

THE COST OF PRIVACY: OPTIMAL RATES OF CONVERGENCE FOR PARAMETER ESTIMATION WITH DIFFERENTIAL PRIVACY

BY T. TONY CAI^{1,*}, YICHEN WANG^{1,†} AND LINJUN ZHANG²

¹*Department of Statistics, The Wharton School, University of Pennsylvania, *tcai@wharton.upenn.edu;*

†wangyc@wharton.upenn.edu

²*Department of Statistics, Rutgers University, linjun.zhang@rutgers.edu*

Privacy-preserving data analysis is a rising challenge in contemporary statistics, as the privacy guarantees of statistical methods are often achieved at the expense of accuracy. In this paper, we investigate the tradeoff between statistical accuracy and privacy in mean estimation and linear regression, under both the classical low-dimensional and modern high-dimensional settings. A primary focus is to establish minimax optimality for statistical estimation with the (ϵ, δ) -differential privacy constraint. By refining the “tracing adversary” technique for lower bounds in the theoretical computer science literature, we improve existing minimax lower bound for low-dimensional mean estimation and establish new lower bounds for high-dimensional mean estimation and linear regression problems. We also design differentially private algorithms that attain the minimax lower bounds up to logarithmic factors. In particular, for high-dimensional linear regression, a novel private iterative hard thresholding algorithm is proposed. The numerical performance of differentially private algorithms is demonstrated by simulation studies and applications to real data sets.

1. Introduction. With the unprecedented availability of data sets containing sensitive personal information, there are increasing concerns that statistical analysis of such data sets may compromise individual privacy. These concerns give rise to statistical methods that provide privacy guarantees at the cost of statistical accuracy, which then motivates us to study the optimal tradeoff between privacy and accuracy in fundamental statistical problems such as mean estimation and linear regression.

A rigorous definition of privacy is a prerequisite for our study. Differential privacy, introduced in [16], is arguably the most widely adopted definition of privacy in statistical data analysis. The promise of a differentially private algorithm is protection of each individual’s privacy from an adversary who has access to the algorithm’s output and possibly even the rest of the data. Differential privacy has gained significant attention in academia [1, 15, 17, 19] and found its way into real world applications developed by Apple [45], Google [23], Microsoft [11] and the U.S. Census Bureau [2].

A usual approach to developing differentially private algorithms is perturbing the output of nonprivate algorithms by random noise [16, 17, 34], and naturally the processed output suffers from some loss of accuracy, which has been extensively observed and studied in the literature [5, 22, 31, 41, 48]. Our paper intends to characterize quantitatively the tradeoff between differential privacy guarantees and statistical accuracy, under the statistical minimax risk framework. Specifically, we study this tradeoff in mean estimation and linear regression problems with the (ϵ, δ) -differential privacy constraint, which is formally defined as follows.

Received February 2020; revised October 2020.

MSC2020 subject classifications. Primary 62F30; secondary 62F12, 62J05.

Key words and phrases. High-dimensional data, differential privacy, mean estimation, linear regression, minimax optimality.

DEFINITION 1 (Differential privacy [16]). A randomized algorithm $M : \mathcal{X}^n \rightarrow \mathcal{R}$ is (ε, δ) -differentially private if for every pair of adjacent data sets $X, X' \in \mathcal{X}^n$ that differ by one individual datum and every (measurable) $S \subseteq \mathcal{R}$,

$$\mathbb{P}(M(X) \in S) \leq e^\varepsilon \cdot \mathbb{P}(M(X') \in S) + \delta,$$

where the probability measure $\mathbb{P}$ is induced by the randomness of M only.

According to the definition, the two parameters ε and δ control the level of privacy against an adversary who attempts to detect the presence of a certain individual in the sample. The privacy constraint becomes more stringent as ε, δ tend to 0.

Our contributions and related literature.

Lower bounds based on tracing attacks. We establish the necessary cost of privacy by proving minimax risk lower bounds with the (ε, δ) -differential privacy constraint. Specifically, we improve existing minimax risk lower bounds for low-dimensional mean estimation and prove new lower bounds for linear regression problems as well as high-dimensional¹ mean estimation. These lower bound results are based on the tracing adversary argument, which originated in the theoretical computer science literature [9, 42]. Early works in this direction were primarily concerned with the accuracy of releasing in-sample quantities, such as k -way marginals, with differential privacy constraints. Some more recent works [20, 27] applied the idea to obtain lower bounds for estimating population quantities such as mean vectors of discrete and Gaussian distributions. Below is a brief summary of our results as compared to existing results; the details are in Sections 3 and 4.

(1) Improved lower bound for low-dimensional mean estimation. In Section 3.1, we show that the minimax squared ℓ_2 risk of sub-Gaussian mean estimation with (ε, δ) -differential privacy is at least $O(\frac{d^2 \log(1/\delta)}{n^2 \varepsilon^2})$ (Theorem 3.1), which improves the $O(\frac{d^2}{n^2 \varepsilon^2})$ minimax lower bound by [27] and matches the deterministic worst case lower bound by [42]. It is further shown that our lower bound is optimal as it can be attained by a differentially private algorithm, the noisy sample mean (Algorithm 3.1; Theorem 3.2).

(2) New lower bounds for linear regression and high-dimensional mean estimation. To the best of our knowledge, our minimax risk lower bounds for high-dimensional mean estimation and linear regression in both low and high dimensions are the first of their kind in the literature. In these three problems, the minimax squared ℓ_2 risk lower bounds are of the order $O(\frac{(s \log d)^2}{n^2 \varepsilon^2})$ (Theorem 3.3), $O(\frac{d^2}{n^2 \varepsilon^2})$ (Theorem 4.1) and $O(\frac{(s \log d)^2}{n^2 \varepsilon^2})$ (Theorem 4.3), respectively, where n, d and s denote the sample size, the dimension and the sparsity of true parameter vector.

For context, there exist several lower bound results for related but different problems: [43] found that the sample complexity lower bound of selecting the top- k largest coordinates of d -dimensional data depends linearly on k and only logarithmically on d ; [5] established an excess empirical risk lower bound of $O(\frac{d}{n \varepsilon^2})$ for (ε, δ) -differentially private empirical risk minimization, by explicitly constructing a worst-case strongly convex and Lipschitz objective function.

¹In computer science literature, the term “high-dimension” refers to settings in which the dimension is allowed to grow with the sample size, and asymptotic dependence on the dimension is of interest. In statistics literature, including this paper, “high-dimension” typically implies that the dimension is greater than the sample size, so sparsity assumptions are often introduced to make the problem feasible.

Differentially private algorithms. We show that the lower bound results are sharp up to logarithmic factors, by constructing differentially private algorithms with rates of convergence matching the corresponding lower bounds.

In low-dimensional problems, the algorithms (Algorithms 3.1 and 4.1) are similar to existing algorithms, such as the noisy Gaussian sample mean [28] or noisy gradient descent [5]. For low-dimensional regression, our contribution is in obtaining an upper bound of the parameter estimation error $\mathbb{E}\|\hat{\beta}_{\text{private}} - \beta_{\text{true}}\|_2^2 = \tilde{O}(\frac{d^2 \log(1/\delta)}{n^2 \varepsilon^2})$ (Theorem 4.2) for the noisy gradient descent algorithm, as opposed to the excess risk bound (or its empirical version) by previous works [4, 5]: $\mathbb{E}[\mathcal{L}_n(\hat{\beta}_{\text{private}}) - \mathcal{L}_n(\hat{\beta}_{\text{nonprivate}})] = O(\frac{\sqrt{d \log(1/\delta)}}{n\varepsilon})$, where $\mathcal{L}_n$ is the least-square objective function of linear regression.

For high-dimensional sparse estimation, our algorithms, to the best of our knowledge, are the first results achieving optimal rates of convergence with the (ε, δ) -differential privacy constraint up to logarithmic factors. The high-dimensional mean estimation algorithm (Algorithms 3.3) is based on a novel application of the “peeling” algorithm first proposed by [21] for reporting top- k coordinates of a vector. The high-dimensional linear regression algorithm (Algorithm 4.2) can be understood as a private version of iterative hard thresholding [7, 25], which roughly speaking, is a projected gradient descent algorithm onto the set of sparse vectors. The focus of our theoretical analysis is again on the parameter estimation error $\mathbb{E}\|\hat{\beta}_{\text{private}} - \beta_{\text{true}}\|_2^2 = \tilde{O}(\frac{(s \log d)^2 \log(1/\delta)}{n^2 \varepsilon^2})$ (Theorems 3.4 and 4.4), as opposed to excess risk results such as $\mathbb{E}[\mathcal{L}_n(\hat{\beta}_{\text{private}}) - \mathcal{L}_n(\beta_{\text{true}})] = O(\frac{s^3 \log d}{n\varepsilon})$ in [30] and $\mathbb{E}[\mathcal{L}_n(\hat{\beta}_{\text{private}}) - \mathcal{L}_n(\hat{\beta}_{\text{nonprivate}})] = O(\frac{\log d + \log(n/\delta)}{(n\varepsilon)^{2/3}})$ in [44].

Other related literature. In theoretical computer science, [41] showed that under strong conditions for privacy parameters, some estimators attain the statistical convergence rates, and hence privacy can be gained for free. [5, 22, 44] proposed differentially private algorithms for convex empirical risk minimization, principal component analysis and high-dimensional sparse regression, and investigated the convergence rates of excess risk.

In the statistics literature, there has also been a series of works that studied differential privacy in statistical estimation. [48] observed that, locally differentially private schemes [29] seem to yield slower convergence rates than the optimal minimax rates in general; [14] developed a framework for deriving statistical minimax rates with the α -local privacy constraint; [39] proved several minimax optimal rates of convergence under α -local differential privacy and exhibited a mechanism that is minimax optimal for linear functionals based on randomized response. It has also been observed that α -local privacy is a much stronger notion of privacy than (ε, δ) -differential privacy that is hardly compatible with high-dimensional problems [14]. As we shall see in this paper, the cost of (ε, δ) -differential privacy in high-dimensional statistical estimation is quite different from that of α -local privacy.

Organization of the paper. The paper is organized as follows. Section 2 formally defines the “cost of privacy” in terms of statistical minimax risk and introduces our technical tools for upper and lower bounding the cost of privacy in various statistical problems. These technical tools are then applied to mean estimation and linear regression problems in Sections 3 and 4, respectively. The numerical performance of the mean estimation and linear regression algorithms are demonstrated by simulated experiments in Section 5 and by real data analysis in Section 6. Section 7 discusses implications of our results in other statistical estimation problems with privacy constraints. The proofs of our theoretical results are in Section 8 and the Supplementary Material [10].

Notation. For real-valued sequences $\{a_n\}$ and $\{b_n\}$, we write $a_n \lesssim b_n$ if $a_n \leq cb_n$ for some universal constant $c \in (0, \infty)$, and $a_n \gtrsim b_n$ if $a_n \geq c'b_n$ for some universal constant $c' \in (0, \infty)$. We say $a_n \asymp b_n$ if $a_n \lesssim b_n$ and $a_n \gtrsim b_n$. In this paper, $c, C, c_0, c_1, c_2, \dots$, refer to universal constants, and their specific values may vary from place to place.

For a positive integer k , we write $[k]$ as shorthand for $\{1, \dots, k\}$. For a vector $\mathbf{v} \in \mathbb{R}^d$ and a subset $S \subseteq [d]$, we use $\mathbf{v}_S$ to denote the restriction of vector $\mathbf{v}$ to the index set S . We write $\text{supp}(\mathbf{v}) := \{j \in [d] : v_j \neq 0\}$. $\|\mathbf{v}\|_p$ denotes the vector ℓ_p norm for $1 \leq p \leq \infty$, with an additional convention that $\|\mathbf{v}\|_0$ denotes the number of nonzero coordinates of $\mathbf{v}$. For a positive definite matrix Σ , we define $\|\mathbf{v}\|_\Sigma = \sqrt{\mathbf{v}^\top \Sigma \mathbf{v}}$. For $\mathbf{v} \in \mathbb{R}^d$ and $R > 0$, let $\Pi_R(\mathbf{v})$ denote the projection of $\mathbf{v}$ onto the ℓ_2 ball $\{\mathbf{u} \in \mathbb{R}^d : \|\mathbf{u}\|_2 \leq R\}$.

2. Problem formulation. In this section, we start with a formal definition of the “cost of privacy” based on the minimax risk with differential privacy constraint in Section 2.1. In Sections 2.2 and 2.3, we provide an overview of our technical tools for upper and lower bounding the cost of privacy.

2.1. The cost of privacy. We quantify the cost of differential privacy in statistical estimation via the minimax risk with differential privacy constraint, defined as follows.

Let $\mathcal{P}$ denote a family of distributions supported on a set $\mathcal{X}$, and let $\boldsymbol{\theta} : \mathcal{P} \rightarrow \Theta \subseteq \mathbb{R}^d$ denote a population quantity of interest. The statistician has access to a data set of n i.i.d. samples, $\mathbf{X} = (\mathbf{x}_1, \dots, \mathbf{x}_n) \in \mathcal{X}^n$, drawn from some distribution $P \in \mathcal{P}$.

With the data, we estimate a population parameter $\boldsymbol{\theta}(P)$ by an estimator $M(\mathbf{X}) : \mathcal{X}^n \rightarrow \Theta$ that belongs to $\mathcal{M}_{\varepsilon, \delta}$, the collection of all (ε, δ) -differentially private procedures. The performance of $M(\mathbf{X})$ is measured by its distance to the truth $\boldsymbol{\theta}(P)$: let $\rho : \Theta \times \Theta \rightarrow \mathbb{R}^+$ be a metric induced by a norm $\|\cdot\|$ on Θ , namely $\rho(\boldsymbol{\theta}_1, \boldsymbol{\theta}_2) = \|\boldsymbol{\theta}_1 - \boldsymbol{\theta}_2\|$, and let $l : \mathbb{R}^+ \rightarrow \mathbb{R}^+$ be an increasing function; the minimax risk of estimating $\boldsymbol{\theta}(P)$ with differential privacy constraint is defined as

$$(2.1) \qquad \inf_{M \in \mathcal{M}_{\varepsilon, \delta}} \sup_{P \in \mathcal{P}} \mathbb{E}[l(\rho(M(\mathbf{X}), \boldsymbol{\theta}(P)))].$$

The quantity characterizes the worst-case performance over $\mathcal{P}$ of the best (ε, δ) -differentially private estimator. The difference between (2.1) and the usual, unconstrained minimax risk

$$(2.2) \qquad \inf_M \sup_{P \in \mathcal{P}} \mathbb{E}[l(\rho(M(\mathbf{X}), \boldsymbol{\theta}(P)))],$$

is the “cost of privacy”. As (2.2) is well understood for mean estimation and linear regression problems, we focus on characterizing the constrained minimax risk (2.1) in this paper. More specifically, we establish upper and lower bounds of (2.1) with technical tools to be introduced in Sections 2.2 and 2.3.

