] \\
&\leq \delta_1 + \sum_{v^* \in \mathcal{V}} \mu(v^*) (e^\epsilon \Pr[(v', y') \leftarrow \text{Exec}^M(x') : y' \in \mathcal{N}(S_{v^*}) \mid v' = v^*] + \delta_2) \\
&\leq \delta_1 + \delta_2 + \sum_{v^* \in \mathcal{V}} \mu(v^*) (e^\epsilon \Pr[(v', y') \leftarrow \text{Exec}^M(x') : y' \in \mathcal{N}(S_{v^*}) \mid v' = v^*]) \\
&\leq \delta_1 + \delta_2 + \sum_{v^* \in \mathcal{V}} e^{\epsilon_1} \Pr[\text{View}^M(x') = v^*] e^{\epsilon_2} \Pr[(v', y') \leftarrow \text{Exec}^M(x') : y' \in \mathcal{N}(S_{v^*}) \mid v' = v^*] \\
&= \delta_1 + \delta_2 + e^{\epsilon_1 + \epsilon_2} \Pr[\text{Exec}^M(x') \in \mathcal{N}(S)].
\end{aligned}$$

□

B.2 Proofs of Lemma 4.1

Proof. For any matching w , we use the notation of the partial sum of the matching: for any $a \in \Omega$, $w(a, \cdot) := \sum_{b \sim a} w(a, b)$. Similarly, for any $b \in \Omega$, we have $w(\cdot, b) := \sum_{a \sim b} w(a, b)$. We also define the size of a matching w as $|w| := \sum_{a, b, a \sim b} w(a, b)$.

We first prove that if there is an (ϵ, δ) -matching, the first statement holds. We prove it by contradiction. Suppose there is a set S violates the first condition, i.e., $\Pr[A \in S] > e^\epsilon \Pr[B \in \mathcal{N}(S)] + \delta$, while there is an (ϵ, δ) -matching w . We have that $\Pr[A \in S] - \sum_{a \in S, b \sim x} w(a, b) \geq \Pr[A \in S] - e^\epsilon \Pr[B \in \mathcal{N}(S)] > \delta$. Then, $1 - \sum_{a, b, a \sim b} w(a, b) = \left(\Pr[A \in S] - \sum_{a \in S, b \sim a} w(a, b) \right) + \left(\Pr[A \in \Omega/S] - \sum_{a \in \Omega/S, b \sim a} w(a, b) \right) > \delta + 0$. So $\sum_{a, b} w(a, b) < 1 - \delta$, which contradicts the condition (d) of the matching.

Now, we prove that if the first statement holds, we can find an (ϵ, δ) -matching w . This direction is more challenging. We first prove it assuming Ω is finite. Let w^* be the “max legal” matching. When we say w^* is legal, we mean that it satisfies the condition (a)-(c) for (ϵ, δ) -matching. Also, the maximality of w^* means the size of w^* is more than or equal to any other legal matching w . We prove that, if $|w^*| < 1 - \delta$, then we can find a set S that violates the condition in the first statement. Let S_0 be the set of all those $a \in \Omega$ that are not “saturated”, i.e., $\sum_{b \sim a} w^*(a, b) < \Pr[A = a]$. Now, for each iteration, we keep expanding the set S_i to S_{i+1} . Define $T_i := \{a \in A \mid \exists b \in \mathcal{N}(S_i), w^*(a, b) > 0\}$. Let $S_{i+1} = S_i \cup T_i$. We keep expanding the set. Since Ω is finite, there exists a maximal S_i that cannot be expanded, i.e., $S_i \cup T_i = S_i$. Notice that for i , all $b \in \mathcal{N}(S_{i+1})$ are saturated. Otherwise, we can find a pair of unsaturated vertices a_0, b_i , such that there exists an alternating path $a_0, b_0, a_1, b_1, \dots, a_i, b_i$ where for all $j \in [i-1]$, $a_j \sim b_j$ and $w^*(a_{j+1}, b_j) > 0$. Now, let $\Delta = \min\{\Pr[A = a] - w^*(a, \cdot), e^\epsilon \Pr[B = b] - w^*(\cdot, b), \min_{j \in [i]} \{w^*(a_{j+1}, b_j)\}\}$. We know $\Delta > 0$. For each $j \in [i-1]$, we simply raise $w^*(a_j, b_j)$ by Δ and decrease $w^*(a_{j+1}, b_j)$ by Δ . This procedure raises the size of w^* by Δ while preserving all legal conditions, which breaks the maximality of w^* .

Now, we argue that this maximal set S_i violates the first statement. We first notice that all $a \in \Omega/S_i$ are saturated. So $\Pr[A \in \Omega/S_i] - \sum_{a \in \Omega/S_i} w^*(a, \cdot) = 0$. Then, $\Pr[A \in S_i] -$

$\sum_{a \in S_i} w^*(a, \cdot) = (\Pr[A \in S_i] - \sum_{a \in S_i} w(a, \cdot)) + (\Pr[A \in \Omega/S_i] - \sum_{a \in \Omega/S_i} w(a, \cdot)) = 1 - |w^*| > \delta$. Now, we know that all $b \in \mathcal{N}(S_i)$ are saturated. Also, $\mathcal{N}(S_i)$ are only saturated by the vertices in S_i , otherwise we can expand S_i more. So we have that $\Pr[A \in S_i] - e^\epsilon \Pr[B \in \mathcal{N}(S_i)] = \Pr[A \in S_i] - \sum_{a \in S_i, b \in \mathcal{N}(S_i), a \sim b} w(a, b) > \delta$. Therefore, S_i violates the condition in the first statement.

Now we discuss the case when Ω is countable infinite. We will try to use the same argument as the finite case, but we need to handle a few additional details. First, is the maximal matching w^* well-defined and attainable in the infinite case? This turns out to be non-trivial. Luckily, our matching problem can be seen as a “locally finite” network flow problem in a countable infinite graph. Due to the Theorem 5.1 in Aharoni et al. [ABG⁺11]⁴, the maximum flow is attainable for any countable infinite and locally finite graph, which implies the maximum matching in our setting is well-defined and attainable in our case. Notice that the axiom of choice is assumed in the proof from Aharoni et al. Also, the expanding procedure could take infinite number of steps to stop. Now, we can simply define the set S_{final} that violates the condition as the union of all S_i for $i \in 0, 1, 2, \dots$. This maximal set S_{final} is well-defined by Zorn’s Lemma. Next, we know that all $y \in \mathcal{N}(S_{\text{final}})$ are saturated. Although the number of vertices in $\mathcal{N}(S_{\text{final}})$ could be infinite, there exists a finite-length alternative path from an unsaturated x to each y in $\mathcal{N}(S_{\text{final}})$, since our expanding procedure goes in round $i \in 0, 1, 2, \dots$. So if there is any unsaturated $y \in \mathcal{N}(S_{\text{final}})$, the maximality of w^* are broken. Every other parts of the proof still hold in the countable infinite case. $\square$

B.3 Proofs for the Stable Compaction Algorithm

Proof of Theorem 5.1: RandBin’s correctness. We prove the correctness statement by induction. We assume that for the first $i - 1$ iteration, $\text{Bin}_1 \dots \text{Bin}_{i-1}$ have the correct real elements and also, no real elements have been truncated.

