
Learner-Private Convex Optimization

Jiaming Xu¹ Kuang Xu² Dana Yang¹

Abstract

Convex optimization with feedback is a framework where a learner relies on iterative queries and feedback to arrive at the minimizer of a convex function. The paradigm has gained significant popularity recently thanks to its scalability in large-scale optimization and machine learning. The repeated interactions, however, expose the learner to privacy risks from eavesdropping adversaries that observe the submitted queries. In this paper, we study how to optimally obfuscate the learner’s queries in convex optimization with first-order feedback, so that their learned optimal value is provably difficult to estimate for the eavesdropping adversary. We consider two formulations of learner privacy: a Bayesian formulation in which the convex function is drawn randomly, and a minimax formulation in which the function is fixed and the adversary’s probability of error is measured with respect to a minimax criterion.

We show that, if the learner wants to ensure the probability of the adversary estimating accurately be kept below $1/L$, then the overhead in query complexity is additive in L in the minimax formulation, but multiplicative in L in the Bayesian formulation. Compared to existing learner-private sequential learning models with binary feedback, our results apply to the significantly richer family of general convex functions with full-gradient feedback. Our proofs are largely enabled by tools from the theory of Dirichlet processes, as well as more sophisticated lines of analysis aimed at measuring the amount of information leakage under a full-gradient oracle.

¹The Fuqua School of Business, Duke University, Durham NC, USA ²Stanford Graduate School of Business, Stanford University, Stanford CA, USA. Correspondence to: Jiaming Xu <jiaming.xu868@duke.edu>, Kuang Xu <kuangxu@stanford.edu>, Dana Yang <xiaoqian.yang@duke.edu>.

1. Introduction

Convex optimization with feedback is a paradigm in which an learner repeatedly queries an external data source in order to identify the optimal solution of a convex function. The interactive nature of the framework is a double-edged sword. On the one hand, the iterative optimization methods offers inherent scalability since the learner is not required to possess the entire function from the start. As such, it has found applications in large-scale distributed machine learning systems, such as Federated Learning (McMahan et al., 2017; McMahan & Ramage, 2017), where a learner interacts with millions of individual users (data providers) in order to perform training. On the other hand, the repeated interactions with external entities exposes the learner to potential adversaries who may steal the learned model by eavesdropping on the queries exchanged during the training process, a woe especially poignant when the system involves a large number of data providers, many of which could be an eavesdropper in disguise ((Juuti et al., 2019), (Kairouz et al., 2019, Section 4.3)).

To address challenges in protecting the learner’s privacy, a recent line of research proposed the framework of Private Sequential Learning, aimed at quantifying the extra query complexities the learner has to suffer in order to ensure the submitted queries provably conceal the learned value (Tsitsiklis et al., 2018; Xu, 2018; Xu et al., 2019). The model is centered around a binary search problem where a learner tries to estimate an unknown value $X^* \in [0, 1]$ by sequentially submitting queries and receiving binary responses, indicating the position of X^* relative to the queries. Meanwhile, an adversary observes all of the learner’s queries but not responses, and tries to use this information to estimate X^* . The learner’s goal is to design a querying strategy with a minimal number of queries so that she can accurately estimate X^* while ensuring that the eavesdropping adversary cannot reliably estimate X^* . Progress has been made towards understanding the optimal querying strategies in this problem, and upper and lower bounds on the query complexity have been developed that differ by additive constants in the case where the learner’s queries are noiseless (Tsitsiklis et al., 2018; Xu et al., 2019), and are order-wise optimal in the case of noisy queries (Xu et al., 2019).

While the original binary search formulation provides valu-

able insights, its simplifying assumption that the learner only has access to binary feedback is a severe restriction when it comes to modeling convex optimization. Indeed, most real-world applications provide the learner access to significantly richer feedback such as a full gradient (e.g., model training in machine learning). We elaborate further on the potential applications of our model in Section 4.

The main purpose of the present paper is to take a step towards closing this gap by studying learner-private optimization with general convex functions and a full-gradient oracle. In a nutshell, our results demonstrate that the most prominent features of the query complexity in the binary search model extend gracefully to the general convex optimization setting. However, to establish that this is the case is far from trivial. A major difficulty stems from the significantly enriched functional class: unlike in a binary search problem where the ground truth is fully described by a scalar (location of X^*), we will see that the private query complexity crucially depends on the shapes of the convex functions in a family, and not just the locations of their minimizers.

This added richness necessitates the development of both new problem formulations and analytical techniques. We propose in this paper two new learner-privacy frameworks: a new minimax formulation, as well as a Bayesian formulation that generalizes earlier Bayesian private sequential learning to a full-gradient oracle. A number of new techniques are developed to analyze query complexity under these formulations: we introduce tools from the theory of Dirichlet processes to construct priors that convey the richness of the model. Tools from nonparametric Bayes theory are deployed for the analysis under such prior distributions. In addition to an enriched functional class, another fundamental challenge lies in the richness of the feedback. Unlike the binary search model, the responses aligns with the location of the query and the shape of the unobserved convex function to a great extent. In the face of a more powerful learner equipped with a full-gradient oracle, we rely on a more sophisticated line of analysis to gauge the amount of information the responses reveal. We will discuss in more detail these ramifications in Section 4.

Relation to private information retrieval (PIR) and private function retrieval (PFR) Our model formulation bears some similarities with the PIR (Abadi et al., 1989; Chor et al., 1995; Gasarch, 2004) and PFR (Mirmohseni & Maddah-Ali, 2018) framework. However, there are major distinctions which result in completely different dynamics between the learner and the adversary. In PIR, the database is assumed to contain a vector $(x_i)_{i \leq N}$. The learner’s goal is to learn the evaluation x_i at some index i by querying the database, while preventing the database (adversary) from learning the value of i . The PFR problem is formulated sim-

ilarly, except that the database is indexed by functions. Note that in PIR/PFC, the private index is assumed to be known to the learner a priori. In contrast, in our framework, the private information X^* is something the learner herself is in the process of discovering. As a result, our problem is posed as a sequential learning problem. It has natural applications in model stealing attack prevention, where eavesdropping adversaries attempt to steal the model parameters by participating in the model training process. The fundamental difference between the two settings also leads to completely different techniques for analysis. For us, privacy is ensured by utilizing the adversary’s lack of knowledge on the responses, which is not the case in PIR/PFC.