2.2. Construction of differentially private algorithms. It is frequently the case that differentially private algorithms are constructed by perturbing the output of a nonprivate algorithm with random noise. Among the most prominent examples are the Laplace and Gaussian mechanisms.

The Laplace and Gaussian mechanisms. As the name suggests, the Laplace and Gaussian mechanisms achieve differential privacy by perturbing an algorithm with Laplace and Gaussian noises, respectively. The scale of such noises is determined by the sensitivity of the algorithm.

DEFINITION 2. For any algorithm f mapping a data set X to $\mathbb{R}^d$, the ℓ_p -sensitivity of f is

$$\Delta_p(f) = \sup_{X, X' \text{ adjacent}} \|f(X) - f(X')\|_p.$$

The sensitivity of an algorithm f characterizes the magnitude of change in the output of f resulted from replacing one element in an input data set; naturally, we introduce some perturbation of comparable scale so that the differentially private version of f is stable regardless of the presence or absence of any individual datum in the data set.

For algorithms with finite ℓ_1 -sensitivity, differential privacy can be attained by adding Laplace noises.

EXAMPLE 2.1 (The Laplace mechanism). For any algorithm f mapping a data set X to $\mathbb{R}^d$ such that $\Delta_1(f) < \infty$, $M_1(X) := f(X) + (\xi_1, \xi_2, \dots, \xi_d)$, where $\xi_1, \xi_2, \dots, \xi_d$ is an i.i.d. sample drawn from $\text{Laplace}(\Delta_1 f / \varepsilon)$ and achieves $(\varepsilon, 0)$ -differential privacy.

Adding Gaussian noises to algorithms with finite ℓ_2 -sensitivity guarantees differential privacy.

EXAMPLE 2.2 (The Gaussian mechanism). For any algorithm f mapping a data set X to $\mathbb{R}^d$ such that $\Delta_2(f) < \infty$, $M_2(X) := f(X) + (\xi_1, \xi_2, \dots, \xi_d)$, where $\xi_1, \xi_2, \dots, \xi_k$ is an i.i.d. sample drawn from $N(0, 2(\Delta_2(f)/\varepsilon)^2 \log(1.25/\delta))$, achieves (ε, δ) -differential privacy.

Although conceptually simple, these mechanism can often lead to complex differentially private algorithms thanks to the post-processing and composition properties of differential privacy.

Post-processing and composition. Conveniently, post-processing a differentially private algorithm preserves privacy.

FACT 2.1 (Post-processing [16, 48]). Let f be an (ε, δ) -differentially private algorithm and g be an arbitrary, deterministic mapping that takes $f(X)$ as an input; then $g(f(X))$ is (ε, δ) -differentially private.

Further, the privacy parameters are additive with respect to compositions of differentially private algorithms.

FACT 2.2 (Composition [16]). For $i = 1, 2$, let f_i be $(\varepsilon_i, \delta_i)$ -differentially private; then $f_1 \circ f_2$ is $(\varepsilon_1 + \varepsilon_2, \delta_1 + \delta_2)$ -differentially private.

The mechanisms and composition theorem reviewed in this section shall later enable us to construct differentially private algorithms for mean estimation and linear regression.

2.3. *Minimax risk lower bounds with differential privacy constraint.* Our technique for proving lower bounds of the minimax risk (2.1) is based on the “tracing adversary” argument originally proposed by [9]. It has proven to be a powerful tool for obtaining lower bounds in the context of releasing sample quantities [42, 43] and for Gaussian mean estimation [20, 27]. In this paper, we refine the tracing adversary technique to prove a sharper lower bound for low-dimensional mean estimation compared to previous works [20, 27] as well as new lower bounds for sparse mean estimation and linear regression problems.

Informally, a tracing adversary (or tracing attack) is an algorithm that attempts to detect the absence/presence of a candidate datum $\tilde{x}$ in a target data set X , by looking at an estimator $M(X)$ computed from the data set. If one can construct a tracing adversary that is powerful given an accurate estimator, an argument by contradiction leads to a lower bound: suppose a differentially private estimator computed from the target data set is sufficiently accurate, the tracing adversary will be able to determine whether a given datum belongs to the data set or not, thereby contradicting with the differential privacy guarantee. The privacy guarantee and the tracing adversary together ensure that a differentially private estimator cannot be “too accurate.” In Sections 3.1, 3.3, 4.1 and 4.3, we shall formally define and analyze such tracing attacks for mean estimation and linear regression problems. For now, we illustrate this general approach with a concrete example of sub-Gaussian mean estimation.

Example: A preliminary lower bound for mean estimation. To illustrate this approach, we consider a tracing attack proposed by [20] and show how its theoretical properties imply a minimax risk lower bound for differentially private mean estimation of d -dimensional sub-Gaussian(σ) distribution. Let $X = \{x_1, x_2, \dots, x_n\}$ be an i.i.d. sample drawn from the d -dimensional product distribution supported on $\{-\sigma, \sigma\}^d$, which is clearly sub-Gaussian(σ), with unknown mean vector $\mu \in [-\sigma, \sigma]^d$. The tracing attack is given by

$$\mathcal{A}_\mu(x, M(X)) = \langle x - \mu, M(X) \rangle.$$

The theoretical properties of this tracing attack are presented in the following lemma.

LEMMA 2.1. *Let $X = \{x_1, x_2, \dots, x_n\}$ be an i.i.d. sample drawn from the d -dimensional product distribution supported on $\{-\sigma, \sigma\}^d$ with unknown mean vector $\mu \in [-\sigma, \sigma]^d$.*

1. *For each $i \in [n]$, let X'_i denote the data set obtained by replacing x_i in X with an independent copy, then for every $\delta > 0$, every $i \in [n]$ and every μ we have*

$$\mathbb{P}(\mathcal{A}_\mu(x_i, M(X'_i)) > \sigma^2 \sqrt{8d \log(1/\delta)}) < \delta.$$

2. *There is a universal constant c_1 such that, if $n < c_1 \sqrt{d/\log(1/\delta)}$, we can find a prior distribution π of μ so that*

$$\mathbb{P}_{X, \mu} \left(\sum_{i \in [n]} \mathcal{A}_\mu(x_i, M(X)) \leq n\sigma^2 \sqrt{8d \log(1/\delta)}, \|M(X) - \bar{X}\|_2 < c_2 \sigma \sqrt{d} \right) < \delta$$

for an appropriate universal constant c_2 .

The lemma is similar in spirit to Lemma 12 in [20] and proved with a new technical argument in Section B.1 of the Supplementary Material [10]. According to the lemma, when $M(X)$ is close to the sample mean $\bar{X}$, the attack takes large values if the candidate datum belongs to the data set X from which $M(X)$ is computed.

We would like to point out that Lemma 2.1 is valid without requiring differential privacy of $M(X)$. For a differentially private $M(X)$ in particular, the lemma makes available a lower bound for $\|M(X) - \bar{X}\|$, as we have sketched informally: if $n < c_1 \sqrt{d/\log(1/\delta)}$ and $M(X)$ is (ε, δ) -differentially private with $0 < \varepsilon < 1$ and $\delta = o(1/n)$, let $\mathcal{C} = \{\sum_{i \in [n]} \mathcal{A}_\mu(x_i, M(X)) \leq n\sigma^2 \sqrt{8d \log(1/\delta)}\}$,

$$\begin{aligned} & \mathbb{P}_{X, \mu}(\|M(X) - \bar{X}\|_2 < c_2 \sigma \sqrt{d}) \\ & \leq \mathbb{P}_{X, \mu}(\mathcal{C} \cap \{\|M(X) - \bar{X}\|_2 < c_2 \sigma \sqrt{d}\}) \\ & \quad + \sum_{i \in [n]} \mathbb{P}_{X, \mu}(\mathcal{A}_\mu(x_i, M(X)) > \sigma^2 \sqrt{8d \log(1/\delta)}) \end{aligned}$$

$$\begin{aligned}
&\leq \mathbb{P}_{X, \mu}(\mathcal{C} \cap \{\|M(X) - \bar{X}\|_2 < c_2 \sigma \sqrt{d}\}) \\
&\quad + n(e^\varepsilon \mathbb{P}(\mathcal{A}_\mu(\mathbf{x}_i, M(X'_i)) > \sigma^2 \sqrt{8d \log(1/\delta)}) + \delta) \\
&\leq \delta + n(e^\varepsilon \delta + \delta) = o(1).
\end{aligned}$$

The second inequality is due to differential privacy; the third inequality uses Lemma 2.1. It follows that, when $n < c_1 \sqrt{d/\log(1/\delta)}$, we have

$$(2.3) \quad \inf_{M \in \mathcal{M}_{\varepsilon, \delta}} \sup_{\mu} \mathbb{E} \|M(X) - \bar{X}\|_2 \geq \inf_{M \in \mathcal{M}_{\varepsilon, \delta}} \mathbb{E}_\pi \mathbb{E}_{X|\mu} \|M(X) - \bar{X}\|_2 \gtrsim \sigma \sqrt{d}.$$

It should be noted, however, that the lower bound result is unsatisfactory in two important ways. First, as formulated in Section 2.1, we are in fact interested in lower bounding a related but distinct quantity, $\inf_{M \in \mathcal{M}_{\varepsilon, \delta}} \sup_{\mu} \mathbb{E} \|M(X) - \mu\|_2$. Second, the sample size range $n < c_1 \sqrt{d/\log(1/\delta)}$ is somewhat artificial; for low-dimensional mean estimation problems, the usual setting of $n \gtrsim d$ is of greater interest. In Section 3.1, we shall resolve these issues and, on the basis of the same tracing attack and Lemma 2.1, establish an optimal lower bound for the mean estimation problem.

3. The cost of privacy in mean estimation. In this section, we study the cost of (ε, δ) -differential privacy in estimating the mean vector of a d -dimensional sub-Gaussian(σ) distribution. Formally, an $\mathbb{R}^d$ -valued random variable $\mathbf{x}$ follows a sub-Gaussian(σ) distribution if for $\mu = \mathbb{E}\mathbf{x}$ and any fixed vector $\|v\|_2 = 1$, we have

$$\mathbb{E} \exp(\lambda \langle \mathbf{x} - \mu, v \rangle) \leq \exp(\lambda^2 \sigma^2) \quad \forall \lambda \in \mathbb{R}.$$

We begin with sharpening the preliminary lower bound (2.3), in Section 3.1. The lower bound is then shown to be optimal via an (ε, δ) -differentially private estimator with convergence rate attaining the lower bound.

We also study the cost of differential privacy in sparse mean estimation, where the unknown mean vector $\mu \in \mathbb{R}^d$ has only a small fraction of nonzero coordinates. This sparse model is useful when the data's dimension d outnumbers the sample size n , rendering the usual sample mean estimator sub-optimal. Instead, if the unknown mean vector is indeed sparse, thresholding the sample mean have been shown to achieve optimal statistical accuracy [13, 26]. We establish in Section 3.3 a minimax risk lower bound for estimating sparse mean with differential privacy constraint, and match this lower bound with a differentially private estimator of the sparse mean in Section 3.4.

3.1. Lower bound of low-dimensional mean estimation. In this section, we prove a sharp lower bound for estimating a d -dimensional sub-Gaussian mean by improving the preliminary lower bound in Section 2.3. We consider the class of d -dimensional sub-Gaussian(σ) distributions with mean vector in $\Theta = \{\mu \in \mathbb{R}^d : \|\mu\|_\infty < \sigma\}$ and denote the class by $\mathcal{P}(\sigma, d, \Theta)$.

The first improvement is a relaxation of the sample size range $n \lesssim \sqrt{d/\log(1/\delta)}$.

LEMMA 3.1. *Let $Y = \{y_1, y_2, \dots, y_n\}$ be sampled with replacement from a set of deterministic vectors $Z = \{z_1, z_2, \dots, z_m\}$ with $n = km$ and $k \geq 1$. There exists a choice of Z with each $z_i \in \{-\sigma, \sigma\}^d$, $m = c_1 \sqrt{d/\log(1/\delta)} \gtrsim 1$ and $k \asymp \log(1/\delta)/\varepsilon$ such that*

$$\mathbb{E} \|M(Y) - \mathbb{E} y_1\|_2 \gtrsim \sigma \sqrt{d}$$

for every (ε, δ) -differentially private M if $0 < \varepsilon < 1$, $n^{-1} \exp(-n\varepsilon) < \delta < n^{-(1+\omega)}$ for some fixed constant $\omega > 0$, and $\log(\delta)/\log(n)$ is nonincreasing in n .

Lemma 3.1 is proved in Section 8.1. In essence, this lemma improves the lower bound (2.3) by extending its range of validity to $n \lesssim \sqrt{d \log(1/\delta)}/\varepsilon$, as the discrete uniform distribution described in the lemma is sub-Gaussian(σ) with $\boldsymbol{\mu} \in \Theta$ thanks to the choice of $\mathbf{z}_i \in \{-\sigma, \sigma\}^d$.

On the basis of Lemma 3.1, we are able to translate the lower bound result to the more interesting large n regime, as described by the following theorem.

THEOREM 3.1. *Let $\mathbf{X} = \{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n\}$ be an i.i.d. sample drawn from some distribution in $\mathcal{P}(d, \sigma, \Theta)$ with mean $\mathbb{E}\mathbf{x}_1 = \boldsymbol{\mu}$. If $0 < \varepsilon < 1$, $n^{-1} \exp(-n\varepsilon) < \delta < n^{-(1+\omega)}$ for some fixed constant $\omega > 0$ with $\log(\delta)/\log(n)$ nonincreasing in n , $d/\log(1/\delta) \gtrsim 1$ and $n \gtrsim \sqrt{d \log(1/\delta)}/\varepsilon$, we have*

$$(3.1) \qquad \inf_{M \in \mathcal{M}_{\varepsilon, \delta}} \sup_{\mathcal{P}(d, \sigma, \Theta)} \mathbb{E} \|\mathbf{M}(\mathbf{X}) - \boldsymbol{\mu}\|_2^2 \gtrsim \sigma^2 \left(\frac{d}{n} + \frac{d^2 \log(1/\delta)}{n^2 \varepsilon^2} \right).$$

The theorem is proved in Section 8.2. The minimax lower bound characterizes the cost of privacy in the mean estimation problem: the cost of privacy dominates the statistical risk when $d \log(1/\delta)/n\varepsilon^2 \gtrsim 1$. This minimax lower bound matches the sample complexity lower bound in [42], which considered the deterministic worst case instead of the i.i.d. statistical setting. [27] studied the Gaussian mean estimation problem but did not obtain a tight bound with respect to δ ; Theorem 3.1 improves the lower bound in [27] by $\log(1/\delta)$. In Section 3.2, we exhibit an algorithm for mean estimation that attains the convergence rate of $\sigma^2(\frac{d}{n} + \frac{d^2 \log(1/\delta)}{n^2 \varepsilon^2})$, showing that the lower bound (3.1) is in fact rate-optimal.

