

Private Optimal Inventory Policy Learning for Feature-Based Newsvendor with Unknown Demand

Tuoyi Zhao,^a Wen-Xin Zhou,^b Lan Wang^{a,*}

^aDepartment of Management Science, Miami Herbert Business School, University of Miami, Coral Gables, Florida 33146; ^bDepartment of Information and Decision Sciences, University of Illinois at Chicago, Chicago, Illinois 60607

*Corresponding author

Contact: txz311@miami.edu,  <https://orcid.org/0009-0006-1866-3920> (TZ); wenxinz@uic.edu,  <https://orcid.org/0000-0002-2761-485X> (W-XZ); lxw611@miami.edu,  <https://orcid.org/0000-0002-3217-0202> (LW)

Received: April 25, 2023

Revised: January 13, 2024

Accepted: April 8, 2024

Published Online in *Articles in Advance*:
October 24, 2024

<https://doi.org/10.1287/mnsc.2023.01268>

Copyright: © 2024 INFORMS

Abstract. The data-driven newsvendor problem with features has recently emerged as a significant area of research, driven by the proliferation of data across various sectors such as retail, supply chains, e-commerce, and healthcare. Given the sensitive nature of customer or organizational data often used in feature-based analysis, it is crucial to ensure individual privacy to uphold trust and confidence. Despite its importance, privacy preservation in the context of inventory planning remains unexplored. A key challenge is the nonsmoothness of the newsvendor loss function, which sets it apart from existing work on privacy-preserving algorithms in other settings. This paper introduces a novel approach to estimating a privacy-preserving optimal inventory policy within the f -differential privacy framework, an extension of the classical (ϵ, δ) -differential privacy with several appealing properties. We develop a clipped noisy gradient descent algorithm based on convolution smoothing for optimal inventory estimation to simultaneously address three main challenges: (i) unknown demand distribution and nonsmooth loss function, (ii) provable privacy guarantees for individual-level data, and (iii) desirable statistical precision. We derive finite-sample high-probability bounds for optimal policy parameter estimation and regret analysis. By leveraging the structure of the newsvendor problem, we attain a faster excess population risk bound compared with that obtained from an indiscriminate application of existing results for general nonsmooth convex loss. Our bound aligns with that for strongly convex and smooth loss function. Our numerical experiments demonstrate that the proposed new method can achieve desirable privacy protection with a marginal increase in cost.

History: Accepted by J. George Shanthikumar, data science.

Funding: This work was supported by the National Science Foundation [Grants DMS-2113409 and DMS 2401268 to W.-X. Zhou, and FRGMS-1952373 to L. Wang].

Supplemental Material: The online appendix and data files are available at <https://doi.org/10.1287/mnsc.2023.01268>.

Keywords: newsvendor • differential privacy • data-driven decision-making • convolution smoothing • regret analysis

1. Introduction

The newsvendor problem, a classical example of the inventory-control problem, is of fundamental importance to operations management. In recent years, there has been growing interest in data-driven feature-based newsvendor problems because of the vast amount of data generated by retail and supply chains, e-commerce, banking, financial, hospitals, and other business domains. The goal of the feature-based newsvendor problem is to estimate the optimal inventory policy based on the historical demand data as well as the observed features (e.g., product characteristics, customer characteristics) associated with the demand. The ability to determine optimal inventory levels based on features (or contextual information) such as location and usage is essential for supply

chain planning. Recently, Ban and Rudin (2019) carefully justified, from a theoretical perspective, the value of incorporating features in the newsvendor problem when such information is available. They proved that ignoring features may lead to estimation bias, which does not diminish as the number of observations gets large.

In a host of applications, feature-based inventory analysis involves sensitive customer or organizational data. Examples include, but are not restricted to, the following:

- *Healthcare.* In hospitals, nurse staffing in the emergency room can be formulated as a feature-based newsvendor problem; see, for example, He et al. (2012), Green et al. (2013), and Ban and Rudin (2019). The features can include the hospital inflow and outflow conditions,

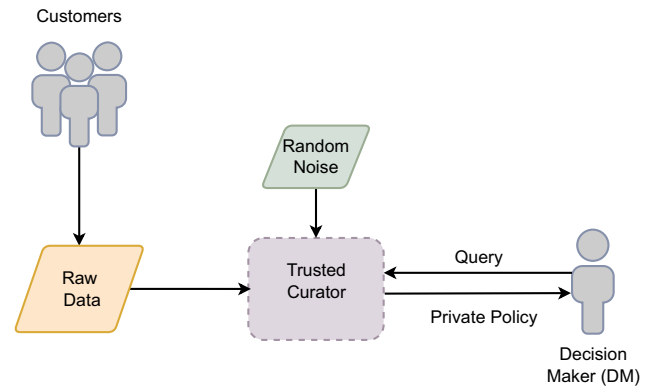
surgical case volume, behavioral health patients' boarding information, doctor staffing information, and nurses' credentials. As another application, surgical procedures require a large number of consumable supplies that need to be kept in hospital inventory and transported to the operating rooms. Görgülü and Sarhangian (2022) formulate the problem of preparing a surgical preference card, a list of items for each surgery, as a newsvendor problem. The features include surgery type, patient disease status, and physicians' past records, among others. Hospitals, in general, would prefer to keep the related information internal and any physician- and patient-level information private.

- *E-commerce.* Companies such as Chewy, an online retailer of pet supplies, or grocery stores with online ordering and delivery service often leverage customer-level data for inventory management. This helps not only coordinate shipping from its local warehouses but also design targeted marketing campaigns (e.g., sending coupons to different groups of customers based on predicted individualized inventory levels).

- *Finance.* In portfolio management, mutual funds hold a certain percentage of their assets in cash to meet redemption demand from investors. The decision on how much cash to reserve can be formulated as a feature-based newsvendor problem. If not enough cash is held, the fund must sell some of its holdings and will incur transaction costs. In the recent Silicon Valley Bank's downfall in March 2023, the bank had to sell its securities to raise cash to meet a wave of withdrawals from customers. A strategy of inventory policy can be planned based on the financial and operational variables and clients' behaviors.

In the context of the applications discussed above, the preservation of privacy is of critical importance. Despite this, systematic studies in the realm of inventory planning are lacking. The protection of individual privacy is essential in maintaining customer trust and confidence and in helping the business avoid financial losses and reputation damage (Williams 2020, Hu et al. 2022, Fainmesser et al. 2023). The prevalence of individual-level data and the increased awareness of privacy concerns motivate us to develop a principled privacy-preserving framework for data-driven feature-based newsvendor problems with unknown demand. We focus on the scenario where there is a trusted curator, such as the company's in-house business analytics team, responsible for processing and analyzing the data. The primary objective of this paper is to develop a data-driven approach that generates valuable outputs. These outputs assist a decision maker (hereafter referred to as "DM") in estimating the optimal inventory level, all while safeguarding the privacy of historical individual-level data. In essence, our approach controls the likelihood of an adversary making harmful inferences about a data subject based on a

Figure 1. (Color online) Illustration of Privacy Protection



differentially private data release, ensuring it remains a small-probability event.

In more detail, we study a feature-based newsvendor problem where the demand distribution is unknown to the DM. The trusted curator has access only to n past records (historical data) $\{d_i, \mathbf{x}_i\}$, $i = 1, \dots, n$, where d_i is the observed demand, and $\mathbf{x}_i$ is the associated vector of features (or covariates). When presented with a new query, the curator utilizes a data-driven iterative algorithm to release a private output. This output guides the DM in determining the optimal inventory level while ensuring that the sensitive information from the historical data remains noninferable from the output. The iterative algorithm introduces a carefully tuned amount of random noise to the statistical outputs, aiming to strike a balance between privacy protection and statistical accuracy. See Figure 1 for an illustration.

Distinct from prior work on privacy-preserving algorithms in other business applications, we confront the challenge of the nonsmoothness of the newsvendor loss function. In this setting, we leverage the recently introduced concept of f -differential privacy (Dong et al. 2022) and propose a noisy clipped gradient descent algorithm based on convolution smoothing for optimal inventory estimation. The new approach simultaneously addresses the three main challenges: (i) unknown demand distribution and nonsmooth loss function, (ii) provable privacy guarantees for individual-level data, and (iii) desirable statistical precision. Importantly, our theoretical and numerical results demonstrate that a reasonable degree of privacy protection can be achieved with minimal sacrifice of data utility, particularly when the size of the historical data set is large.

1.1. Contributions

Our major results and contributions are summarized as follows.

1.1.1. Provable Privacy Protection Guarantee in the f -Differential Privacy Framework. To establish rigorous privacy protection properties, we adopt a recently

introduced novel privacy framework named f -differential privacy (f -DP) (Dong et al. 2022), which generalizes the classical (ϵ, δ) -DP notion (Dwork et al. 2006a, b) with several attractive properties; see Section 4.1 for a more in-depth introduction. The (ϵ, δ) -DP notion was proposed by the computer science community and has become a popular framework for provable privacy protection against arbitrary adversaries while allowing the release of analytical summaries. It provides a statistical hypothesis testing interpretation for differential privacy, thereby making the privacy guarantees easily understandable. Despite its great success, a major shortcoming of (ϵ, δ) -DP is its inability to tightly handle composition (a.k.a. repeated application of the mechanism to the same data set). Using the output of an (ϵ, δ) -DP mechanism, the power of any α -level test is bounded by $e^\epsilon \alpha + \delta$. Recall that the composition of (ϵ_1, δ_1) - and (ϵ_2, δ_2) -DP mechanisms results in an $(\epsilon_1 + \epsilon_2, \delta_1 + \delta_2)$ -DP mechanism. The resulting power bound $e^{\epsilon_1 + \epsilon_2} \alpha + \delta_1 + \delta_2$ of any α -level test no longer tightly characterizes the trade-off between significance level and power. As a fundamental observation, Dong et al. (2022) pointed out that (ϵ, δ) -DP is misparameterized in the sense that the guarantees of the composition of (ϵ_i, δ_i) -DP mechanisms cannot be characterized by any single pair of parameters (ϵ, δ) . Many recent efforts have been devoted to developing relaxations of DP for which composition can be handled exactly. These notions of DP no longer have hypothesis-testing interpretations; rather, they are based on studying divergences that satisfy a certain information processing inequality. We refer to Section 1 of Dong et al. (2022) for an in-depth discussion on this matter.

The main idea of the f -DP is the usage of the so-called trade-off functions as a measure of indistinguishability of two neighboring data sets rather than a few parameters as (ϵ, δ) -DP and other prior relaxations do. It preserves the hypothesis testing interpretation of differential privacy. Furthermore, it captures all the desirable properties of prior differential privacy definitions, in particular, composition, amplification by sampling, and Gaussian mechanism, tightly and analytically. It provides a powerful technique to import existing results proven for the (ϵ, δ) -DP to f -DP.

In the powerful and versatile f -DP framework, we rigorously establish that our data-driven algorithm provides the desired privacy guarantees.

1.1.2. A Computationally Efficient Algorithm to Estimate the Feature-Based Optimal Inventory Policy with Unknown Demand Function. Traditionally, the newsvendor problem is solved based on the assumption that the demand distribution is known up to a small number of parameters. The commonly used data-driven estimation procedures often consist of two steps: the first step estimates the parameters using the observed data, and the second step performs the optimization to estimate

the optimal order quantity. However, in reality, the true demand distribution is hardly ever known to the DM.

For the data-driven feature-based newsvendor problem, Ban and Rudin (2019) proposed a one-step estimation procedure based on empirical risk minimization (ERM) and established its connection to quantile regression (Koenker and Bassett 1978). In terms of computation, Ban and Rudin (2019) reframed the ERM problem as a linear program and utilized existing linear programming solvers. These general-purpose solvers are capable of generating solutions with high precision (low duality gap). However, in the context of machine learning, this is inefficient for two reasons. First, generic toolboxes are often unaware of the problem structure and tend to be too slow or encounter memory issues. Secondly, high precision is not always necessary for machine learning problems, and a duality gap of the order of machine precision may not be required (Bach et al. 2012). More importantly, commonly used algorithms for solving linear programs, such as simplex-based methods and interior point methods, may not be readily adaptable for privacy preservation purposes.

To address the computational and privacy concerns mentioned above, we propose a new approach that utilizes convolution smoothing (Fernandes et al. 2021, He et al. 2023). This approach transforms the nondifferentiable newsvendor loss function into a twice-differentiable, convex, and locally strongly convex surrogate, allowing for fast and scalable gradient-based algorithms for optimization. Additionally, to ensure privacy protection while maintaining computational tractability, inspired by Song et al. (2013), Bassily et al. (2014), and Lee and Kifer (2018), we adopt a noisy optimization approach by adding Gaussian noise to the gradient of the smoothed empirical cost in each iteration. By carefully selecting the scale of the added noise and the number of iterations, we can achieve the desired privacy level along a sequence of outputs. The algorithm is designed for efficient implementation, and in this paper, we provide both privacy protection guarantees and statistical accuracy guarantees for the output of this novel algorithm.

1.1.3. Finite-Sample Performance Bounds and Excess Risk Analysis.

Under a linear demand model with an unknown error distribution and a potentially large number of features, we analyze the convergence of the proposed algorithm and its finite sample performance error bounds. We also derive its regret bound, which is the difference between its expected cost and the optimal cost of the clairvoyant who knows the underlying demand distribution. The regret upper bound is of the order $O(\log(n) \max\{((p + \log n)/\mu n)^2, p/n\})$, where n is the size of the historical data set, p is the number of the features, and μ is the privacy parameter. As we discuss in Section 4.1, $\mu = 0.5$ (or less) indicates a reasonable degree of privacy protection in practice. The term

$O(\log(n)((p + \log n)/\mu n)^2)$ corresponds to the additional regret because of privacy protection, which goes to zero quickly as n gets large for a reasonable choice of μ . The theory and our numerical results suggest that privacy protection can be achieved with a reasonably small additional cost.

The idea of convolution smoothing was initially proposed in the optimization community by Chen and Mangasarian (1995, 1996), where sigmoid functions were used as smooth approximations of the plus function $\max\{x, 0\}$. However, the impact of smoothing on statistical performance, in terms of estimation error bounds or regret bounds, remained largely unknown until recent studies in the context of quantile regression by Fernandes et al. (2021), Tan et al. (2022), and He et al. (2023).

In this paper, we conduct a comprehensive analysis of the noisy clipped gradient descent iterates, as opposed to the hypothesized empirical risk minimizer. By exploring the specific structure of the newsvendor problem, we achieve a faster excess population risk bound compared with the results obtained by indiscriminately applying existing results developed for general nonsmooth convex loss. Our bound matches with what one would obtain when the loss function is both strongly convex and smooth. The combination of carefully selected smoothing and noise-scale parameters allows for control over the trade-off between statistical efficiency and the level of privacy. A key aspect of our analysis is a novel characterization of the local strong convexity and smoothness of the smoothed cost function, which subtly depends on the order of the smoothing parameter. The technical devices we employ in this paper to establish the theoretical framework are distinct from earlier works, such as those presented in Ban and Rudin (2019), and yield sharper results, as elaborated in Section 5. Furthermore, we relax the independent and identically distributed (i.i.d.) error condition in their paper. By allowing the error distribution to be heteroscedastic, we permit the features to influence not only the location of the demand distribution but also its dispersion. Furthermore, we do not require the error distribution to be bounded.

1.2. Notation and Organization

The following general notation will be used throughout the paper. We use $\mathbf{I}_p$ to denote the $p \times p$ identity matrix. For a vector $\mathbf{u} \in \mathbb{R}^p$ ($p \geq 2$), we write $\|\mathbf{u}\|^2 = \mathbf{u}^T \mathbf{u}$. For a positive definite matrix $\mathbf{A}$, we write $\|\mathbf{A}\|_2 = \max_{\mathbf{u}: \|\mathbf{u}\|_2=1} \|\mathbf{A}\mathbf{u}\|_2$ and $\|\mathbf{u}\|_{\mathbf{A}} = \sqrt{\mathbf{u}^T \mathbf{A} \mathbf{u}}$. We use $\mathbb{S}^{p-1}$ to denote the unit sphere in $\mathbb{R}^p$, that is, $\mathbb{S}^{p-1} = \{\mathbf{u} \in \mathbb{R}^p : \|\mathbf{u}\|_2 = 1\}$. For two sequences of positive numbers $\{a_n\}_{n \geq 1}$ and $\{b_n\}_{n \geq 1}$, we write $a_n \lesssim b_n$ if there exists some constant $C > 0$ independent of n such that $a_n \leq Cb_n$ for all n ; we write $a_n \gtrsim b_n$ if $b_n \lesssim a_n$, and we write $a_n \asymp b_n$ if $a_n \lesssim b_n$ and $b_n \lesssim a_n$. For an event or set E , let $\mathbb{1}(E)$ or $\mathbb{1}\{E\}$ denote the indicator function.

