

Rényi Differential Privacy

Ilya Mironov
Google

Abstract—We propose a natural relaxation of differential privacy based on the Rényi divergence. Closely related notions have appeared in several recent papers that analyzed composition of differentially private mechanisms. We argue that the useful analytical tool can be used as a privacy definition, compactly and accurately representing guarantees on the tails of the privacy loss.

We demonstrate that the new definition shares many important properties with the standard definition of differential privacy, while additionally allowing tighter analysis of composite heterogeneous mechanisms.

I. INTRODUCTION

Differential privacy, introduced by Dwork et al. [DMNS06], has been embraced by multiple research communities as a commonly accepted notion of privacy for algorithms on statistical databases. As applications of differential privacy begin to emerge, practical concerns of *tracking* and *communicating* privacy guarantees are coming to the fore.

Informally, differential privacy bounds a shift in the output distribution of a randomized algorithm that can be induced by a small change in its input. The standard definition of ϵ -differential privacy puts a multiplicative bound on the worst-case change in the distribution's density.

Several relaxations of differential privacy explored other measures of closeness between two distributions. The most common such relaxation, the (ϵ, δ) definition, has been a method of choice for evaluating a variety of differentially private algorithms, especially those that rely on the Gaussian additive noise mechanism or whose analysis follows from composition theorems. The additive δ parameter allows suppressing the long tails of the mechanism's distribution where pure ϵ -differential privacy guarantees may not hold.

Compared to the standard definition, (ϵ, δ) -differential privacy offers asymptotically smaller cumulative loss under composition and allows greater flexibility in the choice of privacy-preserving mechanisms.

Despite its notable advantages and numerous applications, the definition of (ϵ, δ) -differential privacy is an imperfect fit for its two most common use cases: the

Gaussian mechanism and a composition rule. We briefly sketch them here and elaborate on these points in the next section.

The first application of (ϵ, δ) -differential privacy was analysis of the Gaussian noise mechanism. In contrast with the Laplace mechanism, whose privacy guarantee is characterized tightly and accurately by ϵ -differential privacy, a single Gaussian mechanism satisfies a *curve* of $(\epsilon(\delta), \delta)$ -differential privacy definitions. Picking any one point on this curve provides a weak lower bound on the mechanism's actual guarantees.

The second common use of (ϵ, δ) -differential privacy is due to applications of advanced composition theorems. The central feature of ϵ -differential privacy is that it is closed under composition; moreover, the ϵ 's of composed mechanisms simply add up, which motivates the concept of a *privacy budget*. By relaxing the guarantee to (ϵ, δ) -differential privacy, advanced composition allows qualitatively incomparable but quantitatively tighter analyses for compositions of (pure) differentially private mechanisms. Iterating this process, however, quickly leads to a combinatorial explosion of parameters, as each application of an advanced composition theorem leads to a wide selection of possibilities for $(\epsilon(\delta), \delta)$ -differentially private guarantees.

In part to address the shortcomings of (ϵ, δ) -differential privacy, several recent works, surveyed in the next section, explored the use of higher-order *moments* as a way of bounding the tails of the privacy loss variable.

Inspired by these theoretical results and their applications, we propose *Rényi differential privacy* as a natural relaxation of differential privacy that is well-suited for expressing guarantees of privacy-preserving algorithms and for composition of heterogeneous mechanisms. Compared to (ϵ, δ) -differential privacy, Rényi differential privacy is a strictly stronger privacy definition. It offers an operationally convenient and quantitatively accurate way of tracking cumulative privacy loss throughout execution of a standalone differentially private mechanism and across many such mechanisms.

The paper presents a self-contained exposition of the new definition, unifying current literature and demon-

strating its applications.

II. DIFFERENTIAL PRIVACY AND ITS FLAVORS

ϵ -DIFFERENTIAL PRIVACY [DMNS06]. We first recall the standard definition of ϵ -differential privacy.

Definition 1 (ϵ -DP). A randomized mechanism $f: \mathcal{D} \mapsto \mathcal{R}$ satisfies ϵ -differential privacy (ϵ -DP) if for any adjacent $D, D' \in \mathcal{D}$ and $S \subset \mathcal{R}$

$$\Pr[f(D) \in S] \leq e^\epsilon \Pr[f(D') \in S].$$

The above definition is contingent on the notion of *adjacent* inputs D and D' , which is domain-specific and is typically chosen to capture the contribution to the mechanism's input by a single individual.

The *Laplace* mechanism is a prototypical instantiation of ϵ -differential privacy, allowing release of an approximate (noisy) answer to an arbitrary query with values in $\mathbb{R}^n$. The mechanism is defined as

$$\mathbf{L}_\epsilon f(x) = f(x) + \Lambda(0, \Delta_1 f / \epsilon),$$

where Λ is the Laplace distribution and ℓ_1 -sensitivity of the query f is

$$\Delta_1 f = \max_{D, D'} \|f(D) - f(D')\|_1$$

taken over all adjacent inputs D and D' .

The basic composition theorem states that if f and g are, respectively, ϵ_1 - and ϵ_2 -DP, then the simultaneous release of $f(D)$ and $g(D)$ satisfies $(\epsilon_1 + \epsilon_2)$ -DP. Moreover, the mechanism g may be selected adaptively, after seeing the output of $f(D)$.

(ϵ, δ) -DIFFERENTIAL PRIVACY [DKM⁺06]. A relaxation of ϵ -differential privacy allows a δ additive term in its defining inequality:

Definition 2 $((\epsilon, \delta)$ -DP). A randomized mechanism $f: \mathcal{D} \mapsto \mathcal{R}$ offers (ϵ, δ) -differential privacy if for any adjacent $D, D' \in \mathcal{D}$ and $S \subset \mathcal{R}$

$$\Pr[f(D) \in S] \leq e^\epsilon \Pr[f(D') \in S] + \delta.$$

The common interpretation of (ϵ, δ) -DP is that it is ϵ -DP "except with probability δ ". Formalizing this statement runs into difficulties similar to the ones addressed by Mironov et al. [MPRV09] for a different relaxation. For any two adjacent inputs, D_1 and D_2 , it is indeed possible to define an ϵ -DP mechanism that agrees with f with all but δ probability. Extending this argument to domains of exponential sizes (for instance, to a boolean hypercube) cannot be done without diluting the

guarantee exponentially [De12]. We conclude that (ϵ, δ) -differential privacy is a *qualitatively* different definition than pure ϵ -DP (unless, of course, $\delta = 0$, which we assume not to be the case through the rest of this section).

Even for the simple case of exactly two feasible databases, the δ additive factor encompasses two very different failure modes. In one, ϵ -DP holds with probability $1 - \delta$, while with probability δ privacy degrades gracefully (e.g., to ϵ_1 -DP with probability $\delta/2$, to ϵ_2 -DP with probability $\delta/4$, etc.). In the other mode, with probability δ the secret becomes completely exposed. The difference between the two failure modes can be quite stark. In the former, there is always some residual deniability; in the latter, the adversary occasionally learns the secret with certainty. Depending on the adversary's tolerance to false positives, plausible deniability may offer adequate protection, but a single (ϵ, δ) -DP privacy statement cannot differentiate between the two alternatives. For a lively parable of the different guarantees offered by ϵ -DP and (ϵ, δ) -DP definitions see McSherry [McS17].

To avoid the worst-case scenario of always violating privacy of a δ fraction of the dataset, the standard recommendation is to choose $\delta = o(1/N)$ or even $\delta = \text{negl}(1/N)$, where N is the number of contributors. This strategy avoids one particularly devastating outcome, but other possibilities of a complete collapse of privacy remain.

The definition of (ϵ, δ) -differential privacy was initially proposed to capture privacy guarantees of the Gaussian mechanism, defined as follows:

$$\mathbf{G}_\sigma f(x) = f(x) + N(0, \sigma^2).$$

Elementary analysis shows that the Gaussian mechanism cannot meet ϵ -DP for any ϵ . Instead, it satisfies a continuum of incomparable (ϵ, δ) -DP guarantees, for all combinations of $\epsilon < 1$ and $\sigma > \sqrt{2 \ln 1.25 / \delta} \Delta_2 f / \epsilon$, where f 's ℓ_2 -sensitivity is defined as

$$\Delta_2 f = \max_{D, D'} \|f(D) - f(D')\|_2$$

taken over all adjacent inputs D and D' .

There are good reasons for choosing the Gaussian mechanism over Laplace: the noise comes from the same Gaussian distribution (closed under addition) as the error that may already be present in the dataset; the variance of the noise is proportional to the query's ℓ_2 sensitivity, which is less than ℓ_1 . Unfortunately, distilling the guarantees of the Gaussian mechanism down to a single number or a small set of numbers using the language of (ϵ, δ) -DP always leaves a possibility of a

complete privacy compromise (that the mechanism itself does not allow).

Another common reason for bringing in (ϵ, δ) -differential privacy is application of advanced composition theorems. Consider the case of k -fold adaptive composition of an (ϵ, δ) -DP mechanism. For any $\delta' > 0$ it holds that the composite mechanism is $(\epsilon', k\delta + \delta')$ -DP, where $\epsilon' = \sqrt{2k \ln(1/\delta')} \epsilon + k\epsilon(e^\epsilon - 1)$ [DRV10]. Note that, similarly to our discussion of the Gaussian mechanism, a single mechanism satisfies a continuum of incomparable (ϵ, δ) -DP guarantees.

Kairouz et al. give a procedure for computing an optimal k -fold composition of an (ϵ, δ) -DP mechanism [KOV15]. Murtagh and Vadhan [MV16] demonstrate that generalizing this result to composition of heterogeneous mechanisms (i.e., satisfying (ϵ_i, δ_i) -DP for different ϵ_i 's) is #P-hard; they describe a PTAS for an approximate solution. None of these works tackle the problem of mechanisms that satisfy several (ϵ, δ) -DP guarantees simultaneously.