3.2. Algorithm for low-dimensional mean estimation. In this section, we show that the minimax lower bound (3.1) can be attained by a differentially private estimator, thereby implying a tight characterization of the cost of privacy in low-dimensional mean estimation.

Let $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n$ be an i.i.d. sample drawn from a sub-Gaussian(σ) distribution on $\mathbb{R}^d$, and we denote $\mathbb{E}\mathbf{x}_1$ by $\boldsymbol{\mu} \in \mathbb{R}^d$. It is further assumed that $\|\boldsymbol{\mu}\|_\infty < c$ for some constant $c = O(1)$. We consider the following simple algorithm based on the Gaussian mechanism, Example 2.2.

The truncation step guarantees that, over a pair of data sets $\mathbf{X}$ and $\mathbf{X}'$ which differ by one single entry, $\|\bar{\mathbf{X}}_R - \bar{\mathbf{X}}_R'\|_2 < 2R\sqrt{d}/n$ and, therefore, the Gaussian mechanism applies. When R is selected so that most of the data is preserved, $\hat{\boldsymbol{\mu}}$ is an accurate estimator of the mean $\boldsymbol{\mu}$.

THEOREM 3.2. *If there exists a constant $T < \infty$ so that $\|\mathbf{x}\|_\infty < T$ with probability one, setting $R = T$ ensures that*

$$\mathbb{E} \|\hat{\boldsymbol{\mu}} - \boldsymbol{\mu}\|_2^2 \lesssim \sigma^2 \left(\frac{d}{n} + \frac{d^2 \log(1/\delta)}{n^2 \varepsilon^2} \right).$$

Otherwise, choosing $R = K\sigma\sqrt{\log n}$ for a sufficiently large K guarantees

$$\mathbb{E} \|\hat{\boldsymbol{\mu}} - \boldsymbol{\mu}\|_2^2 \lesssim \sigma^2 \left(\frac{d}{n} + \frac{d^2 \log(1/\delta) \log n}{n^2 \varepsilon^2} \right).$$

The theorem is proved in Section A.1 of the Supplementary Material [10]. The first case applies to distributions with bounded support, for example, Bernoulli, with the rate of convergence exactly matching the lower bound (3.1). The second case includes unbounded sub-Gaussian distributions such as the Gaussian, where the convergence rate matches the lower bound up to a gap of $O(\log n)$. Overall, the upper and lower bounds suggest that the cost of (ε, δ) -differential privacy in low-dimensional mean estimation is $\tilde{O}(\frac{d^2 \log(1/\delta)}{n^2 \varepsilon^2})$.

Algorithm 3.1: Differentially Private Mean Estimation

Input : Data set $X = \{\mathbf{x}_i\}_{i \in [n]}$, privacy parameters ε, δ , truncation level R .

- 1 Compute $\bar{X}_R$: for $j \in [d]$, $\bar{X}_{R,j} = n^{-1} \sum_{i \in [n]} \Pi_R(x_{ij})$;
- 2 Compute $\hat{\boldsymbol{\mu}} = \bar{X}_R + \mathbf{w}$, where $\mathbf{w} \sim N_d(\mathbf{0}, \frac{4R^2 d \log(1/\delta)}{n^2 \varepsilon^2} \cdot \mathbf{I})$;

Output: $\hat{\boldsymbol{\mu}}$.

It should be noted that Algorithm 3.1 lacks some practicality: the truncation level R is a tuning parameter that needs to be set at the correct level for the convergence rate to hold; we included this somewhat simplistic algorithm here for the theoretical analysis of privacy cost. In Section 5, we consider data-driven and differentially private proxies of the theoretical choice of R and demonstrate their numerical performance. As the focus of this paper is theoretical properties of private estimators, we refer interested readers to [28] and [6] for more practical methods of differentially private mean estimation.

3.3. Lower bound of sparse mean estimation. We consider lower bounding the minimax risk of estimating the mean vector of a sub-Gaussian(σ) distribution when the mean vector is s^* -sparse. Concretely, we index this collection of distributions by the set of mean vectors $\Theta = \{\boldsymbol{\mu} \in \mathbb{R}^d : \|\boldsymbol{\mu}\|_0 \leq s^*, \|\boldsymbol{\mu}\|_\infty < 1\}$, and denote this class of distributions by $\mathcal{P}(\sigma, d, s^*, \Theta)$. Let $X = \{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n\}$ be an i.i.d. sample drawn from a sub-Gaussian(σ) distribution with mean vector $\boldsymbol{\mu} \in \Theta$, we would like to establish a lower bound of $\inf_{M \in \mathcal{M}_{\varepsilon, \delta}} \sup_{\mathcal{P}(\sigma, d, s^*, \Theta)} \mathbb{E} \|M(X) - \boldsymbol{\mu}\|_2^2$ as a function of privacy parameters (ε, δ) as well as d, n, s^* and σ .

As sketched in Section 2.3, our strategy for proving the lower bound requires the existence of a powerful tracing attack. For sparse mean estimation, one reasonable choice of tracing attack is given by

$$(3.2) \quad \mathcal{A}_{\boldsymbol{\mu}, s^*}(\mathbf{x}, M(X)) = \langle (\mathbf{x} - \boldsymbol{\mu})_{\text{supp}(\boldsymbol{\mu})}, M(X) - \boldsymbol{\mu} \rangle.$$

In particular, this attack coincides with the tracing attack proposed by [43] for differentially private top- k selection.

Similar to our lower bound analysis for low-dimensional mean estimation, the key ingredient is to show that the attack typically takes a large value when $\tilde{\mathbf{x}}$ belongs to X and a small value otherwise. This is indeed the case for the tracing attack (3.2), as described by the following lemma.

LEMMA 3.2. *Let $X = \{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n\}$ be an i.i.d. sample drawn from $N_d(\boldsymbol{\mu}, \sigma^2 \mathbf{I})$ with $\boldsymbol{\mu} \in \Theta$. If $s^* = o(d^{1-\omega})$ for some fixed $\omega > 0$, for every (ε, δ) -differentially private estimator M satisfying $\mathbb{E}_{X|\boldsymbol{\mu}} \|M(X) - \boldsymbol{\mu}\|_2^2 = o(1)$ at every $\boldsymbol{\mu} \in \Theta$, the following are true:*

1. *For each $i \in [n]$, let X'_i denote the data set obtained by replacing $\mathbf{x}_i$ in X with an independent copy, then*

$$\mathbb{E} \mathcal{A}_{\boldsymbol{\mu}, s^*}(\mathbf{x}_i, M(X'_i)) = 0, \quad \mathbb{E} |\mathcal{A}_{\boldsymbol{\mu}, s^*}(\mathbf{x}_i, M(X'_i))| \leq \sigma \sqrt{\mathbb{E} \|M(X) - \boldsymbol{\mu}\|_2^2}.$$

2. *There exists a prior distribution of $\boldsymbol{\pi} = \boldsymbol{\pi}(\boldsymbol{\mu})$ supported over Θ such that*

$$\sum_{i \in [n]} \mathbb{E}_{\boldsymbol{\pi}} \mathbb{E}_{X|\boldsymbol{\mu}} \mathcal{A}_{\boldsymbol{\mu}, s^*}(\mathbf{x}_i, M(X)) \gtrsim \sigma^2 s^* \log d.$$

The lemma is proved in Section B.2 of the Supplementary Material [10]. On the basis of Lemma 3.2, we have the following minimax risk lower bound.

THEOREM 3.3. *If $s^* = o(d^{1-\omega})$ for some fixed $\omega > 0$, $0 < \varepsilon < 1$ and $\delta < n^{-(1+\omega)}$ for some fixed $\omega > 0$, we have*

$$(3.3) \quad \inf_{M \in \mathcal{M}_{\varepsilon, \delta}} \sup_{\mathcal{P}(\sigma, d, s^*, \Theta)} \mathbb{E} \|M(X) - \mu\|_2^2 \gtrsim \sigma^2 \left(\frac{s^* \log d}{n} + \frac{(s^* \log d)^2}{n^2 \varepsilon^2} \right).$$

The lower bound is proved in Section B.3 of the Supplementary Material [10]. In this lower bound, it is worth noting that the term due to privacy, similar to the statistical term, only depends logarithmically on the dimension d , suggesting that mean estimation in high dimensions remains viable despite the (ε, δ) -differential privacy constraint. This is in marked contrast with high-dimensional statistical estimation under the (much more demanding) local differential privacy constraint [14, 29], where the minimax risk always depends linearly on d . In the next section, we propose a differentially private estimator that efficiently estimates the sparse mean vector and attains the lower bound (3.3) up to factors of $\log n$.

3.4. Algorithm for sparse mean estimation. Let $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n$ be an i.i.d. sample drawn from a sub-Gaussian(σ) distribution on $\mathbb{R}^d$, with mean $\mathbb{E}\mathbf{x}_1 = \mu \in \mathbb{R}^d$. It is further assumed that $\|\mu\|_0 \leq s^*$ and $\|\mu\|_\infty < c$ for some constant $c = O(1)$.

In this section, we propose a differentially private algorithm for estimating the sparse mean vector μ . At a high level, the algorithm selects the large coordinates of the (truncated) sample mean vector in a differentially private manner, and sets the remaining coordinates to zero. We start with describing and analyzing the differentially private selection step.

The following “peeling” algorithm, developed by [21], is an efficient and differentially private method for selecting the top- s largest coordinates in terms of absolute value. In each of the s iterations, one coordinate is “peeled” from the original vector and added to the output set.

The algorithm is guaranteed to be differentially private when the vector $\mathbf{v} = \mathbf{v}(X)$ has bounded change in value when any single datum in X is modified.

LEMMA 3.3 ([21]). *If for every pair of adjacent data sets Z, Z' we have $\|\mathbf{v}(Z) - \mathbf{v}(Z')\|_\infty < \lambda$, then Algorithm 3.2 is an (ε, δ) -differentially private algorithm.*

Another important property of the Peeling algorithm is its (approximate) accuracy, proved in Section A.2 of the Supplementary Material.

Algorithm 3.2: “Peeling” [21]

Input : vector-valued function $\mathbf{v} = \mathbf{v}(X) \in \mathbb{R}^d$, data X , sparsity s , privacy parameters ε, δ , noise scale λ .

1 Initialize $S = \emptyset$;

2 **for** i in 1 to s **do**

3 Generate $\mathbf{w}_i \in \mathbb{R}^d$ with $w_{i1}, w_{i2}, \dots, w_{id} \stackrel{\text{i.i.d.}}{\sim} \text{Laplace}(\lambda \cdot \frac{2\sqrt{3s \log(1/\delta)}}{\varepsilon})$;

4 Append $j^* = \arg \max_{j \in [d] \setminus S} |v_j| + w_{ij}$ to S ;

5 **end**

6 Set $\tilde{P}_s(\mathbf{v}) = \mathbf{v}_S$;

7 Generate $\tilde{\mathbf{w}}$ with $\tilde{w}_1, \dots, \tilde{w}_d \stackrel{\text{i.i.d.}}{\sim} \text{Laplace}(\lambda \cdot \frac{2\sqrt{3s \log(1/\delta)}}{\varepsilon})$;

Output: $\tilde{P}_s(\mathbf{v}) + \tilde{\mathbf{w}}_S$.

Algorithm 3.3: Differentially Private Sparse Mean Estimation

Input : Data set $\mathbf{X} = \{\mathbf{x}_i\}_{i \in [n]}$, privacy parameters ε, δ , truncation level R , sparsity s .

1 Compute $\bar{\mathbf{X}}_R$: for $j \in [d]$, $\bar{\mathbf{X}}_{R,j} = n^{-1} \sum_{i \in [n]} \Pi_R(x_{ij})$;

2 Compute $\hat{\boldsymbol{\mu}} = \text{Peeling}(\bar{\mathbf{X}}_R, \mathbf{X}, s, \varepsilon, \delta, 2R/n)$;

Output: $\hat{\boldsymbol{\mu}}$.

LEMMA 3.4. *Let S and $\{\mathbf{w}_i\}_{i \in [s]}$ be defined as in Algorithm 3.2. For every $R_1 \subseteq S$ and $R_2 \in S^c$ such that $|R_1| = |R_2|$ and every $c > 0$, we have*

$$\|\mathbf{v}_{R_2}\|_2^2 \leq (1+c)\|\mathbf{v}_{R_1}\|_2^2 + 4(1+1/c) \sum_{i \in [s]} \|\mathbf{w}_i\|_\infty^2.$$

Now returning to the original problem of sparse mean estimation, we construct a differentially estimator of the sparse mean by applying the “peeling” algorithm to a (truncated) sample mean, as follows.

The truncation step ensures that, over a pair of data sets $\mathbf{X}$ and $\mathbf{X}'$ which differ by one single entry, $\|\bar{\mathbf{X}}_R - \bar{\mathbf{X}}'_R\|_\infty < 2R/n$ and, therefore, the privacy guarantee, Lemma 3.3, applies. Algorithm 3.3 further inherits the accuracy of “Peeling” and leads to an accurate estimator of the sparse mean $\boldsymbol{\mu}$, as stated in the following theorem.

THEOREM 3.4. *If $R = K\sigma\sqrt{\log n}$ for a sufficiently large constant K , $s \geq s^*$ and $s \asymp s^*$, then with probability at least $1 - c_1 \exp(-c_2 \log n) - c_1 \exp(-c_2 \log d)$, it holds that*

$$\|\hat{\boldsymbol{\mu}} - \boldsymbol{\mu}\|_2^2 \lesssim \sigma^2 \left(\frac{s^* \log d}{n} + \frac{(s^* \log d)^2 \log(1/\delta) \log n}{n^2 \varepsilon^2} \right).$$

Theorem 3.4 is proved in Section A.3. With the usual choice of $\delta = n^{-(1+\omega)}$, the convergence rate of Algorithm 3.3 attains the lower bound, Theorem 3.3, up to a gap of $\log^2 n$. While the convergence analysis of Algorithm 3.3 requires some theoretical choice of tuning parameters R and s , in Section 5 we discuss data-driven methods of selecting these tuning parameters that achieve reasonably good numerical performance.