The paper is organized as follows. In Section 2, we review the related literature. Section 3 presents the model and introduces the underlying assumptions. In Section 4, we introduce the basics of f -differential privacy, present the new privacy-preserving algorithm for the feature-based newsvendor problem, and provide theoretical justifications for privacy-preserving guarantees. In Section 5, we provide high-probability bounds for the estimated private parameter indexing the optimal inventory policy and the regret analysis. We analyze the performance of our approach through an extensive numerical study and a real data example in Section 6. Section 7 contains some concluding remarks. The technical details are given in the online appendices.

2. Related Review

We briefly review related research in data-driven newsvendor problems and differential privacy for operations management.

2.1. Data-Driven Newsvendor Problem with Unknown Demand

Earlier work on newsvendor problems often assumes the demand distribution is known. There has also been extensive literature on relaxing the known demand distribution assumption but without using any feature information. Ban and Rudin (2019) provided an excellent literature review and broadly characterized these methods into three categories: the Bayesian approach, the minimax approach, and the data-driven approach. Our proposed method is more closely related to the data-driven approach where the DM uses the observed sample to make decisions; see Burnetas and Smith (2000), Godfrey and Powell (2001), Powell et al. (2004), Levi et al. (2007), Kunnumkal and Topaloglu (2008), Huh and Rusmevichientong (2009), and Levi et al. (2015), among others. Related to this line of work, Liyanage and Shanthikumar (2005), Hannah et al. (2010), See and Sim (2010), Beutel and Minner (2012), Ban and Rudin (2019), and Oroojlooyjadid et al. (2020) incorporated feature-based information. Ban and Rudin (2019) provided a systematic study on the benefits of incorporating features, proposed new algorithms, and derived performance bounds. Oroojlooyjadid et al. (2020) considered a deep-learning approach.

Our approach differs from the aforementioned work in several major aspects. First, our algorithm provides individual-level data privacy protection. To the best of our knowledge, this is the first time in the literature of newsvendor problems the issue of privacy is systematically investigated. Second, we provide a theoretical error bound for the T -step output of the proposed algorithm directly. In contrast, the theory of the early work is for the limiting (or theoretical) solution of their proposed algorithms. Third, we substantially relax the

technical conditions on the random error distribution for the theory compared with the earlier work.

2.2. Differential Privacy for Operations Management

In the last decade, privacy preservation has received substantial attention in theoretical computer science, database, and cryptography literature. There exist different notions of privacy. *Differential privacy*, a seminal concept introduced in Dwork et al. (2006a, b), has emerged as the foundation for developing a rigorous framework for provable privacy protection against arbitrary adversaries. The most commonly used form of differential privacy relies on two parameters, $\epsilon \geq 0$ and $0 \leq \delta \leq 1$, and is often also referred to as the (ϵ, δ) -differential privacy. This concept has an intuitive hypothesis interpretation. Suppose an attacker would like to distinguish two neighboring data sets that differ by only one observation. Formulated as a hypothesis testing problem, accepting the null hypothesis means the attacker cannot tell the two data sets apart. Then, for any level α test ($0 < \alpha < 1$) based on the output of a privacy-preserving algorithm satisfying (ϵ, δ) -differential privacy, its power (a.k.a. the probability of rejecting the null hypothesis when the two data sets are different) is upper bounded by $e^\epsilon \alpha + \delta$. Moreover, (ϵ, δ) -differential privacy is immune to postprocessing; that is, combining two differential private algorithms preserves differential privacy. Although (ϵ, δ) -differential privacy provides an elegant formalism for privacy protection, it is known to suffer from the major drawback that it does not tightly handle composition. This makes it challenging to provide a tight analysis of the cumulative privacy loss over multiple computations, thus limiting its applicability to practically useful privacy-preserving algorithms, which often involve injecting privacy protection into different modules and iterative steps.

Although several relaxations of the (ϵ, δ) -differential privacy have been proposed, they do not handle well fundamental primitives associated with differential privacy, such as privacy amplification by subsampling. This motivated us to adopt a recently proposed new notion of *f*-differential privacy (Dong et al. 2022), which extends (ϵ, δ) -differential privacy and overcomes the above limitations. Similar to (ϵ, δ) -differential privacy, *f*-differential privacy characterizes privacy preservation from the hypothesis testing perspective. Rather than using a pair of parameters, (ϵ, δ) , to balance between type I and type II errors, *f*-differential privacy uses a trade-off function. This functional extension of differential privacy avoids the drawbacks mentioned above. We refer to Section 4.1 for more detailed discussions on the properties of *f*-differential privacy. The notion of *f*-differential privacy was recently published as a discussion paper in the leading statistical journal *Journal of the Royal Statistical Society, Series B*. One of the

discussants wrote, “One can expect the latter (*f*-differential privacy) to become a dominant approach in this literature given its appealing intuitive hypothesis-testing interpretation, exact composition property, the central limit role for composition, and computational tractability for approximating privacy losses.”

2.3. Differentially Private Convex Optimization

Our work is also related to the literature on differentially private convex optimization. Differentially private empirical risk minimization (ERM) is a well-studied area. The earlier popular approaches include output perturbation (Chaudhuri and Monteleoni 2008, Wu et al. 2017) and objective perturbation (Chaudhuri and Monteleoni 2008, Chaudhuri et al. 2011, Jain and Thakurta 2014, Abadi et al. 2016, Iyengar et al. 2019, Slavkovic and Molinari 2022).

Motivated by Song et al. (2013) and Bassily et al. (2014), we consider a noisy gradient descent algorithm. Differential privacy with various types of gradient descent algorithms has been studied by Song et al. (2013), Bassily et al. (2014), Wang et al. (2017), Lee and Kifer (2018), Wang (2018), Bassily et al. (2019), and Balle et al. (2020), among others. The methods in Bassily et al. (2014), Bassily et al. (2019), Feldman et al. (2020), and others do not directly apply to our setting. Most of the prior work requires strong convexity and other smoothness conditions that are not satisfied by the newsvendor loss function. In the case for which Lipschitz continuity suffices, the known excess loss rate is suboptimal in our setting, as they do not explore the specific structure for the newsvendor loss as we do. Unlike the earlier literature, we do not assume the gradient is bounded by a constant, and we carefully analyzed a clipped DP gradient descent algorithm. In the statistical literature, Avella-Medina et al. (2023) recently investigated optimization-based approaches for Gaussian differentially private *M*-estimators. However, the objective function in our setting does not satisfy the local strong convexity and smoothness in their paper. However, our proof technique deviates from theirs. We will discuss the main challenges and our proof strategies in Section 5.3. Despite the nonsmooth newsvendor loss, our novel analysis based on restricted strong convexity and smoothness leads to a faster excess population risk rate, which is only obtained in Section 5 of Feldman et al. (2020) when the loss function is both λ -strongly convex and β -smooth. From this perspective, our results and proof techniques bring new insights and results to the differential private convex optimization literature.

Our work is also related to the growing but still limited literature on privacy preservation in operations management. Chen et al. (2022b) addressed privacy preservation for personalized pricing with demand following a generalized linear model. They proposed the new notion of anticipating (ϵ, δ) -differential privacy

that is tailored to the dynamic pricing problem. Lei et al. (2020) and Chen et al. (2022a) considered personalized pricing using the notion of central differential privacy and local differential privacy. The aforementioned papers do not face the challenge of nonsmooth loss function as we have here. However, our work is not a mere application of existing results on private non-smooth convex optimization, nor does it utilize standard arguments such as uniform stability. Furthermore, in contrast to these existing works, we adopt the f -differential privacy framework to study the performance of the new privacy-preserving algorithm and establish its provable privacy protection guarantees.

3. Problem Formulation

3.1. Feature-Based Newsvendor Problem

We consider the classical single-period newsvendor problem setting. The DM needs to determine the ordering level q based on the observable demand d and feature vector $\mathbf{x}$. Both d and $\mathbf{x}$ are random. We assume that the distribution of the demand d is unknown. In the feature-based newsvendor problem, given a realization of the feature vector $\mathbf{x} \in \mathbb{R}^p$, the DM sets the ordering level $q(\mathbf{x})$ to minimize the conditional expected cost function

$$\mathbb{E}\{C(q(\mathbf{x}), d) | \mathbf{x}\} = \mathbb{E}\{[h(q(\mathbf{x}) - d)^+ + b(d - q(\mathbf{x}))^+] | \mathbf{x}\},$$

where h is the per-unit holding cost, b is the per-unit lost-sales penalty cost, $t^+ = \max\{t, 0\}$, and the expectation is taken with respect to the conditional distribution of d given $\mathbf{x}$.

Similarly to Beutel and Minner (2012) and Ban and Rudin (2019), we consider a linear decision function of the form

$$q(\mathbf{x}) = \mathbf{x}^T \boldsymbol{\beta} = \sum_{j=1}^p x_j \beta_j,$$

where the p -dimensional feature vector $\mathbf{x} = (x_1, \dots, x_p)^T$ has $x_1 \equiv 1$, and $\boldsymbol{\beta} = (\beta_1, \dots, \beta_p)^T$ is the coefficient vector, considering the linear decision space is not a restriction in theory or practice. By replacing the features with their transformations (e.g., via series functions), the framework can be adapted to accommodate nonlinear decision rules. More specifically, one may approximate a nonlinear function $\mathbf{x} \mapsto q(\mathbf{x})$ by linear forms $\mathbf{z}(\mathbf{x})^T \boldsymbol{\beta}$, where $\mathbf{x} \mapsto \mathbf{z}(\mathbf{x}) := (z_1(\mathbf{x}), \dots, z_k(\mathbf{x}))^T$ is a vector of approximating functions, and $k = k_n \geq 1$ may increase with n . Then, we denote the transformed features as $\{\mathbf{z}_i = \mathbf{z}(\mathbf{x}_i)\}_{i=1}^n$. Popular choices of the series approximating functions include B-splines (or regression splines), polynomials, Fourier series, and compactly supported wavelets. We refer to Newey (1997) and Chen (2007) for a detailed description of these series functions. Under the linear decision model, we define

the parameter indexing the optimal decision rule as

$$\boldsymbol{\beta}^* = \arg \min_{\boldsymbol{\beta} \in \mathbb{R}^p} C(\boldsymbol{\beta}) := \mathbb{E}\{h(\mathbf{x}^T \boldsymbol{\beta} - d)^+ + b(d - \mathbf{x}^T \boldsymbol{\beta})^+\}, \quad (1)$$

where the expectation is taken with respect to the joint distribution of $(d, \mathbf{x})$. Write $\varepsilon = d - \mathbf{x}^T \boldsymbol{\beta}^*$. Then, the linear decision function (Ban and Rudin 2019) is equivalent to assuming that the conditional $b/(b+h)$ quantile of ε given $\mathbf{x}$ is zero. Unlike Ban and Rudin (2019), we do not assume independence between $\mathbf{x}$ and ε .

3.2. Convolution Smoothing for Empirical Risk Minimization

In the case without features, it is well-known that the optimal decision is given by the $b/(b+h)$ quantile of the demand distribution. It can be estimated by the data-driven sample average approximation (SAA) (Levi et al. 2015), an approach without making any parametric distributional assumption on d . In the setting with features, Ban and Rudin (2019) extended it to the conditional case, proposed a linear programming-based empirical risk minimization algorithm (NV-ERM), and established the connection to conditional quantile regression. More explicitly, one can rewrite $C(q(\mathbf{x}), d) = (b+h)\rho_\tau(d - q(\mathbf{x}))$, where $\tau = b/(b+h)$, and $\rho_\tau(u) = u\{\tau - \mathbb{1}(u < 0)\}$ is referred to as the quantile loss function or the check function corresponding to the τ -th quantile (Koenker and Bassett 1978).

For an arbitrary $\boldsymbol{\beta} \in \mathbb{R}^p$, let $\varepsilon_i(\boldsymbol{\beta}) := d_i - \mathbf{x}_i^T \boldsymbol{\beta}$, $i = 1, \dots, n$. Consider the empirical cumulative distribution function (ECDF) of the $\varepsilon_i(\boldsymbol{\beta})$: $\hat{F}(u; \boldsymbol{\beta}) = (1/n) \sum_{i=1}^n \mathbb{1}\{\varepsilon_i(\boldsymbol{\beta}) \leq u\}$, $u \in \mathbb{R}$. Then, the empirical risk minimization approach of Ban and Rudin (2019) minimizes the following empirical loss function:

$$\hat{C}(\boldsymbol{\beta}) = (b+h) \int_{-\infty}^{\infty} \rho_\tau(u) d\hat{F}(u; \boldsymbol{\beta}), \quad (2)$$

which can be solved via a linear program reformulation. The empirical loss function is nonsmooth and poses significant challenges to developing a privacy protection procedure. To come up with an efficient algorithm to estimate the optimal decision with provable privacy-preserving guarantees, we adopt convolution smoothing to address the challenge associated with the nondifferentiability of the loss function. This aims to simultaneously achieve two goals: (i) to have a feasible algorithm with a privacy-preserving guarantee, and (ii) to have an algorithm with a statistical accuracy guarantee as measured by the accuracy of estimating $\boldsymbol{\beta}^*$ and the regret, which is the cost gap between the estimated policy from the algorithm and the clairvoyant benchmark.

The idea of convolution smoothing originates from Chen and Mangasarian (1995, 1996) in a special case and has been reexamined from a statistical perspective by Fernandes et al. (2021). Specifically, let $\hat{F}_\omega(\cdot; \boldsymbol{\beta})$ be a

smoothed estimator of the distribution function of $\varepsilon_i(\boldsymbol{\beta})$ based on the classical Rosenblatt-Parzen kernel density estimator. That is, for $u \in \mathbb{R}$,

$$\hat{F}_\varpi(u; \boldsymbol{\beta}) = \int_{-\infty}^u \hat{f}_\varpi(t; \boldsymbol{\beta}) dt \quad \text{with} \\ \hat{f}_\varpi(t; \boldsymbol{\beta}) = \frac{1}{n} \sum_{i=1}^n K_\varpi(t - \varepsilon_i(\boldsymbol{\beta})),$$

where $K_\varpi(u) := (1/\varpi)K(u/\varpi)$, $K: \mathbb{R} \rightarrow [0, \infty)$ is a symmetric, nonnegative kernel that integrates to one, and $\varpi = \varpi_n > 0$ is a smoothing parameter. We consider the following smoothed counterpart of $\hat{C}(\boldsymbol{\beta})$:

$$\begin{aligned} \hat{C}_\varpi(\boldsymbol{\beta}) &:= (b+h) \int_{-\infty}^{\infty} \rho_\tau(u) d\hat{F}_\varpi(u; \boldsymbol{\beta}) \\ &= \frac{b+h}{n} \sum_{i=1}^n (\rho_\tau * K_\varpi)(d_i - \mathbf{x}_i^T \boldsymbol{\beta}), \end{aligned} \quad (3)$$

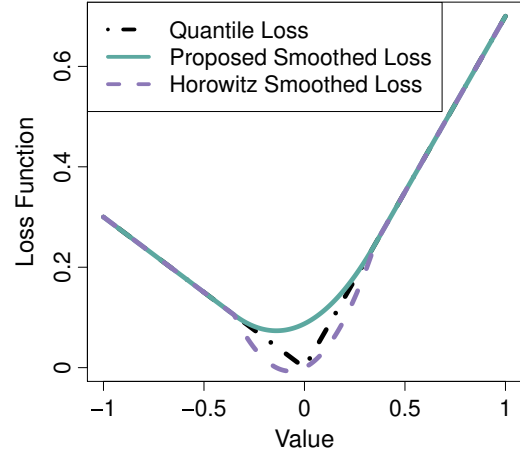
where “ $*$ ” denotes the convolution operator that for any two measurable functions f and g , $(f * g)(u) = \int_{-\infty}^{\infty} f(v)g(u-v)dv$. Therefore, $\hat{C}_\varpi$ is also referred to as the convolution-smoothed loss/cost function. Commonly used kernel functions in optimization and statistics include (i) Gaussian kernel $K(u) = (2\pi)^{-1/2}e^{-u^2/2}$, (ii) Laplacian kernel $K(u) = e^{-|u|}/2$, (iii) logistic kernel $K(u) = e^{-u}/(1+e^{-u})^2$, (iv) uniform kernel $K(u) = (1/2)\mathbb{1}(|u| \leq 1)$, and (v) Epanechnikov kernel $K(u) = (3/4)(1-u^2)\mathbb{1}(|u| \leq 1)$. The following lemma shows that, given any symmetric kernel K and smoothing parameter $\varpi > 0$, the resulting smoothed loss $\rho_\tau * K_\varpi$ provides an upper approximation of ρ_τ with uniform approximation error that scales with ϖ .