(ZERO)-CONCENTRATED DIFFERENTIAL PRIVACY AND THE MOMENTS ACCOUNTANT. The closely related work by Dwork and Rothblum [DR16], followed by Bun and Steinke [BS16], explore privacy definitions—Concentrated Differential Privacy and zero-Concentrated Differential Privacy—that are framed using the language of, respectively, subgaussian tails and the Rényi divergence. The main difference between our approaches is that both Concentrated and zero-Concentrated DP require a bound on *all* positive moments of a privacy loss variable. In contrast, our definition applies to one moment at a time. Although the point-wise bound is weaker, it allows finer-grained analyses.

The work by Abadi et al. [ACG⁺16] on differentially private training of a neural network introduced the *moments accountant* as an internal tool for tracking privacy loss across multiple invocations of the Gaussian mechanism applied to random samples of the training dataset. The paper's results are expressed via a necessarily lossy translation of the accountant's output (bounds on select moments of the privacy loss variable) to the language of (ϵ, δ) -differential privacy.

Taken together, the works on Concentrated DP, zero-Concentrated DP, and the moments accountant point towards accepting Rényi differential privacy as an effective and flexible mechanism for capturing privacy guarantees of a wide variety of algorithms and their combinations.

OTHER RELAXATIONS. We briefly mention two more

relaxations of differential privacy.

Under the indistinguishability-based Computational Differential Privacy (IND-CDP) definition [MPRV09], the test of closeness between distributions on adjacent inputs is computationally bounded (all other definitions considered in this paper hold against an unbounded, information-theoretic adversary). The IND-CDP notion allows much more accurate functionalities in the two-party setting [MMP⁺10]; in the traditional client-server setup there is a natural class of functionalities where the gap between IND-CDP and (ϵ, δ) -DP is minimal [GKY11], and there are (contrived) examples where the computational relaxation permits tasks that are infeasible under information-theoretic definitions [BCV16].

The Pufferfish framework [KM14] applies Bayesian semantics to define privacy guarantees that hold against adversaries with some restricted class of priors. Its notion of similarity between distributions on the inputs that must be protected is the same as in ϵ -DP.

III. RÉNYI DIFFERENTIAL PRIVACY

We describe a generalization of the notion of differential privacy based on the concept of the Rényi divergence. Connection between the two notions has been pointed out before (mostly for one extreme order, known as Kullback-Leibler divergence [DRV10], [DJW13]); our contribution is in systematically exploring the relationship and its practical implications.

The (parameterized) Rényi divergence is classically defined as follows [R61]:

Definition 3 (Rényi divergence). For two probability distributions P and Q defined over $\mathcal{R}$, the Rényi divergence of order $\alpha > 1$ is

$$D_\alpha(P\|Q) = \frac{1}{\alpha - 1} \log \mathbb{E}_{x \sim Q} \left(\frac{P(x)}{Q(x)} \right)^\alpha.$$

(All logarithms are natural.)

For the endpoints of the interval $(1, \infty)$ the Rényi divergence is defined by continuity. Concretely, $D_1(P\|Q)$ is set to be $\lim_{\alpha \rightarrow 1} D_\alpha(P\|Q)$ and can be verified to be equal to Kullback-Leibler divergence (also known as relative entropy):

$$D_1(P\|Q) = \mathbb{E}_{x \sim P} \log \frac{P(x)}{Q(x)}.$$

It is possible, though, that $D_1(P\|Q)$ thus defined is finite whereas $D_\alpha(P\|Q) = +\infty$ for all $\alpha > 1$.

Likewise,

$$D_\infty(P\|Q) = \max_{x \in \text{supp } Q} \log \frac{P(x)}{Q(x)}.$$

For completeness, we reproduce in the Appendix properties of the Rényi divergence important to the sequel: non-negativity, monotonicity, probability preservation, and a weak triangle inequality (Propositions 8–11).

The relationship between the Rényi divergence with $\alpha = \infty$ and differential privacy is immediate. A randomized mechanism f is ϵ -differentially private if and only if its distribution over any two adjacent inputs D and D' satisfies

$$D_\infty(f(D) \| f(D')) \leq \epsilon.$$

It motivates exploring a relaxation of differential privacy based on the Rényi divergence.

Definition 4 ((α, ϵ) -RDP). A randomized mechanism $f: \mathcal{D} \mapsto \mathcal{R}$ is said to have ϵ -Rényi differential privacy of order α , or (α, ϵ) -RDP for short, if for any adjacent $D, D' \in \mathcal{D}$ it holds that

$$D_\alpha(f(D) \| f(D')) \leq \epsilon.$$

Remark 1. Similarly to the definition of differential privacy, a finite value for ϵ -RDP implies that feasible outcomes of $f(D)$ for *some* $D \in \mathcal{D}$ are feasible, i.e., have a non-zero probability of being observed, for *all* inputs from $\mathcal{D}$. Assuming that this is the case, we let the event space be the support of the distribution.

Remark 2. The Rényi divergence can be defined for α smaller than 1, including negative orders. We are not using these orders in our definition of Rényi differential privacy.

The standard definition of differential privacy has been successful as a privacy measure because it simultaneously meets several important criteria. We verify that the relaxed definition inherits many of the same properties. The results of this section are summarized in Table I.

“BAD OUTCOMES” GUARANTEE. A privacy definition is only as useful as its guarantee for data contributors. The simplest such assurance is the “bad outcomes” interpretation. Consider a person, concerned about some adverse consequences, deliberating whether to withhold her record from the database. Let us say that some outputs of the mechanism are labeled as “bad.” The differential privacy guarantee asserts that the probability of observing a bad outcome will not change (either way) by more than a factor of e^ϵ whether anyone’s record is part of the input or not (for appropriately defined “adjacent” inputs). This is an immediate consequence of the definition of differential privacy, where the subset S is the union of bad outcomes.

This guarantee is relaxed for Rényi differential privacy. Concretely, if f is (α, ϵ) -RDP, then by Proposition 10:

$$\begin{aligned} e^{-\epsilon} \Pr[f(D') \in S]^{\alpha/(\alpha-1)} &\leq \Pr[f(D) \in S] \\ &\leq (e^\epsilon \Pr[f(D') \in S])^{(\alpha-1)/\alpha}. \end{aligned}$$

We discuss consequences of this relaxation in Section VII.

ROBUSTNESS TO AUXILIARY INFORMATION. Critical to the adoption of differential privacy as an operationally useful definition is its lack of assumptions on the adversary’s knowledge. More formally, the property is captured by the Bayesian interpretation of privacy guarantees, which compares the adversary’s prior with the posterior.

Assume that the adversary has a prior $p(D)$ over the set of possible inputs $D \in \mathcal{D}$, and observes an output X of an ϵ -differentially private mechanism f . Its posterior satisfies the following guarantee for all pairs of adjacent inputs $D, D' \in \mathcal{D}$ and all $X \in \mathcal{R}$:

$$\frac{p(D|X)}{p(D'|X)} \leq e^\epsilon \frac{p(D)}{p(D')}.$$

In other words, evidence obtained from an ϵ -differentially private mechanism does not move the relative probabilities assigned to adjacent inputs (the odds ratio) by more than e^ϵ .

The guarantee implied by RDP is a probabilistic statement about the change in the posterior. Define the random variable $R(D, D')$ as

$$R(D, D') \sim \frac{p(D|X)}{p(D'|X)} = \frac{p(X|D) \cdot p(D)}{p(X|D') \cdot p(D')}, \quad \text{where } X \sim f(D).$$

It follows immediately from definition that the Rényi divergence of order α between $P = f(D)$ and $Q = f(D')$ bounds the $(\alpha - 1)$ -th moment of the *change* in R :

$$\begin{aligned} \mathbb{E}_P \left[\left\{ \frac{R_{\text{post}}(D, D')}{R_{\text{prior}}(D, D')} \right\}^{\alpha-1} \right] &= \mathbb{E}_Q [P(x)^\alpha Q(x)^{1-\alpha}] = \\ &= \exp[(\alpha - 1)D_\alpha(f(D) \| f(D'))]. \end{aligned}$$

By taking the logarithm of both sides and applying Jensen’s inequality we obtain that

$$\mathbb{E}_{f(D)} [\log R_{\text{post}}(D, D') - \log R_{\text{prior}}(D, D')] \leq D_\alpha(f(D) \| f(D')). \quad (1)$$

(This can also be derived by observing that

$$\mathbb{E}_{f(D)} [\log R_{\text{post}}(D, D') - \log R_{\text{prior}}(D, D')] = D_1(f(D) \| f(D'))$$

and by monotonicity of the Rényi divergence.)

Compare (1) with the guarantee of pure differential privacy, which states that $\log R_{\text{post}}(D, D') - \log R_{\text{prior}}(D, D') \leq \epsilon$ everywhere, not just in expectation.

POST-PROCESSING. A privacy guarantee that can be diminished by manipulating output is unlikely to be useful. Consider a randomized mapping $g: \mathcal{R} \mapsto \mathcal{R}'$. We observe that $D_\alpha(P \| Q) \geq D_\alpha(g(P) \| g(Q))$ by the analogue of the data processing inequality [vEH14, Theorem 9]. It means that if $f(\cdot)$ is (α, ϵ) -RDP, so is $g(f(\cdot))$. In other words, Rényi differential privacy is preserved by post-processing.

PRESERVATION UNDER ADAPTIVE SEQUENTIAL COMPOSITION. The property that makes possible modular construction of differentially private algorithms is self-composition: if $f(\cdot)$ is ϵ_1 -differentially private and $g(\cdot)$ is ϵ_2 -differentially private, then simultaneous release of $f(D)$ and $g(D)$ is $\epsilon_1 + \epsilon_2$ -differentially private. The guarantee even extends to when g is chosen *adaptively* based on f 's output: if g is indexed by elements of $\mathcal{R}$ and $g_X(\cdot)$ is ϵ_2 -differentially private for any $X \in \mathcal{R}$, then publishing (X, Y) , where $X \leftarrow f(D)$ and $Y \leftarrow g_X(D)$, is $\epsilon_1 + \epsilon_2$ -differentially private.