4. The cost of privacy in linear regression. In this section, we consider the Gaussian linear model

$$(4.1) \quad f_{\boldsymbol{\beta}}(y|\mathbf{x}) = \frac{1}{\sqrt{2\pi}\sigma} \exp\left(\frac{-(y - \mathbf{x}^\top \boldsymbol{\beta})^2}{2\sigma^2}\right); \quad \mathbf{x} \sim f_{\mathbf{x}}.$$

Given an i.i.d. sample $(\mathbf{y}, \mathbf{X}) = \{(y_i, \mathbf{x}_i)\}_{i \in [n]}$ drawn from the model, we study the cost of (ε, δ) -differential privacy in estimating the regression coefficients $\boldsymbol{\beta} \in \mathbb{R}^d$. The primary focus is on the high-dimensional setting (Sections 4.3, 4.4) where the dimension d dominates the sample size n , and the regression coefficient $\boldsymbol{\beta}$ is assumed to be sparse; the classical, low-dimensional case of $d = o(n)$ will also be considered (Sections 4.1, 4.2).

4.1. *Lower bound of low-dimensional linear regression.* Let $\mathcal{P}(\sigma, d, \Theta)$ denote the class of distributions $f_{\boldsymbol{\beta}}(\mathbf{y}, \mathbf{x})$, as specified by (4.1), with $\boldsymbol{\beta} \in \Theta = \{\boldsymbol{\beta} \in \mathbb{R}^d : \|\boldsymbol{\beta}\|_2 \leq 1\}$. With an i.i.d. sample $(\mathbf{y}, \mathbf{X}) = \{(y_i, \mathbf{x}_i)\}_{i \in [n]}$ drawn from a distribution in $\mathcal{P}(\sigma, d, \Theta)$, we shall establish a lower bound of $\inf_{M \in \mathcal{M}_{\varepsilon, \delta}} \sup_{\mathcal{P}(\sigma, d, \Theta)} \mathbb{E} \|M(\mathbf{y}, \mathbf{X}) - \boldsymbol{\beta}\|_{\Sigma_{\mathbf{x}}}^2$ via the tracing attack argument.

Consider the attack given by

$$(4.2) \quad \mathcal{A}_{\boldsymbol{\beta}}((\mathbf{y}, \mathbf{x}), M(\mathbf{y}, \mathbf{X})) = \langle M(\mathbf{y}, \mathbf{X}) - \boldsymbol{\beta}, (\mathbf{y} - \mathbf{x}^\top \boldsymbol{\beta}) \mathbf{x} \rangle.$$

Similar to the tracing attacks for mean estimation problems, the attack takes large value when $(y, \mathbf{x})$ belongs to (y, X) and small value otherwise.

LEMMA 4.1. *Let (y, X) be an i.i.d. sample drawn from some distribution in $\mathcal{P}(\sigma, d, \Theta)$ such that $\|\mathbf{x}\|_2 \leq 1$ with probability 1, and $\Sigma_{\mathbf{x}} = \mathbb{E}\mathbf{x}\mathbf{x}^\top$ is diagonal and satisfies $0 < 1/L < d\lambda_{\min}(\Sigma_{\mathbf{x}}) \leq d\lambda_{\max}(\Sigma_{\mathbf{x}}) < L$ for some constant $L = O(1)$. For every (ε, δ) -differentially private estimator M satisfying $\mathbb{E}_{y, X|\beta} \|M(y, X) - \beta\|_2^2 = o(1)$ at every $\beta \in \Theta$, the following are true:*

1. *For each $i \in [n]$, let (y'_i, X'_i) denote the data set obtained by replacing $(y_i, \mathbf{x}_i)$ in (y, X) with an independent copy, then $\mathbb{E}\mathcal{A}_\beta((y_i, \mathbf{x}_i), M(y'_i, X'_i)) = 0$ and*

$$\mathbb{E}|\mathcal{A}_\beta((y_i, \mathbf{x}_i), M(y'_i, X'_i))| \leq \sigma \sqrt{\mathbb{E}\|M(y, X) - \beta\|_{\Sigma_{\mathbf{x}}}^2}.$$

2. *There exists a prior distribution of $\pi = \pi(\beta)$ supported over Θ such that*

$$\sum_{i \in [n]} \mathbb{E}_\pi \mathbb{E}_{y, X|\beta} \mathcal{A}_\beta((y_i, \mathbf{x}_i), M(y, X)) \gtrsim \sigma^2 d.$$

The lemma is proved in Section B.4 of the Supplementary Material [10]. These properties of tracing attack imply a minimax lower bound for (ε, δ) -differentially private estimation of β .

THEOREM 4.1. *If $0 < \varepsilon < 1$ and $\delta < n^{-(1+\omega)}$ for some fixed $\omega > 0$, we have*

$$(4.3) \quad \inf_{M \in \mathcal{M}_{\varepsilon, \delta}} \sup_{\mathcal{P}(\sigma, d, \Theta)} \mathbb{E}\|M(y, X) - \beta\|_{\Sigma_{\mathbf{x}}}^2 \gtrsim \sigma^2 \left(\frac{d}{n} + \frac{d^2}{n^2 \varepsilon^2} \right).$$

The lower bound is proved in Section B.5 of the Supplementary Material [10]. In next section, we show that the lower bound is sharp up to factors of $\log n$ by analyzing a differentially private algorithm for estimating β .

4.2. *Algorithm for low-dimensional linear regression.* For the low-dimensional linear regression problem, we seek a differentially private (approximate) minimizer of the least square objective function

$$\mathcal{L}_n(\beta) = \frac{1}{n} \sum_{i=1}^n (y_i - \mathbf{x}_i^\top \beta)^2.$$

We find this solution via the noisy gradient descent algorithm of [5]. We tailor the convergence analysis to the linear regression problem to obtain convergence in $O(\log n)$ iterations, as opposed to $O(n)$ iterations required by the general-purpose version in [5]. The algorithm and its theoretical properties are described in detail in this section.

The analysis of Algorithm 4.1 relies on some assumptions about $\mathbf{x}$ and β .

(D1) Bounded design: there is a constant $c_x < \infty$ such that $\|\mathbf{x}\|_2 < c_x$ with probability 1.

(D2) Bounded moments of design: $\mathbb{E}\mathbf{x} = \mathbf{0}$ and the covariance matrix $\Sigma_{\mathbf{x}} = \mathbb{E}\mathbf{x}\mathbf{x}^\top$ satisfies $0 < 1/L < d \cdot \lambda_{\min}(\Sigma_{\mathbf{x}}) \leq d \cdot \lambda_{\max}(\Sigma_{\mathbf{x}}) < L$ for some constant $0 < L < \infty$.

(P1) The true parameter vector β satisfies $\|\beta\|_2 < c_0$ for some constant $0 < c_0 < \infty$.

In essence, the assumptions on design require that the rows of design matrix are normalized, and the assumed ℓ_2 bound of β is consistent with the parameter regime in our lower bound analysis, Section 4.1.

Assumptions (D1) and (P1) together guarantee that the algorithm is (ε, δ) -differentially private if the noise level B is sufficiently large.

Algorithm 4.1: Differentially Private Linear Regression

Input : $\mathcal{L}_n(\beta)$, data set $\{(y_i, \mathbf{x}_i)\}_{i \in [n]}$, step size η^0 , privacy parameters ε, δ , noise scale B , number of iterations T , truncation level R , feasibility parameter C , initial value β^0 .

1 **for** t in 0 to $T - 1$ **do**

2 Generate $\mathbf{w}_t \in \mathbb{R}^d$ with $w_{t1}, w_{t2}, \dots, w_{td} \stackrel{\text{i.i.d.}}{\sim} N(0, (\eta^0)^2 2B^2 \frac{\log(2T/\delta)}{n^2(\varepsilon/T)^2})$;

3 Compute $\beta^{t+1} = \Pi_C(\beta^t - (\eta^0/n) \sum_{i=1}^n (\mathbf{x}_i^\top \beta^t - \Pi_R(y_i)) \mathbf{x}_i + \mathbf{w}_t)$;

4 **end**

Output: β^T .

LEMMA 4.2. *If assumptions (D1) and (P1) are true, then Algorithm 4.1 is (ε, δ) -differentially private as long as $B \geq 4(R + c_0 c_x) c_x$ and $C \leq c_0$.*

The lemma is proved in Section A.4 of the Supplementary Material [10]. If (D2) is true as well, we obtain the following theorem which describes the convergence rate of Algorithm 4.1.

THEOREM 4.2. *Let $\{(y_i, \mathbf{x}_i)\}_{i \in [n]}$ be an i.i.d. sample from the linear model (4.1). Suppose assumptions (D1), (D2), (P1) are true. Let the parameters of Algorithm 4.1 be chosen as follows:*

- Set step size $\eta^0 = d/2L$, where L is the constant defined in assumption (D2).
- Set $R = \sigma \sqrt{2 \log n}$, $B = 4(R + c_0 c_x) c_x$ and $C = c_0$, in accordance with Lemma 4.2.
- Number of iterations T . Let $T = (8L^2) \log(c_0^2 n)$, where L is the constant defined in assumption (D2).
- Initialization $\beta^0 = \mathbf{0}$.

If $n \geq K \cdot (Rd^{3/2} \sqrt{\log(1/\delta)} \log n \log \log n / \varepsilon)$ for a sufficiently large constant K , the output of Algorithm 4.1 satisfies

$$(4.4) \quad \|\beta^T - \beta^*\|_{\Sigma_x}^2 \lesssim \sigma^2 \left(\frac{d}{n} + \frac{d^2 \log(1/\delta) \log^3 n}{n^2 \varepsilon^2} \right),$$

with probability at least $1 - c_1 \exp(-c_2 n) - c_1 \exp(-c_2 d) - c_1 \exp(-c_2 \log n)$.

Theorem 4.2 is proved in Section A.5 of the Supplementary Material [10]. For practical application of the algorithm, we note that the theoretical choice of truncation level R , which ensures the privacy protection of the algorithm, depends on the often unknown quantity σ . We provide a data-driven, differentially private alternative to this theoretical choice and demonstrate its numerical performance in Section 5.

4.3. *Lower bound of high-dimensional linear regression.* We next consider the high-dimensional linear regression problem where d potentially dominates sample size n , but the estimand β is sparse. Concretely, let $\mathcal{P}(\sigma, d, s^*, \Theta)$ denote the class of distributions $f_\beta(y, \mathbf{x})$, as specified by (4.1), with $\beta \in \Theta = \{\beta \in \mathbb{R}^d : \|\beta\|_0 \leq s^*, \|\beta\|_2 \leq 1\}$. With an i.i.d. sample $(\mathbf{y}, \mathbf{X}) = \{(y_i, \mathbf{x}_i)\}_{i \in [n]}$ drawn from a distribution in $\mathcal{P}(\sigma, d, s^*, \Theta)$, we consider lower bounding $\inf_{M \in \mathcal{M}_{\varepsilon, \delta}} \sup_{\mathcal{P}(\sigma, d, s^*, \Theta)} \mathbb{E} \|M(\mathbf{y}, \mathbf{X}) - \beta\|_{\Sigma_x}^2$ via the tracing attack argument.

Consider the attack given by

$$(4.5) \quad \mathcal{A}_{\beta, s^*}((y, \mathbf{x}), M(\mathbf{y}, \mathbf{X})) = \langle (M(\mathbf{y}, \mathbf{X}) - \beta)_{\text{supp}(\beta)}, (y - \mathbf{x}^\top \beta) \mathbf{x} \rangle.$$

Similar to the tracing attacks for mean estimation problems, the attack takes large value when $(y, \mathbf{x})$ belongs to $(\mathbf{y}, \mathbf{X})$ and small value otherwise.

LEMMA 4.3. *Let $(\mathbf{y}, \mathbf{X})$ be an i.i.d. sample drawn from some distribution in $\mathcal{P}(\sigma, d, s^*, \Theta)$. Let $S = \text{supp}(\boldsymbol{\beta})$; assume that $\|\mathbf{x}_S\|_2 \leq 1$ and $\mathbf{x}_{S^c} = \mathbf{0}$ with probability 1, and that the restricted covariance matrix $\Sigma_S = \{\mathbb{E}(\mathbf{x}\mathbf{x}^\top)\}_{i,j \in S}$ is diagonal and satisfies $0 < 1/L < s^* \lambda_{\min}(\Sigma_{\mathbf{x}}) \leq s^* \lambda_{\max}(\Sigma_{\mathbf{x}}) < L$ for some constant $L = O(1)$.*

If $s^ = o(d^{1-\omega})$ for some fixed $\omega > 0$, then for every (ε, δ) -differentially private estimator M satisfying $\mathbb{E}_{\mathbf{y}, \mathbf{X}|\boldsymbol{\beta}} \|M(\mathbf{y}, \mathbf{X}) - \boldsymbol{\beta}\|_2^2 = o(1)$ at every $\boldsymbol{\beta} \in \Theta$, the following are true:*

1. *For each $i \in [n]$, let $(\mathbf{y}'_i, \mathbf{X}'_i)$ denote the data set obtained by replacing $(y_i, \mathbf{x}_i)$ in $(\mathbf{y}, \mathbf{X})$ with an independent copy, then $\mathbb{E} \mathcal{A}_{\boldsymbol{\beta}, s^*}((y_i, \mathbf{x}_i), M(\mathbf{y}'_i, \mathbf{X}'_i)) = 0$ and*

$$\mathbb{E} |\mathcal{A}_{\boldsymbol{\beta}, s^*}((y_i, \mathbf{x}_i), M(\mathbf{y}'_i, \mathbf{X}'_i))| \leq \sigma \sqrt{\mathbb{E} \|M(\mathbf{y}, \mathbf{X}) - \boldsymbol{\beta}\|_{\Sigma_{\mathbf{x}}}^2}.$$

2. *There exists a prior distribution of $\boldsymbol{\pi} = \boldsymbol{\pi}(\boldsymbol{\beta})$ over Θ such that*

$$\sum_{i \in [n]} \mathbb{E}_{\boldsymbol{\pi}} \mathbb{E}_{\mathbf{y}, \mathbf{X}|\boldsymbol{\beta}} \mathcal{A}_{\boldsymbol{\beta}, s^*}((y_i, \mathbf{x}_i), M(\mathbf{y}, \mathbf{X})) \gtrsim \sigma^2 s^* \log d.$$

The lemma is proved in Section B.6 of the Supplementary Material [10]. These properties of tracing attack (4.5) imply a minimax lower bound for (ε, δ) -differentially private estimation of $\boldsymbol{\beta}$.