Now, in the i -th iteration, we will read the input array up to index $\text{cnt}[i] + s$. Since s is the upper bound on the additive error of the prefix-sum mechanism, we have $|\text{cnt}[i] - \sum_{j \leq i} \rho_j| \leq s$ and thus, $\text{cnt}[i] + s \geq \sum_{j \leq i} \rho_j$. Then, we know the real elements between $I \left[\sum_{j < i} \rho_j + 1 \right]$ to $I \left[\sum_{j \leq i} \rho_j \right]$ will be visited. By the assumption, they are not moved to the bins before, so they must still be in the buffer. Then, the algorithm obviously sort the buffer and correctly move those real elements to Bin_i . Now, the real elements from $I \left[\sum_{j \leq i} \rho_j + 1 \right]$ to $I \left[\text{cnt}[i] + s \right]$ are in Buf . Again, since $|\text{cnt}[i] - \sum_{j \leq i} \rho_j| \leq s$, we have $\text{cnt}[i] + s - \sum_{j \leq i} \rho_j \leq 2s$. So there will be no more than $2s$ real elements and after another round of oblivious sorting, no real elements will be truncated.

Finally, we have $\rho_i \geq Z/2$, so $\sum_{i \leq B} \rho_i \geq BZ/2 \geq L \geq |I|$. Thus, we will visit the whole input and all the real elements will be moved to the correct bins.

The constant in s . We need to ensure that $s = \frac{15}{\epsilon} \log^2 L \log \frac{1}{\delta}$ is indeed an upper bound on the support of $\mathcal{G}(\epsilon_2, \delta_2)$ and also the additive error of $\text{PrefixSum}^{\epsilon_3, \delta_3}$ on at most L integers. The support of $\mathcal{G}(\epsilon_2, \delta_2)$ is $[0, \frac{2}{\epsilon_2} \ln \frac{2}{\delta_2} + O(1)]$. According to Lemma 2.8 and Theorem 3.6 in [CSS10] and using a union bound, for the release of all L prefix sums, the error will be bounded by $\frac{2\sqrt{2}}{\epsilon_3} (\log L)^{1.5} \ln \frac{2L}{\delta_3}$ with probability $1 - \delta$. We can simply move the utility failure probability δ to the privacy failure probability by releasing all correct prefix sums when any error exceeds the bound. So the prefix sum mechanism will be (ϵ_2, δ_2) private and the error will be bounded by $\frac{2\sqrt{2}}{\epsilon_3} (\log L)^{1.5} \ln \frac{2L}{\delta_3}$ with probability 1. Finally, plugging in $\epsilon_2 = \epsilon_3 = \frac{\epsilon}{3}$ and $\delta_2 = \delta_3 = \frac{\delta}{3}$, it can be verified that the constant

⁴A followup paper by Lochbihler [Loc22] pointed out and fixed some technical issues of the proof in Aharoni et al. [ABG⁺11].

15 is big enough. $\square$

Proof of Theorem 5.2: RandBin's runtime. Before each iteration, Buf has size at most $2s$. For each iteration i , we will move at most $\text{cnt}[i] - \text{cnt}[i-1] + s$ elements to the buffer. Since $|\text{cnt}[i] - \sum_{j \leq i} \rho_j| \leq s$ and $|\text{cnt}[i-1] - \sum_{j \leq i-1} \rho_j| \leq s$, we have $\text{cnt}[i] - \text{cnt}[i-1] \leq \rho_i + 2s$. Then, $\text{cnt}[i] - \text{cnt}[i-1] + s \leq \rho_i + 2s \leq Z + 2s \leq 4s$. Thus, we know the buffer size cannot be more than $6s$, which implies that each oblivious sorting on Buf takes $O(s \log s)$ time. Also, we know $Z = 2s$, so we have $B = \lceil \frac{L}{s} \rceil + 1$. The total time complexity will be $O(Bs \log s) = O(L \log s)$. Since $|I| = \Omega(\frac{1}{\epsilon} \log \frac{1}{\delta})$, we have that $L \leq |I| + O(\frac{1}{\epsilon} \log \frac{1}{\delta}) = O(I)$. Finally, since $s = \Theta(\frac{1}{\epsilon} \log^2 L \cdot \log \frac{1}{\delta})$, the time complexity is $O(|I| (\log \log |I| + \log \frac{1}{\epsilon} + \log \log \frac{1}{\delta}))$. $\square$

Proof of Theorem 5.3: RandBin is (ϵ, δ) -NPDO. Recall that we need to prove the following statement: for any subset S of (view, output) pairs, for any neighboring input pair I, I' , it holds that

$$\Pr[\text{Exec}^{\text{RandBin}}(I) \in S] \leq e^\epsilon \Pr[\text{Exec}^{\text{RandBin}}(I') \in \mathcal{N}(S)] + \delta. \quad (5)$$

where $\text{Exec}^{\text{RandBin}}$ returns the view and the output bin representation, and the neighboring set $\mathcal{N}$ is defined w.r.t. the relation $\sim_B$ (i.e., neighboring of bin representation).

First, observe that the adversary's view is fully determined knowing the noisy input length L and the noisy bin loads cnt . Therefore, henceforth, we may use (L, cnt) to denote the view. Second, observe that the output of the algorithm is fully determined by the input I , the noisy input length L (which determines the parameters Z and B), and the actual random bin load vector ρ . Now, we define the notation $\overline{\text{Exec}}(I)$ to output L, cnt , as well as ρ — notice that here, we have switched the output to just the bin load vector ρ rather than the actual bin representation output. Henceforth, let $\mathcal{Y}$ be the support of the variable ρ , and we abuse notation slightly and use $\mathcal{V}$ to denote the support of the pair (L, cnt) . For any given subset $S \subseteq \mathcal{V} \times \mathcal{Y}$, we define the neighboring set notation $\overline{\mathcal{N}}$ (parametrized by I and I') as follows:

$$\overline{\mathcal{N}}(S) = \{(L, \text{cnt}, \rho') \mid \exists \rho' \text{ s.t. } (L, \text{cnt}, \rho) \in S \text{ and } \text{out}(I, L, \rho) \sim_B \text{out}(I', L, \rho')\}$$

where $\text{out}(I, L, \rho)$ outputs the unique bin representation that is fully determined by the input I , the noisy input length L , and the bin load vector ρ .