We prove a similar statement for the composition of two RDP mechanisms.

Proposition 1. *Let $f: \mathcal{D} \mapsto \mathcal{R}_1$ be (α, ϵ_1) -RDP and $g: \mathcal{R}_1 \times \mathcal{D} \mapsto \mathcal{R}_2$ be (α, ϵ_2) -RDP, then the mechanism defined as (X, Y) , where $X \leftarrow f(D)$ and $Y \leftarrow g(X, D)$, satisfies $(\alpha, \epsilon_1 + \epsilon_2)$ -RDP.*

Proof. Let $h: \mathcal{D} \mapsto \mathcal{R}_1 \times \mathcal{R}_2$ be the result of running f and g sequentially. We write X , Y , and Z for the distributions $f(D)$, $g(X, D)$, and the joint distribution $(X, Y) = h(D)$. X' , Y' , and Z' are similarly defined if

the input is D' . Then

$$\begin{aligned} & \exp [(\alpha - 1) D_\alpha(h(D) \| h(D'))] \\ &= \int_{\mathcal{R}_1 \times \mathcal{R}_2} Z(x, y)^\alpha Z'(x, y)^{1-\alpha} dx dy \\ &= \int_{\mathcal{R}_1} \int_{\mathcal{R}_2} (X(x)Y(x, y))^\alpha (X'(x)Y'(x, y))^{1-\alpha} dy dx \\ &= \int_{\mathcal{R}_1} X(x)^\alpha X'(x)^{1-\alpha} \left\{ \int_{\mathcal{R}_2} Y(x, y)^\alpha Y'(x, y)^{1-\alpha} dy \right\} dx \\ &\leq \int_{\mathcal{R}_1} X(x)^\alpha X'(x)^{1-\alpha} dx \cdot \exp((\alpha - 1)\epsilon_2) \\ &\leq \exp((\alpha - 1)\epsilon_1) \exp((\alpha - 1)\epsilon_2) \\ &= \exp((\alpha - 1)(\epsilon_1 + \epsilon_2)), \end{aligned}$$

from which the claim follows. $\square$

Significantly, the guarantee holds whether the releases of f and g are coordinated or not, or computed over the same or different versions of the input dataset. It allows us to operate with a well-defined notion of a *privacy budget* associated with an individual, which is a finite resource consumed with each differentially private data release.

Extending the concept of the privacy budget, we say that the Rényi differential privacy has a *budget curve* parameterized by the order α . We present examples illustrating this viewpoint in Section VI.

GROUP PRIVACY. Although the definition of differential privacy constraints a mechanism's outputs on pairs of *adjacent* inputs, its guarantee extends, in a progressively weaker form, to inputs that are farther apart. This property has two important consequences. First, the differential privacy guarantee degrades gracefully if our assumptions about one person's influence on the input are (somewhat) wrong. For example, a single family contributing to a survey will likely share many socio-economic, demographic, and health characteristics. Rather than collapsing, the differential privacy guarantee will scale down linearly with the number of family members. Second, the group privacy property allows pre-processing input into a differentially private mechanism, possibly amplifying (in a controlled fashion) one record's impact on the output of the computation.

We define group privacy using a notion of c -stable transformation [McS09]. We say that $g: \mathcal{D} \mapsto \mathcal{D}'$ is c -stable if $g(A)$ and $g(B)$ are adjacent in $\mathcal{D}'$ implies that there exists a sequence of length $c + 1$ so that $D_0 = A, \dots, D_c = B$ and all (D_i, D_{i+1}) are adjacent in $\mathcal{D}$.

The standard notion of differential privacy satisfies the following. If f is ϵ -differentially private and $g: \mathcal{D}' \mapsto \mathcal{D}$

is c -stable, then $f \circ g$ is $c\epsilon$ -differentially private. A similar statement holds for Rényi differential privacy.

Proposition 2. *If $f: \mathcal{D} \mapsto \mathcal{R}$ is (α, ϵ) -RDP, $g: \mathcal{D}' \mapsto \mathcal{D}$ is 2^c -stable and $\alpha \geq 2^{c+1}$, then $f \circ g$ is $(\alpha/2^c, 3^c\epsilon)$ -RDP.*

Proof. We prove the statement for $c = 1$, the rest follows by induction.

Define $h = f \circ g$. Since g is 2-stable, it means that for any adjacent $D, D' \in \mathcal{D}'$ there exist $A \in D$, so that $g(D)$ and A , A and $g(D')$ are adjacent in D .

By Corollary 4 and monotonicity of the Rényi divergence, we have that $h = f \circ g$ satisfies

$$D_{\alpha/2}(h(D) \| h(D')) \leq \frac{\alpha-1}{\alpha-2} D_{\alpha}(h(D) \| h(A)) + D_{\alpha-1}(h(A) \| h(D')) \leq 3\epsilon.$$

□

IV. RDP AND (ϵ, δ) -DP

As we observed earlier, the definition of ϵ -differential privacy coincides with (∞, ϵ) -RDP. By monotonicity of the Rényi divergence, (∞, ϵ) -RDP implies (α, ϵ) -RDP for all finite α .

In turn, an (α, ϵ) -RDP implies $(\epsilon\delta, \delta)$ -differential privacy for any given probability $\delta > 0$.

Proposition 3 (From RDP to (ϵ, δ) -DP). *If f is an (α, ϵ) -RDP mechanism, it also satisfies $(\epsilon + \frac{\log 1/\delta}{\alpha-1}, \delta)$ -differential privacy for any $0 < \delta < 1$.*

Proof. Take any two adjacent inputs D and D' , and a subset of f 's range S . To show that f is (ϵ', δ) -differentially private, where $\epsilon' = \epsilon + \frac{1}{\alpha-1} \log 1/\delta$, we need to demonstrate that $\Pr[f(D) \in S] \leq e^{\epsilon'} \Pr[f(D') \in S] + \delta$. In fact, we prove a stronger statement that $\Pr[f(D) \in S] \leq \max(e^{\epsilon'} \Pr[f(D') \in S], \delta)$.

Recall that by Proposition 10

$$\Pr[f(D) \in S] \leq \{e^{\epsilon} \Pr[f(D') \in S]\}^{1-1/\alpha}.$$

Denote $\Pr[f(D') \in S]$ by Q and consider two cases.

Case I. $e^{\epsilon} Q > \delta^{\alpha/(\alpha-1)}$. Continuing the above,

$$\begin{aligned} \Pr[f(D) \in S] &\leq \{e^{\epsilon} Q\}^{1-1/\alpha} = e^{\epsilon} Q \cdot \{e^{\epsilon} Q\}^{-1/\alpha} \\ &\leq e^{\epsilon} Q \cdot \delta^{-1/(\alpha-1)} \\ &= \exp\left(\epsilon + \frac{\log 1/\delta}{\alpha-1}\right) \cdot Q. \end{aligned}$$

Case II. $e^{\epsilon} Q \leq \delta^{\alpha/(\alpha-1)}$. This case is immediate since

$$\Pr[f(D) \in S] \leq \{e^{\epsilon} Q\}^{1-1/\alpha} \leq \delta,$$

which completes the proof. □

A more detailed comparison between the notions of RDP and (ϵ, δ) -differential privacy that goes beyond these reductions is deferred to Section VII.

V. ADVANCED COMPOSITION THEOREM

The main thesis of this section is that the Rényi differential privacy curve of a composite mechanism is sufficient to draw non-trivial conclusions about its privacy guarantees, similar to the ones given by other advanced composition theorems, such as Dwork et al. [DRV10] or Kairouz et al. [KOV15]. Although our proof is structured similarly to Dwork et al. (for instance, Lemma 1 is a direct generalization of [DRV10, Lemma III.2]), it is phrased entirely in the language of Rényi differential privacy without making any (explicit) use of probability arguments.

Lemma 1. *If P and Q are such that $D_{\infty}(P \| Q) \leq \epsilon$ and $D_{\infty}(Q \| P) \leq \epsilon$, then for $\alpha \geq 1$*

$$D_{\alpha}(P \| Q) \leq 2\alpha\epsilon^2.$$

Proof. If $\alpha \geq 1 + 1/\epsilon$, then

$$D_{\alpha}(P \| Q) \leq D_{\infty}(P \| Q) = \epsilon \leq (\alpha - 1)\epsilon^2.$$

Consider the case when $\alpha < 1 + 1/\epsilon$. We first observe that for any $x > y > 0$, $\lambda = \log(x/y)$, and $0 \leq \beta \leq 1/\lambda$ the following inequality holds:

$$\begin{aligned} x^{\beta+1}y^{-\beta} + x^{-\beta}y^{\beta+1} &= x \cdot e^{\beta\lambda} + y \cdot e^{-\beta\lambda} \\ &\leq x(1 + \beta\lambda + (\beta\lambda)^2) + y(1 - \beta\lambda + (\beta\lambda)^2) \\ &= (1 + (\beta\lambda)^2)(x + y) + \beta\lambda(x - y). \end{aligned} \quad (2)$$

Since all terms of the right hand side of (2) are positive, the inequality applies if λ is an upper bound on $\log x/y$, which we exploit in the argument below.

$$\begin{aligned} &\exp[(\alpha - 1)D_{\alpha}(P \| Q)] \\ &= \int_{\mathcal{R}} P(x)^{\alpha} Q(x)^{1-\alpha} dx \\ &\leq \int_{\mathcal{R}} \{P(x)^{\alpha} Q(x)^{1-\alpha} + Q(x)^{\alpha} P(x)^{1-\alpha}\} dx - 1 \\ &\quad \text{(by nonnegativity of } D_{\alpha}(Q \| P)) \\ &\leq \int_{\mathcal{R}} \{(1 + (\alpha - 1)^2\epsilon^2)(P(x) + Q(x)) + \\ &\quad (\alpha - 1)\epsilon|P(x) - Q(x)|\} dx - 1 \\ &\quad \text{(by (2) for } \beta = \alpha - 1 \leq 1/\epsilon) \\ &= 1 + 2(\alpha - 1)^2\epsilon^2 + (\alpha - 1)\epsilon\|P - Q\|_1. \end{aligned}$$