Relation to data-owner privacy models Similar to Private Sequential Learning, the private convex optimization problem we consider diverges significantly from the existing literature on differentially private iterative learning (Song et al., 2013; Abadi et al., 2016; Agarwal et al., 2018; Jain et al., 2012; Melis et al., 2019), a key difference being that the latter focuses on protecting data owners’ privacy rather than learner’s privacy. To protect data owners’ privacy, the notion of differential privacy (Dwork, 2008) is often adopted and privacy is often achieved by injecting calibrated noise at each iteration of the learning algorithms. In contrast, our work focuses on preventing the adversary inferring the learned model, which is conceptually closer to recent studies of information-theoretically sound obfuscation in sequential decision-making problems (Fanti et al., 2015; Luo et al., 2016; Tsitsiklis & Xu, 2018; Erturk & Xu, 2019; Tang et al., 2020b). See (Xu et al., 2019) for a comprehensive discussion on the distinction between data-owner privacy models and this line of work.

2. The Model: Learner-Private Convex Optimization

We now introduce our model, dubbed Learner-Private Convex Optimization. The emphasis on the learner’s privacy here is to distinguish our model from other forms of private sequential learning, especially those that focus on protecting the privacy of data owners (See proceeding discussion in the Introduction).

Learner Let $\mathcal{F}$ be a family of $\mathbb{R}$ -valued convex functions with domain $[0, 1]$, such that all elements in $\mathcal{F}$ admit a unique minimizer. Suppose there is an unknown *truth* $f^* \in \mathcal{F}$ with the minimizer $X^* := \arg \min_x f^*(x)$. Fix $n \in \mathbb{N}$. Our decision maker is a *learner* who wants to identify X^* by sequentially submitting a total of n queries in $[0, 1]$ to an oracle. For the i th query, q_i , the oracle returns a response r_i that is equal to the gradient of f^* at q_i :

$$r_i = (f^*)'(q_i). \quad (1)$$

If f^* is not differentiable at q_i , then r_i is an arbitrary subgradient of f^* at q_i . We assume that the learner is allowed to introduce outside randomness, in the form of a random seed Y that takes value in a finite discrete alphabet. Formally, we denote by ϕ the *learner's strategy*, which consists of a sequence of mappings $\phi_0, \phi_1, \dots, \phi_{n-1}$ such that the i th query is generated as a function of all previous responses and the random seed:

$$q_i = \phi_{i-1}(r_1, \dots, r_{i-1}, Y). \quad (2)$$

Once the querying process is terminated, the learner constructs an estimator of the optimizer X^* , $\tilde{X}$, based on the n responses. We say that the learner strategy ϕ is ϵ -accurate, if

$$\mathbb{P}_f \left\{ \left| \tilde{X} - x \right| \leq \epsilon/2 \right\} = 1, \quad \forall f \in \mathcal{F}, \quad (3)$$

where x is the minimize of f and the $\mathbb{P}_f$ indicates the induced probability law when the truth f^* is equal to f , and the probability is measured with respect to the randomness in the random seed, Y .

Adversary Meanwhile, an adversary is trying to learn X^* by eavesdropping on the learner's queries: we assume that the adversary observes all n queries submitted by the learner, but not their responses. Denote by $\tilde{X}$ the adversary's estimator, which is a (possibly random) function of $(q_i)_{i=1, \dots, n}$. Wary of such an adversary, the high-level objective of the learner are to (1) generate a query sequence that is largely "uninformative" towards X^* , and (2) at the same minimizing the number of queries needed, n .

We next formalize in what sense a learner's strategy can be private. Generally speaking, a learner strategy is private if we can ensure that the adversary's estimator $\tilde{X}$ is *not accurate*. Importantly, different definitions of the adversary's accuracy will lead to drastically different definitions of privacy, and consequently, distinct algorithms, guarantees and domains of applications. In this paper, we will analyze two privacy metrics, Bayesian and minimax, that parallel the two paradigms in the statistics literature. The Bayesian formulation extends the Bayesian private learning model in (Tsitsiklis et al., 2018), while the minimax formulation is new.

Minimax The truth f^* is a deterministic but unknown function in $\mathcal{F}$. We say that a learner strategy ϕ is (δ, L) -private if

$$\sup_{\tilde{X}} \inf_{f \in \mathcal{F}} \mathbb{P}_f \left\{ \left| \tilde{X} - x \right| \leq \delta/2 \right\} \leq 1/L, \quad (4)$$

where the probability is measured with respect to the internal randomness employed by the learner's querying strategy and that used in the adversary's estimator. In other words, the learner strategy is considered private if the adversary's minimax risk is large.

Bayesian The truth f^* is drawn from a prior distribution π , a probability distribution over $\mathcal{F}$. We say that a learner strategy ϕ is (δ, L) -private if

$$\sup_{\tilde{X}} \mathbb{P} \left\{ \left| \tilde{X} - X^* \right| \leq \delta/2 \right\} \leq 1/L, \quad (5)$$

where the probability is measured with respect to all randomness in the system, including the prior π and any internal randomness employed by the learner's querying strategy and the adversary's estimator.

Private query complexity Finally, we have come to the main metric of interest. In both the minimax and the Bayesian formulations, we define the optimal query complexity, $N(\epsilon, \delta, L)$, as the least number of queries necessary for there to exist an ϵ -accurate learner strategy that is also (δ, L) -private:

$$N(\epsilon, \delta, L) = \min \{ n : \exists \phi \text{ with at most } n \text{ queries, that is } \epsilon\text{-accurate and } (\delta, L)\text{-private} \}.$$

3. Main Results

3.1. Minimax formulation

We will assume that the function class $\mathcal{F}$ satisfies the following assumption:

Assumption 1 (Complexity of $\mathcal{F}$). *Fix $f \in \mathcal{F}$ and interval $I \subset [0, 1]$ that contains the minimizer of f . Then, for every $x \in I$, there exists $g \in \mathcal{F}$ such that g is minimized at x , and the gradient of f and g coincide outside of I .*

Assumption 1 is needed to rule out trivial cases where a learner may exactly pinpoint the location of the minimizer solely by looking at far-away gradients. We show in Section 5 that this richness assumption on $\mathcal{F}$ is in some sense necessary. Examples of function classes that satisfy Assumption 1 include the set of all convex functions on $[0, 1]$, and the set of all piecewise-linear convex functions on $[0, 1]$. The next theorem is our main result for the minimax formulation:

Theorem 1 (Minimax Query Complexity). *Assume that $\mathcal{F}$ satisfies Assumption 1. If $2\epsilon \leq \delta \leq 1/L$, then¹*

$$2L + \log \frac{\delta}{\epsilon} - 2 \leq N(\epsilon, \delta, L) \leq \begin{cases} 2L + \log \frac{\delta}{\epsilon} & \text{if } L \geq \log \frac{1}{\delta} \\ L + \log \frac{1}{\epsilon} & \text{o.w.} \end{cases}.$$

Note that if there were no privacy consideration, the minimax optimal query complexity would be $\log(1/\epsilon)$. Thus under the minimax formulation, a higher level of privacy L leads to an *additive* overhead in the optimal query complexity, that is at most about $2L$.