Now, proving Equation (5) is equivalent to proving that for any $S \subseteq \mathcal{V} \times \mathcal{Y}$, any neighboring inputs I, I' , it holds that

$$\Pr[\overline{\text{Exec}}(I) \in S] \leq e^\epsilon \Pr[\overline{\text{Exec}}(I') \in \overline{\mathcal{N}}(S)] + \delta.$$

We consider these following cases.

Case 1. A substitution transforms I to I' . This is an easy case. Fix any $(L, \text{cnt}, \rho) \in \mathcal{V} \times \mathcal{Y}$ such that $\Pr[\overline{\text{Exec}}(I) = (L, \text{cnt}, \rho)] > 0$. Let's assume that the substituted elements x are placed in Bin_i in the bin representation $\text{out}(I, L, \rho)$. Since I and I' only differ by a substitution, x will be in Bin'_i in the bin representation $\text{out}(I', L, \rho)$ with the same index i . All other bins are exactly the same. Also, the i -th bin will have real bin loads differ at most one. Therefore, by the definition of neighboring relation for bin representation, we know $\text{out}(I, L, \rho) \sim_B \text{out}(I', L, \rho)$. Thus, in this case, given any $S \subseteq \mathcal{V} \times \mathcal{Y}$, $S \subseteq \overline{\mathcal{N}}(S)$. Moreover, the equation $\Pr[\overline{\text{Exec}}(I) = (L, \text{cnt}, \rho)] = \Pr[\overline{\text{Exec}}(I') = (L, \text{cnt}, \rho)]$ trivially holds in this case because $|I| = |I'|$.

Case 2. An insertion or deletion transforms I to I' .

Fix any $(L, \text{cnt}, \rho) \in \mathcal{V} \times \mathcal{Y}$ such that $\Pr[\overline{\text{Exec}}(I) = (L, \text{cnt}, \rho)] > 0$.

(i) Suppose an element x is deleted from I to form I' . Suppose i is the index of the bin in $\text{out}(I, L, \rho)$ that contains x ; observe the choice of B ensures that $i < B$. We write $\rho = (\rho_1, \dots, \rho_B)$. We consider a $\rho' = (\rho'_1, \dots, \rho'_B)$, such $\rho'_{i+1} = \rho_{i+1} - 1$ and $\rho'_j = \rho_j$ for $j \neq i+1$. We first prove that $\text{out}(I, L, \rho) \sim_B \text{out}(I', L, \rho')$. We only need to check the bin load condition and the other two conditions are naturally satisfied. We write $\text{out}(I, L, \rho) = (\text{Bin}_i : i \in [B])$ and $\text{out}(I', L, \rho') = (\text{Bin}'_i : i \in [B])$. It is easy to see that $\text{Bin}_1 = \text{Bin}'_1, \dots, \text{Bin}_{i-1} = \text{Bin}'_{i-1}$. Also, $\text{Bin}_{i+2} = \text{Bin}'_{i+2}, \dots, \text{Bin}_B = \text{Bin}'_B$, because $\rho'_{i+1} = \rho_{i+1} - 1$ offsets the influence of the deletion of x for all following bins. Thus, only bin i and $i+1$ will be different in $\text{out}(I, L, \rho)$ and $\text{out}(I', L, \rho')$. For example, suppose $I = a \perp bcd \perp \dots$, $I' = abcd \perp \dots$ where a, b, c, d are the real elements and $\perp$ denotes the filler elements. Assuming $\rho = (3, 3, \dots)$ and $\rho' = (3, 2, \dots)$, the corresponding bin representation will be $\underline{a \perp b} \ \underline{cd \perp}$ and $\underline{abc} \ \underline{d \perp \perp}$, where the first two bins will have different bin loads and the differences are at most one for each bin. More formally, since $\rho_i = \rho'_i$ and I' is missing the element x , Bin_i has the particular element x , while Bin'_i has the element $I[\sum_{j \leq i} y_j + 1]$ instead (which will be the same as $I'[\sum_{j \leq i} y_j]$). Thus, the real element numbers in Bin_i and Bin'_i may differ by at most one. Also, Bin'_{i+1} simply misses the first element in Bin_{i+1} and every other elements are the same. So the real element numbers in Bin_{i+1} and Bin'_{i+1} may differ by at most one. Henceforce, we prove that $\text{out}(I, L, \rho) \sim_B \text{out}(I', L, \rho')$. For convenience, we define a mapping $\varphi_L : \mathcal{Y} \rightarrow \mathcal{Y}$, such that it maps every ρ to ρ' as the construction above. Notice that φ_L relies on L because the bin size Z and the bin number B relies on a fixed L , and then we restrict $\text{dom}(\varphi_L)$ to those ρ that have exactly B bins of size Z .

We then prove $\varphi_L(\rho)$ is injective by contradiction. Suppose there are $\rho^0, \rho^1 \in \text{dom}(\varphi_L)$ such that $\rho^0 \neq \rho^1$ but $\varphi_L(\rho^0) = \varphi_L(\rho^1)$. Assume x is in Bin_{i_0} in $\text{out}(I, L, \rho^0)$ and in Bin_{i_1} in $\text{out}(I, L, \rho^1)$, and $i_0 < i_1$ without loss of generality. We know that the prefix $\rho^0_{1 \dots i_0}$ are the same as the prefix $\varphi_L(\rho^0)_{1 \dots i_0}$ by the construction of φ_L . Also, the prefix $\rho^1_{1 \dots i_1}$ are the same as the prefix $\varphi_L(\rho^1)_{1 \dots i_1}$. Thus, the prefix $\rho^0_{1 \dots i_0}$ are the same as the prefix $\rho^1_{1 \dots i_0}$ and by the definition of out , x should be in Bin_{i_0} of $\text{out}(I, L, \rho^1)$. This is a contradiction.

(ii) Suppose an element x is inserted into I to form I' . Consider the bin representation $\text{out}(I, L, \rho)$. Suppose x is in Bin_i . We can define an injective mapping φ similarly as the mapping in part (i): $\varphi_L(\rho) = \rho'$ such that $\rho'_{i+1} = \rho_{i+1} + 1$ and $\rho'_j = \rho_j$ for $j \neq i+1$. We can prove φ_L is injective and the corresponding bin representation are neighbors by similar arguments in part (i).

It remains to show the (ϵ, δ) -differentially private inequality. We can use the same proof for insertion and deletion, because $\|\rho - \varphi_L(\rho)\|_1 = 1$ in both cases.

We abused the notation before and now we make a clear distinction: the notations L, cnt, ρ denote the random variables, while the notations ℓ, cnt, ρ denote the particular point in the sample space.