Property	Differential Privacy	Rényi Differential Privacy
Change in probability of outcome S	$\Pr[f(D) \in S] \leq e^\epsilon \Pr[f(D') \in S]$ $\Pr[f(D) \in S] \geq e^{-\epsilon} \Pr[f(D') \in S]$	$\Pr[f(D) \in S] \leq (e^\epsilon \Pr[f(D') \in S])^{(\alpha-1)/\alpha}$ $\Pr[f(D) \in S] \geq e^{-\epsilon} \Pr[f(D') \in S]^{\alpha/(\alpha-1)}$
Change in odds between D and D'	$\frac{R_{\text{post}}(D, D')}{R_{\text{prior}}(D, D')} \leq e^\epsilon$ always	$\mathbb{E} \left[\left\{ \frac{R_{\text{post}}(D, D')}{R_{\text{prior}}(D, D')} \right\}^{\alpha-1} \right] \leq \exp[(\alpha-1)\epsilon]$
Change in log odds between D and D'	$ \Delta \log R(D, D') \leq \epsilon$ always	$\mathbb{E}[\Delta \log R(D, D')] \leq \epsilon$
Post-processing	f is ϵ -DP (or (α, ϵ) -RDP) $\Rightarrow g \circ f$ is ϵ -DP (or (α, ϵ) -RDP, resp.)	
Adaptive sequential composition (basic)	f, g are ϵ -DP (or (α, ϵ) -RDP) $\Rightarrow (f, g)$ is 2ϵ -DP (resp., $(\alpha, 2\epsilon)$ -RDP)	
Group privacy, pre-processing	f is ϵ -DP (or (α, ϵ) -RDP), g is 2^c -stable $\Rightarrow f \circ g$ is $2^c\epsilon$ -DP (resp., $(\alpha/2^c, 3^c\epsilon)$ -RDP)	

TABLE I
SUMMARY OF PROPERTIES SHARED BY DIFFERENTIAL PRIVACY AND RDP.

Taking the logarithm of both sides and using that $\log(1+x) < x$ for positive x we find that

$$D_\alpha(P\|Q) \leq 2(\alpha-1)\epsilon^2 + \epsilon\|P-Q\|_1. \quad (3)$$

Observe that

$$\begin{aligned} \|P-Q\|_1 &= \int |P(x) - Q(x)| \, dx \\ &= \int_{\mathcal{R}} \min(P(x), Q(x)) \left| \frac{\max(P(x), Q(x))}{\min(P(x), Q(x))} - 1 \right| \, dx \\ &\leq \min(2, e^\epsilon - 1) \leq 2\epsilon^2. \end{aligned}$$

Plugging the bound on $\|P-Q\|_1$ into (3) completes the proof.

The claim for $\alpha = 1$ follows by continuity. $\square$

Proposition 4. *Let $f: \mathcal{D} \mapsto \mathcal{R}$ be an adaptive composition of n mechanisms all satisfying ϵ -differential privacy. Let D and D' be two adjacent inputs. Then for any $S \subset \mathcal{R}$:*

$$\Pr[f(D) \in S] \leq \exp \left(2\epsilon \sqrt{n \log 1 / \Pr[f(D') \in S]} \right) \cdot \Pr[f(D') \in S].$$

Proof. By applying Lemma 1 to the Rényi differential privacy curve of the underlying mechanisms and Proposition 1 to their composition, we find that for all $\alpha \geq 1$

$$D_\alpha(f(D)\|f(D')) \leq 2\alpha n \epsilon^2.$$

Denote $\Pr[f(D') \in S]$ by Q and consider two cases.

Case I: $\log 1/Q \geq \epsilon^2 n$. Choosing with some foresight $\alpha = \sqrt{\log 1/Q} / (\epsilon \sqrt{n}) \geq 1$, we have by Proposition 10 (probability preservation):

$$\begin{aligned} \Pr[f(D) \in S] &\leq \left\{ \exp[D_\alpha(f(D)\|f(D'))] \cdot Q \right\}^{1-1/\alpha} \\ &\leq \exp(2(\alpha-1)n\epsilon^2) \cdot Q^{1-1/\alpha} \\ &< \exp \left(\epsilon \sqrt{n \log 1/Q} - (\log Q)/\alpha \right) \cdot Q \\ &= \exp \left(2\epsilon \sqrt{n \log 1/Q} \right) \cdot Q. \end{aligned}$$

Case II: $\log 1/Q < \epsilon^2 n$. This case follows trivially, since the right hand side of the claim is larger than 1:

$$\exp \left(2\epsilon \sqrt{n \log 1/Q} \right) \cdot Q \geq \exp(2 \log 1/Q) \cdot Q = 1/Q > 1. \quad \square$$

The notable feature of Proposition 4 is that its privacy guarantee—bounded probability gain—comes in the form that depends on the event’s probability. We discuss this type of guarantee in Section VII.

The following corollary gives a more conventional (ϵ, δ) variant of advanced composition.

Corollary 1. *Let f be the composition of the n ϵ -differentially private mechanisms. Let $0 < \delta < 1$ be such that $\log(1/\delta) \geq \epsilon^2 n$. Then f satisfies (ϵ', δ) -differential privacy where*

$$\epsilon' = 4\epsilon \sqrt{2n \log(1/\delta)}.$$

Proof. Let D and D' be two adjacent inputs, and S be some subset of the range of f . To argue (ϵ', δ) -differential privacy of f , we need to verify that

$$\Pr[f(D) \in S] \leq e^{\epsilon'} \Pr[f(D') \in S] + \delta.$$

In fact, we prove a somewhat stronger statement, namely that $\Pr[f(D) \in S] \leq \max(e^{\epsilon'} \Pr[f(D') \in S], \delta)$.

By Proposition 4

$$\Pr[f(D) \in S] \leq \exp\left(2\epsilon\sqrt{n \log 1/\Pr[f(D') \in S]}\right) \cdot \Pr[f(D') \in S].$$

Denote $\Pr[f(D') \in S]$ by Q and consider two cases.

Case I: $8 \log 1/\delta > \log 1/Q$. Then

$$\begin{aligned} \Pr[f(D) \in S] &\leq \exp\left(2\epsilon\sqrt{n \log 1/Q}\right) \cdot Q \\ &< \exp\left(2\epsilon\sqrt{8n \log 1/\delta}\right) \cdot Q \\ &\quad (\text{by } 8 \log 1/\delta > \log 1/Q) \\ &= \exp(\epsilon') \cdot Q. \end{aligned}$$

Case II: $8 \log 1/\delta \leq \log 1/Q$. Then

$$\begin{aligned} \Pr[f(D) \in S] &\leq \exp\left(2\epsilon\sqrt{n \log 1/Q}\right) \cdot Q \\ &\leq \exp\left(2\sqrt{\log 1/\delta \cdot \log 1/Q}\right) \cdot Q \\ &\quad (\text{since } \log(1/\delta) \geq \epsilon^2 n) \\ &\leq \exp\left(\sqrt{1/2} \log 1/Q\right) \cdot Q \\ &\quad (\text{by } 8 \log 1/\delta \leq \log 1/Q) \\ &= Q^{1-1/\sqrt{2}} \leq Q^{1/8} \\ &\leq \delta. \end{aligned} \quad (\text{ditto})$$

□

Remark 3. The condition $\log(1/\delta) \geq \epsilon^2 n$ corresponds to the so-called “high privacy” regime of the advanced composition theorem [KOV15], where $\epsilon' < (1 + \sqrt{2}) \log(1/\delta)$. Since δ is typically chosen to be small, say, less than 1%, it covers the case of $\epsilon' < 11$. In other words, if $\log(1/\delta) > \epsilon^2 n$, this and other composition theorems are unlikely to yield strong bounds.

VI. BASIC MECHANISMS

In this section we analyze Rényi differential privacy of three basic mechanisms and their self-composition: randomized response, Laplace and Gaussian noise addition. The results are summarized in Table II and plotted for select parameters in Figures 1 and 2.

A. Randomized response

Let f be a predicate, i.e., $f: \mathcal{D} \mapsto \{0, 1\}$. The Randomized Response mechanism for f is defined as

$$\mathbf{RR}_p f(D) = \begin{cases} f(D) & \text{with probability } p \\ 1 - f(D) & \text{with probability } 1 - p \end{cases}.$$

The following statement can be verified by direct application of the definition of Rényi differential privacy:

Proposition 5. *Randomized Response mechanism $\mathbf{RR}_p(f)$ satisfies*

$$\left(\alpha, \frac{1}{\alpha-1} \log(p^\alpha(1-p)^{1-\alpha} + (1-p)^\alpha p^{1-\alpha})\right) \text{-RDP}$$

if $\alpha > 1$, and

$$\left(\alpha, (2p-1) \log \frac{p}{1-p}\right) \text{-RDP}$$

if $\alpha = 1$.

B. Laplace noise

Through the rest of this section we assume that $f: \mathcal{D} \mapsto \mathbb{R}$ is a function of sensitivity 1, i.e., for any two adjacent $D, D' \in \mathcal{D}$: $|f(D) - f(D')| \leq 1$.

Define the Laplace mechanism for f of sensitivity 1 as

$$\mathbf{L}_\lambda f(D) = f(D) + \Lambda(0, \lambda),$$

where $\Lambda(\mu, \lambda)$ is Laplace distribution with mean μ and scale λ , i.e., its probability density function is $\frac{1}{2\lambda} \exp(-|x - \mu|/\lambda)$.