¹Here and subsequently log refers to logarithm with base 2.

Remark 1 (Multidimensional Extensions). *By considering a separable class of functions, and using the ℓ_∞ norm to measure the error of the learner and the adversary’s estimators, Theorem 1 can be extended to d dimensions. The upper and lower bounds of the query complexity take the same form, with L replaced with $L^{1/d}$. See the supplementary material for the precise statement and proof.*

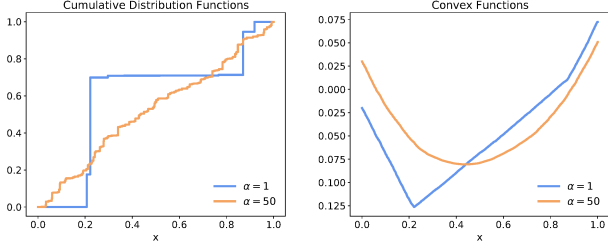


Figure 1. The left figure exemplifies realizations of F following the Dirichlet Process with base function $\lambda_{[0,1]}$ and different concentration parameters α . The right figure shows the corresponding convex functions f^* , with $\gamma_+ = 0.5$ and $\gamma_- = -0.5$.

3.2. Bayesian formulation

In the Bayesian formulation, we seek a function class and prior distribution that are sufficiently rich to capture real-world data, while at the same time amenable to analysis. A good candidate in this respect is the so-called Dirichlet process, a family of measures over non-decreasing functions, which we will use to model the gradient function of f^* . Dirichlet processes are fundamental objects in non-parametric Bayes theory and widely used in Bayesian isotonic regression for modeling monotone functions (Lavine & Mockus, 1995; Bornkamp & Ickstadt, 2009; Neelon & Dunson, 2004). We begin by defining a Dirichlet process:

Definition 1 (Dirichlet Process). *Given a base probability measure μ_0 on $\mathcal{X}$ and a concentration parameter $\alpha > 0$. A random probability measure μ over $\mathcal{X}$ is said to follow the Dirichlet process $DP(\mu_0, \alpha)$, if for any finite partition of $\mathcal{X} = \cup_{i=1}^n \mathcal{X}_i$,*

$$(\mu(\mathcal{X}_1), \dots, \mu(\mathcal{X}_n)) \sim \text{Dir}((\alpha\mu_0(\mathcal{X}_1), \dots, \alpha\mu_0(\mathcal{X}_n))),$$

where $\text{Dir}(c)$ denotes the Dirichlet distribution over the n -dimensional simplex Δ^{n-1} with density

$$g_{\text{Dir}(c)}(x_1, \dots, x_n) \propto \prod_{i=1}^n x_i^{c_i-1}, \quad x \in \Delta^{n-1}. \quad (6)$$

We now construct the prior distribution of f^* using a Dirichlet process. The prior is parameterized by two quantities:

1. a concentration parameter $\alpha > 0$, which controls the dispersion of the distribution of the minimizer;

2. a probability distribution η over $[0, 1]$, which captures the range of gradients of f^* . We assume that η admits a density that is bounded from above and away from 0 (e.g., $\text{Unif}[0, 1]$).

Definition 2 (Bayesian Prior using Dirichlet Process). *Fix α and η . Denote by $\lambda_{[0,1]}$ the Lebesgue measure restricted to $[0, 1]$. Then, the prior π corresponds to the following procedure for generating f^* :*²

1. Sample γ_+ from η . Set $\gamma_- = -\gamma_+$.
2. Sample μ from the Dirichlet process with concentration parameter α and base distribution $\lambda_{[0,1]}$. Let F be the cumulative distribution function of μ .
3. Set $f^*(x) = \gamma_-x + \int_0^x (\gamma_+ - \gamma_-)F(t)dt$, for $x \in [0, 1]$.

Note that $(f^*(x))' = \gamma_+ (2F(x) - 1)$ and thus the minimizer X^* of f^* corresponds to the median of F , or more precisely the smallest x for which $F(x) \geq 1/2$. By construction, F is a monotone simple function that consists of countably many points of discontinuity that are dense on $[0, 1]$. Its level of discreteness is modeled through the concentration parameter α . For a small α , the increase of F from 0 to 1 is mostly from a few abrupt jumps, and the convex function f^* resembles a piece-wise linear function with finitely many pieces; as α grows, the increase of F becomes more gradual, and f^* starts to concentrate around a smooth quadratic function. See Figure 1 for some realizations of the distribution function F and the corresponding convex function f^* for different value of α .³

The following theorem is our main result for the Bayesian formulation.

Theorem 2 (Bayesian Query Complexity). *Fix $\alpha > 0$. Suppose that $2\epsilon \leq \delta < \frac{1}{2LH_\alpha}$, with $H_\alpha = (3 + 2e^{-1})\alpha + 14$. Then*

$$c_1 L \log \frac{\delta}{\epsilon} \leq N(\epsilon, \delta, L) \leq L \log \frac{\delta}{\epsilon} + c_2 L + \log \frac{1}{\delta L},$$

where c_1, c_2 are positive constants that only depend on α such that $c_1 \rightarrow 1$ as $\alpha \rightarrow 0$.

The above theorem shows that, in the Bayesian formulation, the query complexity overhead due to privacy constraints scales *multiplicatively* with respect to the privacy level L .

²Note that in this definition we have restricted the gradients to lie in $[-1, 1]$ and the function f^* to have zero intercept. Both restrictions are without loss of generality, since any constant offset will not change the location of a minimizer and similarly our results will carry through if one wishes to incorporate a different gradient scaling factor.

³To plot the convex functions together, we shift them by some constants on the y -axis. This shift is irrelevant to the optimization task since the response only contains gradient information.

Note that this is substantially higher than the minimax setting where such overhead is only additive in L . When $\alpha \rightarrow 0$, F converges to a step function and our query complexity bounds recover the existing ones in the binary search problem (Xu, 2018), showing that $N(\epsilon, \delta, L) \sim L \log \frac{1}{\epsilon}$ as $\epsilon \rightarrow 0$ for fixed δ, L .