For any measurable $S \subseteq \mathcal{V} \times \mathcal{Y}$ and $\ell \in \mathbb{Z}_+, \rho \in \mathcal{Y}$, denote $S_{\ell, \rho} := \{\text{cnt} : (\ell, \text{cnt}, \rho) \in S\}$. We also denote $\varphi(S) = \{(\ell, \text{cnt}, \varphi_\ell(\rho)) \mid (\ell, \text{cnt}, \rho) \in S\}$.

Observe that following properties:

1. First, since $d_E(I, I') \leq 1$, for all $\ell \in \mathbb{Z}_+$, $\Pr[L = \ell] \leq e^{\epsilon_1} \Pr[L' = \ell] + \delta_1$.
2. Because of the truncated geometric distribution $\text{Geom}(\epsilon_2, \delta_2)$ and the choice of Z , we have $\Pr[\rho \notin \text{dom}(\varphi_\ell) \mid L = \ell] \leq \delta_2$.
3. For $\rho \in \text{dom}(\varphi_\ell)$, since ρ and ρ' in both executions have the same distribution, we have $\Pr[\rho = \rho \mid L = \ell] \leq e^{\epsilon_2} \cdot \Pr[\rho' = \varphi_\ell(\rho) \mid L' = \ell]$, noting that $\|\rho - \varphi_\ell(\rho)\|_1 = 1$.
4. For $\rho \in \text{dom}(\varphi_\ell)$, again we have $\|\rho - \varphi_\ell(\rho)\|_1 = 1$. Then, the (ϵ_3, δ_3) -differentially private prefix sum mechanism implies that $\Pr[\text{cnt} \in S_{\ell, y} \mid \rho = \rho, L = \ell] \leq e^{\epsilon_3} \cdot \Pr[\text{cnt}' \in S_{\ell, y} \mid \rho' = \rho]$.

$$\varphi_\ell(\rho), L' = \ell] + \delta_3.$$

Then, we have the following inequalities:

$$\begin{aligned}
& \Pr[\overline{\text{Exec}}(I) = (L, \text{cnt}, \rho) \in S | L = \ell] \\
& \leq \sum_{\rho \in \text{dom}(\varphi_\ell)} \Pr[\rho = \rho | L = \ell] \cdot \Pr[\text{cnt} \in S_{\ell, \rho} | \rho = \rho, L = \ell] + \delta_2 \\
& \leq \sum_{\rho \in \text{dom}(\varphi_\ell)} \Pr[\rho = \rho | L = \ell] \cdot (e^{\epsilon_3} \cdot \Pr[\text{cnt}' \in S_{\ell, \rho} | \rho' = \varphi_\ell(\rho), L' = \ell] + \delta_3) + \delta_2 \\
& \leq \sum_{\rho \in \text{dom}(\varphi_\ell)} e^{\epsilon_3 + \epsilon_2} \Pr[\rho' = \varphi_\ell(\rho) | L' = \ell] \cdot \Pr[\text{cnt}' \in S_{\ell, \rho} | \rho' = \varphi_\ell(\rho), L' = \ell] + \delta_3 + \delta_2 \\
& = \sum_{\rho' \in \text{range}(\varphi_\ell)} e^{\epsilon_3 + \epsilon_2} \Pr[\rho' = \rho' | L' = \ell] \cdot \Pr[\text{cnt}' \in S_{\ell, \varphi_\ell^{-1}(\rho')} | \rho' = \rho', L' = \ell] + \delta_3 + \delta_2 \\
& = e^{\epsilon_3 + \epsilon_2} \cdot \Pr[\overline{\text{Exec}}(I') = (L', \text{cnt}', \rho') \in \varphi(S) | L' = \ell] + \delta_3 + \delta_2,
\end{aligned}$$

where the last equality holds because for all $\rho' \in \text{range}(\varphi_\ell)$, $\text{cnt}' \in S_{\ell, \varphi_\ell^{-1}(\rho')}$ iff $(\ell, \text{cnt}', \rho') \in \varphi(S)$.

Since $\Pr[L = \ell] \leq e^{\epsilon_1} \cdot \Pr[L' = \ell] + \delta_1$ for all ℓ , using a standard DP-composition argument, we have

$$\Pr[\overline{\text{Exec}}(I) = (L, \text{cnt}, \rho) \in S] \leq e^{\epsilon_1 + \epsilon_2 + \epsilon_3} \cdot \Pr[\overline{\text{Exec}}(I') = (L', \text{cnt}', \rho') \in \varphi(S)] + \delta_3 + \delta_2 + \delta_1.$$

Finally, since φ_ℓ is injective and we know for all $(\ell, \text{cnt}, \rho) \in S$, $(\ell, \text{cnt}, \varphi_\ell(\rho)) \in \mathcal{N}(S)$, we have that

$$\Pr[\overline{\text{Exec}}(I') = (L', \text{cnt}', \rho') \in \varphi(S)] \leq \Pr[\overline{\text{Exec}}(I') = (L', \text{cnt}', \rho') \in \overline{\mathcal{N}}(S)],$$

and we finish the proof. $\square$

Proof of Theorem 5.4: CompactBin's correctness. We need to show that no filler element will be output before the last real element and no real element will be truncated. We know fact that for every $i \in [B]$, $|\sum_{j \leq i} R_j - \text{cnt}[i]| \leq s$ since s is the upper bound on error of the prefix sum mechanism.

Let i^* be the smallest i such that $\text{cnt}[i] - s \geq 0$. If all $\text{cnt}[i] < s$, then let $i^* = B + 1$. First, we observe that during the first $i^* - 1$ iteration, no element will be output. Also, at the end of the iteration $i^* - 1$, Buf contains $\sum_{j \leq i^* - 1} R_j$ real elements. We know that $\sum_{j \leq i^* - 1} R_j \leq \text{cnt}[i^* - 1] + s \leq 2s$, so no real elements get truncated.

Starting from the i^* iteration, for each iteration i , after reading the i -th bin into Buf, the algorithm has read $\sum_{j \leq i} R_j$ real elements. Then, if we know at the end of the iteration, the output array does not have more than $\sum_{j \leq i} R_j$ elements, then all the elements that in the output tape are still real. We simply have that for $i \geq i^*$, the output array will be appended to the length of $\text{cnt}[i] - s$ at the end of the iteration and $\text{cnt}[i] - s \leq \sum_{j \leq i} R_j$. Also, we know that after the output step, Buf contains exactly $\sum_{j \leq i} R_j - (\text{cnt}[i] - s) \leq 2s$ real elements. So no real elements will be truncated.