To derive the RDP budget curve for the exponential mechanism we compute the Rényi divergence for Laplace distribution and its offset.

Proposition 6. *For any $\alpha \geq 1$ and $\lambda > 0$:*

$$D_\alpha(\Lambda(0, \lambda) \parallel \Lambda(1, \lambda)) = \frac{1}{\alpha-1} \log \left\{ \frac{\alpha}{2\alpha-1} \exp\left(\frac{\alpha-1}{\lambda}\right) + \frac{\alpha-1}{2\alpha-1} \exp\left(\frac{-\alpha}{\lambda}\right) \right\}.$$

Proof. For continuous distributions P and Q defined over the real interval with densities p and q

$$D_\alpha(P \parallel Q) = \frac{1}{\alpha-1} \log \int_{-\infty}^{\infty} p(x)^\alpha q(x)^{1-\alpha} dx.$$

To compute the integral for $p(x) = \frac{1}{2\lambda} \exp(-|x|/\lambda)$ and $q(x) = \frac{1}{2\lambda} \exp(-|x-1|/\lambda)$, we evaluate it separately

Mechanism	Differential Privacy	Rényi Differential Privacy for α
Randomized Response	$\left \log \frac{p}{1-p} \right $	$\alpha > 1: \frac{1}{\alpha-1} \log (p^\alpha (1-p)^{1-\alpha} + (1-p)^\alpha p^{1-\alpha})$ $\alpha = 1: (2p-1) \log \frac{p}{1-p}$
Laplace Mechanism	$1/\lambda$	$\alpha > 1: \frac{1}{\alpha-1} \log \left\{ \frac{\alpha}{2\alpha-1} \exp(\frac{\alpha-1}{\lambda}) + \frac{\alpha-1}{2\alpha-1} \exp(-\frac{\alpha}{\lambda}) \right\}$ $\alpha = 1: 1/\lambda + \exp(-1/\lambda) - 1 = .5/\lambda^2 + O(1/\lambda^3)$
Gaussian Mechanism	∞	$\alpha/(2\sigma^2)$

TABLE II
SUMMARY OF RDP PARAMETERS FOR BASIC MECHANISMS.

over the intervals $(-\infty, 0]$, $[0, 1]$ and $[1, +\infty]$.

$$\begin{aligned}
& \int_{-\infty}^{+\infty} p(x)^\alpha q(x)^{1-\alpha} dx = \\
& \quad \frac{1}{2\lambda} \int_{-\infty}^0 \exp(\alpha x/\lambda + (1-\alpha)(x-1)/\lambda) dx \\
& \quad + \frac{1}{2\lambda} \int_0^1 \exp(-\alpha x/\lambda + (1-\alpha)(x-1)/\lambda) dx \\
& \quad + \frac{1}{2\lambda} \int_1^{+\infty} \exp(-\alpha x/\lambda - (1-\alpha)(x-1)/\lambda) dx \\
& = \frac{1}{2} \exp((\alpha-1)/\lambda) \\
& \quad + \frac{1}{2(2\alpha-1)} (\exp((\alpha-1)/\lambda) - \exp(-\alpha/\lambda)) \\
& \quad + \frac{1}{2} \exp(-\alpha/\lambda) \\
& = \frac{\alpha}{2\alpha-1} \exp((\alpha-1)/\lambda) + \frac{\alpha-1}{2\alpha-1} \exp(-\alpha/\lambda),
\end{aligned}$$

from which the claim follows. $\square$

Since the Laplace mechanism is additive, the Rényi divergence between $\mathbf{L}_\lambda f(D)$ and $\mathbf{L}_\lambda f(D')$ depends only on α and the distance $|f(D) - f(D')|$. Proposition 6 implies the following:

Corollary 2. *If real-valued function f has sensitivity 1, then the Laplace mechanism $\mathbf{L}_\lambda f$ satisfies $(\alpha, \frac{1}{\alpha-1} \log \left\{ \frac{\alpha}{2\alpha-1} \exp(\frac{\alpha-1}{\lambda}) + \frac{\alpha-1}{2\alpha-1} \exp(-\frac{\alpha}{\lambda}) \right\})$ -RDP.*

Predictably,

$$\begin{aligned}
& \lim_{\alpha \rightarrow +\infty} D_\alpha(\Lambda(0, \lambda) \| \Lambda(1, \lambda)) = \\
& D_\infty(\Lambda(0, \lambda) \| \Lambda(1, \lambda)) = 1/\lambda.
\end{aligned}$$

This is, of course, consistent with the Laplace mechanism satisfying $1/\lambda$ -differential privacy. The other extreme evaluates to the following expression $\lim_{\alpha \rightarrow 1} D_\alpha(\Lambda(0, \lambda) \| \Lambda(1, \lambda)) = 1/\lambda + \exp(-1/\lambda) - 1$, which is well approximated by $.5/\lambda^2$ for large λ .

C. Gaussian noise

Assuming, as before, that f is a real-valued function, the Gaussian mechanism for approximating f is defined as

$$\mathbf{G}_\sigma f(D) = f(D) + N(0, \sigma^2),$$

where $N(0, \sigma^2)$ is normally distributed random variable with standard deviation σ^2 and mean 0.

The following statement is a closed-form expression of the Rényi divergence between a Gaussian and its offset (for a more general version see [vEH14], [LV87]).

Proposition 7. $D_\alpha(N(0, \sigma^2) \| N(\mu, \sigma^2)) = \alpha\mu^2/(2\sigma^2)$.

Proof. By direct computation we verify that

$$\begin{aligned}
& D_\alpha(N(0, \sigma^2) \| N(\mu, \sigma^2)) \\
& = \frac{1}{\alpha-1} \log \int_{-\infty}^{\infty} \frac{1}{\sigma\sqrt{2\pi}} \exp(-\alpha x^2/(2\sigma^2)) \\
& \quad \exp(-(1-\alpha)(x-\mu)^2/(2\sigma^2)) dx \\
& = \frac{1}{\alpha-1} \log \left\{ \frac{\sigma\sqrt{2\pi}}{\sigma\sqrt{2\pi}} \exp[(\alpha^2 - \alpha)\mu^2/(2\sigma^2)] \right\} \\
& = \alpha\mu^2/(2\sigma^2).
\end{aligned}$$

$\square$

The following corollary is immediate:

Corollary 3. *If f has sensitivity 1, then the Gaussian mechanism $\mathbf{G}_\sigma f$ satisfies $(\alpha, \alpha/(2\sigma^2))$ -RDP.*

Observe that the RDP budget curve for the Gaussian mechanism is particularly simple—a straight line (Figure 1). Recall that the (adaptive) composition of RDP mechanisms satisfies Rényi differential privacy with the budget curve that is the sum of the mechanisms' budget curves. It means that a composition of Gaussian mechanisms will behave, privacy-wise, “like” a Gaussian mechanism. Concretely, a composition of n Gaussian mechanisms each with parameter σ will have the RDP curve of a Gaussian mechanism with parameter $\sigma/\sqrt{n}$.

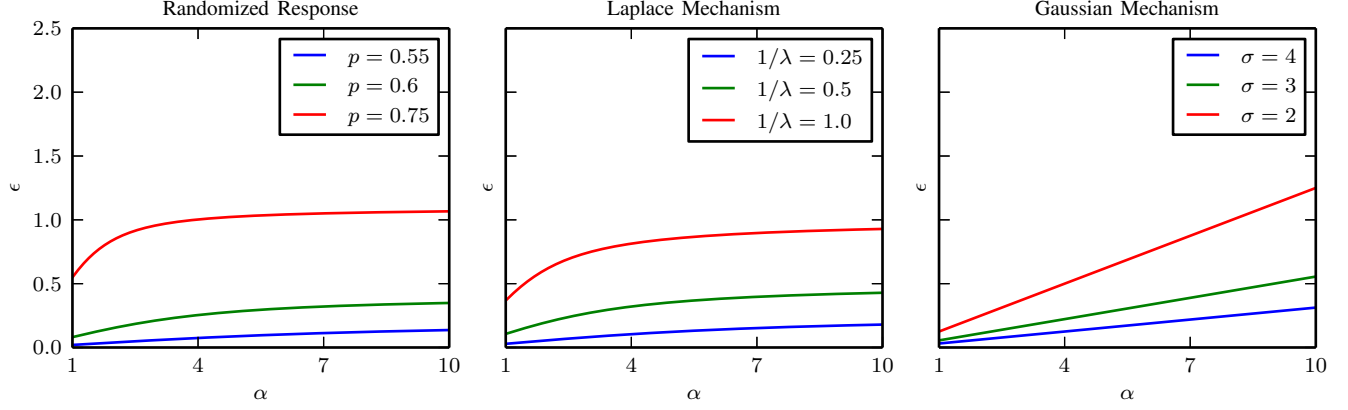


Fig. 1. (α, ϵ) -Rényi differential privacy budget curve for three basic mechanisms with varying parameters.