THEOREM 4.3. *If $s^* = o(d^{1-\omega})$ for some fixed $\omega > 0$, $0 < \varepsilon < 1$ and $\delta < n^{-(1+\omega)}$ for some fixed $\omega > 0$, we have*

$$(4.6) \quad \inf_{M \in \mathcal{M}_{\varepsilon, \delta}} \sup_{\mathcal{P}(\sigma, d, s^*, \Theta)} \mathbb{E} \|M(\mathbf{y}, \mathbf{X}) - \boldsymbol{\beta}\|_{\Sigma_{\mathbf{x}}}^2 \gtrsim \sigma^2 \left(\frac{s^* \log d}{n} + \frac{(s^* \log d)^2}{n^2 \varepsilon^2} \right).$$

The lower bound is proved in Section B.7 of the Supplementary Material [10]. Similar to the cost of privacy in high-dimensional mean estimation, the lower bound here depends only logarithmically on dimension d . We show in the next section that this lower bound is achieved up to factors of $\log n$ by an (ε, δ) -differentially private algorithm.

4.4. *Algorithm for high-dimensional linear regression.* When the dimension of $\boldsymbol{\beta}$ exceeds the sample size, directly minimizing $\mathcal{L}_n(\boldsymbol{\beta}) = n^{-1} \sum_{i=1}^n (y_i - \mathbf{x}_i^\top \boldsymbol{\beta})^2$ no longer leads to an accurate estimate of $\boldsymbol{\beta}$, as seen from the rank deficiency of $\nabla^2 \mathcal{L}_n(\boldsymbol{\beta}) = n^{-1} \mathbf{X}^\top \mathbf{X}$. As a consequence, the differentially private, noisy gradient algorithm for the low-dimensional setting is no longer applicable.

To leverage the sparsity of $\boldsymbol{\beta}$, we recall the “peeling” algorithm for sparse mean estimation in Section 3.4, and arrive at the following modification of Algorithm 4.1.

If the “Peeling” step is replaced by nonprivate, exact projection of the gradient step onto $\{\mathbf{v} \in \mathbb{R}^d : \|\mathbf{v}\|_0 \leq s\}$, we recover the well-known iterative hard thresholding algorithm [7, 25] for high-dimensional sparse regression.

The analysis of Algorithm 4.2 requires some assumptions similar to their low-dimensional counterparts in Section 4.2, as follows.

(P1') The true parameter vector $\boldsymbol{\beta}$ satisfies $\|\boldsymbol{\beta}\|_2 < c_0$ for some constant $0 < c_0 < \infty$ and $\|\boldsymbol{\beta}\|_0 \leq s^* = o(n)$.

(D1') Bounded design: for every index set $I \subseteq [d]$ with $|I| = o(n)$, there is a constant $c_{\mathbf{x}} < \infty$ such that $\sqrt{|I|} \|\mathbf{x}_I\|_\infty < c_{\mathbf{x}}$ with probability 1.

(D2') Bounded moments of design: $\mathbb{E} \mathbf{x} = \mathbf{0}$ and for every index set $I \subseteq [d]$ with $|I| = o(n)$, the (restricted) covariance matrix $\Sigma_I = \mathbb{E} \mathbf{x}_I \mathbf{x}_I^\top$ satisfies $0 < 1/L < |I| \cdot \lambda_{\min}(\Sigma_I) \leq |I| \cdot \lambda_{\max}(\Sigma_I) < L$ for some constant $0 < L < \infty$.

Algorithm 4.2: Differentially Private Sparse Linear Regression

Input : $\mathcal{L}_n(\beta)$, data set $(\mathbf{y}, \mathbf{X}) = \{(y_i, \mathbf{x}_i)\}_{i \in [n]}$, step size η^0 , privacy parameters ϵ, δ , noise scale B , number of iterations T , truncation level R , feasibility parameter C , sparsity s , initial value β^0 .

1 **for** t in 0 to $T - 1$ **do**

2 Compute $\beta^{t+0.5} = \beta^t - (\eta^0/n) \sum_{i=1}^n (\mathbf{x}_i^\top \beta^t - \Pi_R(y_i)) \mathbf{x}_i$;

3 $\beta^{t+1} = \Pi_C(\text{Peeling}(\beta^{t+0.5}, (\mathbf{y}, \mathbf{X}), s, \epsilon/T, \delta/T, \eta^0 B/n))$.

4 **end**

Output: β^T .

These assumptions can be understood as restricted versions of their counterparts, (P1), (D1) and (D2), in the low-dimensional case, Section 4.2. When assumptions (P1') and (D1') hold, the algorithm is guaranteed to be (ϵ, δ) -differentially private as long as the noise level B is chosen properly.

LEMMA 4.4. *If assumption (P1') and (D1') are true, then Algorithm 4.2 is (ϵ, δ) -differentially private as long as $B \geq 4(R + c_0 c_x) c_x / \sqrt{s}$.*

The lemma is proved in Section A.6. With assumption (D2') in addition, we can obtain the following convergence result for Algorithm 4.2.

THEOREM 4.4. *Let $\{(y_i, \mathbf{x}_i)\}_{i \in [n]}$ be an i.i.d. sample from the linear model (4.1). Suppose assumptions (P1'), (D1') and (D2') are true. Let $R = \sigma \sqrt{2 \log n}$, $C = c_0$ and $B = 4(R + c_0 c_x) c_x / \sqrt{s}$ in accordance with Lemma 4.4, and $\beta^0 = \mathbf{0}$. Then there exists some absolute constant ρ such that, if $s = \rho L^4 s^*$, $\eta^0 = s/6L$, $T = \rho L^2 \log(8c_0^2 L n)$ and $n \geq K \cdot (R(s^*)^{3/2} \log d \sqrt{\log(1/\delta)} \log n / \epsilon)$ for a sufficiently large constant K , the bound*

$$(4.7) \quad \|\beta^T - \beta\|_{\Sigma_x}^2 \lesssim \sigma^2 \left(\frac{s^* \log d}{n} + \frac{(s^* \log d)^2 \log(1/\delta) \log^3 n}{n^2 \epsilon^2} \right)$$

holds with probability at least $1 - c_1 \exp(-c_2 \log(d/s^ \log n)) - c_1 \exp(-c_2 n) - c_1 \exp(-c_2 \log n)$.*

The theorem is proved in Section 8.3. This convergence rate attains the corresponding lower bound (4.6) up to factors of $\log n$, for the usual choice of $\delta = n^{-(1+\omega)}$. For selecting tuning parameters R and s in Algorithm 4.2, we demonstrate in Section 5 data-driven and differentially private alternatives to the theoretical choices required by Theorem 4.4.

5. Simulation studies. In this section, we perform simulation studies of our algorithms to evaluate their numerical performance and demonstrate the cost of privacy in various estimation problems. The data are generated as follows.

Mean estimation $\mathbf{x}_1, \dots, \mathbf{x}_n$ are independently drawn from $N_d(\boldsymbol{\mu}, \mathbf{I}_d)$. Over repetitions of the experiments, the coordinates of $\boldsymbol{\mu}$ are sample i.i.d. from Uniform $(-10, 10)$ for the low-dimensional problem; in the high-dimensional case, the first s^* coordinates of $\boldsymbol{\mu}$ are sampled i.i.d. from Uniform $(-10, 10)$ and the other coordinates are set to 0.

Linear regression The data $(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_n, y_n)$ are generated from the linear model $y_i = \mathbf{x}_i^\top \beta + \epsilon_i$. The entries of design matrix are sampled i.i.d from the uniform distribution over $(-1/\sqrt{d}, 1/\sqrt{d})$ so that the row normalization assumption (D1) in Section 4 is satisfied;

$\epsilon_1, \dots, \epsilon_n$ is an i.i.d sample from $N(0, 1)$. β is sampled uniformly from the unit sphere $\{v \in \mathbb{R}^d : \|v\|_2 = 1\}$ for the low-dimensional problem; in the high-dimensional problem, the vector of first s coordinates is sampled uniformly from the unit sphere $\{v \in \mathbb{R}^{s^*} : \|v\|_2 = 1\}$, and the other coordinates are set to 0.

We shall carry out three sets of experiments with the simulated data:

- Compare the performance of our algorithms under different choices of R , the truncation tuning parameter.
- Compare the performance of the high-dimensional algorithms under different choices of s , the sparsity tuning parameter.
- Compare our algorithms with their nonprivate counterparts, and with other differentially private algorithms in the literature.

5.1. Tuning of truncation level. For each of our four algorithms, we consider three methods of determining the truncation tuning parameter R :

- No truncation.
- The theoretical choice: R is set to be the theoretical value of $4\sigma\sqrt{\log n}$.
- Data-driven: compute differentially private estimates of the data set’s 2.5% and 97.5% percentiles by Algorithm 1’ in [31] (see “Extension to distributions supported on $(-\infty, \infty)$,” page 6), and truncate the data set at these levels.

As shown in Figure 1 below, the data-driven method incurs comparable errors to the no truncation case and the theoretical choice of R , suggesting that it is a viable method for choosing R in practice. It should be cautioned that the optimistic performance of constant quantile

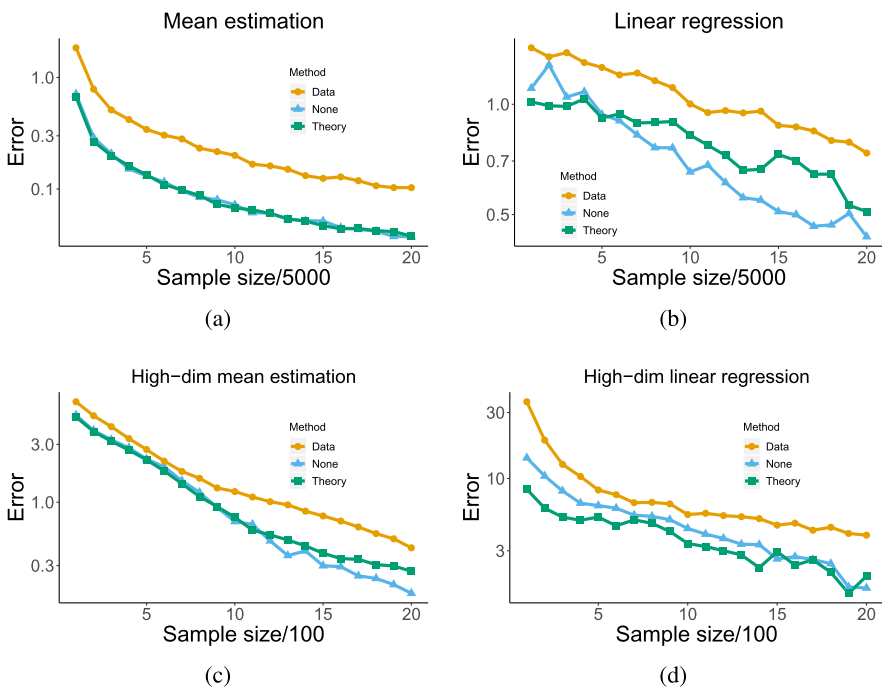


FIG. 1. Average ℓ_2 -error over 100 repetitions plotted against sample size n , with privacy level set at $(0.5, 10/n^{1.1})$. (a) and (b): mean estimation and linear regression with $d = 20$ and n from 5000 to 100,000. (c) and (d): high-dimensional mean estimation and linear regression with n increasing from 100 to 2000, $d = n$ and $s = 20$.

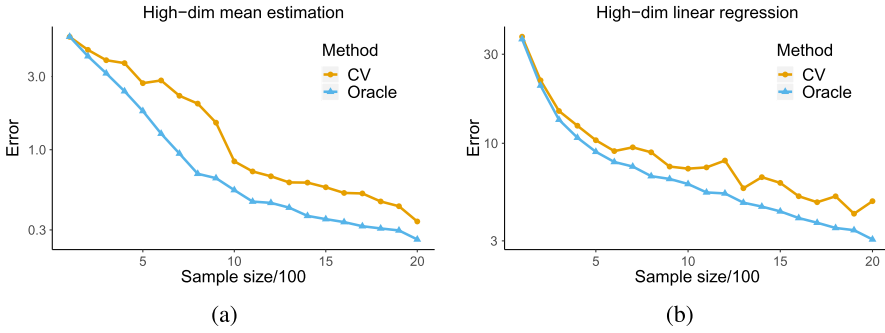


FIG. 2. Average ℓ_2 -error over 50 repetitions plotted against sample size n , with privacy level set at $(0.5, 10/n^{1.1})$. (a) and (b): high-dimensional mean estimation and linear regression with n increasing from 100 to 2000, $d = n$ and $s = 20$.

truncation benefits from the symmetry and light-tailedness of the Gaussian distribution; it may not be applicable to all types of data distribution.

5.2. Tuning of s . Our algorithms for high-dimensional problems require a sparsity tuning parameter s . We compare their performances when supplied with the true sparsity s^* and when s is chosen by 5-fold cross validation. The cross-validation error is first computed over a uniform grid of values from $s^*/2$ to $2s^*$. We then truncate these cross-validation errors with Algorithm 1' in [31], so that the truncated cross-validation errors have bounded sensitivity. With bounded sensitivity, the exponential mechanism [34] can be applied to the (truncated) cross-validation errors to select a value of s in a differentially private manner.

Informed by the previous section on tuning R , the truncation tuning parameters for experiments in this section are selected by the data-driven method. In each problem, as the Figures 2 and 3 show, selecting s by cross validation leads to errors comparable with their counterparts when the algorithms are supplied with the true sparsity s^* .

5.3. Comparisons with other algorithms. We compare our algorithms with their nonprivate counterparts, as well as other differentially private algorithms in the literature. For the low-dimensional problems, we consider Algorithm 4 in [28] for mean estimation and Algorithm 1 in [40] for regression. For the high-dimensional problems, we compare with the method in [44].

There are significant gaps in performance between our algorithms and those in [40, 44]. It is important to note, however, that the primary strength of Algorithm 1 in [40] is its ability to produce accurate test statistics with differential privacy, and the algorithm by [44] is primarily targeted at minimizing the excess empirical risk, so these numerical experiments may not be fully reflective of their advantages.

To further understand the improved numerical performance, we report here some observations from the numerical experiments. For the private Johnson–Lindenstrauss projection algorithm in [40], we observed that the ridge regression subroutine of the algorithm is frequently activated even when n is very large, resulting in a ridge regression solution with regularization parameter of order $O(\log(1/\delta)/\epsilon)$ and leading to significant bias. For the private Frank–Wolfe algorithm in [44], the solution is often nonsparse with large values outside the true support of β , while our algorithm guarantees a sparse solution by construction and converges to the nonprivate solution as n grows.