4. Discussion

In this section, we examine some real-world applications of our privacy model and discuss some of the most salient features of our main results and modeling assumptions.

Motivating examples A learner naturally suffers from privacy breaches if the learning process involves interactions with third-party users. An example would be the aforementioned Federated Learning framework. A typical Federated Learning model training process can be posed as iterative optimization of some unknown function. Iterations of model updates are generated from the feedback from a large number of users (see e.g. the *Federated Averaging* algorithm (McMahan et al., 2017)). Since the model updates (queries) are broadcasted to the participating users, the learner is exposed to eavesdropping attacks. Due to the high cost of large-scale model training, it is of great importance to protect the learner from such privacy breaches, and do so at a minimal cost (Kairouz et al., 2019).

Another potential application is pricing optimization, where the goal is to learn the optimal release price of a product by conducting market experiments at test price points (queries). See (Xu et al., 2019; Tsitsiklis et al., 2018) for more detailed discussions on the Federated learning and pricing optimization examples.

Given the close connection between convex and monotone functions, our work can also be applied to learning monotone functions, for example to clinical dose-response studies (Ramgopal et al., 1993; Bornkamp & Ickstadt, 2009). In dose-response analysis, the potency curve $\mu(x)$ is a monotone function that models the treatment effectiveness as a function of the dosage. An important problem is to estimate the minimum effective dose (MED)

$$\text{MED} = \min_x \{x : \mu(x) > \mu(0) + \Delta\}$$

for some threshold Δ . Note that the MED is the minimizer X^* of some unknown convex function f^* (e.g. $f^*(x) = \int_0^t \mu(t)dt - [\mu(0) + \Delta]x$). We also remark that the Dirichlet process is widely used in isotonic regression for modeling monotone functions (Lavine & Mockus, 1995; Bornkamp & Ickstadt, 2009), as we do when modeling the gradient of the convex function.

Applying the Bayesian and Minimax privacy criteria

Our results show that the two privacy criteria lead to distinct query complexity scalings, so it would be instructive to

understand in what application domain each metric is most applicable. We expect the Bayesian formulation to be most relevant in data-driven machine learning and optimization with feedback such as in Federated Learning and pricing optimization; the aforementioned dose-response analysis is also a natural application of the Bayesian formulation due to the close connection between potency curves and convex functions. The minimax formulation is a new metric proposed in this paper. One interesting application is in law and criminal justice, where a prosecutor should have to prove that the accuracy of any conclusion drawn from evidence holds up *regardless* of the value of a certain hidden parameter (Young et al., 2001). Other potential applications include autonomous driving, where the performance guarantee of an estimator needs to be valid in the worst case, for the sake of public safety.

Comparisons with private sequential learning As mentioned in the Introduction, our convex optimization framework generalizes the Private Sequential Learning (PSL) model. Recall that in the PSL framework, the responses are binary and only indicate whether the minimizer is to the left or right of a given query; this is equivalent, in our setting, to returning only the sign of the gradient. There are several major differences that distinguish the convex optimization framework from the PSL model. First and foremost, the learner now has access to the entire gradient instead of only its sign. A most direct implication of this enriched information structure is that, when analyzing the amount of information leakage of a learner strategy, we will have to keep track of the distributions over target functions, as opposed to only the minimizers, as was the case in PSL. Moreover, when the learner has access to full gradients, it is in principle possible for the learner to gather information about the minimizer’s precise location even from queries that are submitted far away from the minimizer, which was not possible within bisection search. For instance, if the underlying target function is known to be quadratic, then two queries placed anywhere are sufficient to uncover the minimizer. To address these complexities, our goal is to precisely measure the amount of information about the minimizer that the learner and adversary may obtain from a given sequence of queries. We will do so both by developing more sophisticated information theoretic arguments, and by exploiting structural properties of the Dirichlet process.

Open questions Our results leave open a number of questions. For the Bayesian query complexity in one dimension, there remains a gap between the leading constants in the upper and lower bounds, in the regime where α is bounded away from zero. Generalizing the main theorems to a multi-dimensional setting, where $x \in \mathbb{R}^d$, $d \geq 2$, is also interesting and practically relevant. We take a first step in this direction by extending our results to multi-dimensional

separable functions (see supplementary material), while the general case with non-separable objective functions remains open and appears to be challenging. Our problem formulation only considers first-order feedback. An interesting direction is to consider convex optimization with more general types of feedback, e.g., bandit feedback (Agarwal et al., 2013).

A different notion of minimax privacy in (Tang et al., 2020a) A recent work (Tang et al., 2020a) also aims to extend the private sequential learning model of (Tsitsiklis et al., 2018) to convex optimization. They use a different notion of minimax privacy criteria that bear some superficial similarities to ours. However, the definition of privacy in (Tang et al., 2020a) contains crucial errors that render it vacuous, in the sense that there cannot exist any private learner strategy satisfying that definition. To be precise, here is Definition 2 of (Tang et al., 2020a): fix $\epsilon, \delta \in (0, 1)$. A learner strategy is said to be (ϵ, δ) -private if for any adversary estimator $\tilde{X}$ and any truth $f \in \mathcal{F}$,

$$\mathbb{P}_f(\text{err}(\tilde{X}, f) \leq \epsilon) \leq \delta, \quad (7)$$

where $\text{err}(\cdot, \cdot)$ is a certain error function which measures the discrepancy between the adversary estimator and the true minimizer. For instance, in our example $\text{err}(\tilde{X}, f) = |\tilde{X} - \arg \min f(x)|$.

The problem with this privacy definition is that it can never be satisfied by any learner strategy. Indeed, for any $f \in \mathcal{F}$ with minimizer x^* , there always exists an adversary estimator that trivially yields zero estimation error with probability one: simply set $\tilde{X} = x^*$, without even taking into account the queries. Under this trivial estimator, we automatically have $\mathbb{P}_f(\text{err}(\tilde{X}, f) = 0) = 1$, so (7) cannot possibly hold uniformly across all adversary estimators and all f . Unfortunately, this would further suggest that the analysis and conclusions in (Tang et al., 2020a) contain errors as well.

5. Proof of Main Results

We present in this section the proof sketch of our main results and defer the full proof to the supplementary material due to the space constraint.