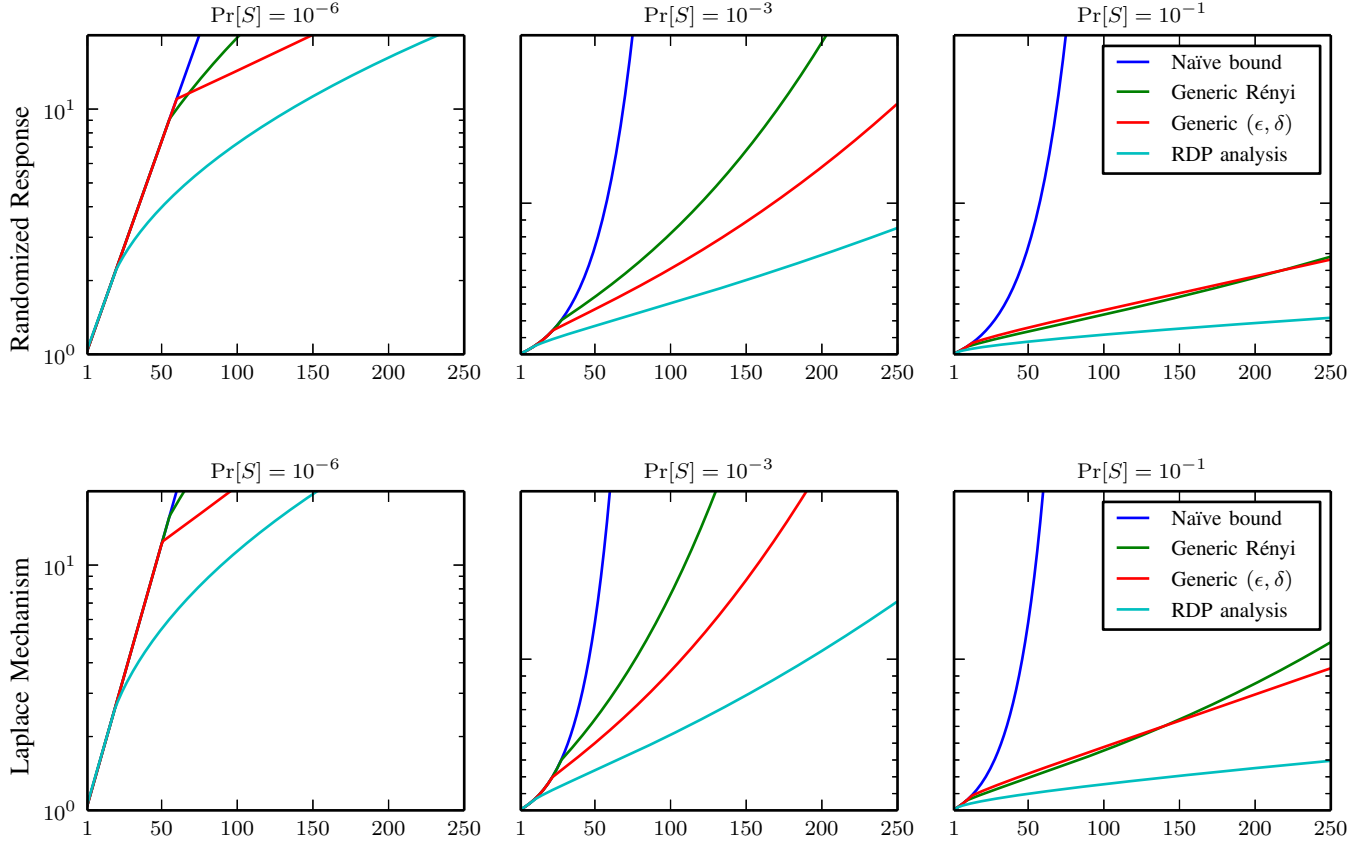


Fig. 2. Various privacy guarantees of the randomized response with parameter $p = 51\%$ (top row) and the Laplace mechanism with parameter $\lambda = 20$ (bottom row) under self-composition. The x -axis is the number of compositions (1–250). The y -axis, in log scale, is the upper bound on the multiplicative increase in probability of event S , where S 's initial mass is either 10^{-6} (left), 10^{-3} (center), or $.1$ (right). The four plot lines are the “naïve” $n\epsilon$ bound (blue); optimal choice (ϵ, δ) in the standard advanced composition theorem (red); generic bound of Proposition 4 (blue); optimal choice of α in Proposition 10 (cyan).

D. Privacy of basic mechanisms under composition

The “bad outcomes” interpretation of Rényi differential privacy ties the probabilities of seeing the same outcome under runs of the mechanism applied to adjacent inputs. The dependency of the upper bound on the increase in probability on its initial value is complex, especially compared to the standard differential privacy guarantee. The main advantage of this more involved analysis is that for most parameters the bound becomes tighter.

In this section we compare numerical bounds for several analyses of self-composed mechanisms (see Figure 2), presented as three sets of graphs, where $\Pr[f(D) \in S]$ takes values 10^{-6} , 10^{-3} , and 10^{-1} .

Each of the six graphs in Figure 2 (three probability values $\times$ {randomized response, Laplace}) plots bounds in logarithmic scale on the relative increase in probability of S (i.e., $\Pr[f(D') \in S] / \Pr[f(D) \in S]$) offered by four analyses. The first, “naïve”, bound follows from the basic composition theorem for differential privacy and, as expected, is very pessimistic for all but a handful of parameters. A tighter, advanced composition theorem [DRV10], gives a choice of δ , from which one computes ϵ' so that the n -fold composition satisfies (ϵ', δ) -differential privacy. The second curve plots the bound for the optimal (tightest) choice of (ϵ', δ) . Two other bounds come from Rényi differential privacy analysis: our generic advanced composition theorem (Proposition 4) and the bound of Proposition 10 for the optimal combination of (α, ϵ) from the RDP curve of the composite mechanism.

Several observations are in order. The RDP-specific analysis for both mechanisms is tighter than all generic bounds whose only input is the mechanism’s differential privacy parameter. On the other hand, our version of the advanced composition bound (Proposition 4) is consistently outperformed by the standard (ϵ, δ) -form of the composition theorem, where δ is chosen *optimally*. We elaborate on this distinction in the next section.

VII. DISCUSSION

Rényi differential privacy is a natural relaxation of the standard notion of differential privacy that preserves many of its essential properties. It can most directly be compared with (ϵ, δ) -differential privacy, with which it shares several important characteristics.

PROBABILISTIC PRIVACY GUARANTEE. The standard “bad outcomes” guarantee of ϵ -differential privacy is independent of the probability of a bad outcome: it

may increase only by a factor of $\exp(\epsilon)$. Its relaxation, (ϵ, δ) -differential privacy, allows for an additional δ term, which allows for a complete privacy compromise with probability δ .

In stark contrast, Rényi differential privacy even with very weak parameters never allows a total breach of privacy with no residual uncertainty. The following analysis quantifies this assurance.

Let f be (α, ϵ) -RDP with $\alpha > 1$. Recall that for any two adjacent inputs D and D' , and an arbitrary prior p the odds function $R(D, D') \sim p(D)/p(D')$ satisfies $\mathbb{E} \left[\left\{ \frac{R_{\text{post}}(D, D')}{R_{\text{prior}}(D, D')} \right\}^{\alpha-1} \right] \leq \exp((\alpha-1)\epsilon)$. By Markov’s inequality $\Pr[R_{\text{post}}(D, D') > \beta R_{\text{prior}}(D, D')] < e^\epsilon / \beta^{1/(\alpha-1)}$. For instance, if $\alpha = 2$, the probability that the ratio between two posteriors increases by more than the β factor drops off as $O(1/\beta)$.

BASELINE-DEPENDENT GUARANTEES. The Rényi differential privacy bound gets weaker for less likely outcomes. For instance, if f is a $(10.0, .1)$ -RDP mechanism, an event of probability .5 under $f(D)$ can be as large as .586 and as small as .419 under $f(D')$. For smaller events the range is (in relative terms) wider. If the probability under $f(D)$ is .001, then $\Pr[f(D') \in S] \in [.00042, 0.00218]$. For $\Pr[f(D) \in S] = 10^{-6}$ the range is wider still: $\Pr[f(D') \in S] \in [.195 \cdot 10^{-6}, 4.36 \cdot 10^{-6}]$.

Compared to the pure ϵ -differential privacy this type of guarantee is conceptually weaker and more onerous in application: in order to decide whether the increased risk is tolerable, one is required to estimate the baseline risk first.

However, in comparison with (ϵ, δ) -differential privacy Rényi differential privacy leads to qualitatively simpler and occasionally quantitatively stronger bounds. The reason is that (ϵ, δ) -differential privacy often arises as a result of some analysis that implicitly comes with an ϵ - δ tradeoff. Finding an optimal value of (ϵ, δ) given the baseline risk may be non-trivial, especially in closed form. Contrast the following two, basically equivalent, statements of advanced composition theorems (Proposition 4 and its Corollary 1):

Let $f: \mathcal{D} \mapsto \mathcal{R}$ be an adaptive composition of n mechanisms all satisfying ϵ -differential privacy for $\epsilon \leq 1$. Let D and D' be two adjacent inputs. Then for any $S \subset \mathcal{R}$:

$$\Pr[f(D') \in S] \leq \exp \left(2\epsilon \sqrt{n \log 1 / \Pr[f(D) \in S]} \right) \cdot \Pr[f(D) \in S].$$