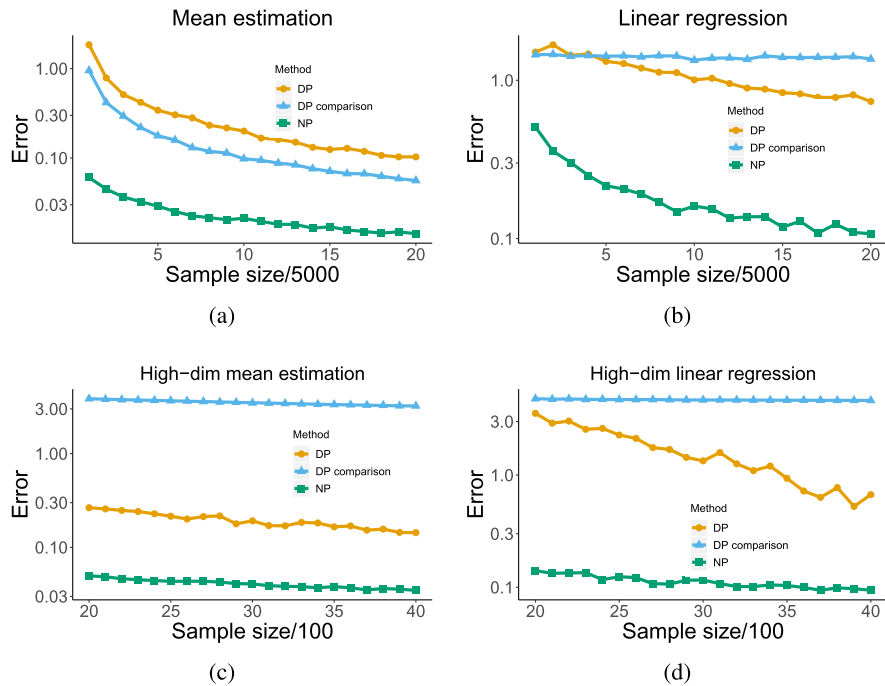


FIG. 3. Average ℓ_2 -error over 100 repetitions plotted against sample size n , with privacy level set at $(0.5, 10/n^{1.1})$. (a) and (b): mean estimation and linear regression with fixed $d = 20$ and n increasing from 5000 to 100,000. (c) and (d): high-dimensional mean estimation and linear regression with n increasing from 2000 to 4000, $d = 2n$ and $s = 20$.

6. Data analysis. In this section, we demonstrate the numerical performance of the differentially private algorithms on real data sets.

6.1. SNP array of adults with schizophrenia. We analyze the SNP array data of adults with schizophrenia, collected by [33], to illustrate the performance of our high-dimensional sparse mean estimator. In the data set, there are 387 adults with schizophrenia, 241 of which are labeled as “average IQ” and 146 of which are labeled as “low IQ.” The SNP array is obtained by genotyping the subjects with the Affymetrix Genome-Wide Human SNP 6.0 platform. For our analysis, we focus on the 2000 SNPs with the highest minor allele frequencies (MAFs); the full data set is available at <https://www.ncbi.nlm.nih.gov/geo/query/acc.cgi?acc=GSE106818>.

Privacy-preserving data analysis is very much relevant for this data set and genetic data in general, because as [24] shows, an adversary can infer the absence/presence of an individual’s genetic data in a large data set by cross-referencing summary statistics, such as MAFs, from multiple genetic data sets. As MAFs can be calculated from the mean of an SNP array, differentially-private estimators of the mean allow reporting the MAFs without compromising any individual’s privacy.

The data set takes the form of a 387×2000 matrix. The entries of the matrix take values 0, 1 or 2, representing the number of minor allele(s) at each SNP and, therefore, the MAF of each SNP location in this sample can be obtained by computing the mean of the rows in this matrix. Sparsity is introduced by considering the difference in MAFs of the two IQ groups: the MAFs of the two groups are likely to differ at a small number of SNP locations among the 2000 SNPs considered. For m ranging from 10 to 120, we subsample m subjects from each of the two IQ groups, say $\{\mathbf{x}_{11}, \mathbf{x}_{12}, \dots, \mathbf{x}_{1m}\}$ and $\{\mathbf{x}_{21}, \mathbf{x}_{22}, \dots, \mathbf{x}_{2m}\}$, and apply our sparse mean estimator to $\{\mathbf{x}_{11} - \mathbf{x}_{21}, \mathbf{x}_{12} - \mathbf{x}_{22}, \dots, \mathbf{x}_{1m} - \mathbf{x}_{2m}\}$ with $s = 20$ and privacy parameters

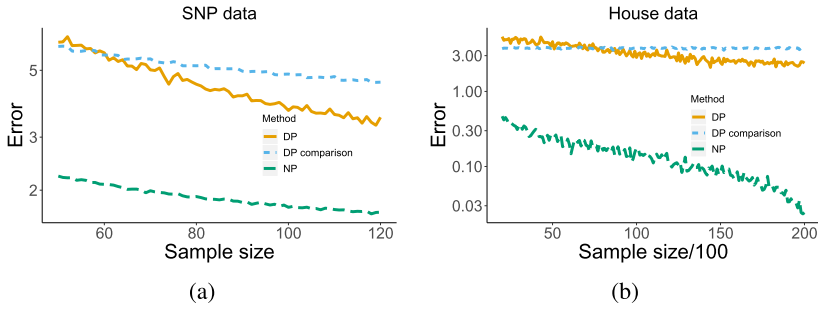


FIG. 4. (a): The estimate of $\mathbb{E}[\|\hat{\mu} - \mu\|_2]$ for the differentially private sparse mean estimator as sample size increases from 50 to 120, with $s = 20$. (b): The estimate of $\mathbb{E}[\|\hat{\beta} - \beta\|_2]$ for the differentially private OLS estimator, compared with its differentially private counterpart, as sample size increases from 2000 to 20,000.

$(\epsilon, \delta) = (0.5, 10/n^{1.1})$. The error of this estimator is then calculated by comparing with the mean of the entire sample. This procedure is repeated 100 times to obtain Figure 4(a), which displays the estimate of $\mathbb{E}[\|\hat{\mu} - \mu\|_2]$ as m increases from 50 to 120. We also plotted the corresponding curve for the method in [44] for comparison.

6.2. Housing prices in California. For the linear regression problem, we analyze a housing price data set with economic and demographic covariates, constructed by [37] and available for download at <http://lib.stat.cmu.edu/datasets/houses.zip>. In this data set, each subject is a block group in California in the 1990 Census; there are 20,640 block groups in this data set. The response variable is the median house value in the block group; the covariates include the median income, median age, total population, number of households and the total number of rooms of all houses in the block group. In general, summary statistics such as mean or median do not have any differential privacy guarantees, so the absence of information on individual households in the data set does not preclude an adversary from extracting sensitive individual information from the summary statistics. Privacy-preserving methods are still desirable in this case.

For m ranging from 100 to 20,600, we subsample m subjects from the data set to compute the differentially private OLS estimate, with privacy parameters $(\epsilon, \delta) = (0.5, 10/n^{1.1})$. The error of this estimator is then calculated by comparing with the nonprivate OLS estimator computed using the entire sample. This procedure is repeated 100 times to obtain Figure 4(b), which displays the estimate of $\mathbb{E}[\|\hat{\beta} - \beta\|_2]$ as m increases from 2000 to 20,000. The design matrix is standardized before applying the algorithm. The corresponding curve for the method in [40] is also plotted for comparison.

7. Discussion. Our paper investigates the tradeoff between statistical accuracy and privacy, by providing minimax lower bounds with differential privacy constraint and proposing differentially private algorithms with rates of convergence attaining the lower bounds up to logarithmic factors. For the lower bounds, we considered a technique based on tracing adversary and illustrated its utility by establishing minimax lower bounds for differentially private mean estimation and linear regression. These lower bounds are shown to be tight up to logarithmic factors via analysis of differentially private algorithms with matching rates of convergence.

Beyond the theoretical results, numerical performance of the private algorithms are demonstrated in simulations and real data analysis. The results suggest that the proposed algorithms have robust performance with respect to various choices of tuning parameters, achieve accuracy comparable to or better than that of existing differentially private algorithms, and can compute efficiently for sample sizes and dimensions up to tens of thousands.

The numerical results corroborate the cost of privacy delineated in the theorems by exhibiting shrinking but nonvanishing gaps of accuracy between the private algorithms and their non-private counterparts. The theoretical and numerical results together can inform practitioners of differential privacy the necessary sacrifice of accuracy at a prescribed level of privacy, or the appropriate choice of privacy parameters if a given level of accuracy is desired.

There are many promising avenues for future research. It is of significant interest to study the optimal tradeoff of privacy and accuracy in statistical problems beyond mean estimation and linear regression. Examples include covariance/precision matrix estimation, graphical model recovery, nonparametric regression and principal component analysis. Along the way, it is of importance to further develop general approaches of designing privacy-preserving algorithms, as well as more general lower bound techniques than those presented in this work.

One natural extension is uncertainty quantification with privacy constraints, which is largely unexplored in the statistics literature. Notably, [28] established the rate-optimal length of differentially private confidence intervals for the (one-dimensional) Gaussian mean. The technical tools developed in our paper may provide insights for constructing optimal statistical inference procedures in the context of, say, high-dimensional sparse mean estimation and linear regression.

Yet another intriguing direction of research is the cost of other notions of privacy, such as concentrated differential privacy [18], Rényi differential privacy [35] and Gaussian differential privacy [12]. These notions of privacy have found important applications such as stochastic gradient Langevin dynamics, stochastic Monte Carlo sampling [47] and deep learning [8].

8. Proofs. In this section, we prove the lower bound of low-dimensional mean estimation, Lemma 3.1 and Theorem 3.1, and the upper bound of high-dimensional linear regression, Theorem 4.4.

8.1. *Proof of Lemma 3.1.*

PROOF OF LEMMA 3.1. Let $k = (C/2) \log(\frac{1}{n\delta})/\varepsilon$, with the value of $0 < C < 1$ to be chosen later. By the assumed regime of δ as a function of n , we have $k \asymp \log(1/\delta)/\varepsilon$ and $k < n/2$. We assume that k divides n without the loss of generality.

For an arbitrary $M \in \mathcal{M}_{\varepsilon,\delta}$, we define $M_k(\mathbf{Z}) \equiv M(\mathbf{Y})$. Because M is (ε, δ) -differentially private, M_k is also differentially private by post-processing. To lower bound $\mathbb{E}[\|\mathbf{M}(\mathbf{Y}) - \mathbb{E}\mathbf{y}_1\|_2|\mathbf{Z}]$, we observe that it suffices to find some appropriate distribution of $\mathbf{Z}$ so that $\mathbb{E}\|M_k(\mathbf{Z}) - \bar{\mathbf{Z}}\|_2$ can be lower bounded: as $\|\mathbf{M}(\mathbf{Y}) - \mathbb{E}\mathbf{y}_1\|_2 = \|M_k(\mathbf{Z}) - \bar{\mathbf{Z}}\|_2$ by construction, there must be a realization of $\mathbf{Z}$ such that $\mathbb{E}[\|\mathbf{M}(\mathbf{Y}) - \mathbb{E}\mathbf{y}_1\|_2|\mathbf{Z}]$ is also lower bounded.

Since the sample size of $\mathbf{Z}$ does satisfy the assumption of the preliminary lower bound (2.3), the bound applies provided that M_k is a differentially private algorithm with respect to $\mathbf{Z}$. To this end, we consider the group privacy lemma.

LEMMA 8.1 (Group privacy [42]). *For every $m \geq 1$, if M is (ε, δ) -differentially private, then for every pair of data sets $\mathbf{X} = \{\mathbf{x}_k\}_k$ and $\mathbf{Z} = \{\mathbf{z}_k\}_k$ satisfying $\sum_i \mathbb{1}(\mathbf{x}_i \neq \mathbf{z}_i) \leq m$, and every measurable set S ,*

$$\mathbb{P}(M(\mathbf{X}) \in S) \leq e^{\varepsilon m} \mathbb{P}(M(\mathbf{Z}) \in S) + \frac{e^{\varepsilon m} - 1}{e - 1} \cdot \delta.$$

The group privacy lemma suggests that, to characterize the privacy parameters of M_k , it suffices to upper-bound the number of changes in $\mathbf{Y}$ incurred by replacing one element

of $\mathbf{Z}$. Let m_i denote the number of times that $\mathbf{z}_i$ appears in a sample of size n drawn with replacement from $\mathbf{Z}$, then our quantity of interest here is simply $\max_{i \in [n/k]} m_i$.

To analyze $\max_{i \in [n/k]} m_i$, we first show that δ as a function of n must satisfy one of the following two statements:

1. There is a fixed constant τ such that $\frac{n}{(n/k) \log(n/k)} \leq C\tau/\varepsilon$ when n is sufficiently large.
2. Case 1 fails to hold: we have $k > (C/2)\tau/\varepsilon \log(n/k)$ for any constant τ , as long as n is sufficiently large.

The dichotomy is made possible by the assumption that $\log(\delta)/\log(n)$ is nonincreasing in n and $\delta < n^{-(1+\omega)}$ for some fixed $\omega > 0$. Under this assumption, we have either $\lim_{n \rightarrow \infty} \log(\delta)/\log(n) = c < -1$ or $\lim_{n \rightarrow \infty} \log(\delta)/\log(n) = -\infty$.

For the first case, we have $k = (C/2) \log(\frac{1}{n\delta})/\varepsilon = (C/2\varepsilon) \cdot (\log(1/\delta) - \log(n)) \in ((C/2\varepsilon) \cdot c_1 \log n, (C/2\varepsilon) \cdot c_2 \log n)$ for some $c_1, c_2 > 0$ when n is sufficiently large. Therefore,

$$\frac{n}{(n/k) \log(n/k)} = \frac{k}{\log(n/k)} \leq \frac{c_2 \cdot (C/2\varepsilon) \cdot \log n}{\log(2n\varepsilon/(C \cdot c_1 \cdot \log n))} \leq C\tau/\varepsilon,$$

for some $\tau > 0$. This corresponds to the first statement.

For the second case $\lim_{n \rightarrow \infty} \log(\delta)/\log(n) = -\infty$, we then have for any constant $c_3 > 0$ and sufficiently large n , $\log(1/\delta) > c_3 \log(n)$. Then $k = (C/2) \log(\frac{1}{n\delta})/\varepsilon = (C/2\varepsilon) \cdot (\log(1/\delta) - \log(n)) \geq (C/2\varepsilon) \cdot (c_3 - 1) \log n$. Consequently, we have

$$\frac{k}{\log(n/k)} > \frac{(c_3 - 1) \cdot (C/2\varepsilon) \cdot \log n}{\log(n)} \leq (c_3 - 1) \cdot (C/2\varepsilon)$$

for sufficiently large n . Since c_3 can take value of any positive number, this corresponds to the second statement.

Case 1. As $(m_1, m_2, \dots, m_{n/k})$ follows a uniform multinomial distribution, we consider a useful result from [38], stated below.

LEMMA 8.2 ([38]). *If $(y_1, y_2, \dots, y_d)$ follows a uniform multinomial(ℓ) distribution, and $\frac{\ell}{d \log d} \leq c$ for some constant c not depending on ℓ and d , then for every $\zeta > 0$,*

$$\mathbb{P}\left(\max_{i \in [d]} y_i > (r_c + \zeta) \log d\right) = o(1),$$

where r_c is the unique root of $1 + y(\log c - \log y + 1) - c = 0$ that is strictly greater than c .