Lemma 1. Let K be a symmetric, nonnegative kernel function with $\kappa_1 := \int_{-\infty}^{\infty} |u|K(u)du < \infty$. For any $\varpi > 0$, it holds uniformly over $u \in \mathbb{R}$ that $\rho_\tau(u) \leq (\rho_\tau * K_\varpi)(u) \leq \rho_\tau(u) + \kappa_1\varpi/2$.

From Lemma 1, we see that uniformly over $\boldsymbol{\beta} \in \mathbb{R}^p$, $\hat{C}(\boldsymbol{\beta}) \leq \hat{C}_\varpi(\boldsymbol{\beta}) \leq \hat{C}(\boldsymbol{\beta}) + 0.5\kappa_1(b+h)\varpi$. The analysis of the discrepancy between the respective minimizers necessitates a more nuanced examination and additional assumptions on the data-generating process. In the following subsection, we elucidate the assumptions necessary for the analysis of the privacy-preserving algorithm proposed for feature-based newsvendor problems. We also discuss their connections with the assumptions imposed in Ban and Rudin (2019).

Remark 1. To address the nondifferentiability (at the origin) of the check loss, Horowitz (1998) proposed a more direct approach by replacing the indicator function $\mathbb{1}(u < 0)$ in $\rho_\tau(u)$ with $G(-u/\varpi)$, where $G(\cdot)$ is a smooth, nondecreasing function that takes values between zero and one, and $\varpi > 0$ is a smoothing parameter. However, Horowitz’s smoothing gains

Figure 2. (Color online) Illustration of Two Smoothed Check/Quantile Loss Functions



smoothness at the cost of convexity, which inevitably raises optimization challenges, particularly when dealing with a large number of features is large. On the contrary, provided a nonnegative kernel is used, the convolution-smoothed loss $\rho_\tau * K_\varpi$ remains convex, as its second derivative $(\rho_\tau * K_\varpi)^{(2)}(u) = K_\varpi(u) = K(u/\varpi)/\varpi$ is everywhere nonnegative. See Figure 2 below for a visualization of Horowitz’s and convolution-smoothed check losses.

Using an approximate Hessian matrix, Chen et al. (2019) proposed a Newton-type algorithm to solve Horowitz’s smoothed empirical loss minimization problem. In their convergence analysis, they assumed the initial estimator to be consistent, albeit at a suboptimal rate. In contrast, our analysis imposes minimal assumptions on the initial value, as demonstrated in Theorems 4 and 5 in Section 5.3. To achieve differential privacy in their algorithm, it may be necessary to inject noise into both the gradient and the approximate Hessian of the empirical loss function. However, the theoretical analysis of their algorithm in the presence of inconsistent initial estimates is currently unknown.

Remark 2. The choice of the smoothing technique has an impact on the optimality of the algorithm. Section A of Bassily et al. (2014) considered an example in the setting of hinge loss with the Huberization method (quadratic smoothing) and argued that it does not allow one to get the optimal excess risk bounds. An alternative popular smoothing method for minimizing a nondifferentiable convex objective $\boldsymbol{\beta} \mapsto \hat{C}(\boldsymbol{\beta})$ is through Moreau-Yosida inf-convolution (or envelope) (Nesterov 2005). Recall that $\rho_\tau(u) = \tau u \mathbb{1}(u > 0) - (1 - \tau)u \mathbb{1}(u \leq 0)$. Its Moreau-Yosida envelope is given by

$$\rho_\tau^\gamma(u) = \min_{v \in \mathbb{R}} \left\{ \rho_\tau(v) + \frac{1}{2\gamma}(u-v)^2 \right\},$$

where $\gamma > 0$ is the regularization parameter. It follows

from standard calculations that

$$\rho'_\tau(u) = \begin{cases} -(1-\tau)u - \frac{1}{2}(1-\tau)^2\gamma & \text{if } u < -(1-\tau)\gamma, \\ \frac{1}{2\gamma}u^2, & \text{if } u \in [-(1-\tau)\gamma, \tau\gamma], \\ \tau u - \frac{1}{2}\tau^2\gamma & \text{if } u > \tau\gamma \end{cases} \quad (4)$$

and $(\rho'_\tau)'(u) = \tau \mathbb{1}(u > \tau\gamma) - (1-\tau) \mathbb{1}\{u < -(1-\tau)\gamma\} + u \mathbb{1}\{-(1-\tau)\gamma \leq u \leq \tau\gamma\}/\gamma$. When the convex loss is non-smooth, that is, not everywhere differentiable, Bassily et al. (2019) proposed a variant of the noisy stochastic gradient descent (SGD) algorithm and established upper bounds on the expected excess risk. Applying their general results to the check loss ρ_τ , it follows that the expected excess risk of the output from their (ϵ, δ) -DP algorithm is of order

$$O\left(\sqrt{\frac{p}{n}} + \sqrt{\log(1/\delta)} \frac{p}{\epsilon n}\right).$$

In Section 5.2, by proving a form of restricted strong convexity that holds with high probability, we will show that the expected excess risk of the proposed noisy GD estimator satisfies a faster rate, which is of the order

$$O\left(\log(n) \left\{ \frac{p}{n} + \left(\frac{p + \log n}{\mu n} \right)^2 \right\}\right),$$

where $\mu > 0$ is the privacy parameter for the Gaussian mechanism.

3.3. Assumptions

Suppose we observe an i.i.d. sample $\{(d_i, \mathbf{x}_i)\}_{i=1}^n$ from $(d, \mathbf{x})$ that follows a linear demand model $d = \mathbf{x}^\top \boldsymbol{\beta}^* + \varepsilon$, where the observation noise ε is such that the conditional τ -th quantile of ε given $\mathbf{x}$ is zero, and $\tau = b/(b+h)$. Moreover, we assume that the conditional density function of $\varepsilon|\mathbf{x}$, denoted by $f_{\varepsilon|\mathbf{x}}$, exists and satisfies some regularity conditions described below. Compared with the Linear Model (18) considered in Ban and Rudin (2019), here, we do not assume the independence between the observation noise ε and the random feature $\mathbf{x}$. Under the above model, the τ -th conditional quantile of d given $\mathbf{x}$ is $\mathbf{x}^\top \boldsymbol{\beta}^*$. Specifically, let $F_d(\cdot|\mathbf{x})$ be the conditional distribution function of d given $\mathbf{x}$. Then, the conditional τ -th conditional quantile of d given $\mathbf{x}$ is formally defined as $Q_d(\tau|\mathbf{x}) = \inf\{u : F_d(u|\mathbf{x}) \geq \tau\}$. We write $\mathbf{x} = (x_1, \mathbf{x}_-^\top)^\top$, where $x_1 \equiv 1$ and $\mathbf{x}_- = (x_2, \dots, x_p)^\top$ consists of the remaining random features. For theoretical analysis, we assume without loss of generality that $\mu_j := E(x_j) = 0$ for $j = 2, \dots, p$. Otherwise, it suffices to work with the demeaned model $Q_d(\tau|\mathbf{x}) = \beta_1^* + \sum_{j=2}^p (x_j - \mu_j) \beta_j^*$, where $\beta_1^* = \beta_1^* + \sum_{j=2}^p \mu_j \beta_j^*$.

To facilitate the analysis, we adopt the following technical conditions.

Condition 1 (Kernel Function). Let $K(\cdot)$ be a symmetric, nonnegative, and Lipschitz continuous kernel function; that is, $K(u) = K(-u)$, $K(u) \geq 0$ for all u and $\int_{-\infty}^{\infty} K(u) du = 1$. Moreover, assume $\kappa_u := \sup_{u \in \mathbb{R}} K(u)$ and $\kappa_\ell := \int_{-\infty}^{\infty} |u|^\ell K(u) du$, $\ell = 1, 2$ are bounded.

Condition 2 (Feature Distribution). The random predictor $\mathbf{x}_- \in \mathbb{R}^{p-1}$ is sub-Gaussian: there exists $v_1 \geq 1$ such that $E e^{\lambda \mathbf{u}^\top \mathbf{w}_-} \leq e^{\lambda^2 v_1^2 / 2}$ for all $\lambda \in \mathbb{R}$ and $\mathbf{u} \in \mathbb{S}^{p-2}$, where $\mathbf{w}_- = \mathbf{S}^{-1/2} \mathbf{x}_-$, and $\mathbf{S} = E(\mathbf{x}_- \mathbf{x}_-^\top)$ is positive definite.

Under Condition 2, the $p \times p$ matrix $\Sigma = E(\mathbf{x} \mathbf{x}^\top)$ is also positive definite. Let $\mathbf{w} = \Sigma^{-1/2} \mathbf{x} = (1, \mathbf{w}_-^\top)^\top$ denote the standardized feature vector satisfying $E(\mathbf{w} \mathbf{w}^\top) = \mathbf{I}_p$ and $E(\mathbf{w}) = (1, \mathbf{0}_{p-1}^\top)^\top$.

Condition 3 (Observational Noise $\varepsilon = d - \mathbf{x}^\top \boldsymbol{\beta}^*$). There exist constants $l_0 > 0$ and $f_u \geq f_l > 0$ such that $|f_{\varepsilon|\mathbf{x}}(u) - f_{\varepsilon|\mathbf{x}}(v)| \leq l_0 |u - v|$, $f_{\varepsilon|\mathbf{x}}(u) \leq f_u$ for all $u, v \in \mathbb{R}$ almost surely (over $\mathbf{x}$), and

$$\inf_{t \in [0, 1], \mathbf{v} \in \mathbb{S}^{p-1}} E\{f_{\varepsilon|\mathbf{x}}(t\langle \mathbf{w}, \mathbf{v} \rangle) \langle \mathbf{w}, \mathbf{v} \rangle^2\} \geq f_l. \quad (5)$$

Below, we discuss how Conditions 2 and 3 are comparable to assumptions 1 and 2 in Ban and Rudin (2019). For the random feature vector $\mathbf{x} = (1, \mathbf{x}_-^\top)^\top \in \mathbb{R}^p$, Ban and Rudin (2019) assumed that the coordinates of $\mathbf{x}_-$ are normalized with mean zero and standard deviation one, and $\|\mathbf{x}\|_2 \leq C\sqrt{p}$ for some constant $C > 0$. Condition 2, on the other hand, requires feature vector $\mathbf{x}$ to be sub-Gaussian, extending the concept of sub-Gaussian random variables to higher dimensions through one-dimensional marginals. Examples of such sub-Gaussian vectors include (i) Gaussian and Bernoulli random vectors, (ii) spherical random vectors, (iii) random vectors uniformly distributed on the Euclidean ball centered at the origin with radius $\sqrt{p}$, and (iv) random vectors uniformly distributed on the unit cube $[-1, 1]^p$. We refer to chapter 3.4 in Vershynin (2018) for detailed discussions on multivariate sub-Gaussian distributions. For sub-Gaussian feature vectors $\mathbf{x}_i$'s, from theorem 2.1 in Hsu et al. (2012), it follows that $\max_{1 \leq i \leq n} \|\mathbf{x}_i\|_2 \lesssim \sqrt{p + \log n}$ with high probability. For the observation noise ε , Ban and Rudin (2019) assumed that its density function, denoted by f_ε , is bounded away from zero on some compact interval $[\bar{D}, \bar{D}]$. In Condition 3, it is worth noticing that the constant f_l may depend on both the conditional distribution of ε given $\mathbf{x}$, and the distribution of $\mathbf{x} \in \mathbb{R}^p$. To see this, define $\iota_\delta = \inf\{\iota > 0 : E\langle \mathbf{w}, \mathbf{u} \rangle^2 \mathbb{1}_{\{|\langle \mathbf{w}, \mathbf{u} \rangle| > \iota\}} \leq \delta \text{ for all } \mathbf{u} \in \mathbb{S}^{p-1}\}$ for $\delta \in (0, 1]$. It is easy to see that $\delta \mapsto \iota_\delta$ is nondecreasing, and $\iota_\delta \leq (m_q/\delta)^{1/(q-2)}$ for any $q > 2$, where $m_q := \sup_{\mathbf{u} \in \mathbb{S}^{p-1}} E|\langle \mathbf{w}, \mathbf{u} \rangle|^q$. Then, a sufficient condition for (5) is $\min_{|t| \leq \iota_\delta} f_{\varepsilon|\mathbf{x}}(t) \geq c_\delta > 0$ almost surely (over $\mathbf{x}$) for

some $0 < \delta < 1$ because

$$\inf_{t \in [0, 1], \mathbf{v} \in \mathbb{S}^{p-1}} \mathbb{E}\{f_{\varepsilon}(\mathbf{x}(t\langle \mathbf{w}, \mathbf{v} \rangle)\langle \mathbf{w}, \mathbf{v} \rangle^2) \geq c_{\delta}\}$$

$$\inf_{\mathbf{v} \in \mathbb{S}^{p-1}} \mathbb{E}\{\langle \mathbf{w}, \mathbf{v} \rangle^2 \mathbb{1}(|\langle \mathbf{w}, \mathbf{v} \rangle| \leq \iota_{\delta})\} \geq (1 - \delta)c_{\delta}.$$

4. A Privacy-Preserving Algorithm for the Feature-Based Newsvendor Problem

4.1. Preliminaries on f -Differential Privacy

We will first review the (ϵ, δ) -differential privacy concept introduced in Dwork et al. (2006a, b). Let $\mathcal{S}$ denote a data set consisting of observations $\{\mathbf{z}_1, \dots, \mathbf{z}_n\}$, where $\mathbf{z}_i = (d_i, \mathbf{x}_i)$, $i = 1, \dots, n$. A pair of data sets $\mathcal{S}$ and $\mathcal{S}'$ are said to be neighboring data sets if they differ in only one data point.

Definition 1. ((ϵ, δ) -Differential Privacy, or (ϵ, δ) -DP). A randomized algorithm M is (ϵ, δ) -differentially private if, for any neighboring data sets $\mathcal{S}$ and $\mathcal{S}'$ and any event $\mathcal{E}$, we have

$$P(M(\mathcal{S}) \in \mathcal{E}) \leq e^{\epsilon} P(M(\mathcal{S}') \in \mathcal{E}) + \delta,$$

where $\epsilon \geq 0$ and $0 \leq \delta \leq 1$ are constants.

An intuitive way to understand the concept of (ϵ, δ) -differential privacy is via the lens of a hypothesis testing problem for distinguishing two neighboring data sets $\mathcal{S}$ and $\mathcal{S}'$:

H_0 : the underlying data set is $\mathcal{S}$ versus

H_1 : the underlying data set is $\mathcal{S}'$.

Consider any given test procedure ϕ based on the output of the randomized algorithm M , and denote its type I error and type II error by α_{ϕ} and β_{ϕ} , respectively. It can be shown that for a test ϕ based on the output of any (ϵ, δ) -differentially private algorithm, its power is bounded by $\min\{e^{\epsilon}\alpha_{\phi} + \delta, 1 - e^{-\epsilon}(1 - \alpha_{\phi} - \delta)\}$. If ϵ and δ are both small, then any α -level ($0 < \alpha < 1$) test will be nearly powerless.

Dong et al. (2022) extended Dwork et al. (2006a, b) by introducing the trade-off function to characterize the trade-off between the type I and type II errors.

Definition 2 (Trade-Off Function (Dong et al. 2022)). For any two probability distributions P and Q , the trade-off function $T(P, Q) : [0, 1] \rightarrow [0, 1]$ is defined as $T(P, Q)(\alpha) = \inf\{\beta_{\phi} : \alpha_{\phi} \leq \alpha\}$, where the infimum is taken over all measurable rejection rules to distinguish P and Q .

The greater the trade-off function is, the harder it is to distinguish the two distributions via hypothesis testing. Dong et al. (2022) showed that a function $f : [0, 1] \rightarrow [0, 1]$ is a trade-off function if and only if f is convex, continuous, nonincreasing, and $f(x) \leq 1 - x$ for all $x \in [0, 1]$. In the following, for any two functions f and g defined on $[0, 1]$, we write $g \geq f$ if $g(x) \geq f(x)$,

$\forall x \in [0, 1]$. We abuse the notation a little by identifying $\mathcal{S}$ and $\mathcal{S}'$ with their respective probability distributions.