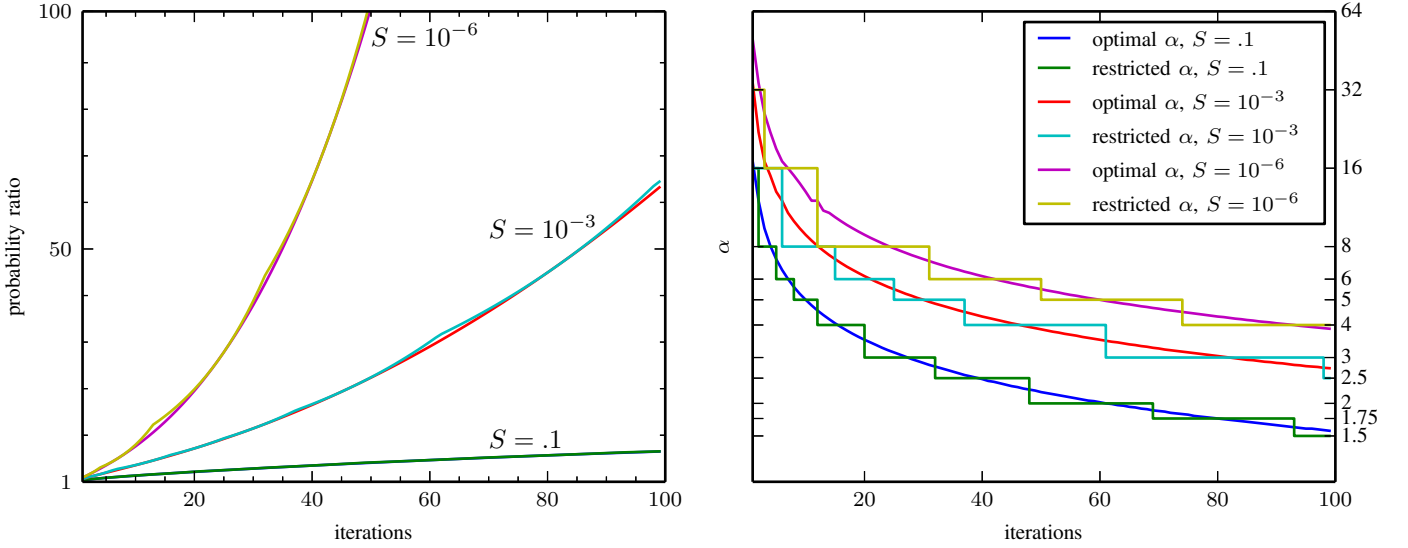


Fig. 3. Left: Bounds on the ratio $\Pr[f(D') \in S] / \Pr[f(D) \in S]$ for $\Pr[f(D) \in S] \in \{.1, 10^{-3}, 10^{-6}\}$ for up to 100 iterations of a mixed mechanism (randomized response with $p = .52$, Laplace with $\lambda = 20$ and Gaussian with $\sigma = 10$). Each bound is computed twice: once for an optimal choice of α and once for α restricted to $\{1.5, 1.75, 2, 2.5, 3, 4, 5, 6, 8, 16, 32, 64, +\infty\}$. The curves for two choices of α are nearly identical. Right: corresponding values of α in log scale.

(by Proposition 4) or

$$\Pr[f(D') \in S] \leq \exp\left(4\epsilon\sqrt{2n\log 1/\delta}\right) \cdot \Pr[f(D) \in S] + \delta,$$

where $0 < \epsilon, \delta < 1$ such that $\log(1/\delta) \geq \epsilon^2 n$ (by Corollary 1).

Given some value of baseline risk $\Pr[f(D) \in S]$, which formulation is easier to interpret? We argue that it is the former, since the (ϵ, δ) form has a free parameter (δ) that ought to be optimized in order to extract a tight bound that Proposition 4 gives directly.

The use of (ϵ, δ) bounds gets even more complex if we consider a composition of heterogeneous mechanisms. It brings us to the last point of comparison between (ϵ, δ) - and Rényi differential privacy measures.

KEEPING TRACK OF ACCUMULATED PRIVACY LOSS. A finite privacy budget associated with an individual is an intuitive and appealing concept, to which ϵ -differential privacy gives a rigorous mathematical expression. Cumulative loss of differential privacy over the course of a mechanism run, a protocol, or one's lifetime can be tracked easily thanks to the additivity property of differential privacy. Unfortunately, doing so naively likely

exaggerates privacy loss, which grows sublinearly in the number of queries with all but negligible probability (via advanced composition theorems).

Critically, applying advanced composition theorems breaks the convenient abstraction of privacy as a non-negative real number. Instead, the guarantee comes in the (ϵ, δ) form that effectively corresponds to a single point on an implicitly defined curve. Composition of multiple, heterogeneous mechanisms makes applying the composition rule optimally much more challenging, as one may choose various (ϵ, δ) points to represent their privacy (in the analysis, not evaluation!). It begs the question of how to represent the privacy guarantee of a complex mechanism: distilling it to a single number throws away valuable information, while publishing the entire (ϵ, δ) curve shifts the problem to the aggregation step. (See Kairouz et al. [KOV15] for an optimal bound on composition of homogeneous mechanisms and Murtagh and Vadhan [MV16] for hardness results and an approximation scheme for composition of mechanisms with heterogeneous privacy guarantees.)

Rényi differential privacy offers an alternative thanks to its composition rule: RDP curves for composed mechanisms simply add up. Importantly, α 's of (α, ϵ) -Rényi differential privacy do not change. If RDP statements

are reported for a common set of α 's (which includes $+\infty$, to keep track of pure differential privacy), RDP of the aggregate is the sum of the reported vectors. Since the composition theorem of Proposition 4 only considers the RDP curve of a mechanism, it means that, qualitatively, the sublinear loss of privacy as the function of the number of queries will be preserved.

For an example of this approach we tabulate the bound on privacy loss for an iterative mechanism consisting of three basic mechanisms: randomized response, Gaussian, and Laplace. Its RDP curve is given, in the closed form, by application of the basic composition rule to RDP curves of the underlying mechanisms (Table II). The privacy guarantee is presented in Figure 3 for three values of the baseline risk: .1, .001, and 10^{-6} . For each set of parameters two curves are plotted: one for an optimal value of α from $(1, +\infty]$, the other for an optimal α restricted to the set of 13 values $\{1.5, 1.75, 2, 2.5, 3, 4, 5, 6, 8, 16, 32, 64, +\infty\}$. The two curves are nearly identical, which illustrates our thesis that reporting RDP curves for a restricted set of α 's preserves tightness of privacy analysis.

VIII. CONCLUSIONS AND OPEN QUESTIONS

We put forth the proposition that Rényi divergence yields useful insight into analysis of differentially private mechanisms. Among our findings

- Rényi differential privacy (RDP) is a natural generalization of pure differential privacy.
- RDP shares, with some adaptations, many properties that make differential privacy a useful and versatile tool.
- RDP analysis of Gaussian noise is particularly simple.
- A composition theorem can be proved based solely on the properties of RDP, which implies that RDP packs sufficient information about a composite mechanism as to enable its analysis without consideration of its components.
- Furthermore, an RDP curve may be sampled in just a few points to provide useful guarantees for a wide range of parameters. If these points are chosen consistently across multiple mechanisms, this information can be used to estimate aggregate privacy loss.

Naturally, multiple questions remain open. Among those

- As Lemma 1 demonstrates, the RDP curve of a differentially private mechanism is severely constrained. Making fuller use of these constraints is

a promising direction, in particular towards formal bounds on tightness of RDP guarantees from select α values.

- Proposition 10 (probability preservation) is not tight when $D_\alpha(P\|Q) \rightarrow 0$. We expect that $P(A) \rightarrow Q(A)$ but the bound does not improve beyond $P(A)^{(\alpha-1)/\alpha}$.

ACKNOWLEDGMENTS

We would like to thank Cynthia Dwork, Kunal Talwar, Salil Vadhan, and Li Zhang for numerous fruitful discussions, and Mark Bun and Thomas Steinke for sharing a draft of [BS16].

REFERENCES

- [ACG⁺16] Martín Abadi, Andy Chu, Ian Goodfellow, H Brendan McMahan, Ilya Mironov, Kunal Talwar, and Li Zhang. Deep learning with differential privacy. In *Proceedings of the 2016 ACM SIGSAC Conference on Computer and Communications Security*, pages 308–318. ACM, 2016.
- [BCV16] Mark Bun, Yi-Hsiu Chen, and Salil P. Vadhan. Separating computational and statistical differential privacy in the client-server model. In Martin Hirt and Adam D. Smith, editors, *Theory of Cryptography—14th International Conference, TCC 2016-B, Part I*, pages 607–634, 2016.
- [BS16] Mark Bun and Thomas Steinke. Concentrated differential privacy: Simplifications, extensions, and lower bounds. In Martin Hirt and Adam D. Smith, editors, *Theory of Cryptography—14th International Conference, TCC 2016-B, Part I*, pages 635–658, 2016.
- [De12] Anindya De. Lower bounds in differential privacy. In Ronald Cramer, editor, *Theory of Cryptography—9th Theory of Cryptography Conference, TCC 2012*, pages 321–338, 2012.
- [DJW13] John C. Duchi, Michael I. Jordan, and Martin J. Wainwright. Local privacy and statistical minimax rates. In *2013 IEEE 54th Annual Symposium on Foundations of Computer Science*, pages 429–438. IEEE, October 2013.
- [DKM⁺06] Cynthia Dwork, Krishnaram Kenthapadi, Frank McSherry, Ilya Mironov, and Moni Naor. Our data, ourselves: Privacy via distributed noise generation. In *Advances in Cryptography—Eurocrypt '06*, pages 486–503. Springer, 2006.
- [DMNS06] Cynthia Dwork, Frank McSherry, Kobbi Nissim, and Adam D. Smith. Calibrating noise to sensitivity in private data analysis. In Shai Halevi and Tal Rabin, editors, *Third Theory of Cryptography Conference, TCC 2006*, pages 265–284. Springer, 2006.
- [DR16] Cynthia Dwork and Guy N. Rothblum. Concentrated differential privacy. *CoRR*, abs/1603.01887, 2016.
- [DRV10] Cynthia Dwork, Guy N. Rothblum, and Salil Vadhan. Boosting and differential privacy. In Luca Trevisan, editor, *51st Annual IEEE Symposium on Foundations of Computer Science, FOCS 2010*, pages 51–60. IEEE, October 2010.