It follows that $\mathbb{P}(\max_i m_i \leq r \log n) = 1 - o(1)$, where r is the unique root of $1 + x(\log(C\tau/\varepsilon) - \log x + 1) - (C\tau/\varepsilon) = 0$ that is greater than $C\tau/\varepsilon$. Such a root exists, because $f_{C,\tau,\varepsilon}(x) := 1 + x(\log(C\tau/\varepsilon) - \log x + 1) - (C\tau/\varepsilon)$ is strictly concave and achieves the global maximum value of 1 at $x = C\tau/\varepsilon$.

Let $\mathcal{E} := \{\max_i m_i \leq r \log n\}$. Under $\mathcal{E}$, Lemma 8.1 implies that M_k is an $(\varepsilon r \log n, \delta e^{\varepsilon r \log n})$ -differentially private algorithm. We may essentially repeat the lower bound argument leading to the preliminary lower bound (2.3), as follows. Let $\mathbf{Z} = \{\mathbf{z}_1, \mathbf{z}_2, \dots, \mathbf{z}_{n/k}\}$ be sampled i.i.d., from the data distribution specified in Lemma 2.1 including the prior distribution on $\boldsymbol{\mu} = \mathbb{E} \mathbf{z}_1$, so that Lemma 2.1 applies to $\mathbf{Z}$. For every $i \in [n/k]$, let $\mathcal{C} = \{\sum_{i \in [n/k]} \mathcal{A}_{\boldsymbol{\mu}}(\mathbf{z}_i, M_k(\mathbf{Z})) \leq (n/k)\sigma^2 \sqrt{8d \log(1/\delta)}\}$. We have

$$\begin{aligned} & \mathbb{P}(\|M_k(\mathbf{Z}) - \bar{\mathbf{Z}}\|_2 < c\sigma\sqrt{d}) \\ & \leq \mathbb{P}(\mathcal{E}^c) + \mathbb{P}(\mathcal{C} \cap \{\|M_k(\mathbf{Z}) - \bar{\mathbf{Z}}\|_2 < c\sigma\sqrt{d}\}) + \mathbb{P}(\mathcal{C}^c) \\ & \leq \mathbb{P}(\mathcal{E}^c) + \mathbb{P}(\mathcal{C} \cap \{\|M_k(\mathbf{Z}) - \bar{\mathbf{Z}}\|_2 < c\sigma\sqrt{d}\}) \end{aligned}$$

$$\begin{aligned}
 & + \sum_{i \in [n/k]} \mathbb{P}(\mathcal{A}_\mu(\mathbf{z}_i, \mathbf{M}_k(\mathbf{Z})) > \sigma^2 \sqrt{8d \log(1/\delta)}) \\
 & \leq o(1) + \delta + n(e^{\varepsilon r \log n} \delta + \delta e^{\varepsilon r \log n}) = o(1) + \delta + 2n^{-\tau + \varepsilon r}.
 \end{aligned}$$

This probability is always bounded away from 1, because $\varepsilon r < \tau$ with appropriately chosen C : since $f_{C, \tau, \varepsilon}(\tau/\varepsilon) = (\tau/\varepsilon)(1 + \log C - C) + 1$ and $0 < \varepsilon < 1$, for every $\tau > 0$ there is a sufficiently small $0 < C < 1$ such that $f_{C, \tau, \varepsilon}(\tau/\varepsilon) < 0$. Since $f_{C, \tau, \varepsilon}(C\tau/\varepsilon) = 1$ is the global maximum, we have $r < \tau/\varepsilon$, or equivalently $\varepsilon r < \tau$, as desired.

Case 2. Each m_i is a sum of n independent Bernoulli(k/n) random variables. Chernoff's inequality implies that

$$\mathbb{P}\left(\max_i m_i > \frac{1}{\varepsilon} \log\left(\frac{1}{2n\delta}\right)\right) \leq \frac{n}{k} \cdot \exp\left(-\frac{(1/2C - 1)}{3}k\right).$$

Recall that $k = (C/2) \log(\frac{1}{2n\delta})/\varepsilon$ by construction. By the assumption of Case 2, we have $k > (C/2)\tau/\varepsilon \log(n/k)$ for any constant τ , as long as n is sufficiently large. Then we have

$$\mathbb{P}\left(\max_i m_i > \frac{1}{\varepsilon} \log\left(\frac{1}{2n\delta}\right)\right) \leq \left(\frac{n}{k}\right)^{1 - (1 - C/2)\tau/3\varepsilon}.$$

The probability can be made arbitrarily small by fixing $C = 1/2$ and choosing large τ . Now that we have a high-probability bound for $\max_i m_i$, the group privacy lemma and union bound, as in Case 1, imply that $\mathbb{P}(\mathcal{C}) = o(1)$ and, therefore, by Lemma 2.1,

$$\begin{aligned}
 & \mathbb{P}(\|\mathbf{M}_k(\mathbf{Z}) - \bar{\mathbf{Z}}\|_2 < c\sigma\sqrt{d}) \\
 & \leq \mathbb{P}(\mathcal{E}^c) + \mathbb{P}(\mathcal{C} \cap \{\|\mathbf{M}_k(\mathbf{Z}) - \bar{\mathbf{Z}}\|_2 < c\sigma\sqrt{d}\}) + \mathbb{P}(\mathcal{C}^c) = o(1).
 \end{aligned}$$

In each of the two cases, we found that $\mathbb{P}(\|\mathbf{M}_k(\mathbf{Z}) - \bar{\mathbf{Z}}\|_2 < c\sigma\sqrt{d}) = o(1)$. The proof is now complete by the reduction from $\mathbb{E}[\|\mathbf{M}(\mathbf{Y}) - \mathbb{E}\mathbf{y}_1\|_2 | \mathbf{Z}]$ to $\mathbb{E}\|\mathbf{M}_k(\mathbf{Z}) - \bar{\mathbf{Z}}\|_2$. $\square$

8.2. Proof of Theorem 3.1. It suffices to prove the second term of the minimax lower bound, as the first term is the statistical minimax lower bound for sub-Gaussian mean estimation.

For $i \in [n]$, consider $\mathbf{x}_i = \mathbf{0} \in \mathbb{R}^d$ with probability $1 - \alpha$ and $\mathbf{x}_i = \mathbf{y}_i$ with probability α , where $\mathbf{y}_i$ follows the discrete uniform distribution specified in Lemma 3.1. When $n \gtrsim \sqrt{d \log(1/\delta)}/\varepsilon$, there exists some $0 < \alpha < 1$ such that $\alpha n \asymp \sqrt{d \log(1/\delta)}/\varepsilon$. The distribution of $\mathbf{x}_i$ is indeed sub-Gaussian(σ) with $\boldsymbol{\mu} \in \Theta$.

Consider the random index set $\mathcal{S} = \{i \in [n] : \mathbf{x}_i \neq \mathbf{0}\}$. For every $M \in \mathcal{M}_{\varepsilon, \delta}$, we have

$$\mathbb{E}[\|\mathbf{M}(\mathbf{X}) - \boldsymbol{\mu}\|_2] \geq \sum_{\mathcal{S}=\mathcal{S} \subseteq [n], |\mathcal{S}| \leq n\alpha} \mathbb{E}[\|\mathbf{M}(\mathbf{X}) - \boldsymbol{\mu}\|_2 | \mathcal{S} = \mathcal{S}] \mathbb{P}(\mathcal{S} = \mathcal{S}).$$

Now for each fixed $\mathcal{S}$, define $\tilde{M}(\mathbf{X}_{\mathcal{S}}) = \alpha^{-1} M(\{\mathbf{x}_i : i \in \mathcal{S}\} \cup \{\mathbf{0}\}^{n-|\mathcal{S}|})$. We note that $\tilde{M}(\mathbf{X}_{\mathcal{S}})$ is an (ε, δ) -differentially private algorithm with respect to $\mathbf{X}_{\mathcal{S}} = \{\mathbf{x}_i : i \in \mathcal{S}\}$, by observing that modifying any single datum in $\mathbf{X}_{\mathcal{S}}$ incurs the same privacy loss to M as it does to $\tilde{M}$. By construction, it also holds that $\boldsymbol{\mu} = \mathbb{E}\mathbf{x}_1 = \alpha \mathbb{E}\mathbf{y}_1$. We then have

$$\begin{aligned}
 \mathbb{E}[\|\mathbf{M}(\mathbf{X}) - \boldsymbol{\mu}\|_2 | \mathcal{S} = \mathcal{S}] & \geq \mathbb{E}[\|\alpha \tilde{M}(\mathbf{X}_{\mathcal{S}}) - \alpha \mathbb{E}\mathbf{y}_1\|_2 | \mathcal{S} = \mathcal{S}] \\
 & \geq \alpha \mathbb{E}[\|\tilde{M}(\mathbf{X}_{\mathcal{S}}) - \mathbb{E}\mathbf{y}_1\|_2] \gtrsim \alpha \sigma \sqrt{d} \asymp \sigma \frac{d \sqrt{\log(1/\delta)}}{n\varepsilon}.
 \end{aligned}$$

For the last inequality, we invoked the lower bound proved in Lemma 3.1, since the sample size of $\mathbf{X}_{\mathcal{S}}$ is at most $\alpha n \asymp \sqrt{d \log(1/\delta)}/\varepsilon$. The proof is complete.

8.3. Proof of Theorem 4.4.

PROOF OF THEOREM 4.4. Let $\hat{\beta} = \arg \min_{\|\beta\|_2 \leq c_0, \|\beta\|_0 \leq s^*} \mathcal{L}_n(\beta)$. While global strong convexity and smoothness are no longer possible when $d > n$, because β^t and $\hat{\beta}$ are sparse, we have following fact known as restricted strong convexity (RSC) and restricted smoothness (RSM) [3, 32, 36].

FACT 8.1. Under assumptions of Theorem 4.4, it holds with probability at least $1 - c_1 \exp(-c_2 n)$ that

$$(8.1) \quad \frac{1}{8Ls} \|\beta^t - \hat{\beta}\|_2^2 \leq \langle \nabla \mathcal{L}_n(\beta^t) - \nabla \mathcal{L}_n(\hat{\beta}), \beta^t - \hat{\beta} \rangle \leq \frac{4L}{s} \|\beta^t - \hat{\beta}\|_2^2.$$

Under the event $\mathcal{E}_1 = \{\Pi_R(y_i) = y_i, \forall i \in [n]\}$, Fact 8.1 implies that $\mathcal{L}_n(\beta^t) - \mathcal{L}_n(\hat{\beta})$ decays exponentially fast in t .

LEMMA 8.3. Under assumptions of Theorem 4.4 and event $\mathcal{E}_1$, (8.1) implies that there exists an absolute constant ρ such that

$$(8.2) \quad \mathcal{L}_n(\beta^{t+1}) - \mathcal{L}_n(\hat{\beta}) \leq \left(1 - \frac{1}{\rho L^2}\right) (\mathcal{L}_n(\beta^t) - \mathcal{L}_n(\hat{\beta})) + c_3 \left(\sum_{i \in [s]} \|\mathbf{w}_i^t\|_\infty^2 + \|\tilde{\mathbf{w}}_{S^{t+1}}^t\|_2^2 \right),$$

for every t , where $\mathbf{w}_1^t, \mathbf{w}_2^t, \dots, \mathbf{w}_s^t$ are the Laplace noise vectors added to $\beta^t - (\eta^0/n) \times \sum_{i=1}^n (\mathbf{x}_i^\top \beta^t - \Pi_R(y_i)) \mathbf{x}_i$ when the support of β^{t+1} is iteratively selected by “Peeling,” S^{t+1} is the support of β^{t+1} and $\tilde{\mathbf{w}}^t$ is the noise vector added to the selected s -sparse vector.

We take Lemma 8.3, which is proved in Section A.7 of the Supplementary Material [10], to prove Theorem 4.4. We iterate (8.2) over t and notate $\mathbf{W}_t = c_3 (\sum_{i \in [s]} \|\mathbf{w}_i^t\|_\infty^2 + \|\tilde{\mathbf{w}}_{S^{t+1}}^t\|_2^2)$ to obtain

$$(8.3) \quad \begin{aligned} \mathcal{L}_n(\beta^T) - \mathcal{L}_n(\hat{\beta}) &\leq \left(1 - \frac{1}{\rho L^2}\right)^T (\mathcal{L}_n(\beta^0) - \mathcal{L}_n(\hat{\beta})) + \sum_{k=0}^{T-1} \left(1 - \frac{1}{\rho L^2}\right)^{T-k-1} \mathbf{W}_k \\ &\leq \left(1 - \frac{1}{\rho L^2}\right)^T 8Lc_0^2 + \sum_{k=0}^{T-1} \left(1 - \frac{1}{\rho L^2}\right)^{T-k-1} \mathbf{W}_k. \end{aligned}$$

The second inequality is a consequence of the upper inequality in (8.1) and the ℓ_2 bounds of β^0 and $\hat{\beta}$. We can also bound $\mathcal{L}_n(\beta^T) - \mathcal{L}_n(\hat{\beta})$ from below by the lower inequality in (8.1):

$$(8.4) \quad \mathcal{L}_n(\beta^T) - \mathcal{L}_n(\hat{\beta}) \geq \mathcal{L}_n(\beta^T) - \mathcal{L}_n(\beta^*) \geq \frac{1}{16Ls} \|\beta^T - \beta^*\|_2^2 - \langle \nabla \mathcal{L}_n(\beta^*), \beta^* - \beta^T \rangle.$$

Now (8.3) and (8.4) imply that, with $T = (\rho L^2) \log(8c_0^2 Ln)$,

$$(8.5) \quad \begin{aligned} \frac{1}{16Ls} \|\beta^T - \beta^*\|_2^2 &\leq \|\nabla \mathcal{L}_n(\beta^*)\|_\infty \sqrt{s + s^*} \|\beta^* - \beta^T\|_2 \\ &+ \frac{1}{n} + \sum_{k=0}^{T-1} \left(1 - \frac{1}{\rho L^2}\right)^{T-k-1} \mathbf{W}_k. \end{aligned}$$

To further bound $\|\beta^T - \beta^*\|_2^2$, we observe that under $\mathcal{E}_1$ and two other events,

$$\mathcal{E}_2 = \left\{ \max_t W_t \leq K \frac{R^2 (s^*)^3 \log^2 d \log(1/\delta) \log^2 n}{n^2 \varepsilon^2} \right\},$$

$$\mathcal{E}_3 = \left\{ \|\nabla \mathcal{L}_n(\beta^*)\|_\infty \leq 4\sigma \|\mathbf{x}\|_\infty \sqrt{\frac{\log d}{n}} \right\},$$

(8.5) and assumptions (D1'), (D2') yield

$$\|\beta^T - \beta\|_{\Sigma_x}^2 \lesssim \sigma^2 \left(\frac{s^* \log d}{n} + \frac{(s^* \log d)^2 \log(1/\delta) \log^3 n}{n^2 \varepsilon^2} \right).$$

It remains to show that the events $\mathcal{E}_1, \mathcal{E}_2, \mathcal{E}_3$ occur simultaneously with high probability. We have $\mathbb{P}(\mathcal{E}_1^c) \leq c_1 \exp(-c_2 \log n)$ because $y_1, y_2, \dots, y_n \stackrel{\text{i.i.d.}}{\sim} N(0, \sigma^2)$ and $R \asymp \sigma \sqrt{\log n}$.