Definition 3 (f -Differential Privacy (Dong et al. 2022)). A randomized algorithm M is said to be f -differentially private if, for any neighboring data sets $\mathcal{S}$ and $\mathcal{S}'$,

$$T(\mathcal{S}, \mathcal{S}') \geq f$$

for some trade-off function f .

If $f = T(P, Q)$ for some distributions P and Q , then a mechanism M is f -DP if distinguishing any two neighboring data sets based on the output of M is at least as difficult as distinguishing P and Q based on a single draw. This functional perspective avoids some of the pitfalls associated with (ϵ, δ) -differential privacy. f -DP is a generalization of (ϵ, δ) -DP. A result of Wasserman and Zhou (2010) indicates that a mechanism M is (ϵ, δ) -DP if and only if M is $f_{\epsilon, \delta}$ -DP, where $f_{\epsilon, \delta}(\alpha) = \max\{0, 1 - e^{\epsilon}\alpha_{\phi} - \delta, e^{-\epsilon}(1 - \alpha_{\phi} - \delta)\}$.

Definition 4 (Gaussian Differential Privacy (Dong et al. 2022)). A randomized algorithm M is said to satisfy μ -Gaussian differential privacy if for any neighboring data sets $\mathcal{S}$ and $\mathcal{S}'$

$$T(\mathcal{S}, \mathcal{S}') \geq G_{\mu},$$

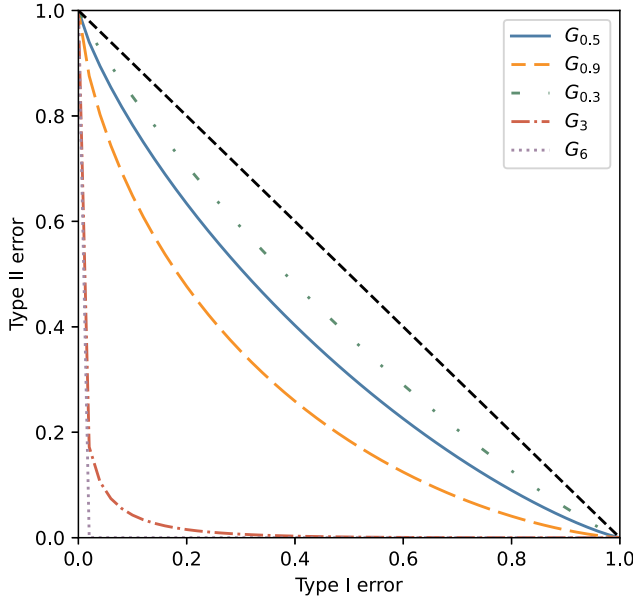
where $G_{\mu}(\alpha) = \Phi(\Phi^{-1}(1 - \alpha) - \mu)$, and Φ is the distribution function of the standard normal distribution.

GDP provides a parametric family of f -DP that guarantees and enjoys many desirable properties. As a rule of thumb, $\mu \leq 0.5$ guarantees a reasonable amount of privacy, $\mu = 1$ is borderline private, and $\mu > 6$ promises almost no privacy guarantee. The trade-off of type I and type II errors is illustrated in Figure 3.

4.2. Proposed Differentially Private Algorithm

To obtain a differentially private counterpart of $\hat{\beta}_{\varpi} = \arg \min_{\beta} \hat{C}_{\varpi}(\beta)$, we draw inspiration from works Song et al. (2013) and Bassily et al. (2014) and utilize a noisy optimization approach that involves adding Gaussian noise during each iteration of the gradient descent method. Minimizing the loss function $\hat{C}(\beta)$ used in Ban and Rudin (2019) poses a challenge, as it is convex but not differentiable everywhere, and subgradient methods commonly employed for such cases exhibit slow (sub-linear) convergence, leading to computational instability. In this work, we propose a differentially private algorithm that leverages convolution smoothing and noisy gradient descent. Its privacy protection guarantee is established in Theorem 1 in Section 4.3. We further investigate its finite-sample performance in Section 5 and demonstrate that the proposed approach achieves a balanced trade-off between statistical accuracy and computational stability, owing to the effectiveness of convolution smoothing.

Figure 3. (Color online) Trade-Off Functions for GDP with Privacy Parameter $\mu = 0.3, 0.5, 0.9, 3$, and 6 , Respectively



Given a nonnegative kernel function $K(\cdot)$ introduced earlier, we define

$$\bar{K}(u) = \int_{-\infty}^u K(v)dv \quad \text{and} \quad \bar{K}_\varpi(u) = \bar{K}(u/\varpi)$$

so that $\bar{K}'_\varpi(u) = K_\varpi(u)$. If the $p \times p$ matrix $\Sigma = E(\mathbf{x}\mathbf{x}^T)$ was known, we consider the following noisy smoothed gradient descent method, starting at iteration 0 with an initial estimate $\boldsymbol{\beta}^{(0)}$. At iteration $t = 0, 1, 2, \dots, T-1$, we update the solution as follows:

$$\boldsymbol{\beta}^{(t+1)} = \boldsymbol{\beta}^{(t)} - \frac{\eta_0}{n} \Sigma^{-1/2} \left[\sum_{i=1}^n \{ \bar{K}_\varpi(\mathbf{x}_i^T \boldsymbol{\beta}^{(t)} - d_i) - \tau \} w_B(\mathbf{w}_i) + \sigma \mathbf{g}_t \right], \quad (6)$$

where $\mathbf{w}_i = \Sigma^{-1/2} \mathbf{x}_i$ denotes the standardized covariates vector in the sense that $E(\mathbf{w}_i \mathbf{w}_i^T) = \mathbf{I}_p$, $w_B(\mathbf{u}) = \mathbf{u} / \max\{1, \|\mathbf{u}\|_2/B\}$, and $\mathbf{g}_t \in \mathbb{R}^p$ ($t = 0, 1, \dots, T-1$) are independent $\mathcal{N}(\mathbf{0}, \mathbf{I}_p)$ vectors. Here, $\eta_0 > 0$ is the step size, $T \geq 1$ is a prespecified number of iterations, $B \geq 1$ is a truncation parameter, and σ is a positive constant adjusting the level of noise injected into the gradient perturbation. We summarize this procedure in Algorithm 1.

The presence of Σ in Algorithm 1 is primarily for theoretical convenience, as it allows us to establish upper bounds on the estimation error $E\{\mathbf{x}^T(\hat{\boldsymbol{\beta}}_\varpi - \boldsymbol{\beta}^*)\}^2$ by expressing it as $\|\Sigma^{1/2}(\hat{\boldsymbol{\beta}}_\varpi - \boldsymbol{\beta}^*)\|_2^2$, where the expectation is only taken with respect to $\mathbf{x}$ that is independent of $\hat{\boldsymbol{\beta}}_\varpi$. When Σ is unknown, we consider the following update:

$$\boldsymbol{\beta}^{(t+1)} = \boldsymbol{\beta}^{(t)} - \frac{\eta_0}{n} \left[\sum_{i=1}^n \{ \bar{K}_\varpi(\mathbf{x}_i^T \boldsymbol{\beta}^{(t)} - d_i) - \tau \} w_B(\mathbf{x}_i) + \sigma \mathbf{g}_t \right] \quad (7)$$

with a slight abuse of notation.

In both (6) and (7), the use of covariate clipping/truncation guarantees a bounded ℓ_2 -sensitivity of the gradient function, which is the key to achieving differential privacy. For the initial value, one can take $\boldsymbol{\beta}^{(0)}$ to be either $\mathbf{0}$ or a random guess that is uniformly distributed over the unit sphere.

Algorithm 1 (Private ERM via Noisy Smoothed Gradient Descent)

Input: data set $\{(d_i, \mathbf{x}_i)\}_{i=1}^n$, probability level $\tau \in (0, 1)$, bandwidth $\varpi > 0$, initial value $\boldsymbol{\beta}^{(0)}$, step size $\eta_0 > 0$, noise scale $\sigma > 0$, truncation level $B \geq 1$, number of iterations $T \geq 1$.

- 1: **for** $t = 0, 1, \dots, T-1$ **do**
- 2: Generate standard multivariate normal vector $\mathbf{g}_t \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_p)$;
- 3: Compute clipped/truncated covariates $\bar{\mathbf{w}}_i = \mathbf{w}_i \min\{1, B/\|\mathbf{w}_i\|_2\}$ for $i = 1, \dots, n$;
- 4: Compute $\boldsymbol{\beta}^{(t+1)} = \boldsymbol{\beta}^{(t)} - (\eta_0/n) \cdot \Sigma^{-1/2} [\sum_{i=1}^n \{ \bar{K}_\varpi(\mathbf{x}_i^T \boldsymbol{\beta}^{(t)} - d_i) - \tau \} \bar{\mathbf{w}}_i + \sigma \mathbf{g}_t]$;

5: **end for**

Output: $\boldsymbol{\beta}^{(T)}$.

4.3. Privacy-Protection Guarantee of the Proposed Algorithm

We first review some useful results about f -DP.

Proposition 1 (Theorem 1 in Dong et al. (2022)). Define the Gaussian mechanism that operates on a statistic θ as $M(S) = \theta(S) + \xi$, where $\xi \sim \mathcal{N}(0, \text{sens}(\theta)^2/\mu^2)$, and the sensitivity $\text{sens}(\theta) = \sup_{S, S'} |\theta(S) - \theta(S')|$ with the supremum over all neighborhood data sets S and S' . Then, M is μ -GDP.

A nice property of f -DP is that the composition of private mechanisms is closed and tight in the f -DP framework. A two-step composition can be written as $M(S) = (y_1, M_2(S, y_1))$, where $y_1 = M_1(S)$ with $M_1: X \rightarrow Y_1$ being the first mechanism, and $M_2: X \times Y_1 \rightarrow Y_2$ is the second mechanism. In general, given a sequence of mechanisms $M_i: X \times Y_1 \times \dots \times Y_{i-1} \rightarrow Y_i$, $i = 1, \dots, n$, we consider the n -fold composed mechanism: $M: X \rightarrow Y_1 \times \dots \times Y_n$.

Definition 5. The tensor product of two trade-off functions $f = T(P, Q)$ and $g = T(P', Q')$ is $f \otimes g = T(P \times P', Q \times Q')$.

Proposition 2 (Theorem 4 in Dong et al. (2022)). Let $M_i(\cdot, y_1, \dots, y_{i-1})$ be f_i -DP for all $y_1 \in Y_1, \dots, y_{i-1} \in Y_{i-1}$. Then, the n -fold composed mechanism $M: X \rightarrow Y_1 \times \dots \times Y_n$ is $f_1 \otimes \dots \otimes f_n$ -DP.

Corollary 1 (Corollary 2 in Dong et al. (2022)). The n -fold composition of μ_i -GDP mechanisms is $\sqrt{\mu_1^2 + \dots + \mu_n^2}$ -GDP.

Remark 3. Dong et al. (2022) revealed that the Gaussian mechanism in Proposition 1 satisfies $G_\mu(\alpha) = \inf_{\text{neighboring } S, S'} T(M(S), M(S'))(\alpha)$. For all possible type I error rate α , the infimum could be achieved at two neighboring data sets that satisfy $|\theta(S) - \theta(S')| = \text{sens}(\theta)$. This implies that the characterization by GDP is precise in the point-wise sense. GDP offers the tightest possible privacy bound of the Gaussian mechanism.

Theorem 1 below establishes the Gaussian differential privacy property of the proposed algorithm. For any given sample size n , number of iterations T , and noise scale σ , the output of the algorithm $\beta^{(T)}$ satisfies the Gaussian differential privacy property with privacy level μ , as outlined in the theorem.

Theorem 1 (Privacy Protection Guarantees). *Given an initial estimate $\beta^{(0)} \in \mathbb{R}^p$ and a data set $Z_n = \{(d_i, \mathbf{x}_i)\}_{i=1}^n$, consider the noisy gradient descent iterates $\{\beta^{(t)}\}_{t=0, \dots, T}$ defined in (6). Given a privacy level $\mu > 0$, if $\sigma \geq 2\bar{\tau}BT^{1/2}/\mu$ with $\bar{\tau} = \max(\tau, 1 - \tau)$, then the final output $\beta^{(T)}$ is μ -GDP.*

Proof of Theorem 1. Consider two data sets Z_n and Z'_n that differ by one datum, say $(d_1, \mathbf{x}_1) \in Z_n$ versus $(d'_1, \mathbf{x}'_1) \in Z'_n$. At the first iteration, note that

$$\begin{aligned} & \|\Sigma^{1/2}\beta^{(1)}(Z_n) - \Sigma^{1/2}\beta^{(1)}(Z'_n)\|_2 \\ &= \frac{\eta_0}{n} \left\| \{\bar{K}_\varpi(\langle \mathbf{x}_1, \beta^{(t)} \rangle - d_1) - \tau\} w_B(\mathbf{w}_1) \right. \\ & \quad \left. - \{\bar{K}_\varpi(\langle \mathbf{x}'_1, \beta^{(t)} \rangle - d'_1) - \tau\} w_B(\mathbf{w}'_1) \right\|_2 \\ &\leq 2 \max(1 - \tau, \tau) B \frac{\eta_0}{n} = 2\bar{\tau}B \frac{\eta_0}{n}. \end{aligned}$$

By Proposition 1, $\Sigma^{1/2}\beta^{(1)}$ is $(T^{-1/2}\mu)$ -GDP as long as $\sigma \geq 2\bar{\tau}BT^{1/2}/\mu$. After postprocessing (by a deterministic map), $\beta^{(1)}$ is also $(T^{-1/2}\mu)$ -GDP.

By definition, the second iterate $\beta^{(2)} = \beta^{(2)}(Z_n)$ takes $\beta^{(1)}$ as input in addition to the data set. Together, Proposition 2 and Corollary 1 show that the twofold composed (joint) mechanism $(\beta^{(1)}, \beta^{(2)})$ is $\sqrt{\mu^2/T + \mu^2/T}$ -GDP. Using the same argument repeatedly, we conclude that the T -fold composed mechanism $(\beta^{(1)}, \dots, \beta^{(T)})$ is μ -GDP, and hence, so is $\beta^{(T)}$. $\square$

5. Theoretical Performance

Theorem 1 in Section 4.3 establishes the privacy protection guarantees of the proposed new algorithm. This section provides a statistical analysis of the privacy-preserving coefficient estimate $\beta^{(T)}$ under Conditions 1–3 from Section 3.3. Section 5.1 provides upper bounds for the finite-sample bias of the estimated optimal policy both in high probability and under expectation. Section 5.2 provides the regret analysis. To prove these bounds, the main technical challenge is that the empirical cost function after convolution smoothing does not

satisfy the local strong convexity as required in Avella-Medina et al. (2023). We provide main proof strategies and important intermediate results in Section 5.3.

5.1. Performance Bound on Estimation Error

The parameter indexing the clairvoyant optimal policy, where the demand distribution is known a priori, is given by (1): $\beta^* = \arg \min_{\beta: q(\mathbf{x}) = \mathbf{x}^T \beta} \mathbb{E}\{C(q(\mathbf{x}), d) | \mathbf{x}\}$. Given a feature vector $\mathbf{x}$, the clairvoyant optimal to-order-quantity is $\mathbf{x}^T \beta^*$. Given an estimate $\hat{\beta}$ of the unknown parameter β^* based on a random, independent sample $\{(d_1, \mathbf{x}_1), \dots, (d_n, \mathbf{x}_n)\}$ of size n drawn from the linear demand model, we use the following two metrics to evaluate its performance: (i) the estimation error $\|\hat{\beta} - \beta^*\|_2$ under the vector ℓ_2 -norm or its variant, and (ii) the excess (population) risk $C(\hat{\beta}) - C(\beta^*)$, where $C(\cdot)$ is defined in (1).

With a predetermined bandwidth $\varpi > 0$, step size $\eta_0 > 0$, truncation level $B \geq 1$, number of iterations $T \geq 1$, and noise scale $\sigma = 2\bar{\tau}BT^{1/2}/\mu$, let $\beta^{(T)}$ be the μ -GDP estimator of β^* obtained from Algorithm 1, where $\bar{\tau} = \max(\tau, 1 - \tau)$. Our first theorem provides upper bounds for the estimation error $\|\beta^{(T)} - \beta^*\|_\Sigma$ both in high probability and under expectation, where $\|\cdot\|_\Sigma$ denotes the Σ -induced norm; that is, $\|\mathbf{u}\|_\Sigma = \sqrt{\mathbf{u}^T \Sigma \mathbf{u}}$ for $\mathbf{u} \in \mathbb{R}^p$.