- [GKY11] Adam Groce, Jonathan Katz, and Arkady Yerukhimovich. Limits of computational differential privacy in the client/server setting. In Yuval Ishai, editor, *Theory of Cryptography—8th Theory of Cryptography Conference, TCC 2011*, pages 417–431, 2011.
- [KM14] Daniel Kifer and Ashwin Machanavajjhala. Pufferfish: A framework for mathematical privacy definitions. *ACM Transactions on Database Systems (TODS)*, 39(1):3, 2014.
- [KOV15] Peter Kairouz, Sewoong Oh, and Pramod Viswanath. The composition theorem for differential privacy. In *Proceedings of the 32nd International Conference on Machine Learning*, pages 1376–1385, 2015.
- [LPR13] Vadim Lyubashevsky, Chris Peikert, and Oded Regev. On ideal lattices and learning with errors over rings. *Journal of the ACM*, 60(6):1–35, November 2013.
- [LSS14] Adeline Langlois, Damien Stehlé, and Ron Steinfeld. GGHLite: More efficient multilinear maps from ideal lattices. In Phong Q. Nguyen and Elisabeth Oswald, editors, *Advances in Cryptology—EUROCRYPT 2014*, volume 8441 of *Lecture Notes in Computer Science*, pages 239–256. Springer Berlin Heidelberg, 2014.
- [LV87] Friedrich Liese and Igor Vajda. *Convex Statistical Distances*. Teubner, 1987.
- [McS09] Frank D. McSherry. Privacy integrated queries: an extensible platform for privacy-preserving data analysis. In Carsten Binnig and Benoit Dageville, editors, *Proceedings of the ACM SIGMOD International Conference on Management of Data, SIGMOD 2009*, pages 19–30, 2009.
- [McS17] Frank D. McSherry. How many secrets do you have? <https://github.com/frankmcsherry/blog/blob/master/posts/2017-02-08.md>, February 2017.
- [MMP⁺10] Andrew McGregor, Ilya Mironov, Toniann Pitassi, Omer Reingold, Kunal Talwar, and Salil Vadhan. The limits of two-party differential privacy. In Luca Trevisan, editor, *51st Annual IEEE Symposium on Foundations of Computer Science, FOCS 2010*, pages 81–90. IEEE, 2010.
- [MMR09] Yishay Mansour, Mehryar Mohri, and Afshin Ros-tamizadeh. Multiple source adaptation and the Rényi divergence. In *UAI '09 Proceedings of the Twenty-Fifth Conference on Uncertainty in Artificial Intelligence*, pages 367–374. AUAI Press, June 2009.
- [MPRV09] Ilya Mironov, Omkant Pandey, Omer Reingold, and Salil P. Vadhan. Computational differential privacy. In Shai Halevi, editor, *Advances in Cryptology—CRYPTO 2009*, pages 126–142, 2009.
- [MV16] Jack Murtagh and Salil Vadhan. The complexity of computing the optimal composition of differential privacy. In Eyal Kushilevitz and Tal Malkin, editors, *Theory of Cryptography—13th International Conference, TCC 2016-A, Part I*, pages 157–175, 2016.
- [R61] Alfréd Rényi. On measures of entropy and information. In *Proceedings of the fourth Berkeley symposium on mathematical statistics and probability*, volume 1, pages 547–561, 1961.
- [Sha11] Ofer Shayevitz. On Rényi measures and hypothesis testing. In *2011 IEEE International Symposium on Information Theory Proceedings*, pages 894–898. IEEE, July 2011.
- [vEH14] Tim van Erven and Peter Harremoës. Rényi divergence and Kullback-Leibler divergence. *IEEE Transactions*

on Information Theory, 60(7):3797–3820, July 2014. arxiv.org/abs/1206.2459.

APPENDIX

For comprehensive exposition of properties of the Rényi divergence we refer to two recent papers [vEH14], [Sha11]. Here we recall and re-prove several facts useful for our analysis.

Proposition 8 (Non-negativity). *For $1 \leq \alpha$ and arbitrary distributions P, Q*

$$D_\alpha(P||Q) \geq 0.$$

Proof. Assume that $\alpha > 1$. Define $\phi(x) = x^{1-\alpha}$ and $g(x) = Q(x)/P(x)$. Then

$$\begin{aligned} D_\alpha(P||Q) &= \frac{1}{\alpha-1} \log \mathbb{E}_P[\phi(g(x))] \\ &\geq \frac{1}{\alpha-1} \log \phi(\mathbb{E}_P[g(x)]) \\ &= 0 \end{aligned}$$

by the Jensen inequality applied to the convex function ϕ . The case of $\alpha = 1$ follows by letting ϕ to be $\log 1/x$. $\square$

Proposition 9 (Monotonicity). *For $1 \leq \alpha < \beta$ and arbitrary P, Q*

$$D_\alpha(P||Q) \leq D_\beta(P||Q).$$

Proof (following [vEH14]). Assume that $\alpha > 1$. Observe that the function $x \mapsto x^{\frac{\alpha-1}{\beta-1}}$ is concave. By Jensen’s inequality

$$\begin{aligned} D_\alpha(P||Q) &= \frac{1}{\alpha-1} \log \mathbb{E}_P \left(\frac{P(x)}{Q(x)} \right)^{\alpha-1} \\ &= \frac{1}{\alpha-1} \log \mathbb{E}_P \left(\frac{P(x)}{Q(x)} \right)^{(\beta-1) \frac{\alpha-1}{\beta-1}} \\ &\leq \frac{1}{\alpha-1} \log \left\{ \mathbb{E}_P \left(\frac{P(x)}{Q(x)} \right)^{\beta-1} \right\}^{\frac{\alpha-1}{\beta-1}} \\ &= D_\beta(P||Q). \end{aligned}$$

The case of $\alpha = 1$ follows by continuity. $\square$

The following proposition appears in Langlois et al. [LSS14], generalizing Lyubashevsky et al. [LPR13].

Proposition 10 (Probability preservation [LSS14]). *Let $\alpha > 1$, P and Q be two distributions defined over $\mathcal{R}$ with identical support, $A \subset \mathcal{R}$ be an arbitrary event. Then*

$$P(A) \leq (\exp[D_\alpha(P||Q)] \cdot Q(A))^{(\alpha-1)/\alpha}.$$

Proof. The result follows by application of the Hölder's inequality, which states that for real-valued functions f and g , and real $p, q > 1$, such that $1/p + 1/q = 1$,

$$\|fg\|_1 \leq \|f\|_p \|g\|_q.$$

By setting $p = \alpha$, $q = \alpha/(\alpha-1)$, $f(x) = P(x)/Q(x)^{1/q}$, $g(x) = Q(x)^{1/q}$, and applying Hölder's, we have

$$\begin{aligned} \int_A P(x) dx &\leq \left(\int_A P(x)^\alpha Q(x)^{1-\alpha} dx \right)^{\frac{1}{\alpha}} \left(\int_A Q(x) dx \right)^{\frac{\alpha-1}{\alpha}} \\ &\leq \exp[D_\alpha(P\|Q)]^{(\alpha-1)/\alpha} Q(A)^{(\alpha-1)/\alpha}, \end{aligned}$$

completing the proof. $\square$

The most salient feature of the bound is its (often non-monotone) dependency on α : as α approaches 1, $D_\alpha(P\|Q)$ shrinks (by monotonicity of the Rényi divergence) but the power to which it is raised goes to 0, pushing the result in the opposite direction. Several of our proofs proceed by finding the optimal, or approximately optimal, α minimizing the bound.

The Rényi divergence is not a metric: it is not symmetric and it does not satisfy the triangle inequality. A weaker variant of the triangle inequality tying together the Rényi divergence of different orders does hold. Its general version is presented below.

Proposition 11 (Weak triangle inequality). *Let P, Q, R be distributions on $\mathcal{R}$. Then for $\alpha \geq 1$ and for any $p, q > 1$ satisfying $1/p + 1/q = 1$ it holds that*

$$D_\alpha(P\|Q) \leq \frac{\alpha - 1/p}{\alpha - 1} D_{p\alpha}(P\|R) + D_{q(\alpha-1/p)}(R\|Q).$$

Proof. By Hölder's inequality we have:

$$\begin{aligned} \exp[(\alpha - 1)D_\alpha(P\|Q)] &= \int_{\mathcal{R}} P(x)^\alpha Q(x)^{1-\alpha} dx \\ &= \int_{\mathcal{R}} \frac{P(x)^\alpha}{R(x)^{\alpha-1/p}} \frac{R(x)^{\alpha-1/p}}{Q(x)^{\alpha-1}} dx \\ &\leq \left\{ \int_{\mathcal{R}} \frac{P(x)^{p\alpha}}{R(x)^{p\alpha-1}} dx \right\}^{1/p} \left\{ \int_{\mathcal{R}} \frac{R(x)^{q\alpha-q/p}}{Q(x)^{q\alpha-q}} dx \right\}^{1/q} \\ &= \exp[(\alpha - 1/p)D_{p\alpha}(P\|R)] \cdot \exp[(\alpha - 1)D_{q\alpha-q/p}(R\|Q)]. \end{aligned}$$

By taking the logarithm and dividing both sides by $\alpha - 1$ we establish the claim. $\square$

Several important special cases of the weak triangle inequality can be obtained by fixing parameters p and q (compare it with [MMR09, Lemma 12] and [LSS14, Lemma 4.1]):

Corollary 4. *For P, Q, R with common support we have*

- 1) $D_\alpha(P\|Q) \leq \frac{\alpha-1/2}{\alpha-1} D_{2\alpha}(P\|R) + D_{2\alpha-1}(R\|Q)$.
- 2) $D_\alpha(P\|Q) \leq \frac{\alpha}{\alpha-1} D_\infty(P\|R) + D_\alpha(R\|Q)$.
- 3) $D_\alpha(P\|Q) \leq D_\alpha(P\|R) + D_\infty(R\|Q)$.
- 4) $D_\alpha(P\|Q) \leq \frac{\alpha-\alpha/\beta}{\alpha-1} D_\beta(P\|R) + D_\beta(R\|Q)$, for some explicit $\beta = 2\alpha - .5 + O(1/\alpha)$.

Proof. All claims follow from the weak triangle inequality (Proposition 11) where p and q are chosen, respectively, as

- 1) $p = q = 2$.
- 2) $p \rightarrow \infty$ and $q = p/(p-1) \rightarrow 1$.
- 3) $q \rightarrow \infty$ and $p = q/(q-1) \rightarrow 1$.
- 4) such that $\alpha p = \alpha q - 1$ and $1/p + 1/q = 1$.

In the last case $\beta = p\alpha = 2\alpha - .5 + O(1/\alpha)$. $\square$