At the end of the algorithm, we append the whole Buf to the output. Since this Buf is sorted, the first filler element will be after the last real element. $\square$

Proof of Theorem 5.5: **CompactBin** is (ϵ, δ) -NPDO. First, due to Theorem 5.4 and the definition of $\sim_B$, we know that **CompactBin** is strongly-NP. Also, notice that the view of the adversary fully depends on the (ϵ_1, δ_1) -DP prefix sum mechanism result **cnt**. By the definition of $\sim_B$, we know there are at most two bins that have difference at most one given the neighboring input. Notice that we set $\epsilon_1 = \epsilon/2, \delta_1 = \delta/(2(1 + e^{\epsilon_1}))$. With group privacy theorem, it is straightforward that releasing **cnt** to the adversary is (ϵ, δ) -DP and thus **CompactBin** is (ϵ, δ) -DO. Therefore, we have **CompactBin** is (ϵ, δ) -NPDO. $\square$

Proof of Theorem 5.6: **CompactBin**'s runtime. We have B iterations. Before each iteration, the **Buf** has size at most $2s$. When we read the i -th bin into **Buf**, its size will be at most $2s + Z$. So performing oblivious sorting on **Buf** takes $O((Z + s) \log(Z + s))$ time. In total, the time complexity will be $O(B(Z + s) \log(Z + s))$. $\square$

B.4 Proofs for the privacy amplification theorem by DO-shuffle

In Section 6, we mentioned that it is suffice to prove the mechanism M'_1 that provides the adversary with auxiliary information is (ϵ, δ) -NPDO w.r.t the input neighboring relation of $\sim_H$ (neighboring by Hamming distance at most 1) and output neighboring relation of $\sim_S$ (neighboring by swapping two messages).

Equivalent view of M_1 . Without loss of generality, fix a pair of neighboring input configurations $\mathbf{x} = (x_1, x_2, \dots, x_n)$ and $\mathbf{x}' = (x'_1, x_2, \dots, x_n)$ that only differ in the first client's input. Feldman et al. [FMT21] defined the following equivalent view of M_1 . For client 1, its message are generated as following: for $\mathbf{x}$, the message y_1 will be sampled from $\mathcal{R}(x_1)$; for $\mathbf{x}'$, y_1 is sampled from $\mathcal{R}(x'_1)$. For each client i from 2 to n , it behaves the same given the input being $\mathbf{x}$ or $\mathbf{x}'$. That is, y_i is sampled from $\mathcal{R}(x_i)$. Since $\mathcal{R}$ is an ϵ_0 -DP mechanism, for any $x, x' \in \mathcal{X}$ where $\mathcal{X}$ is the input domain and any $y \in \mathcal{Y}$ where $\mathcal{Y}$ is the randomizer's output domain, we have $\frac{\Pr[\mathcal{R}(x)=y]}{\Pr[\mathcal{R}(x')=y]} \geq e^{-\epsilon_0}$. Thus, we can “decompose” the distribution of $\mathcal{R}(x')$ as $\mathcal{R}(x') = e^{-\epsilon_0} \mathcal{R}(x) + (1 - e^{-\epsilon_0}) \mathcal{Q}_x(x')$. It means that the distribution of $\mathcal{R}(x')$ can be seen as a mixture of $\mathcal{R}(x)$ and some appropriate leftover distribution $\mathcal{Q}_x(x')$. Similarly, for all x_i where $i = 2, \dots, n$, we can decompose the output distribution of $\mathcal{R}(x_i)$ as

$$\mathcal{R}(x_i) = \begin{cases} \mathcal{R}(x_1) & w.p. \frac{1}{2e^{\epsilon_0}}, \\ \mathcal{R}(x'_1) & w.p. \frac{1}{2e^{\epsilon_0}}, \\ \mathcal{Q}_{\{x_1, x'_1\}}(x_i), & w.p. 1 - \frac{1}{e^{\epsilon_0}}. \end{cases}$$

Here, $\mathcal{Q}_{\{x_1, x'_1\}}(x_i)$ is also some appropriate left-over distribution after the decomposition. The properties of $\mathcal{Q}_{\{x_1, x'_1\}}(x_i)$ are irrelevant in this proof, so we only need to know $\mathcal{Q}_{\{x_1, x'_1\}}(x_i)$ is a legal and well-defined distribution. In other words, with probability $1/e^{\epsilon_0}$, each client from 2 to n will clone either $\mathcal{R}(x_1)$ or $\mathcal{R}(x'_1)$ with equal probability, and with the remaining probability, it will sample from the leftover distribution $\mathcal{Q}_{\{x_1, x'_1\}}(x_i)$.

Definition of M'_1 . Recall that in M_1 , the adversary's view is empty. We will define an augmented algorithm M'_1 that is almost identical as M_1 except that we add some extra information to the adversary's view. We define the view of the mechanism M'_1 to contain the random variables (U, V, T) . Here, V is the set of the indices of those clients that sample their messages from $\mathcal{R}(x_1)$ or $\mathcal{R}(x'_1)$. V always contains index 1 and it may contain some indices from 2 to n , depending on those clients'

decision on sampling from the corresponding decomposed distribution. U is the multi-set of the messages sent by those clients in V . T contains the messages sent by all the non-clones (i.e., clients not in V) along with their indices.

Later in our proof, we will consider a probability space conditioned on the adversary's view (U, V, T) . Thus, it may be helpful to think of M'_1 as first sampling the variables (U, V, T) , and then sampling the output $(y_1, \dots, y_n)$ accordingly. Specifically, consider the following equivalent way of sampling the experiment M'_1 .

- First, decide out of the $n - 1$ clients, how many clients sample from $\mathcal{R}(x_1)$ or $\mathcal{R}(x'_1)$. We sample this quantity, denoted as c , from $\text{Bin}(n - 1, e^{-\epsilon_0})$.
- Then, conditioned on c clients sampling from $\mathcal{R}(x_1)$ or $\mathcal{R}(x'_1)$, we sample the number of the clients cloning $\mathcal{R}(x_1)$, denoted as H , from $\text{Bin}(c, 1/2)$. Let $H' = c - H$ be the number of clients cloning $\mathcal{R}(x'_1)$.
- We now denote (m, m') be the total numbers of clients (including client 1) that are sampling from $\mathcal{R}(x_1)$ and $\mathcal{R}(x'_1)$, respectively. When the input is $\mathbf{x}$, $(m, m') = (H + 1, H')$. When the input is $\mathbf{x}'$, $(m, m') = (H, H' + 1)$.
- Generate U, V as following: Sample m messages independently from $\mathcal{R}(x_1)$ and sample m' messages independently from $\mathcal{R}(x'_1)$. Let U be the multi-set of these $m + m'$ messages. Let V be a subset of the indices of $m + m'$ clients picked uniformly at random from $[n]$ that must include client 1.
- Generate T as following: for each $i \notin V$, sample a message y_i from $\mathcal{Q}_{\{x_1, x'_1\}}(x_i)$ and add the index-message pair (i, y_i) to T .
- Finally, generate the output (each client's message) in the following way.
 1. If the input is $\mathbf{x}$, sample a subset of m clients from V uniformly at random conditioned on containing client 1. If the input is $\mathbf{x}'$, the condition will be changed to exclude client 1.
 2. The clients in this set will be assigned messages from the m messages that are sampled earlier from $\mathcal{R}(x_1)$ (for instance, in a uniformly random order), and the remaining clients in V will be assigned messages from the remaining m' messages sampled from $\mathcal{R}(x'_1)$.
 3. For a client $i \notin V$, its message is already determined by T .