Theorem 2. *In addition to Conditions 1–3, assume $\kappa_l := \min_{|u| \leq 1} K(u) > 0$. Let the triplet of parameters (ϖ, B, T) and sample size n satisfy*

$$\begin{aligned} \varpi &\asymp \left(\frac{p + \log n}{n} \right)^{1/4}, \quad B \asymp \sqrt{p + \log n}, \\ T &\asymp \log n \quad \text{and} \quad n \gtrsim T^{1/2} \frac{p + \log n}{\mu f_l}. \end{aligned} \quad (8)$$

Moreover, let the step size η_0 satisfy $0 < \eta_0 \leq 1/\max(2f_u, f_l + \bar{\tau})$. Then, the μ -GDP estimated coefficient $\beta^{(T)}$ obtained from noisy gradient descent initialized at any $\beta^{(0)} \in \beta^* + \Theta_\Sigma(1)$ satisfies

$$\|\beta^{(T)} - \beta^*\|_\Sigma \leq C_0 \left(\eta_0 T^{1/2} \frac{p + \log n}{\mu n} + \frac{1}{f_l} \sqrt{\frac{p \log n}{n}} \right)$$

with probability at least $1 - \frac{C_1}{n^2}$

and

$$\mathbb{E}\|\beta^{(T)} - \beta^*\|_\Sigma \leq C_2 \sqrt{\log n} \left(\eta_0 \frac{p + \log n}{\mu n} + \frac{1}{f_l} \sqrt{\frac{p}{n}} \right),$$

where C_0, C_1 , and C_2 are positive constants independent of (n, p) .

As a benchmark, we use $\hat{\beta}_\varpi$ to denote the nonprivate empirical (smoothed) risk minimizer; that is, $\hat{\beta}_\varpi = \arg \min_{\beta \in \mathbb{R}^p} \hat{C}_\varpi(\beta)$, with $\hat{C}_\varpi(\cdot)$ defined in (3). From the proof of Theorem 2, we see that the second term on the

right-hand side upper bound is the bound on the estimation error of β_ϖ , which is of order $\sqrt{p/n}$ with a properly chosen smoothing parameter ϖ (He et al. 2023). The first term quantifies the “cost of privacy” of the noisy gradient descent algorithm for solving the feature-based newsvendor problem. For sufficiently small values of μ , the obtained upper bound (on $\|\cdot\|_\Sigma$ -error) matches the minimax lower bound, up to logarithmic factors, for (μ, δ) -DP estimation of β^* under a linear model with normal errors; see theorem 4.1 in Cai et al. (2021). By corollary 1 in Dong et al. (2022), an algorithm is μ -GDP if and only if $(\mu, \delta(\mu))$ -DP, where $\delta(\mu) = \Phi(-1 + \mu/2) - e^\mu \Phi(-1 - \mu/2)$.

5.2. Regret Analysis

We next provide a finite-sample analysis of the regret $C(\beta^{(T)}) - C(\beta^*)$, where $C(\beta) = \mathbb{E}\{\hat{C}(\beta)\}$ is as in (1). The regret is the difference between the expected cost obtained with the estimated privacy-preserving optimal policy and the clairvoyant optimal expected cost with known demand distribution but without privacy protection.

Without loss of generality (up to a constant scale), we consider the regret of $\beta^{(T)}$, defined as $Q(\beta^{(T)}) - Q(\beta^*) = Q(\beta^{(T)}) - \inf_{\beta \in \mathbb{R}^p} Q(\beta)$, where $Q(\beta) = (b+h)^{-1} C(\beta) = \mathbb{E}\{\rho_\tau(d - \mathbf{x}^T \beta)\}$ with $\tau = b/(b+h)$ as the population cost (without smoothing). Recall that the coefficient β^* indexing the optimal inventory policy satisfies the first-order condition $\nabla Q(\beta^*) = \mathbf{0}$. By the mean value theorem and Condition 3 that $\sup_{u \in \mathbb{R}} f_{\varepsilon|\mathbf{x}}(u) \leq f_u$, it can be shown that $Q(\beta) - Q(\beta^*) = Q(\beta) - Q(\beta^*) - \langle \nabla Q(\beta^*), \beta - \beta^* \rangle \leq 0.5f_u \|\beta - \beta^*\|_\Sigma^2$ for any $\beta \in \mathbb{R}^p$. This, combined with Theorem 2, implies the following high-probability upper bound on the regret along with an expected regret bound.

Theorem 3 (High-Probability Bound for Excess Risk). *Under the same set of assumptions in Theorem 2, we have*

$$Q(\beta^{(T)}) - Q(\beta^*) \lesssim \log(n) \left\{ \frac{1}{f_u} \left(\frac{p + \log n}{\mu n} \right)^2 + \frac{f_u p}{f_l^2 n} \right\} \quad (9)$$

with probability at least $1 - C_1 n^{-2}$.

Corollary 2 (Excess Population Risk Bound). *Under the conditions of Theorem 3, the excess population risk bound satisfies*

$$\mathbb{E}\{Q(\beta^{(T)}) - Q(\beta^*)\} \lesssim \log(n) \{f_u^{-1}(p + \log n)^2 / (\mu n)^2 + (f_u / f_l^2) \cdot p/n\}. \quad (10)$$

For fixed cost parameters b and h , the Bound (9) indicates that if the number of relevant features p is small relative to the number of observations in the sense that $p \log n = o(n)$, the expected cost of the μ -GDP estimated decision converges to that of the optimal decision at a fast rate $O(p/n + \mu^{-2}(p/n)^2)$, up to logarithmic factors.

Without privacy guarantees, Ban and Rudin (2019) obtained a similar performance bound, which implies consistency but under a stronger condition on the number of features; that is, $p^2 = o(n)$.

5.3. Proof Strategy and Key Intermediate Results

The statistical analysis of the noisy gradient descent iterates $\{\beta^{(t)}\}_{t=1, \dots, T}$ depends crucially on the landscape of the empirical loss function, as highlighted in recent research (Cai et al. 2021, Avella-Medina et al. 2023). Unlike many commonly used loss functions in statistical learning, such as squared loss, Huber loss (and its smoothed variants), and logistic loss, the quantile loss ρ_τ exhibits piecewise linearity, which lacks local strong convexity. Instead, its “curvature energy” is concentrated in a single point. Notably, the local strong convexity and smoothness of $\hat{Q}_\varpi(\cdot)$ are intricately influenced by the bandwidth used in the analysis.

In this study, we present the key findings for analyzing the landscape of $\hat{Q}_\varpi(\cdot)$. Specifically, we illustrate that by conditioning on a series of “good events” related to the empirical smoothed cost function, the proposed noisy gradient descent iterates exhibit favorable convergence properties. Moreover, we demonstrate that these good events will occur with high probability under Conditions 1–3. We conclude this section with a supplementary analysis of the initialization of the algorithm.

Given a kernel function $K(\cdot)$ and bandwidth $\varpi > 0$, we define the empirical smoothed loss

$$\hat{Q}_\varpi(\beta) = (b+h)^{-1} \hat{C}_\varpi(\beta) = \frac{1}{n} \sum_{i=1}^n \underbrace{(\rho_\tau * K_\varpi)}_{=: \ell_\varpi}(d_i - \mathbf{x}_i^T \beta),$$

where “ $*$ ” is the convolution operator. Its gradient and Hessian are given, respectively, by

$$\nabla \hat{Q}_\varpi(\beta) = \frac{1}{n} \sum_{i=1}^n \{\bar{K}_\varpi(\mathbf{x}_i^T \beta - d_i) - \tau\} \mathbf{x}_i \quad \text{and}$$

$$\nabla^2 \hat{Q}_\varpi(\beta) = \frac{1}{n} \sum_{i=1}^n K_\varpi(d_i - \mathbf{x}_i^T \beta) \mathbf{x}_i \mathbf{x}_i^T.$$

Let $Q(\beta) = \mathbb{E}\{\hat{Q}(\beta)\}$ and $Q_\varpi(\beta) = \mathbb{E}\{\hat{Q}_\varpi(\beta)\}$ be the population quantile and smoothed quantile losses, respectively. Note that although the parameter β^* indexing the theoretically optimal inventory policy satisfies the moment condition $\nabla Q(\beta^*) = \mathbf{0}$, in general, $\nabla Q_\varpi(\beta^*) \neq \mathbf{0}$. Therefore, we use

$$b^* := \|\Sigma^{-1/2} \nabla Q_\varpi(\beta^*)\|_2 = \|\mathbb{E}\{\bar{K}_\varpi(\mathbf{x}^T \beta^* - d) - \tau\} \mathbf{w}\|_2$$

to quantify the smoothing bias. Together, Condition 1 and the Lipschitz continuity of $f_{\varepsilon|\mathbf{x}}(\cdot)$ ensure that $b^* \leq 0.5l_0\kappa_2\varpi^2$; see Lemma 1.3 in the supplementary material.

For any $r > 0$, define the local ellipses centered at the origin and β^* , respectively, as

$$\Theta_\Sigma(r) = \{\delta \in \mathbb{R}^p : \|\delta\|_\Sigma \leq r\} \quad \text{and} \\ \Theta_\Sigma^*(r) = \{\beta \in \mathbb{R}^p : \|\beta - \beta^*\|_\Sigma \leq r\}.$$

Moreover, for every $\beta \in \mathbb{R}^p$, we write

$$\widehat{D}_\varpi(\delta) = \widehat{Q}_\varpi(\beta) - \widehat{Q}_\varpi(\beta^*) \quad \text{and} \\ D_\varpi(\delta) = \mathbb{E}\{\widehat{D}_\varpi(\delta)\} \quad \text{for } \delta = \beta - \beta^*. \quad (11)$$

Given parameters $B, R \geq 1$ and $\delta_0, \delta_1 > 0$, define the “good events”

$$\mathcal{E}_0(B) = \left\{ \max_{1 \leq i \leq n} \|\mathbf{w}_i\|_2 \leq B \right\}, \quad (12)$$

$$\mathcal{E}_1(\delta_0, \delta_1) = \{|\widehat{D}_\varpi(\delta) - D_\varpi(\delta)| \leq \delta_0 \|\beta\|_\Sigma \\ \text{for all } \delta \in \Theta_\Sigma(1) \setminus \Theta_\Sigma(1/n)\} \cap \{|\nabla \widehat{Q}_\varpi(\beta) \\ - \nabla Q_\varpi(\beta)|\|_{\Sigma^{-1}} \leq \delta_1 \text{ for all } \beta \in \Theta_\Sigma^*(1)\}, \quad (13)$$

$$\mathcal{E}_2(R) = \{|\nabla^2 \widehat{Q}_\varpi(\beta) - \nabla^2 Q_\varpi(\beta)|\|_{\Sigma^{-1}} \leq f_u \\ \text{for all } \beta \in \Theta_\Sigma^*(R)\}. \quad (14)$$

Here we write $\|\mathbf{A}\|_{\Sigma^{-1}} = \|\Sigma^{-1/2} \mathbf{A} \Sigma^{-1/2}\|_2$ for any $p \times p$ matrix $\mathbf{A}$. In the following, we will restrict our analysis to the intersection of the above events.

Proposition 3 (Restricted Strong Convexity and Smoothness). *Let $0 < \varpi \leq f_l/(2l_0\kappa_1)$, and set $\phi_1 = 0.5(f_l - l_0\kappa_1\varpi) \geq 0.25f_l > 0$. Then, conditioned on the event $\mathcal{E}_1(\delta_0, \delta_1) \cap \mathcal{E}_2(R)$, we have*

$$\widehat{Q}_\varpi(\beta) - \widehat{Q}_\varpi(\beta^*) \\ \geq \begin{cases} \phi_1 \|\beta - \beta^*\|_\Sigma^2 & \text{for all } \beta \in \Theta_\Sigma^*(1) \setminus \Theta_\Sigma^*(1/n) \\ -(\delta_0 + b^*) \|\beta - \beta^*\|_\Sigma & \\ (\phi_1 - \delta_0 - b^*) \|\beta - \beta^*\|_\Sigma & \text{for all } \beta \in \Theta_\Sigma^*(1)^c, \end{cases} \quad (15)$$

$$\widehat{Q}_\varpi(\beta^*) - \widehat{Q}_\varpi(\beta) - \langle \nabla \widehat{Q}_\varpi(\beta), \beta^* - \beta \rangle \\ \geq \phi_1 \|\beta - \beta^*\|_\Sigma^2 - (\delta_0 + \delta_1) \|\beta - \beta^*\|_\Sigma \\ \text{for all } \beta \in \Theta_\Sigma^*(1) \setminus \Theta_\Sigma^*(1/n), \quad (16)$$

and

$$\widehat{Q}_\varpi(\beta_2) - \widehat{Q}_\varpi(\beta_1) - \langle \nabla \widehat{Q}_\varpi(\beta_1), \beta_2 - \beta_1 \rangle \\ \leq f_u \|\beta_2 - \beta_1\|_\Sigma^2 \quad \text{for all } \beta_1, \beta_2 \in \Theta_\Sigma^*(R). \quad (17)$$

Note that the Lower Bound (16) implies a restricted strong convexity (RSC) for $\widehat{Q}_\varpi(\beta)$ when $\beta \in \Theta_\Sigma^*(1)/\Theta_\Sigma^*(n^{-1} \vee r_1)$ with $r_1 = (\delta_0 + \delta_1)/\phi_1$. The Upper Bound (17) is related to the local strong smoothness of the empirical cost, which no longer holds without convolution smoothing. In addition, we define a “good” event on which the smoothed empirical loss $\widehat{Q}_\varpi(\cdot)$ satisfies

a refined RSC property. Given a radius $r > 0$ and a curvature parameter $\phi_2 \in (0, f_l)$, define

$$\mathcal{E}_3(r, \phi_2) = \left\{ \inf_{\beta_1 \in \beta^* + \Theta_\Sigma(r/2), \beta_2 \in \beta_1 + \Theta_\Sigma(r)} \frac{\widehat{Q}_\varpi(\beta_1) - \widehat{Q}_\varpi(\beta_2) - \langle \nabla \widehat{Q}_\varpi(\beta_2), \beta_1 - \beta_2 \rangle}{\|\beta_1 - \beta_2\|_\Sigma^2} \geq \phi_2 \right\}.$$

Now, we are ready to present the following general upper bound on the estimation error conditioning on the above good events.

Theorem 4. *Assume Conditions 1 and 3 hold and $\beta^{(0)} \in \Theta_\Sigma^*(1)$, and let (ϖ, η_0) satisfy $0 < \varpi \leq f_l/(2l_0\kappa_1)$ and $0 < \eta_0 \leq 1/\max(2f_u, f_l + \bar{\tau})$, where $\bar{\tau} = \max(\tau, 1 - \tau)$. Set $R = 2$, and let $\delta_0, \delta_1 > 0$ be such that*

$$r_0 := (\delta_0 + b^*)/\phi_1 < 1 \quad \text{and} \quad r_1 := (\delta_0 + \delta_1)/\phi_1 \leq 1, \quad (18)$$

where $\phi_1 = 0.5(f_l - l_0\kappa_1\varpi)$. Moreover, let $\Delta = \phi_1 - \delta_0 - b^* \in (0, f_l/2)$ and $\epsilon = \eta_0\phi_1 \in (0, 1/2)$. For any $z \geq 0$, let the sample size satisfy

$$n \geq \sigma \frac{p^{1/2} + \sqrt{2(\log T + z)}}{\Delta}. \quad (19)$$

Then, conditioned on the event $\mathcal{E}_0(B) \cap \mathcal{E}_1(\delta_0, \delta_1) \cap \mathcal{E}_2(2)$, the noisy gradient descent iterate $\beta^{(T)}$ with $T \geq 2\log(n)/\log((1 - \epsilon)^{-1})$ satisfies

$$\|\beta^{(T)} - \beta^*\|_\Sigma \leq r^* \\ := \sqrt{\frac{1}{n^2} + (1 + 1/\epsilon) \left\{ (2p/\epsilon + 3z) \left(\frac{\eta_0\sigma}{n} \right)^2 + (r_0 \vee r_1)^2 \right\}} \quad (20)$$

with probability (over the i.i.d. normal vectors $\{\mathbf{g}_t\}_{t=0}^{T-1}$) at least $1 - 2e^{-z}$. Moreover, the nonprivate empirical (smoothed) risk minimizer satisfies $\|\beta_\varpi - \beta^*\|_\Sigma \leq r_0$.

Let $r = r^* + r_0$. Conditioned further on $\mathcal{E}_3(r, \phi_2)$, $\beta^{(T)}$ with $T \geq 3\log(n)/\log((1 - \epsilon)^{-1})$ satisfies

$$\|\beta^{(T)} - \widehat{\beta}_\varpi\|_\Sigma \leq r^+ := \sqrt{\frac{r^2}{n} + (1 + 1/\epsilon)(2p/\epsilon + 3z) \left(\frac{\eta_0\sigma}{n} \right)^2} \quad (21)$$

with probability (over the i.i.d. normal vectors $\{\mathbf{g}_t\}_{t=0}^{T-1}$) at least $1 - 2e^{-z}$.