5.1. Proof sketch under the Bayesian setting

Proof of the upper bound in Theorem 2. The upper bound is established by analyzing a constructive algorithm. The key challenge is that the prior distribution on X^* is always non-uniform under the Dirichlet process model. In particular, we can no longer simply apply the replicated search strategy from (Xu et al., 2019), since the non-uniform distribution of X^* provides the adversary with additional prior information.

To address this difficulty, our key algorithmic idea is to find L intervals that occupy the same prior mass, while at the same time are at least δ -separated from each other. One of these intervals contains the true value X^* . On each of the other $L - 1$ intervals, we sample a proxy for X^* according to the conditional distribution of X^* restricted to the interval.

Let ν denote the distribution of X^* . For an interval $I \subset [0, 1]$, write ν_I for the probability distribution of ν conditioned on I , i.e., $\frac{d\nu_I}{d\nu}(x) = \mathbb{1}\{x \in I\}/\nu(I)$. We design the following multi-phase querying strategy to attain the desired upper bound.

Algorithm 1 Querying Strategy under the Bayesian Setting

- 1: Recursively query the median of the posterior distribution of X^* , until it is supported on an interval I with $\nu(I) \in [2\delta LH_\alpha, 4\delta LH_\alpha]$.
 - 2: Let κ_j be the j/L quantile of ν_I for $j = 0, 1, \dots, L$ and let $I_j = [\kappa_{j-1}, \kappa_j]$ for $j \in [L]$. Query $\kappa_1, \dots, \kappa_{L-1}$ and identify j^* for which $f'(\kappa_{j^*-1}) \leq 0$ and $f'(\kappa_{j^*}) > 0$ so that I_{j^*} contains X^* .
 - 3: Query the median m_j of ν_{I_j} for $j \in [L]$. If $f'(m_{j^*}) > 0$, let $J_j = [\kappa_{j-1}, m_j]$ for all j ; otherwise let $J_j = [m_j, \kappa_j]$.
 - 4: For all $j \neq j^*$, sample $X_j \sim \nu_{J_j}$ independently. Denote $X_{j^*} = X^*$. For $j = 1, \dots, L$, run the regular bisection search on J_j to locate X_j up to ϵ -accuracy.
-

Phase 1 runs the median-based bisection search, which is equivalent to the regular bisection search on $U = F_\nu(X^*) \sim \text{Unif}[0, 1]$, where F_ν is the CDF of ν . Note that this step is always possible under the assumption $2\delta LH_\alpha \leq 1$. Phase 2 divides I into L subintervals $I_1, \dots, I_L$ with equal ν -probability and determines I_{j^*} containing X^* . Phase 3 is the key to ensure adequate separation between the subintervals $\{J_j\}_{j \in [L]}$. Phase 4 serves to achieve the ϵ -accuracy while obfuscating the adversary.

The querying strategy outlined in Algorithm 1 is clearly ϵ -accurate by design. We now show that it is also (δ, L) -private. The high-level proof idea is to consider an adversary who has access to $X_1, \dots, X_L$. Using a genie-aided argument, we argue that this adversary is stronger than the one who only has access to the query sequence. We then establish that the conditional distribution of X^* given $X_1, \dots, X_L$ is uniform on the X_j 's. Moreover, phase 3 of the querying strategy ensures that the X_j 's are all δ -separated. Therefore even with the additional knowledge of $X_1, \dots, X_L$, the adversary cannot estimate X^* accurately with probability higher than $1/L$.

Proof of Privacy: Since the adversary only has access to the query sequence q , any adversary's estimator $\tilde{X}$ must be a (random) function of q , that is $\tilde{X} \equiv \tilde{X}(q)$. Meanwhile by the design of our querying strategy, q can be com-

pletely reconstructed from $X_1, \dots, X_L$. To see that, note that $I, \{I_j\}, \{J_j\}$ and all the queries in phase 4 are deterministic functions of $X_1, \dots, X_L$. Therefore there is a mapping $\tilde{\psi}$ such that $\tilde{X}(q) = \tilde{\psi}(X_1, \dots, X_L)$. Thus,

$$\begin{aligned} & \mathbb{P} \left\{ \left| \tilde{X} - X^* \right| \leq \frac{\delta}{2} \right\} \\ &= \mathbb{E} \left[\mathbb{P} \left\{ \left| \tilde{X}(q) - X^* \right| \leq \frac{\delta}{2} \mid q \right\} \right] \\ &\leq \mathbb{E} \left[\sup_{\tilde{\psi}} \mathbb{P} \left\{ \left| \tilde{\psi}(X_1, \dots, X_L) - X^* \right| \leq \frac{\delta}{2} \mid X_1, \dots, X_L \right\} \right] \\ &\leq \mathbb{E} \left[\sup_{\tilde{x} \in [0,1]} \mathbb{P} \left\{ \left| \tilde{x} - X^* \right| \leq \frac{\delta}{2} \mid X_1, \dots, X_L \right\} \right]. \quad (8) \end{aligned}$$

We claim that

- (i) $X^* \mid X_1, \dots, X_L \sim \text{Unif}\{X_1, \dots, X_L\}$.
- (ii) With probability 1, $|X_i - X_j| > \delta$ for all $i \neq j$.

Assuming the two claims hold (the proofs are deferred to the supplementary material),

$$\begin{aligned} & \sup_{\tilde{x} \in [0,1]} \mathbb{P} \left\{ \left| \tilde{x} - X^* \right| \leq \frac{\delta}{2} \mid X_1, \dots, X_L \right\} \\ &= \sup_{\tilde{x} \in [0,1]} \frac{1}{L} \sum_{j \leq L} \mathbb{1} \left\{ \left| \tilde{x} - X_j \right| \leq \frac{\delta}{2} \right\} \leq \frac{1}{L}, \end{aligned}$$

where the equality is from (i) and the inequality is from (ii). Continuing (8), we have $\mathbb{P}\{|\tilde{X} - X^*| \leq \delta/2\} \leq 1/L$. Thus our strategy is (δ, L) -private.

Finally, the number of queries needed follows from a straightforward bookkeeping calculation, which we defer to the supplementary material. $\square$

Proof of the lower bound in Theorem 2. For the lower bound, the challenge lies in tracking and quantifying the amount of information the learner gains from the responses. Compared to the binary search model, the full gradient responses can potentially reveal too much information to the learner. To tackle this challenge, our key proof strategy is to find an event on which the learner cannot gather information on X^* too rapidly. The proof follows the following main steps.