A key lemma from Feldman et al. [FMT21] shows that releasing (m, m') , i.e., the number of clients (including client 1) sampling from $\mathcal{R}(x_1)$ and $\mathcal{R}(x'_1)$, is (ϵ, δ) differentially private, where ϵ is much smaller than ϵ_0 . Formally, they prove the following result.

Lemma B.1 ([FMT21]). *Let $c \xleftarrow{\$} \text{Bin}(n - 1, e^{-\epsilon_0})$, $H \xleftarrow{\$} \text{Bin}(c, 1/2)$ and $H' = c - H$. For any $\delta > 0$ and $\epsilon_0 \leq \log\left(\frac{n}{16 \log(2/\delta)}\right)$, the joint distributions on $(H + 1, H')$ and $(H, H' + 1)$ are (ϵ, δ) -close where $\epsilon = O\left(\frac{(1 - e^{-\epsilon_0})e^{\epsilon_0/2}\sqrt{\log(1/\delta)}}{\sqrt{n}}\right)$.*

Here, two distributions $\mathcal{D}$ and $\mathcal{D}'$ are (ϵ, δ) -close means that for any subset S in the sample space, $\Pr[x \xleftarrow{\$} \mathcal{D} : x \in S] \leq e^\epsilon \Pr[x' \xleftarrow{\$} \mathcal{D}' : x' \in S] + \delta$ and vice versa.

Notice that conditioned on (m, m') , whose distribution is (ϵ, δ) -close, the generation of the view (U, V, T) is the same regardless of the input being $\mathbf{x}$ or $\mathbf{x}'$. Thus, by post-processing theorem and Lemma B.1, we can easily show that M'_1 is (ϵ, δ) -DO. Also, suppose there are t clients that are corrupted by the adversary. Notice that all honest clients run $\mathcal{R}()$ locally. Then, we can simply just

consider the honest client set, which has size $n - t$ and the aforementioned argument still holds, as long as we plug $n - t$ into Lemma B.1. We have:

Corollary B.2. *For any $\delta > 0$ and $\epsilon_0 \leq \log\left(\frac{n-t}{16\log(2/\delta)}\right)$, M'_1 is (ϵ, δ) -DO with $\epsilon = O\left(\frac{(1-e^{-\epsilon_0})e^{\epsilon_0/2}\sqrt{\log(1/\delta)}}{\sqrt{n-t}}\right)$ when at most t clients are corrupted by the adversary.*

Now, we only need to prove the following key lemma and the proof is finished. In this proof, we will use techniques from Gordon et al. [GKLX22] in a non-black box way.

Lemma B.3. M'_1 is $(0, 0)$ -NP.

Proof. Given any view of the adversary (U, V, T) that has non-zero probability when the input is $\mathbf{x}$ or $\mathbf{x}'$, it suffices to give a $(0, 0)$ -matching on the output space of M'_1 conditioned on (U, V, T) as in Lemma 4.1.

If client i is sampling from $\mathcal{R}(x_1)$ or $\mathcal{R}(x'_1)$, we call that client i is *cloning* $\mathcal{R}(x_1)$ or $\mathcal{R}(x'_1)$. Recall that V is the subset of clients (including client 1) that clone either $\mathcal{R}(x_1)$ or $\mathcal{R}(x'_1)$, and U is the multi-set of the messages they generate.

Without loss of generality, we assume that all clients in V are honest; otherwise, we remove dishonest clients from V (and also their messages from U). Denote $\ell = |V|$ and we can rename the clients such that $V = [\ell]$ without loss of generality. For clients $i \notin V$, their messages in the output are already determined by T . For the rest of the proof, since we will always conditioned on T , we will focus on the part of the output space $\mathcal{Y}^\ell$ corresponding to the clients in V .

To recap the setup, we have fixed some neighboring inputs $\mathbf{x}$ and $\mathbf{x}'$ and conditioned on some subset $V = [\ell]$ of clients that clone either $\mathcal{R}(x_1)$ or $\mathcal{R}(x'_1)$, and their multi-set U of output messages (such that the view (U, V) has non-zero probability in both executions). We use SEQ_U to denote the multi-set of the $\ell!$ permutations of messages in U , where each permutation can be interpreted as an output in $\mathcal{Y}^\ell$. Since U itself is a multi-set, then so is SEQ_U ; for the purpose of distinguishing repeated elements in U , we can artificially attach labels on them which will help to resolve any ambiguity. Observe that conditioning on (U, V) , only outputs in SEQ_U can have non-zero probabilities in either executions. So we only need to build the mappings $w(\mathbf{y}, \mathbf{y}')$ for $\mathbf{y}, \mathbf{y}' \in \text{SEQ}_U$.

Conditional Probability Calculations. Before we build the mapping, we need to figure out the conditional probabilities of the outputs. For each $\mathbf{y} \in \text{SEQ}_U$, we will calculate $\Pr[M'_1(\mathbf{x}) = \mathbf{y} | \text{View}^{M'_1}(\mathbf{x}) = (U, V, T)]$. For each $\mathbf{y}' \in \text{SEQ}_U$, we will calculate $\Pr[M'_1(\mathbf{x}') = \mathbf{y}' | \text{View}^{M'_1}(\mathbf{x}') = (U, V, T)]$. For simplicity, we use $\Pr_{\mathbf{x}}[\cdot]$ to denote probabilities associated with the execution given input $\mathbf{x}$, and $\Pr_{\mathbf{x}'}[\cdot]$ to denote those associated given input $\mathbf{x}'$. So we write the previous conditional probabilities as $\Pr_{\mathbf{x}}[\mathbf{y} | U, V, T], \Pr_{\mathbf{x}'}[\mathbf{y}' | U, V, T]$ for short.

We use μ and μ' to denote the distribution for $\mathcal{R}(x_1)$ and $\mathcal{R}(x'_1)$, respectively. Also, we define $\mu(\cdot)$ and $\mu'(\cdot)$ as their probability density functions. For $i \in V$ such that $i \neq 1$, client i will clone either $\mathcal{R}(x_1)$ or $\mathcal{R}(x'_1)$ with $\frac{1}{2}$ probability (conditioned on only V); in other words, it will generate a message sampled from the mixture distribution $\omega = \frac{1}{2}\mu + \frac{1}{2}\mu'$. We also denote $\omega(\cdot)$ as its probability density function.