Remark 4. Under local strong convexity and smoothness conditions, Avella-Medina et al. (2023) established statistical convergence guarantees for private M-estimators obtained via noisy gradient descent. Let $\mathcal{L}_n : \mathbb{R}^p \rightarrow \mathbb{R}$ be a general (empirical) loss function of interest and $\Theta \subseteq \mathbb{R}^p$ be the parameter space. As high-level conditions, Avella-Medina et al. (2023) assumed that $\mathcal{L}_n$ is locally τ_1 -strongly convex and τ_2 -smooth; that is,

$$\mathcal{L}_n(\beta_1) - \mathcal{L}_n(\beta_2) \geq \langle \nabla \mathcal{L}_n(\beta_2), \beta_1 - \beta_2 \rangle + \tau_1 \|\beta_1 - \beta_2\|_2^2, \\ \forall \beta_1, \beta_2 \in \{\beta \in \mathbb{R}^p : \|\beta - \beta^*\|_2 \leq r\}$$

for some $r > 0$, and $\mathcal{L}_n(\boldsymbol{\beta}_1) - \mathcal{L}_n(\boldsymbol{\beta}_2) \leq \langle \nabla \mathcal{L}_n(\boldsymbol{\beta}_2), \boldsymbol{\beta}_1 - \boldsymbol{\beta}_2 \rangle + \tau_2 \|\boldsymbol{\beta}_1 - \boldsymbol{\beta}_2\|_2^2$, $\forall \boldsymbol{\beta}_1, \boldsymbol{\beta}_2 \in \Theta$. To our knowledge, it remains uncertain whether the aforementioned local strong assumption holds with high probability for either the empirical newsvendor loss $\hat{Q}$ or its convolution-smoothed counterpart $\hat{Q}_\varpi$. Therefore, a more delicate argument is required to analyze the convergence of noisy gradient descent iterates obtained from Algorithm 1. Our proof of Theorem 4 crucially relies on the structural properties of $\hat{Q}_\varpi$ stated in Proposition 3. In Proposition 4 below, we will show that the event conditioned on in Proposition 3 holds with high probability.

Remark 5. The proof of the error bound $\|\hat{\boldsymbol{\beta}}_\varpi - \boldsymbol{\beta}^*\|_\Sigma \leq r_0$ in Theorem 4, which holds conditioning on event $\mathcal{E}_1(\delta_0, \delta_1)$, extends the argument in He et al. (2023). The main contribution of Theorem 4 is to establish finite-sample performance bounds for the noisy gradient descent iterates $\{\boldsymbol{\beta}^{(t)}\}_{t=1}^T$ in a sequential manner, which involves a more intricate analysis compared with that for the nonprivate empirical (smoothed) risk minimizer $\hat{\boldsymbol{\beta}}_\varpi$. More specifically, the analysis conducted in He et al. (2023) necessitates that the empirical loss satisfies only Condition (15), whereas our approach requires a more comprehensive version of the Restricted Strong Convexity Property (16). We also establish a connection between statistical theory and algorithmic complexity, demonstrating that to achieve a statistically efficient estimator as shown in Theorem 2, the computational complexity is of order $O(np \log(n))$. In contrast, the conventional interior-point method commonly used for solving the LP reformulation of empirical check loss minimization demands a significantly higher average-case computational complexity of $O_p(n^{1.25} p^3 \log n)$ (Portnoy and Koenker 1997).

The convergence result stated in Theorem 4 relies on the assumption that the initial value $\boldsymbol{\beta}^{(0)}$ falls within the neighborhood $\Theta_\Sigma^*(1)$, which we term as the *tightening region*. In each iteration of the noisy gradient descent, the current estimate contracts toward the true parameter, progressively moving closer to the region of near-optimal convergence.

In general, let us define $R_0 := \|\boldsymbol{\beta}^{(0)} - \hat{\boldsymbol{\beta}}_\varpi\|_\Sigma$ to be the distance between the initial value $\boldsymbol{\beta}^{(0)}$ and the nonprivate empirical risk minimizer $\hat{\boldsymbol{\beta}}_\varpi$. The following result presents the number of iterations necessary for the noisy gradient descent to enter the tightening region.

Theorem 5. Assume Conditions 1 and 3 hold, and let (ϖ, η_0) satisfy $0 < \varpi \leq f_l/(2l_0\kappa_1)$ and $0 < \eta_0 \leq \min\{1, 1/(2f_u)\}$. Without loss of generality, assume $R_0 = \|\boldsymbol{\beta}^{(0)} - \hat{\boldsymbol{\beta}}_\varpi\|_\Sigma > 1$, and let

$$\Delta = \phi_1 - \delta_0 - b^* \in (0, f_l/2), \quad r_0 = (\delta_0 + b^*)/\phi_1 \in (0, 1),$$

where $\phi_1 = (f_l - l_0\kappa_1\varpi)/2$. Given $z \geq 0$, let the number of iterations T_0 and sample size n satisfy

$$T_0 \geq \frac{R_0^2}{\eta_0 \Delta} \quad \text{and} \quad n \geq 2B_{T_0}\sigma$$

$$\max \left\{ \frac{2R_0 + (\bar{\tau}B + 1/4)(T_0 + 1)\eta_0}{\Delta}, \frac{e-1}{4-e}(\bar{\tau}B + 1/4)T_0 \right\}, \quad (22)$$

where $B_{T_0} = \sqrt{p} + \sqrt{2(\log T_0 + z)}$. Then, conditioned on $\mathcal{E}_0(B) \cap \mathcal{E}_1(\delta_0, \delta_1) \cap \mathcal{E}_2(R)$ with $R = 2R_0 + r_0$, the noisy gradient descent iterate $\boldsymbol{\beta}^{(T_0)}$ satisfies

$$\hat{Q}_\varpi(\boldsymbol{\beta}^{(T_0)}) - \hat{Q}_\varpi(\hat{\boldsymbol{\beta}}_\varpi) \leq \Delta \quad \text{and} \quad \|\boldsymbol{\beta}^{(T_0)} - \boldsymbol{\beta}^*\|_\Sigma \leq 1$$

with probability (over normal vectors $\{\mathbf{g}_t\}_{t=0}^{T_0-1}$) at least $1 - e^{-z}$.

The aforementioned high-level findings demonstrate that, given a sequence of “good events” associated with the empirical smoothed cost function, the proposed noisy gradient descent iterates exhibit desirable convergence properties. To complement this deterministic analysis, we further provide probabilistic bounds, which subsequently yield finite-sample performance bounds as presented in Sections 5.1 and 5.2.

Proposition 4. Assume Conditions 1–3 hold. Given $R > 0$, for any $z > 0$, we have that with probability at least $1 - 5e^{-z}$,

$$\max_{1 \leq i \leq n} \|\Sigma^{-1/2} \mathbf{x}_i\|_2 \leq C_0 v_1 \sqrt{p + \log(n) + z},$$

$$\sup_{\boldsymbol{\delta} \in \Theta_\Sigma(1) \setminus \Theta_\Sigma(1/n)} \frac{|\hat{D}_\varpi(\boldsymbol{\delta}) - D_\varpi(\boldsymbol{\delta})|}{\|\boldsymbol{\delta}\|_\Sigma} \leq C_1 v_1 \sqrt{\frac{p + \log_2(n) + z}{n}},$$

$$\sup_{\boldsymbol{\beta} \in \Theta_\Sigma^*(1)} \|\nabla \hat{Q}_\varpi(\boldsymbol{\beta}) - \nabla Q_\varpi(\boldsymbol{\beta})\|_{\Sigma^{-1}} \leq C_2 v_1 \sqrt{\frac{p \log(n/\varpi) + z}{n}} + \frac{f_u \varpi + 2\kappa_u}{n}$$

and

$$\sup_{\boldsymbol{\beta} \in \Theta_\Sigma^*(R)} \|\nabla^2 \hat{Q}_\varpi(\boldsymbol{\beta}) - \nabla^2 Q_\varpi(\boldsymbol{\beta})\|_{\Sigma^{-1}}$$

$$\leq C_2 v_1^2 \left\{ \sqrt{\frac{p \log(n/\varpi) + z}{n\varpi}} + \frac{p \log(n/\varpi) + z}{n\varpi} \right\}$$

$$+ C_2' R \frac{v_1(p + \log n + z)^{1/2} + l_0 m_3 \varpi^2}{n^2},$$

provided that $n \gtrsim v_1^4(p + z)$, where $C_0, C_1, C_2, C_2' > 0$ are absolute constants. Moreover, $b^* = \|\nabla Q_\varpi(\boldsymbol{\beta}^*)\|_{\Sigma^{-1}} \leq 0.5l_0\kappa_2\varpi^2$, and $\|\nabla^2 Q_\varpi(\boldsymbol{\beta}^*)\|_{\Sigma^{-1}} \leq f_u$.

Proposition 5. In addition to Conditions 1–3, assume

$$\kappa_l := \min_{|u| \leq 1} K(u) > 0 \quad \text{and} \quad \inf_{|u| \leq 1} \frac{1}{2u} \int_{-u}^u f_{\varepsilon|\mathbf{x}}(v) dv \geq f'_l \quad \text{almost surely}$$

for some $f'_l > 0$. Let $r_{\text{loc}} = \varpi / (16 \max\{m_4, 3\}^{1/4} v_1)$. Then, for any $z > 0$, we have that with probability at least $1 - e^{-z}$,

$$\widehat{Q}_{\varpi}(\boldsymbol{\beta}_1) - \widehat{Q}_{\varpi}(\boldsymbol{\beta}_2) - \langle \nabla \widehat{Q}_{\varpi}(\boldsymbol{\beta}_2), \boldsymbol{\beta}_1 - \boldsymbol{\beta}_2 \rangle \geq \phi_2 \|\boldsymbol{\beta}_1 - \boldsymbol{\beta}_2\|_{\Sigma}^2 \quad (23)$$

holds uniformly over $\boldsymbol{\beta}_1 \in \boldsymbol{\beta}^* + \Theta_{\Sigma}(r_{\text{loc}}/2)$ and $\boldsymbol{\beta}_2 \in \boldsymbol{\beta}_1 + \Theta_{\Sigma}(r_{\text{loc}})$, provided that the “effective sample size” $n\varpi$ satisfies $n\varpi \gtrsim m_4^{1/2} v_1^2(p+z)$, where $\phi_2 > 0$ is a constant depending only on (κ_l, f'_l) .

Let the sample size n and bandwidth $\varpi = \varpi_n > 0$ satisfy $n\varpi \gtrsim p \log n$ and $\varpi \lesssim \{(p + \log n)/n\}^{1/4}$. Then, Proposition 4 implies $b^* \lesssim \sqrt{(p + \log n)/n}$ and event $\mathcal{E}_0(B) \cap \mathcal{E}_1(\delta_0, \delta_1) \cap \mathcal{E}_2(2)$ with (B, δ_0, δ_1) satisfying

$$B \asymp \sqrt{p + \log n}, \quad \delta_0 \asymp \sqrt{\frac{p + \log n}{n}} \quad \text{and} \quad \delta_1 \asymp \sqrt{\frac{p \log n}{n}}$$

occurs with high probability. Combining this with Theorem 4 implies the finite-sample performance bounds in Theorem 2.

Moreover, under the additional assumptions stated in Proposition 5, there exist some curvature parameter $\phi_2 > 0$ and a local radius $r_{\text{loc}} \asymp \varpi$ such that the event $\mathcal{E}_3(r_{\text{loc}}, \phi_2)$ also occurs with high probability. This further implies that the μ -GDP estimate $\boldsymbol{\beta}^{(T)}$, obtained from noisy gradient descent initialized at any $\boldsymbol{\beta}^{(0)} \in \boldsymbol{\beta}^* + \Theta_{\Sigma}(1)$, satisfies with probability at least $1 - Cn^{-1}$ that

$$\|\boldsymbol{\beta}^{(T)} - \widehat{\boldsymbol{\beta}}_{\varpi}\|_{\Sigma} \lesssim \frac{\sqrt{\log n} p + \log n}{(f_l \wedge f'_l)^2 \mu n}.$$

Here, we implicitly assume that both the smoothing parameter ϖ and the number of iterations T are chosen appropriately. The above bound, in turn, implies that the excess (smoothed) empirical risk is bounded with high probability by

$$\begin{aligned} \widehat{C}_{\varpi}(\boldsymbol{\beta}^{(T)}) - \widehat{C}_{\varpi}(\widehat{\boldsymbol{\beta}}_{\varpi}) &= \widehat{C}_{\varpi}(\boldsymbol{\beta}^{(T)}) - \min_{\boldsymbol{\beta} \in \mathbb{R}^p} \widehat{C}_{\varpi}(\boldsymbol{\beta}) \\ &\lesssim \log n \left(\frac{p + \log n}{\mu n} \right)^2. \end{aligned}$$

The above rate also matches (up to a logarithm factor) the one in Bassily et al. (2014) for Lipschitz and strongly convex loss functions (after adjusting for scaling differences).

6. Numerical and Empirical Studies

In this section, we utilize synthetic data as well as real-world data to showcase the empirical performance of the proposed privacy-preserving feature-driven policy.

We compare its performance against that of the theoretically optimal policy, which assumes known demand but lacks privacy protection measures. For the sake of simplicity and consistency, we employ the Gaussian kernel in all of our numerical experiments.

6.1. Synthetic Data

We consider the linear demand model $d = \mathbf{x}^T \boldsymbol{\theta}^* + \varepsilon$, where $\boldsymbol{\theta}^* = (1.5, 1, -2.5, -1.5, 3)^T \in \mathbb{R}^5$ and $\mathbf{x}^T = (1, \mathbf{z}^T)^T$. The feature vector $\mathbf{z} \in \mathbb{R}^4$ is generated from a centered multivariate normal distribution with covariance matrix $\Sigma = (0.5^{|j-k|})_{1 \leq j, k \leq 4}$. Independent of the feature vector $\mathbf{x}$, the observation noise variable ε follows one of the following three distributions: (i) standard norm distribution $\mathcal{N}(0, 1)$, (ii) t -distribution with three degrees of freedom (t_3), and (iii) Gaussian mixture distribution $0.9\mathcal{N}(0, 1) + 0.1\mathcal{N}(0, 100)$.

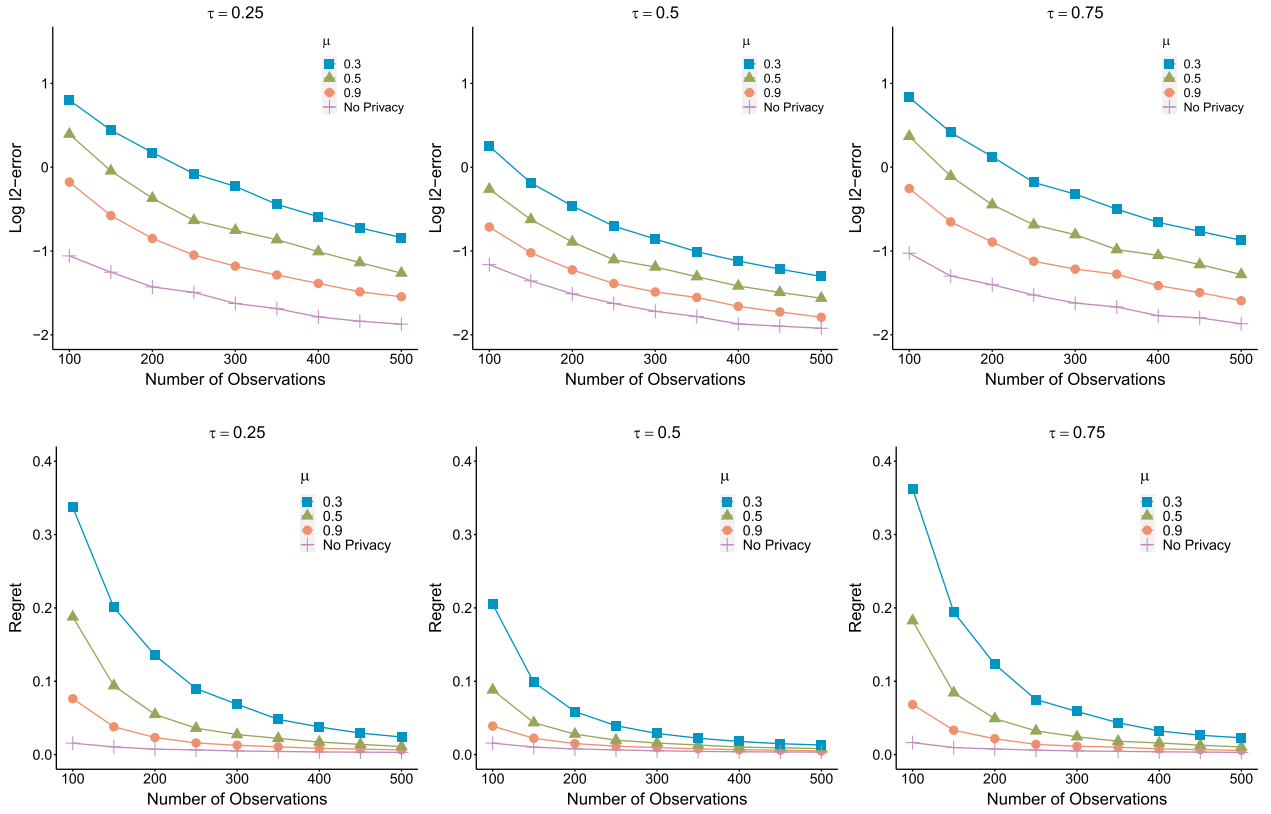
In all of our numerical experiments, we set $b + h = 1$ so that with the distribution of ε known as a priori information, the optimal quantity to order can be determined by the τ -th quantile of the conditional distribution of d given $\mathbf{x}$, where $\tau = b$. Specifically, the clairvoyant optimal policy is $\mathbf{x}^T \boldsymbol{\beta}^*$ with $\boldsymbol{\beta}^* = \boldsymbol{\theta}^* + (Q_{\varepsilon}(\tau), 0, 0, 0, 0)^T$, where $Q_{\varepsilon}(\cdot)$ denotes the quantile function of ε .