Step 1: quantify the learner's information. We adopt the notion of "learner's intervals", $I_0, I_1, \dots$. Here, $I_0 = [0, 1]$ and I_i is the smallest interval that the learner knows to contain X^* after the first i queries.

Step 2: analyze the conditional distribution of X^* over the learner's interval. This is the key step of the proof. We want

to find a "good" event $\mathcal{B}$ on which the conditional distribution is uniform and hence the learner does not possess too much information on the location of X^* . To this end, we crucially exploit the stick-breaking characterization of the Dirichlet Process, which we describe next.

Given base distribution μ_0 and scaling parameter $\alpha > 0$, draw $\{X_k\}_{k=1}^\infty$ i.i.d. from μ_0 , and independently draw $\{V_k\}_{k=1}^\infty$ i.i.d. from $\text{Beta}(1, \alpha)$. From a stick of unit length, break off the first stick of length V_1 ; break off V_2 fraction of the remaining stick and repeat. In other words, denote by β_k the length of the k 'th stick. We have

$$\beta_k = V_k \cdot \prod_{j \leq k-1} (1 - V_j)$$

and $\sum_{k=1}^\infty \beta_k = 1$. Let $\mu = \sum_{k=1}^\infty \beta_k \delta_{X_k}$ be the discrete distribution supported on $\{X_k\}_{k=1}^\infty$, where δ_{X_k} denotes the point mass distribution at X_k . Then μ with the distribution function of F follows the Dirichlet process $\text{DP}(\mu_0, \alpha)$.

Here is an intuitive explanation on how the stick-breaking process helps us prove the uniformity of the conditional distribution of X^* . Under our prior construction, X^* is at the median of $F \sim \text{DP}(\lambda_{[0,1]}, \alpha)$, where we recall that $\lambda_{[0,1]}$ is the Lebesgue measure on $[0, 1]$. Therefore, X^* occurs at one of the stick-breaking locations X_k . Even though the X_k 's are distributed i.i.d. uniformly in $[0, 1]$, X^* itself does not follow the uniform distribution since the index i that corresponds to X^* is random. The key observation is that the conditional distribution of X^* is uniform conditional on the event $\mathcal{A}$ that the length of the longest stick is at least $1/2$. To prove uniformity, we first show that on the event $\mathcal{A}$, the median X^* must occur at the X_k that corresponds to the longest stick. Moreover, by independence of the stick lengths $\{\beta_k\}_{k \geq 1}$ and the locations $\{X_k\}_{k \geq 1}$, the distribution of the location corresponding to the longest stick is uniform in $[0, 1]$. Furthermore, the posterior distribution of X^* remains uniform as queries are sequentially submitted. The following Lemma 1 contains the precise statement on uniformity.

Some notation is necessary before stating Lemma 1. Firstly, denote by $\beta_{(1)}, \beta_{(2)}, \dots$ the order statistics of the lengths of the sticks in the stick-breaking process corresponding to F . Let

$$\mathcal{A} = \{\beta_{(1)} \geq 1/2\} = \cup_{z \geq 1/2} \mathcal{A}_z, \quad \mathcal{A}_z \triangleq \{\beta_{(1)} = z\}.$$

Let $J \subset [0, 1]$ be an arbitrary fixed interval. Write $[q_-, q_+] = I_i \cap J$. Let the event $\mathcal{B} = \mathcal{B}(z, J, y, i, \rho^{(i)}, \rho_-, \rho_+)$ encode the random instances of F, Y and the first i responses, defined as

$$\mathcal{B} = \left\{ \mathcal{A}_z, X^* \in J, Y = y, r^{(i)} = \rho^{(i)}, F(q_\pm) = \rho_\pm \right\}.$$

See Figure 2 for an example of F and some quantities in the definition of $\mathcal{B}$.

Lemma 1. For all $z \geq 1/2$, $J, y, i, \rho^{(i)}, \rho_- < 1/2, \rho_+ > 1/2$, we have

$$\mathcal{L}(X^* | \mathcal{B}) = \text{Unif}[q_-, q_+],$$

where $\mathcal{L}(\cdot)$ denotes the (conditional) distribution.

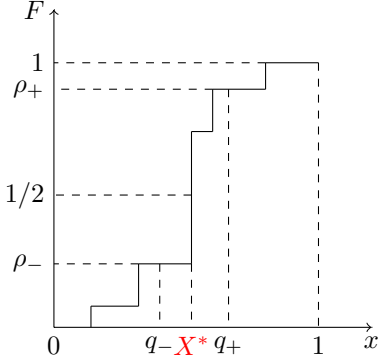


Figure 2. Conditional on $X^* \in J$ and the responses to the first i queries, the range of X^* is narrowed down to $I_i \cap J = [q_-, q_+]$. Further conditioning on $F(q_-) = \rho_-$ and $F(q_+) = \rho_+$, we show that F restricted to $[q_-, q_+]$ also follows a Dirichlet process after appropriate scaling.

The proof of Lemma 1 crucially utilizes the self-similarity property of the Dirichlet process. In short, it ensures that the values of F inside of $[q_-, q_+]$ conditional on information outside of $[q_-, q_+]$ also follows a scaled Dirichlet process. Thus the learner cannot gain too much information about the location of X^* in $[q_-, q_+]$.

Step 3: use Lemma 1 to control the speed at which the learner's interval shrinks. Divide $[0, 1]$ into $2/\delta$ subintervals $J_1, \dots, J_{2/\delta}$ of length $\delta/2$, and let J^* denote the subinterval of contains X^* . In this step, by integrating over instances of $\mathcal{B}$, and letting J range over the $2/\delta$ subintervals, we prove the following lemma.