We first compute the probabilities corresponding to the execution on $\mathbf{x}$. Conditioned on V , for each $\mathbf{y} \in \text{SEQ}_U$, we have

$$\Pr_{\mathbf{x}}[\mathbf{y} | V, T] = \mu(y_1)\omega(y_2) \cdots \omega(y_\ell) = \frac{\mu(y_1)}{\omega(y_1)} \prod_{u \in U} \omega(u).$$

By enumerating all $\ell!$ sequences in SEQ_U , we have $\Pr_{\mathbf{x}}[U | V, T] = \sum_{\mathbf{y} \in \text{SEQ}_U} \Pr_{\mathbf{x}}[\mathbf{y} | V, T]$.

Observe that the above probability is the sum of $\ell!$ terms, each of which corresponds to a sequence in SEQ_U . Hence, for each $\mathbf{y} \in \text{SEQ}_U$,

$$\Pr_{\mathbf{x}}[\mathbf{y} \mid (U, V, T)] = \frac{\Pr_{\mathbf{x}}[\mathbf{y} \mid V, T]}{\Pr_{\mathbf{x}}[U \mid V, T]} = \frac{\frac{\mu(y_1)}{\omega(y_1)} \prod_{u \in U} \omega(u)}{\sum_{\mathbf{z} \in \text{SEQ}_U} \frac{\mu(z_1)}{\omega(z_1)} \prod_{u \in U} \omega(u)} = \frac{1}{(\ell-1)!} \frac{\mu(y_1)/\omega(y_1)}{\sum_{z \in U} \mu(z)/\omega(z)},$$

where the last equality holds, because we can group the $\ell!$ possible sequences in SEQ_U according to the message sent from client 1, i.e., z_1 . Hence, there are ℓ groups and each group has $(\ell-1)!$ sequences.

By the same calculation, for the execution on $\mathbf{x}'$, we have:

$$\Pr_{\mathbf{x}'}[\mathbf{y} \mid (U, V, T)] = \frac{1}{(\ell-1)!} \frac{\mu'(y_1)/\omega(y_1)}{\sum_{z \in U} \mu'(z)/\omega(z)}.$$

Building the (0,0)-matching. We are now ready to build the (0,0)-matching $w : \mathcal{Y}^\ell \times \mathcal{Y}^\ell \rightarrow [0, 1]$. Notice that we only need to assign non-zero value to $w(\mathbf{y}, \mathbf{y}')$ where $\mathbf{y}, \mathbf{y}' \in \text{SEQ}_U$, because only those sequences will have non-zero probability in the execution conditioned on (U, V, T) .

Given a vector $\mathbf{y} \in \mathcal{Y}^\ell$ and some index $k \in [\ell]$, we use $\mathbf{y}_{1 \leftrightarrow k}$ as the sequence produced by interchanging coordinates between the 1st and the k th positions. For each $\mathbf{y} \in \text{SEQ}_U$, it has ℓ neighbors: $\mathbf{y}_{1 \leftrightarrow 1}, \mathbf{y}_{1 \leftrightarrow 2}, \dots, \mathbf{y}_{1 \leftrightarrow \ell}$. We let

$$w(\mathbf{y}, \mathbf{y}_{1 \leftrightarrow k}) = \frac{1}{(\ell-1)!} \frac{\mu(y_1)/\omega(y_1)}{\sum_{z \in U} \mu(z)/\omega(z)} \times \frac{\mu'(y_k)/\omega(y_k)}{\sum_{z \in U} \mu'(z)/\omega(z)}$$

It's easy to check that this matching satisfies all the conditions in Lemma [4.1](#):

- (a) We know that $\mathbf{y} \sim \mathbf{y}_{1 \leftrightarrow k}$, so $w(\mathbf{y}, \mathbf{y}')$ will only be non-zero if $\mathbf{y} \sim \mathbf{y}'$.
- (b) For each $\mathbf{y} \in \text{SEQ}_U$, we have that

$$\sum_{k \in [\ell]} w(\mathbf{y}, \mathbf{y}_{1 \leftrightarrow k}) = \frac{1}{(\ell-1)!} \frac{\mu(y_1)/\omega(y_1)}{\sum_{z \in U} \mu(z)/\omega(z)} \frac{\sum_{k \in [\ell]} \mu'(y_k)/\omega(y_k)}{\sum_{z \in U} \mu'(z)/\omega(z)} = \frac{1}{(\ell-1)!} \frac{\mu(y_1)/\omega(y_1)}{\sum_{z \in U} \mu(z)/\omega(z)}.$$

The last equality holds because $y_1, \dots, y_\ell$ is a permutation of U .

So we get that $\sum_{k \in [\ell]} w(\mathbf{y}, \mathbf{y}_{1 \leftrightarrow k}) = \Pr_{\mathbf{x}}[\mathbf{y} \mid (U, V, T)]$.

- (c) For all $\mathbf{y}' = (y'_1, \dots, y'_\ell) \in \text{SEQ}_U$, it also has ℓ neighbors $\mathbf{y}'_{1 \leftrightarrow 1}, \mathbf{y}'_{1 \leftrightarrow 2}, \dots, \mathbf{y}'_{1 \leftrightarrow \ell}$ and only those neighbors will satisfy $w(\mathbf{y}'_{1 \leftrightarrow k}, \mathbf{y}') > 0$. We simply check that

$$\sum_{k \in [\ell]} w(\mathbf{y}'_{1 \leftrightarrow k}, \mathbf{y}') = \frac{1}{(\ell-1)!} \frac{\mu'(y'_1)/\omega(y'_1)}{\sum_{z \in U} \mu'(z)/\omega(z)} \frac{\sum_{k \in [\ell]} \mu(y'_k)/\omega(y'_k)}{\sum_{z \in U} \mu(z)/\omega(z)} = \frac{1}{(\ell-1)!} \frac{\mu'(y'_1)/\omega(y'_1)}{\sum_{z \in U} \mu'(z)/\omega(z)}.$$

The last equality holds because $y'_1, \dots, y'_\ell$ is a permutation of U .

So we get that $\sum_{k \in [\ell]} w(\mathbf{y}'_{1 \leftrightarrow k}, \mathbf{y}') = \Pr_{\mathbf{x}'}[\mathbf{y}' \mid (U, V, T)]$.

- (d) We simply have that

$$\sum_{\mathbf{y}, \mathbf{y}' \in \text{SEQ}_U} w(\mathbf{y}, \mathbf{y}') = \sum_{\mathbf{y} \in \text{SEQ}_U} \sum_{k \in [\ell]} w(\mathbf{y}, \mathbf{y}_{1 \leftrightarrow k}) = \sum_{\mathbf{y} \in \text{SEQ}_U} \Pr_{\mathbf{x}}[\mathbf{y} \mid (U, V, T)] = 1.$$

All conditions hold for the (0,0)-matching, so we know that $\mathbf{M}'_1$ is (0,0)-NP.