For $\mathcal{E}_2$, we invoke Lemma A.1 in the Supplementary Material [10]. For each iterate t , the individual coordinates of $\tilde{\mathbf{w}}^t, \mathbf{w}_i^t$ are sampled i.i.d. from the Laplace distribution with scale $\eta^0 \cdot \frac{2B\sqrt{3s\log(T/\delta)}}{n\varepsilon/T}$, where the noise scale $B \lesssim R/\sqrt{s}$ and $T \asymp \log n$ by our choice. If $n \geq K \cdot (R(s^*)^{3/2} \log d \sqrt{\log(1/\delta) \log n / \varepsilon})$ for a sufficiently large constant K , Lemma A.1 and the union bound imply that, with probability at least $1 - c_1 \exp(-c_2 \log(d/(s^* \log n)))$, $\max_t W_t$ is bounded by $K \frac{R^2 (s^*)^3 \log^2 d \log(1/\delta) \log^2 n}{n^2 \varepsilon^2}$ for some appropriate constant K . For $\mathcal{E}_3$, under assumptions (D1') and (D2'), it is a standard probabilistic result (see, e.g., [46], pages 210–211) that $\mathbb{P}(\mathcal{E}_3^c) \leq 2e^{-2 \log d}$. $\square$

Acknowledgments. We would like to thank the Associate Editor and referees for their helpful suggestions and comments that lead to a great improvement of the paper. We also thank Chi-Yun Wu for advice on Section 6.

Funding. The research of T. Cai was supported in part by NSF Grants DMS-1712735 and DMS-2015259 and NIH Grants R01-GM129781 and R01-GM123056. The research of L. Zhang was supported in part by NSF Grant NSF DMS-2015378.

SUPPLEMENTARY MATERIAL

Supplement to “The cost of privacy: Optimal rates of convergence for parameter estimation with differential privacy” (DOI: [10.1214/21-AOS2058SUPP](https://doi.org/10.1214/21-AOS2058SUPP); .pdf). In the Supplementary Material, we prove all the main theorems and the technical lemmas.

REFERENCES

- [1] ABADI, M., CHU, A., GOODFELLOW, I., MCMAHAN, H. B., MIRONOV, I., TALWAR, K. and ZHANG, L. (2016). Deep learning with differential privacy. In *ACM CCS 2016* 308–318. ACM, New York.
- [2] ABOWD, J. M. (2016). The challenge of scientific reproducibility and privacy protection for statistical agencies. In *Census Scientific Advisory Committee*.
- [3] AGARWAL, A., NEGAHBAN, S. and WAINWRIGHT, M. J. (2010). Fast global convergence of gradient methods for high-dimensional statistical recovery. In *Advances in Neural Information Processing Systems* 37–45.
- [4] BASSILY, R., FELDMAN, V., TALWAR, K. and THAKURTA, A. G. (2019). Private stochastic convex optimization with optimal rates. In *Advances in Neural Information Processing Systems* 11282–11291.
- [5] BASSILY, R., SMITH, A. and THAKURTA, A. (2014). Private empirical risk minimization: Efficient algorithms and tight error bounds. In *55th Annual IEEE Symposium on Foundations of Computer Science—FOCS 2014* 464–473. IEEE Computer Soc., Los Alamitos, CA. [MR3344896 https://doi.org/10.1109/FOCS.2014.56](https://doi.org/10.1109/FOCS.2014.56)

- [6] BISWAS, S., DONG, Y., KAMATH, G. and COINPRESS, J. U. (2020). Practical private mean and covariance estimation. Preprint. Available at [arXiv:2006.06618](https://arxiv.org/abs/2006.06618).
- [7] BLUMENSATH, T. and DAVIES, M. E. (2009). Iterative hard thresholding for compressed sensing. *Appl. Comput. Harmon. Anal.* **27** 265–274. [MR2559726 https://doi.org/10.1016/j.acha.2009.04.002](https://doi.org/10.1016/j.acha.2009.04.002)
- [8] BU, Z., DONG, J. LONG, Q. and SU, W. J. (2019). Deep learning with gaussian differential privacy. Preprint. Available at [arXiv:1911.11607](https://arxiv.org/abs/1911.11607).
- [9] BUN, M., ULLMAN, J. and VADHAN, S. (2014). Fingerprinting codes and the price of approximate differential privacy. In *STOC'14—Proceedings of the 2014 ACM Symposium on Theory of Computing* 1–10. ACM, New York. [MR3238925 https://doi.org/10.1145/2591796.2591877](https://doi.org/10.1145/2591796.2591877)
- [10] CAI, T. T., WANG, Y. and ZHANG, L. (2021). Supplement to “The cost of privacy: Optimal rates of convergence for parameter estimation with differential privacy.” <https://doi.org/10.1214/21-AOS2058SUPP>
- [11] DING, B., KULKARNI, J. and YEKHANIN, S. (2017). Collecting telemetry data privately. In *NeurIPS 2017* 3571–3580.
- [12] DONG, J. ROTH, A. and SU, W. J. (2019). Gaussian differential privacy. Preprint. Available at [arXiv:1905.02383](https://arxiv.org/abs/1905.02383).
- [13] DONOHO, D. L. (1994). Statistical estimation and optimal recovery. *Ann. Statist.* **22** 238–270. [MR1272082 https://doi.org/10.1214/aos/1176325367](https://doi.org/10.1214/aos/1176325367)
- [14] DUCHI, J. C., JORDAN, M. I. and WAINWRIGHT, M. J. (2018). Minimax optimal procedures for locally private estimation. *J. Amer. Statist. Assoc.* **113** 182–201. [MR3803452 https://doi.org/10.1080/01621459.2017.1389735](https://doi.org/10.1080/01621459.2017.1389735)
- [15] DWORK, C. and FELDMAN, V. (2018). Privacy-preserving prediction. Preprint. Available at [arXiv:1803.10266](https://arxiv.org/abs/1803.10266).
- [16] DWORK, C., MCSHERRY, F., NISSIM, K. and SMITH, A. (2006). Calibrating noise to sensitivity in private data analysis. In *Theory of Cryptography. Lecture Notes in Computer Science* **3876** 265–284. Springer, Berlin. [MR2241676 https://doi.org/10.1007/11681878_14](https://doi.org/10.1007/11681878_14)
- [17] DWORK, C. and ROTH, A. (2013). The algorithmic foundations of differential privacy. *Found. Trends Theor. Comput. Sci.* **9** 211–487. [MR3254020 https://doi.org/10.1561/04000000042](https://doi.org/10.1561/04000000042)
- [18] DWORK, C. and ROTHBLUM, G. N. (2016). Concentrated differential privacy. Preprint. Available at [arXiv:1603.01887](https://arxiv.org/abs/1603.01887).
- [19] DWORK, C., SMITH, A., STEINKE, T. and ULLMAN, J. (2017). Exposed! A survey of attacks on private data. *Annu. Rev. Stat. Appl.* **4** 61–84.
- [20] DWORK, C., SMITH, A., STEINKE, T., ULLMAN, J. and VADHAN, S. (2015). Robust traceability from trace amounts. In *2015 IEEE 56th Annual Symposium on Foundations of Computer Science—FOCS 2015* 650–669. IEEE Computer Soc., Los Alamitos, CA. [MR3473333 https://doi.org/10.1109/FOCS.2015.46](https://doi.org/10.1109/FOCS.2015.46)
- [21] DWORK, C., SU, W. J. and ZHANG, L. (2018). Differentially private false discovery rate control. Preprint. Available at [arXiv:1807.04209](https://arxiv.org/abs/1807.04209).
- [22] DWORK, C., TALWAR, K., THAKURTA, A. and ZHANG, L. (2014). Analyze Gauss: Optimal bounds for privacy-preserving principal component analysis. In *STOC'14—Proceedings of the 2014 ACM Symposium on Theory of Computing* 11–20. ACM, New York. [MR3238926 https://doi.org/10.1145/2591796.2591883](https://doi.org/10.1145/2591796.2591883)
- [23] ERLINGSSON, Ú., PIHUR, V. and RAPPOR, A. K. (2014). Randomized aggregatable privacy-preserving ordinal response. In *ACM CCS 2014* 1054–1067. ACM, New York.
- [24] HOMER, N., SZELINGER, S., REDMAN, M., DUGGAN, D., TEMBE, W., MUEHLING, J., PEARSON, J. V., STEPHAN, D. A., NELSON, S. F. et al. (2008). Resolving individuals contributing trace amounts of DNA to highly complex mixtures using high-density SNP genotyping microarrays. *PLoS Genet.* **4** e1000167. <https://doi.org/10.1371/journal.pgen.1000167>
- [25] JAIN, P., TEWARI, A. and KAR, P. (2014). On iterative hard thresholding methods for high-dimensional m -estimation. In *NeurIPS 2014* 685–693.
- [26] JOHNSTONE, I. M. (1994). On minimax estimation of a sparse normal mean vector. *Ann. Statist.* **22** 271–289. [MR1272083 https://doi.org/10.1214/aos/1176325368](https://doi.org/10.1214/aos/1176325368)
- [27] KAMATH, G., LI, J., SINGHAL, V. and ULLMAN, J. (2018). Privately learning high-dimensional distributions. Preprint. Available at [arXiv:1805.00216](https://arxiv.org/abs/1805.00216).
- [28] KARWA, V. and VADHAN, S. (2017). Finite sample differentially private confidence intervals. Preprint. Available at [arXiv:1711.03908](https://arxiv.org/abs/1711.03908).
- [29] KASIVISWANATHAN, S. P., LEE, H. K., NISSIM, K., RASKHODNIKOVA, S. and SMITH, A. (2011). What can we learn privately? *SIAM J. Comput.* **40** 793–826. [MR2823508 https://doi.org/10.1137/090756090](https://doi.org/10.1137/090756090)
- [30] KIFER, D., SMITH, A. and THAKURTA, A. G. (2012). Private convex empirical risk minimization and high-dimensional regression. In *COLT 2012* 25.1–25.40.

- [31] LEI, J. (2011). Differentially private m -estimators. In *NeurIPS* 2011 361–369.
- [32] LOH, P.-L. and WAINWRIGHT, M. J. (2015). Regularized M -estimators with nonconvexity: Statistical and algorithmic theory for local optima. *J. Mach. Learn. Res.* **16** 559–616. [MR3335800](#)
- [33] LOWTHER, C., MERICO, D., COSTAIN, G., WASERMAN, J., BOYD, K., NOOR, A., SPEEVAK, M., STAVROPOULOS, D. J., WEI, J. et al. (2017). Impact of IQ on the diagnostic yield of chromosomal microarray in a community sample of adults with schizophrenia. *Gen. Med.* **9** 105.
- [34] MCSHERRY, F. and TALWAR, K. (2007). Mechanism design via differential privacy. In *FOCS* 2007 **7** 94–103.
- [35] MIRONOV, I. (2017). Renyi differential privacy. In *CSF* 2017 263–275. IEEE, New York.
- [36] NEGAHBAN, S., YU, B., WAINWRIGHT, M. J. and RAVIKUMAR, P. K. (2009). A unified framework for high-dimensional analysis of m -estimators with decomposable regularizers. In *Advances in Neural Information Processing Systems* 1348–1356.
- [37] PACE, R. K. and BARRY, R. (1997). Sparse spatial autoregressions. *Statist. Probab. Lett.* **33** 291–297.
- [38] RAAB, M. and STEGER, A. (1998). “Balls into bins”—A simple and tight analysis. In *Randomization and Approximation Techniques in Computer Science (Barcelona, 1998). Lecture Notes in Computer Science* **1518** 159–170. Springer, Berlin. [MR1729169](#) https://doi.org/10.1007/3-540-49543-6_13
- [39] ROHDE, A. and STEINBERGER, L. (2020). Geometrizing rates of convergence under local differential privacy constraints. *Ann. Statist.* **48** 2646–2670. [MR4152116](#) <https://doi.org/10.1214/19-AOS1901>
- [40] SHEFFET, O. (2017). Differentially private ordinary least squares. In *Proceedings of the 34th International Conference on Machine Learning* **70** 3105–3114. [JMLR.org](#).
- [41] SMITH, A. (2011). Privacy-preserving statistical estimation with optimal convergence rates [extended abstract]. In *STOC’11—Proceedings of the 43rd ACM Symposium on Theory of Computing* 813–821. ACM, New York. [MR2932032](#) <https://doi.org/10.1145/1993636.1993743>
- [42] STEINKE, T. and ULLMAN, J. (2017). Between pure and approximate differential privacy. *J. Priv. Confid.* **7**.
- [43] STEINKE, T. and ULLMAN, J. (2017). Tight lower bounds for differentially private selection. In *58th Annual IEEE Symposium on Foundations of Computer Science—FOCS* 2017 552–563. IEEE Computer Soc., Los Alamitos, CA. [MR3734260](#) <https://doi.org/10.1109/FOCS.2017.57>
- [44] TALWAR, K., THAKURTA, A. G. and ZHANG, L. (2015). Nearly optimal private lasso. In *NeurIPS* 2015 3025–3033.
- [45] APPLE DIFFERENTIAL PRIVACY TEAM (2017). Privacy at scale.
- [46] WAINWRIGHT, M. J. (2019). *High-Dimensional Statistics: A Non-asymptotic Viewpoint. Cambridge Series in Statistical and Probabilistic Mathematics* **48**. Cambridge Univ. Press, Cambridge. [MR3967104](#) <https://doi.org/10.1017/9781108627771>
- [47] WANG, Y.-X., FIENBERG, S. and SMOLA, A. (2015). Privacy for free: Posterior sampling and stochastic gradient Monte Carlo. In *ICML* 2493–2502.
- [48] WASSERMAN, L. and ZHOU, S. (2010). A statistical framework for differential privacy. *J. Amer. Statist. Assoc.* **105** 375–389. [MR2656057](#) <https://doi.org/10.1198/jasa.2009.tm08651>