Lemma 2. For all i , we have that

$$\mathbb{E} \left(\log \frac{|I_{i+1} \cap J^*|}{|I_i \cap J^*|} \middle| \mathcal{A} \right) \geq -\mathbb{P} \{q_{i+1} \in J^* \mid \mathcal{A}\}. \quad (9)$$

Step 4: In this step, we apply Lemma 2 to obtain the desired lower bound on the optimal query complexity. Let n be the total number of queries submitted by the learner. By writing

$\log |I_n \cap J^*|$ as a telescoping sum, we have that

$$\begin{aligned} & \mathbb{E} (\log |I_n \cap J^*| \mid \mathcal{A}) \\ &= \log |I_0 \cap J^*| + \sum_{i=0}^{n-1} \mathbb{E} \left(\log \frac{|I_{i+1} \cap J^*|}{|I_i \cap J^*|} \middle| \mathcal{A} \right) \\ &= \log \frac{\delta}{2} + \sum_{i=0}^{n-1} \mathbb{E} \left(\log \frac{|I_{i+1} \cap J^*|}{|I_i \cap J^*|} \middle| \mathcal{A} \right) \\ &\geq \log \frac{\delta}{2} - \mathbb{E} (\text{number of queries in } J^* \mid \mathcal{A}). \end{aligned} \quad (10)$$

From the accuracy requirement, we must have $|I_n| \leq \epsilon$ with probability 1. Therefore

$$\mathbb{E} (|I_n \cap J^*| \mid \mathcal{A}) \leq \mathbb{E} (|I_n| \mid \mathcal{A}) \leq \epsilon/2,$$

so that by Jensen's inequality,

$$\mathbb{E} (\log |I_n \cap J^*| \mid \mathcal{A}) \leq \log \mathbb{E} (|I_n \cap J^*| \mid \mathcal{A}) \leq \log \frac{\epsilon}{2}.$$

Combining the last display with (10) yields

$$\mathbb{E} (\text{number of queries in } J^* \mid \mathcal{A}) \geq \log \frac{\delta}{\epsilon}. \quad (11)$$

Consider an adversary who adopts the *proportional-sampling* strategy (Xu, 2018). That is, suppose the adversary's estimator $\tilde{X}$ is sampled from the empirical distribution of the queries. For this particular $\tilde{X}$,

$$\begin{aligned} \frac{1}{L} &\geq \mathbb{P} \left\{ \tilde{X} \in [X^* - \delta/2, X^* + \delta/2] \right\} \\ &= \frac{\mathbb{E}(\text{number of queries in } [X^* - \delta/2, X^* + \delta/2])}{n}, \end{aligned}$$

which gives a lower bound on the total number of queries:

$$n \geq L \mathbb{E}(\text{number of queries in } [X^* - \delta/2, X^* + \delta/2]).$$

Since $J^* \subset [X^* - \delta/2, X^* + \delta/2]$, combining the last display with (11) yields that

$$n \geq L \mathbb{E}(\text{number of queries in } J^*) \geq \mathbb{P}(\mathcal{A}) L \log \frac{\delta}{\epsilon}.$$

We have thus arrived at the desired query complexity lower bound with

$$c_1 = \mathbb{P}(\mathcal{A}) = \mathbb{P} \{ \beta_{(1)} > 1/2 \} \geq \mathbb{P} \{ \beta_1 > 1/2 \},$$

where $\beta_1 \sim \text{Beta}(1, \alpha)$ is the length of the first stick from the stick-breaking characterization of the Dirichlet process. $\square$

5.2. Proof sketch under the Minimax setting

Since the response contains the full gradient information, the key challenge in the analysis is to track the amount of information available to the learner. Note that aside from the directional information $\mathbb{1}\{X^* \geq q\}$, the response for a query q contains additional information on $(f^*)'(q)$. The key insight in the proof under the minimax setting, is that under the Assumption 1 on the richness of the family of functions, only the directional information is relevant to the learning task. Therefore, it suffices to only track the learner’s knowledge with the directional information from the responses.

Starting with the upper bound, we design a querying strategy that is ϵ -accurate, (δ, L) -private, and submits at most $\max\{2L + \log(\delta/\epsilon), L + \log(1/\epsilon)\}$ queries. In particular, our querying strategy only utilizes the directional information of the gradient responses. Firstly, note that since the gradient responses contain the binary directional information, the learner can always check whether an interval contains X^* by querying the two endpoints. We refer to a pair of queries at q and $q + \epsilon$ as a *guess*. The key privacy-ensuring mechanism is to check L guesses that are δ apart from each other. By doing so, the learner manually plants L possible locations for X^* that an adversary cannot rule out without observing the responses, thus achieving (δ, L) -privacy.

To prove the lower bound, we need to show that a querying strategy that only utilizes the directional information can be optimal. Firstly, let us give a heuristic argument of why only the gradient information is relevant to learning X^* under Assumption 1. Given $(f^*)'(a) < 0$ and $(f^*)'(b) > 0$, under Assumption 1, X^* can be anywhere between a and b regardless of the value of the gradients $(f^*)'(a)$, $(f^*)'(b)$. We should point out that the richness assumption is necessary. For example, suppose $\mathcal{F}$ is the family of convex polynomial functions with fixed degree d . Then the learner can solve for the X^* by submitting d distinct queries at arbitrary locations, making both learning and obfuscation trivial.

The lower bound proof contains two main ingredients.

- (a) Step 1: Rigorously justify the claim that under Assumption 1, the learner does not benefit from the additional gradient information aside from the one-bit directional response. In particular, we show that the learner cannot search faster than the bisection method on any interval $I \subset [0, 1]$. Therefore, for each interval of length δ , it takes at least $\log(\delta/\epsilon)$ queries in I to achieve ϵ -accuracy, in the worst case.
- (b) Step 2: Relate the adversary’s statistical performance to the size of the information set (Tsitsiklis et al., 2018) of

a query sequence q , defined as

$$\mathcal{I}(q) = \{x \in [0, 1] : \exists f \in \mathcal{F} \text{ and } y, \\ \text{s.t. } x = \arg \min f, \text{ and } q(f, y) = q\}.$$

The information set contains all possible values of X^* that could lead to the query sequence q . We show that to ensure the adversary achieves δ -accuracy with probability at most $1/L$, there must be some q for which the δ -covering number of $\mathcal{I}(q)$ is at least L . Note that from the ϵ -accuracy requirement, each member of $\mathcal{I}(q)$ is sandwiched between a pair of queries in q that are at most ϵ -apart. Therefore, q contains at least L such pairs of queries, contributing a total of $2L$ queries.

After performing these two steps, some challenges remain. The functions associated with q (in step 2) may not coincide with the worst-case instances that arise from step 1. Therefore, the remaining task is to combine the two lower bounds $\log(\delta/\epsilon)$ and $2L$. To this end, we show the existence of some interval I , such that for some f minimized in I , the learner must pay not only the $\log(\delta/\epsilon)$ queries for accuracy, but also the $2L$ queries for privacy. The high-level idea behind the proof is to divide q into two sub-sequences q_{before} , q_{after} , before and after the $2L$ queries (in step 2) are submitted. The key observation is that q_{before} is shared by a large class of functions whose minimizers lie in some δ -length interval I . For all these functions, the cost of $2L$ queries would have been committed in q_{before} . For at least one of them, an extra cost of $\log(\delta/\epsilon)$ queries must be paid in q_{after} to achieve ϵ -accuracy.