□

C Discussion: Generalizing to Uncountable Sample Spaces

In this section, we describe how Lemma 4.1 may be generalized to uncountable sample spaces. We show how Lemma 4.1 can be derived from a couple of conjectures that can be viewed as uncountable variants of some well-known primal-dual results on bipartite matching and vertex cover. Proving these conjectures is an interesting open question.

We first introduce generalized bipartite matching on measure spaces. Note that we use notation from measure theory that can capture both discrete and uncountable spaces.

Definition C.1 (Generalized Bipartite Matching). Consider the following bipartite matching problem.

Bipartite graph. Suppose (Ω_1, μ_1) and (Ω_2, μ_2) are measure spaces, each of which represents a vertex set in a bipartite graph. Suppose $\mathcal{E} \subseteq \Omega_1 \times \Omega_2$ is a (measurable) subset representing the edges of the bipartite graph, where for each $x \in \Omega_1$, $\mathcal{E}(x) := \{y \in \Omega_2 : (x, y) \in \mathcal{E}\}$ is a measurable subset denoting the *neighbors* of x in Ω_2 , and $\mathcal{E}(y) \subseteq \Omega_1$ is defined analogously for $y \in \Omega_2$.

Vertex capacities. Let $p : \Omega_1 \rightarrow \mathbb{R}_+$ and $q : \Omega_2 \rightarrow \mathbb{R}_+$ be measurable functions denoting vertex capacities. We consider the special case that $\int_{x \in \Omega_1} p(x) d\mu_1(x) < +\infty$ and $\int_{y \in \Omega_2} q(y) d\mu_2(y) < +\infty$.

Feasible matching. A feasible matching is a measurable function $w : \mathcal{E} \rightarrow \mathbb{R}_+$ satisfying the vertex capacity constraints:

- For all $x \in \Omega_1$, $\int_{y \in \mathcal{E}(x)} w(x, y) d\mu_2(y) \leq p(x)$.
- For all $y \in \Omega_2$, $\int_{x \in \mathcal{E}(y)} w(x, y) d\mu_1(x) \leq q(y)$.

Maximum matching. The goal is to “find” a feasible matching w that maximizes the objective function:

$$\int_{(x,y) \in \mathcal{E}} w(x, y) d\mu_1(x) d\mu_2(y).$$

Interpreting Lemma 4.1 with Definition C.1. Definition C.1 is related to the second statement in Lemma 4.1. The edges in $\mathcal{E}$ corresponds to neighboring elements in the relation $\sim$. For the discrete spaces, the capacity for a vertex a in the first space is $p(a) = \Pr[A = a]$, and the capacity for b in the second space is $q(b) = e^\epsilon \cdot \Pr[B = b]$. Statement 2(d) means that there is a feasible matching w with objective value at least $1 - \delta$.

As in the finite case, the dual problem for Definition C.1 is the generalized vertex cover problem.

Definition C.2 (Generalized Vertex Cover Problem). Consider the bipartite graph in Definition C.1 between measure space (Ω_1, μ_1) and (Ω_2, μ_2) with edges in $\mathcal{E} \subseteq \Omega_1 \times \Omega_2$.

Vertex costs. The measurable functions $p : \Omega_1 \rightarrow \mathbb{R}_+$ and $q : \Omega_2 \rightarrow \mathbb{R}_+$ are interpreted as vertex costs.

Feasible vertex cover. A feasible vertex cover consists of a pair of measurable functions $\alpha : \Omega_1 \rightarrow [0, 1]$ and $\beta : \Omega_2 \rightarrow [0, 1]$ such that for all $(x, y) \in \mathcal{E}$, $\alpha(x) + \beta(y) \geq 1$.

Minimum vertex cover. The goal is to “find” a feasible vertex cover (α, β) that minimizes the cost objective:

$$\int_{x \in \Omega_1} \alpha(x) p(x) d\mu_1(x) + \int_{y \in \Omega_2} \beta(y) q(y) d\mu_2(y).$$

One can easily check that weak duality holds for the problems in Definitions C.1 and C.2. Observe that $w \equiv 0$ is feasible for the primal problem and $\alpha \equiv \beta \equiv 1$ is feasible for the dual problem. The first conjecture is that strong duality also holds.

Conjecture C.3 (Strong Duality). *There exist primal and dual feasible solutions in Definitions C.1 and C.2 with the same objective value, and hence, each is an optimum solution in its corresponding problem.*

Ease of proof. We actually believe that Conjecture C.3 is more like a claim, because strong duality for finite LP can be proved via Farkas' Lemma, which can be derived from the hyperplane separation theorem. Hence, it should be a fairly straightforward exercise to apply the corresponding Hahn-Banach separation theorem to general spaces, where one might need to add appropriate assumptions on the measure spaces (Ω_1, μ_1) and (Ω_2, μ_2) .

Conjecture C.4 (Integrality Gap for Vertex Cover). *For the vertex cover problem in Definition C.2, there exists an integral feasible solution $\alpha : \Omega_1 \rightarrow \{0, 1\}$ and $\beta : \Omega_2 \rightarrow \{0, 1\}$ that attains the minimum.*

Ease of proof. If $\mathcal{E}$ is finite, then this is a well-known result by considering the vertex cover polytope for bipartite graphs. However, we are not aware of any result in the literature that can readily extend the result to general uncountable bipartite graphs.

We next show that Lemma 4.1 is an easy corollary from Conjectures C.3 and C.4.

Deriving Lemma 4.1. We consider the more difficult direction of statement 1 implies statement 2. Proceeding with proof by contradiction, we assume that the primal objective of the maximum matching is strictly less than $1 - \delta$. By strong duality from Conjecture C.3, there exists a feasible dual solution whose objective is also strictly less than $1 - \delta$.

By Conjecture C.4, we may assume that the dual solution (α, β) is integral. Define $S := \{x \in \Omega_1 : \alpha(x) = 0\}$ and $T := \{y \in \Omega_2 : \beta(y) = 1\}$. Since (α, β) is a feasible vertex cover, it follows that all the neighbors in S are contained in T .

Recall that the cost functions satisfy $p(a) = \Pr[A = a]$ and $q(b) = e^\epsilon \Pr[B = b]$ (which are interpreted as probability density functions in the continuous case). Hence, the objective value of (α, β) is:

$$\Pr[A \notin S] + e^\epsilon \cdot \Pr[B \in T] \geq 1 - \Pr[A \in S] + e^\epsilon \cdot \Pr[B \in \mathcal{N}(S)].$$

Finally, since the value of (α, β) is also strictly less than $1 - \delta$, we immediately have

$$\Pr[A \in S] > e^\epsilon \cdot \Pr[B \in \mathcal{N}(S)] + \delta, \text{ obtaining the required contradiction.}$$