For the hyperparameters in the noisy gradient descent method, we set $T = 10$, $B = 2$, and $\sigma = \lceil 2\tau BT^{1/2}/\mu \rceil$, where $\mu \in \{0.3, 0.5, 0.9\}$ is privacy level, and $\tau \in \{0.25, 0.5, 0.75\}$. The step size η_0 is chosen via a backtracking line search. As suggested in He et al. (2023), the bandwidth ϖ is taken to be $\sqrt{\tau(1-\tau)} \cdot \{(p + \log n)/n\}^{2/5}$. The final output $\boldsymbol{\beta}^{(T)}$ is μ -GDP according to Theorem 1. We fix $p = 5$ and let the sample size increase from 100 to 500. Figures 4–6 present plots of the logarithmic ℓ_2 -error and regret versus the sample size under different error distributions and privacy levels, averaged over 300 repetitions. The regret of $\boldsymbol{\beta}^{(T)}$, defined as $\mathbb{E}[C(\mathbf{x}^T \boldsymbol{\beta}^{(T)}, d)] - \mathbb{E}[C(\mathbf{x}^T \boldsymbol{\beta}^*, d)]$, where the expectation is taken over the joint distribution of $(d, \mathbf{x})$, is evaluated using an additional data set of size one million. Figures 7–9 display the box plots of regrets under different error distributions and privacy levels based on 300 repetitions with a sample size of 400. Table 1 reports the corresponding average regrets and standard deviations.

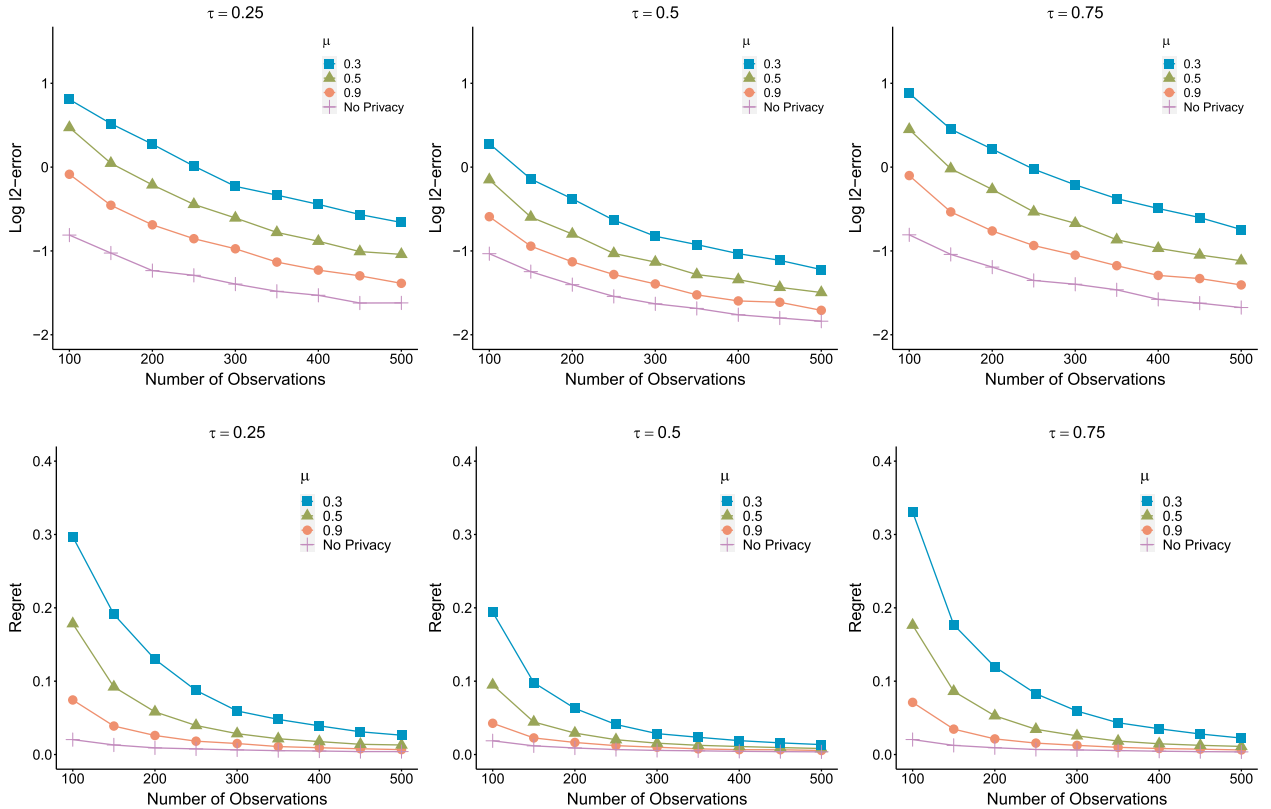
From Figures 4–6, we see that both the estimation errors and regrets decrease as the number of observations grows, as expected. The spacing of these error curves further illustrates the impact of privacy. From Table 1, we observe that the variability of regrets is low when the sample size is reasonably large. These numerical results also highlight the robustness of newsvendor loss minimization against heavy-tailed error distributions.

6.2. Real Data Example

We demonstrate the effectiveness of the proposed privacy-preserving algorithm using the restaurant data from Buttler et al. (2022). This data set comprises

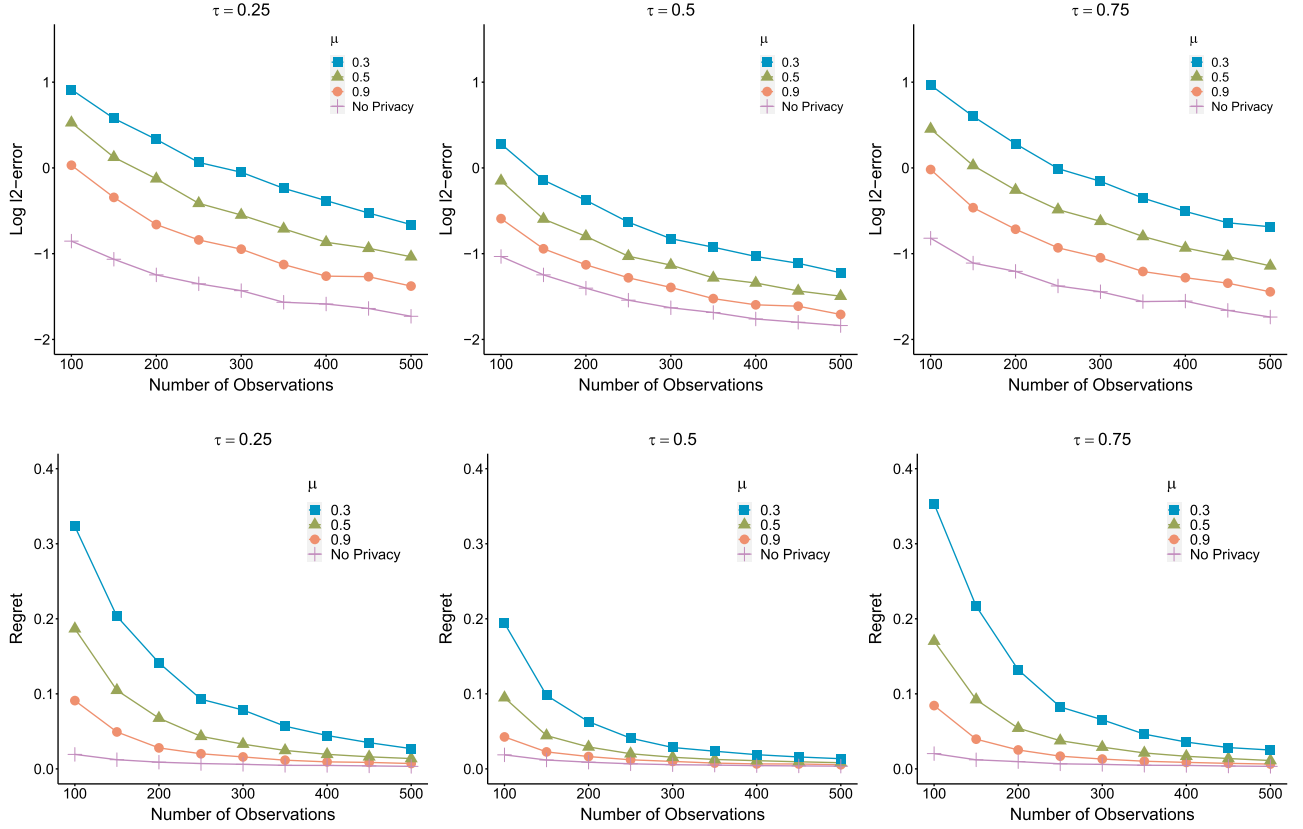
Figure 4. (Color online) Estimation Errors and Regrets of Different Estimators When $\varepsilon \sim \mathcal{N}(0,1)$ 

Note. Plots of logarithmic ℓ_2 estimation error and regret vs. the number of observations, averaged over 300 replications when $\varepsilon \sim \mathcal{N}(0,1)$.

Figure 5. (Color online) Estimation Errors and Regrets of Different Estimators When $\varepsilon \sim t_3$ 

Note. Plots of logarithmic ℓ_2 estimation error and regret vs. the number of observations, averaged over 300 replications when $\varepsilon \sim t_3$.

Figure 6. (Color online) Estimation Errors and Regrets of Different Estimators When $\varepsilon \sim 0.9\mathcal{N}(0, 1) + 0.1\mathcal{N}(0, 100)$



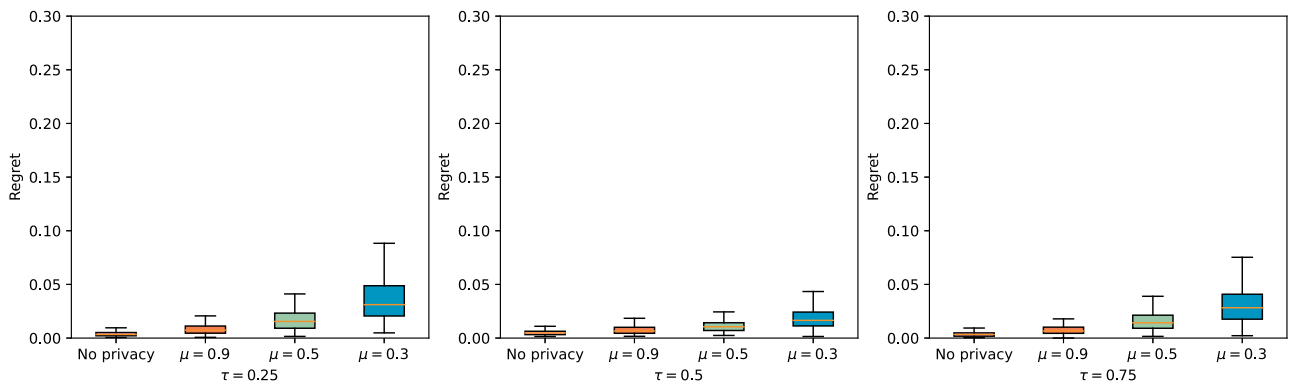
Note. Plots of logarithmic ℓ_2 estimation error and regret vs. the number of observations, averaged over 300 replications when $\varepsilon \sim 0.9\mathcal{N}(0, 1) + 0.1\mathcal{N}(0, 100)$.

demand data for main ingredients at a casual restaurant in Stuttgart over approximately 750 days. The restaurant manager needs to decide on the amount of ingredients to defrost overnight to prepare meals, considering that leftover ingredients result in holding costs. Therefore, we formulate the problem of determining the optimal amount of ingredients to defrost as a newsvendor problem. It is worth noting that during the data collection period, the store manager's strategy

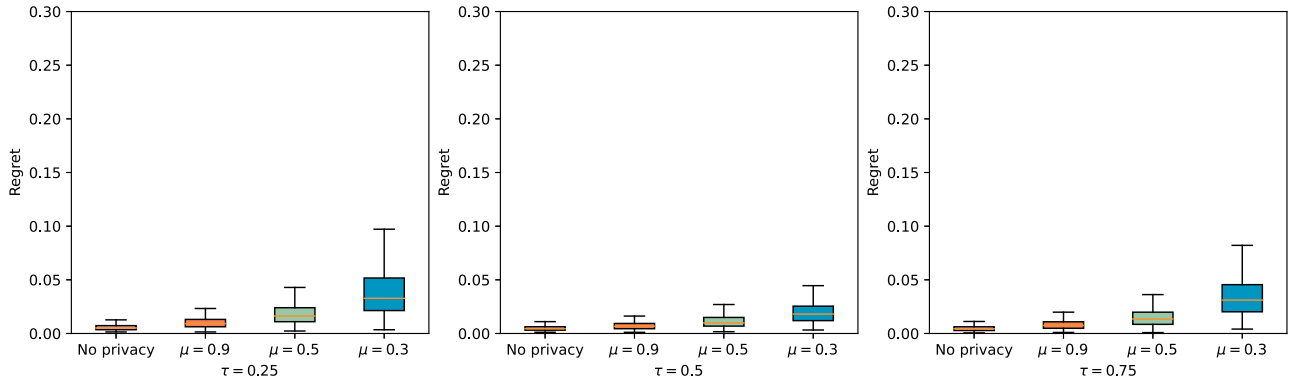
was to maintain a service level of nearly 100%, rendering the issue of censored demand negligible.

In our analysis, we utilize the private algorithm to determine the optimal strategy for defrosting the amount of lamb (which is the ingredient with the highest demand) to minimize costs and maximize performance. We compare the outcomes obtained from the private algorithm with those from the standard nonprivate algorithm. Our model incorporates three distinct

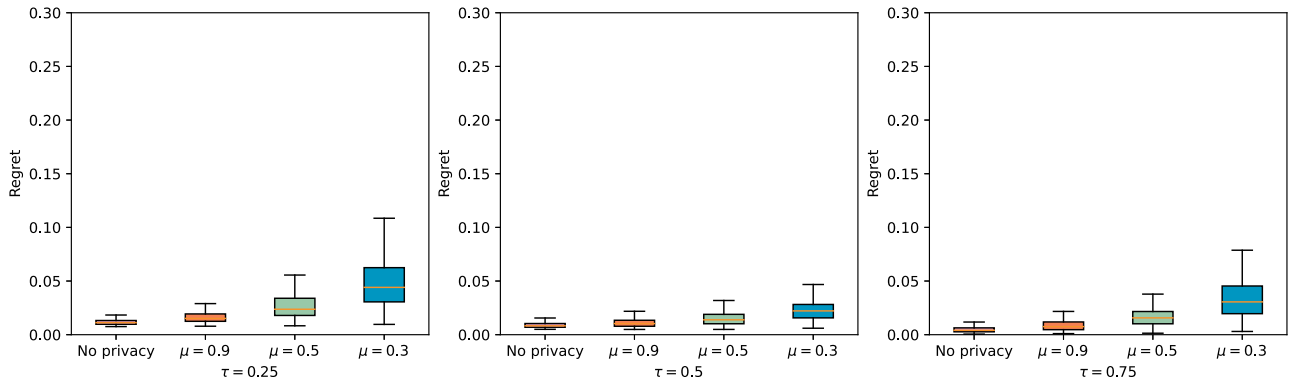
Figure 7. (Color online) Regrets of Different Estimators When Sample Size Equals 400 and $\varepsilon \sim \mathcal{N}(0, 1)$



Note. Box plots of regret with different privacy levels over 300 replications when the sample size is 400.