Acknowledgements

J. Xu is supported by the NSF Grants IIS-1838124, CCF-1850743, and CCF-1856424. D. Yang is supported by the NSF Grant CCF-1850743 and IIS-1838124. We thank the anonymous reviewers for the constructive comments, and Niva Ran and Benjamin Ran for suggesting ideas that inspired the algorithm used in the upper bound of the Bayesian formulation of the problem.

References

- Abadi, M., Feigenbaum, J., and Kilian, J. On hiding information from an oracle. *Journal of computer and system sciences*, 39(1):21–50, 1989.
- Abadi, M., Chu, A., Goodfellow, I., McMahan, B., Mironov, I., Talwar, K., and Zhang, L. Deep learning with differential privacy. In *Proceedings of the 2016 ACM SIGSAC Conference on Computer and Communications Security*, pp. 308–318. ACM, 2016.
- Agarwal, A., Foster, D. P., Hsu, D., Kakade, S. M., and Rakhlin, A. Stochastic convex optimization with bandit

- feedback. *SIAM Journal on Optimization*, 23(1):213–240, 2013.
- Agarwal, N., Suresh, A. T., Yu, F. X. X., Kumar, S., and McMahan, B. cpsgd: Communication-efficient and differentially-private distributed SGD. In *Advances in Neural Information Processing Systems*, pp. 7564–7575, 2018.
- Bornkamp, B. and Ickstadt, K. Bayesian nonparametric estimation of continuous monotone functions with applications to dose–response analysis. *Biometrics*, 65(1): 198–205, 2009.
- Chor, B., Goldreich, O., Kushilevitz, E., and Sudan, M. Private information retrieval. In *Proceedings of IEEE 36th Annual Foundations of Computer Science*, pp. 41–50. IEEE, 1995.
- Dwork, C. Differential privacy: A survey of results. In Agrawal, M., Du, D., Duan, Z., and Li, A. (eds.), *Theory and Applications of Models of Computation*, pp. 1–19, Berlin, Heidelberg, 2008. Springer Berlin Heidelberg. ISBN 978-3-540-79228-4.
- Erturk, M. S. and Xu, K. Dynamically protecting privacy, under uncertainty. *arXiv preprint arXiv:1911.08875*, 2019.
- Fanti, G., Kairouz, P., Oh, S., and Viswanath, P. Spy vs. spy: Rumor source obfuscation. In *ACM SIGMETRICS Performance Evaluation Review*, volume 43, pp. 271–284. ACM, 2015.
- Gasarch, W. A survey on private information retrieval. *Bulletin of the EATCS*, 82(72-107):113, 2004.
- Jain, P., Kothari, P., and Thakurta, A. Differentially private online learning. In *Conference on Learning Theory*, pp. 24–1, 2012.
- Juuti, M., Szyller, S., Marchal, S., and Asokan, N. Prada: protecting against dnn model stealing attacks. In *2019 IEEE European Symposium on Security and Privacy (EuroS&P)*, pp. 512–527. IEEE, 2019.
- Kairouz, P., McMahan, H. B., Avent, B., Bellet, A., Bennis, M., Bhagoji, A. N., Bonawitz, K., Charles, Z., Cormode, G., Cummings, R., et al. Advances and open problems in federated learning. *arXiv preprint arXiv:1912.04977*, 2019.
- Lavine, M. and Mockus, A. A nonparametric bayes method for isotonic regression. *Journal of Statistical Planning and Inference*, 46(2):235–248, 1995.
- Luo, W., Tay, W. P., and Leng, M. Infection spreading and source identification: A hide and seek game. *IEEE Transactions on Signal Processing*, 64(16):4228–4243, 2016.
- McMahan, B. and Ramage, D. Federated learning: Collaborative machine learning without centralized training data. 2017. <https://ai.googleblog.com/2017/04/federated-learning-collaborative.html>. Date accessed: July 31, 2020.
- McMahan, B., Moore, E., Ramage, D., Hampson, S., and y Arcas, B. A. Communication-efficient learning of deep networks from decentralized data. In *Artificial Intelligence and Statistics*, pp. 1273–1282. PMLR, 2017.
- Melis, L., Song, C., De Cristofaro, E., and Shmatikov, V. Exploiting unintended feature leakage in collaborative learning. In *2019 IEEE Symposium on Security and Privacy (SP)*, pp. 691–706. IEEE, 2019.
- Mirmohseni, M. and Maddah-Ali, M. A. Private function retrieval. In *2018 Iran Workshop on Communication and Information Theory (IWCIT)*, pp. 1–6. IEEE, 2018.
- Neelon, B. and Dunson, D. B. Bayesian isotonic regression and trend analysis. *Biometrics*, 60(2):398–406, 2004.
- Ramgopal, P., Laud, P., and Smith, A. Nonparametric bayesian bioassay with prior constraints on the shape of the potency curve. *Biometrika*, 80(3):489–498, 1993.
- Song, S., Chaudhuri, K., and Sarwate, A. D. Stochastic gradient descent with differentially private updates. In *2013 IEEE Global Conference on Signal and Information Processing*, pp. 245–248. IEEE, 2013.
- Tang, W., Ho, C.-J., and Liu, Y. Optimal query complexity of secure stochastic convex optimization. *Advances in Neural Information Processing Systems*, 33, 2020a.
- Tang, W., Wang, W., Fanti, G., and Oh, S. Privacy-utility tradeoffs in routing cryptocurrency over payment channel networks. In *Abstracts of the 2020 SIGMETRICS/Performance Joint International Conference on Measurement and Modeling of Computer Systems*, pp. 81–82, 2020b.
- Tsitsiklis, J. N. and Xu, K. Delay-predictability trade-offs in reaching a secret goal. *Operations Research*, 66(2): 587–596, 2018.
- Tsitsiklis, J. N., Xu, K., and Xu, Z. Private sequential learning. *arXiv preprint arXiv:1805.02136*, 2018.
- Xu, J., Xu, K., and Yang, D. Optimal query complexity for private sequential learning against eavesdropping. *arXiv preprint arXiv:1909.09836*, 2019.
- Xu, K. Query complexity of Bayesian private learning. In *Advances in Neural Information Processing Systems*, pp. 2431–2440, 2018.

Young, W., Cameron, N., and Tinsley, Y. Juries in criminal trials. Technical report, New Zealand Law Commission Report, 2001.