Figure 8. (Color online) Regrets of Different Estimators When Sample Size Equals 400 and $\varepsilon \sim t_3$ 

Note. Box plots of regret with different privacy levels over 300 replications when the sample size equals 400.

Figure 9. (Color online) Regrets of Different Estimators When Sample Size Equals 400 and $\varepsilon \sim 0.9\mathcal{N}(0, 1) + 0.1\mathcal{N}(0, 100)$ 

Note. Box plots of regret with different privacy levels over 300 replications when the sample size is 400.

features: (i) calendric features, which include binary variables indicating holidays or nonholidays extracted from the date; (ii) lag features, incorporating demand information from the previous periods, such as demand from exactly one week ago and exactly two weeks ago; and (iii) weather features, encompassing rain and temperature data. We assume a per-unit (per-kilogram) holding cost h for lamb of \$30. We consider four different values for the lost-sales penalty cost b : \$50, \$70, \$90, and \$120 per kilogram. These values correspond to gross profit margins (excluding labor costs) of roughly 62.5%, 70%, 75%, and 80%, respectively, which are

similar to the gross profit margin of a financially viable restaurant, estimated to be around 70%. We use a training data set of $n = 552$ past demand observations to train our model, and we evaluate its performance by measuring the out-of-sample error on a separate testing data set consisting of 184 observations.

The algorithm hyperparameters are set as $T = 10$, $B = 2$ and $\sigma = \lceil 2\tau BT^{1/2}/\mu \rceil$ with privacy level $\mu \in \{0.3, 0.5, 0.9\}$. We conduct 100 random partitions of the data set into training and testing data and summarize the average out-of-sample cost across these partitions. Table 2 presents the out-of-sample cost of our private estimator

Table 1. Synthetic Data Analysis

	Nonprivate	$\mu = 0.9$	$\mu = 0.5$	$\mu = 0.3$
$\varepsilon \sim \mathcal{N}(0, 1)$	0.004 (0.002)	0.009 (0.006)	0.017 (0.011)	0.038 (0.026)
$\varepsilon \sim t_3$	0.012 (0.003)	0.017 (0.006)	0.027 (0.012)	0.052 (0.034)
$\varepsilon \sim 0.9\mathcal{N}(0, 1) + 0.1\mathcal{N}(0, 100)$	0.006 (0.003)	0.01 (0.005)	0.019 (0.011)	0.04 (0.026)

Note. Synthetic data analysis: mean and standard deviation (in the parentheses) of the regret for different estimators in three different models over 300 replications, when the sample size equals 40.

Table 2. Restaurant Data Analysis

	Nonprivate	$\mu = 0.9$	$\mu = 0.5$	$\mu = 0.3$
$b = 50$	313.08	315.87	316.71	317.49
$b = 70$	365.67	365.75	367.09	369.32
$b = 90$	405.45	405.22	407.47	410.43
$b = 120$	452.45	453.07	456.21	459.89

Note. Restaurant data analysis: average out-of-sample cost of the private and nonprivate estimators.

(at different privacy levels) and the naive nonprivate estimator for various choices of b . The average cost of our private algorithm is at most 2% higher than the cost of the nonprivate algorithm. This indicates that the proposed algorithm can be effectively used by the restaurant to predict future demands while maintaining a reasonable level of privacy protection, albeit at a slightly higher cost.

7. Concluding Remarks and Discussions

In this paper, we investigate the learning of privacy-preserving optimal policies for feature-based newsvendor problems with unknown demand. We consider the problem within the framework of f -differential privacy, a recently proposed approach that extends the classical (ϵ, δ) -differential privacy with several appealing features. To address the challenge of nonsmoothness associated with the newsvendor loss function, we propose a new noisy gradient algorithm based on convolution smoothing. We provide privacy-preserving guarantees for the T -step output of the proposed algorithm and establish rigorous finite-sample high-probability bounds for estimation error and regret. Importantly, we demonstrate that a reasonable level of privacy protection can be achieved without sacrificing performance compared with the clairvoyant policy with known demand distribution but without privacy protection.

A future endeavor is to find proper conditions on the mini-batch size m under which the noisy SGD estimators are consistent. If m is fixed, there will be a nonvanishing noise term in noisy SGD. Thus, we might not have consistent noisy SGD estimators unless $m \rightarrow \infty$, and the cost of privacy might not be negligible unless we also have that $m^2/n \rightarrow \infty$. These problems deserve further attention in future research. From the practical perspective, the choice of appropriate mini-batch size is critical and nontrivial to ensure high-quality performance. In the newsvendor problem, when the data size is not on the scale of millions, the full gradient descent method remains computationally efficient and exhibits fast geometric convergence.

Acknowledgments

The authors are grateful to the department editor, associate editor, and two anonymous referees for their extensive and constructive comments.

References

- Abadi M, Chu A, Goodfellow I, McMahan HB, Mironov I, Talwar K, Zhang L (2016) Deep learning with differential privacy. *Proc. 2016 ACM SIGSAC Conf. Comput. Comm. Security* (Association for Computing Machinery, New York), 308–318.
- Avella-Medina M, Bradshaw C, Loh PL (2023) Differentially private inference via noisy optimization. *Ann. Statist.* 51(5):2067–2092.
- Bach F, Jenatton R, Mairal J, Obozinski G (2012) Optimization with sparsity-inducing penalties. *Foundations Trends Machine Learning* 4(1):1–106.
- Balle B, Kairouz P, McMahan B, Thakkar O, Guha Thakurta A (2020) Privacy amplification via random check-ins. *Proc. 34th Internat. Conf. Neural Inform. Processing Systems (NIPS '20)* (Curran Associates Inc., Red Hook, NY), 4623–4634.
- Ban GY, Rudin C (2019) The big data newsvendor: Practical insights from machine learning. *Oper. Res.* 67(1):90–108.
- Bassily R, Smith A, Thakurta A (2014) Private empirical risk minimization: Efficient algorithms and tight error bounds. *IEEE 55th Annual Sympos. Foundations Comput. Sci. (IEEE, Piscataway, NJ)*, 464–473.
- Bassily R, Feldman V, Talwar K, Guha Thakurta A (2019) Private stochastic convex optimization with optimal rates. *Proc. 33rd Internat. Conf. Neural Inform. Processing Systems* (Curran Associates Inc., Red Hook, NY), 11282–11291.
- Beutel AL, Minner S (2012) Safety stock planning under causal demand forecasting. *Internat. J. Production Econom.* 140(2):637–645.
- Burnetas AN, Smith CE (2000) Adaptive ordering and pricing for perishable products. *Oper. Res.* 48(3):436–443.
- Buttler S, Philippi A, Stein N, Pibernik R (2022) A meta analysis of data-driven newsvendor approaches. *ICLR 2022 Workshop Setting Up ML Evaluation Standards Accelerate Progress* (ICLR, Appleton, WI).
- Cai TT, Wang Y, Zhang L (2021) The cost of privacy: Optimal rates of convergence for parameter estimation with differential privacy. *Ann. Statist.* 49(5):2825–2850.
- Chaudhuri K, Monteleoni C (2008) Privacy-preserving logistic regression. *Proc. 21st Internat. Conf. Neural Inform. Processing Systems (NIPS'08)* (Curran Associates Inc., Red Hook, NY), 289–296.
- Chaudhuri K, Monteleoni C, Sarwate AD (2011) Differentially private empirical risk minimization. *J. Machine Learning Res.* 12(3):1069–1109.
- Chen X (2007) Large sample sieve estimation of semi-nonparametric models. Heckman JJ, Leamer EE, eds. *Handbook of Econometrics*, vol. 6 (Elsevier, Amsterdam), 5549–5632.
- Chen C, Mangasarian OL (1995) Smoothing methods for convex inequalities and linear complementarity problems. *Math. Programming* 71(1):51–69.
- Chen C, Mangasarian OL (1996) A class of smoothing functions for nonlinear and mixed complementarity problems. *Comput. Optim. Appl.* 5(2):97–138.
- Chen X, Liu W, Zhang Y (2019) Quantile regression under memory constraint. *Ann. Statist.* 47(6):3244–3273.
- Chen X, Miao S, Wang Y (2022a) Differential privacy in personalized pricing with nonparametric demand models. *Oper. Res.* 71(2):581–602.
- Chen X, Simchi-Levi D, Wang Y (2022b) Privacy-preserving dynamic personalized pricing with demand learning. *Management Sci.* 68(7):4878–4898.
- Dong J, Roth A, Su WJ (2022) Gaussian differential privacy. *J. Roy. Statist. Soc. B* 84(1):3–37.
- Dwork C, McSherry F, Nissim K, Smith A (2006b) Calibrating noise to sensitivity in private data analysis. Halevi S, Rabin T, eds. *Theory of Cryptography* (Springer, Berlin, Heidelberg), 265–284.
- Dwork C, Kenthapadi K, McSherry F, Mironov I, Naor M (2006a) Our data, ourselves: Privacy via distributed noise generation.

- Vaudenay S, ed. *Adv. Cryptology EUROCRYPT 2006*, Lecture Notes in Computer Science, vol. 4004 (Springer, Berlin, Heidelberg), 486–503.
- Fainmesser IP, Galeotti A, Momot R (2023) Digital privacy. *Management Sci.* 69(6):3157–3173.
- Feldman V, Koren T, Talwar K (2020) Private stochastic convex optimization: Optimal rates in linear time. *Proc. 52nd Annual ACM SIGACT Sympos. Theory Comput.* (Association for Computing Machinery, New York), 439–449.
- Fernandes M, Guerre E, Horta E (2021) Smoothing quantile regressions. *J. Bus. Econom. Statist.* 39(1):338–357.
- Godfrey GA, Powell WB (2001) An adaptive, distribution-free algorithm for the newsvendor problem with censored demands, with applications to inventory and distribution. *Management Sci.* 47(8):1101–1112.
- Görgülü B, Sarhangian V (2022) A newsvendor approach to design of surgical preference cards. *Service Sci.* 14(3):213–230.
- Green LV, Savin S, Savva N (2013) “Nursevendor problem”: Personnel staffing in the presence of endogenous absenteeism. *Management Sci.* 59(10):2237–2256.
- Hannah LA, Powell WB, Blei DM (2010) Nonparametric density estimation for stochastic optimization with an observable state variable. *Proc. 23rd Internat. Conf. Neural Inform. Processing Systems - Volume 1 (NIPS'10)* (Curran Associates Inc., Red Hook, NY), 820–828.
- He B, Dexter F, Macario A, Zenios S (2012) The timing of staffing decisions in hospital operating rooms: Incorporating workload heterogeneity into the newsvendor problem. *Manufacturing Service Oper. Management* 14(1):99–114.
- He X, Pan X, Tan KM, Zhou WX (2023) Smoothed quantile regression with large-scale inference. *J. Econometrics* 232(2):367–388.
- Horowitz JL (1998) Bootstrap methods for median regression models. *Econometrica* 66(6):1327–1351.
- Hsu D, Kakade S, Zhang T (2012) A tail inequality for quadratic forms of subgaussian random vectors. *Electronic Comm. Probab.* 17(52):1–6.
- Hu M, Momot R, Wang J (2022) Privacy management in service systems. *Manufacturing Service Oper. Management* 24(5):2761–2779.
- Huh WT, Rusmevichientong P (2009) A nonparametric asymptotic analysis of inventory planning with censored demand. *Math. Oper. Res.* 34(1):103–123.
- Iyengar R, Near JP, Song D, Thakkar O, Thakurta A, Wang L (2019) Toward practical differentially private convex optimization. *2019 IEEE Sympos. Security Privacy (SP)* (IEEE, Piscataway, NJ), 299–316.
- Jain P, Thakurta AG (2014) (Near) dimension independent risk bounds for differentially private learning. Xing EP, Jebara T, eds. *Proc. 31st Internat. Conf. Machine Learn.*, Proceedings of Machine Learning Research, vol. 32, no. 1 (PMLR, New York), 476–484.
- Koenker R, Bassett G (1978) Regression quantiles. *Econometrica* 46(1):33–50.
- Kunnumkal S, Topaloglu H (2008) Using stochastic approximation methods to compute optimal base-stock levels in inventory control problems. *Oper. Res.* 56(3):646–664.
- Lee J, Kifer D (2018) Concentrated differentially private gradient descent with adaptive per-iteration privacy budget. *Proc. 24th ACM SIGKDD Internat. Conf. Knowledge Discovery Data Mining* (Association for Computing Machinery, New York), 1656–1665.
- Lei Y, Miao S, Momot R (2024) Privacy-preserving personalized revenue management. *Management Sci.* 70(7):4875–4892.
- Levi R, Perakis G, Uichanco J (2015) The data-driven newsvendor problem: New bounds and insights. *Oper. Res.* 63(6):1294–1306.
- Levi R, Roundy RO, Shmoys DB (2007) Provably near-optimal sampling-based policies for stochastic inventory control models. *Math. Oper. Res.* 32(4):821–839.
- Liyanage LH, Shanthikumar JG (2005) A practical inventory control policy using operational statistics. *Oper. Res. Lett.* 33(4):341–348.
- Nesterov Y (2005) Smooth minimization of non-smooth functions. *Math. Programming* 103(1):127–152.
- Newey W (1997) Convergence rates and asymptotic normality for series estimators. *J. Econometrics* 79(1):147–168.
- Oroojlooyjadid A, Snyder LV, Takáč M (2020) Applying deep learning to the newsvendor problem. *IIEE Trans.* 52(4):444–463.
- Portnoy S, Koenker R (1997) The Gaussian hare and the Laplacian tortoise: Computability of squared-error vs. absolute-error estimators. *Statist. Sci.* 12(4):279–300.
- Powell W, Ruszczyński A, Topaloglu H (2004) Learning algorithms for separable approximations of discrete stochastic optimization problems. *Math. Oper. Res.* 29(4):814–836.
- See CT, Sim M (2010) Robust approximation to multiperiod inventory management. *Oper. Res.* 58(3):583–594.
- Slavkovic A, Molinari R (2022) Perturbed M-estimation: A further investigation of robust statistics for differential privacy. Carriquiry AL, Tanur JM, Eddy WF, eds. *Statistics in the Public Interest*, Springer Series in the Data Sciences (Springer International Publishing, Cham, Switzerland), 337–361.
- Song S, Chaudhuri K, Sarwate AD (2013) Stochastic gradient descent with differentially private updates. *2013 IEEE Global Conf. Signal Inform. Processing* (IEEE, Piscataway, NJ), 245–248.
- Tan KM, Wang L, Zhou WX (2022) High-dimensional quantile regression: Convolution smoothing and concave regularization. *J. Roy. Statist. Soc. Ser. B. Statist. Methodology* 84(1):205–233.
- Vershynin R (2018) *High-Dimensional Probability: An Introduction with Applications in Data Science*, Cambridge Series in Statistical and Probabilistic Mathematics (Cambridge University Press, Cambridge, UK).
- Wang YX (2018) Revisiting differentially private linear regression: Optimal and adaptive prediction & estimation in unbounded domain. Preprint, submitted July 7, <https://arxiv.org/abs/1803.02596>.
- Wang D, Ye M, Xu J (2017) Differentially private empirical risk minimization revisited: Faster and more general. *Proc. 31st Internat. Conf. Neural Inform. Processing Systems (NIPS'17)* (Curran Associates Inc., Red Hook, NY), 2719–2728.
- Wasserman L, Zhou S (2010) A statistical framework for differential privacy. *J. Amer. Statist. Assoc.* 105(489):375–389.
- Williams EA (2020) Where insights meet privacy: Privacy-preserving machine learning. *Forbes* (July 2), <https://www.forbes.com/sites/forbestechcouncil/2020/07/02/where-insights-meet-privacy-privacy-preserving-machine-learning/?sh=13a873987332>.
- Wu X, Li F, Kumar A, Chaudhuri K, Jha S, Naughton J (2017) Bolt-on differential privacy for scalable stochastic gradient descent-based analytics. *Proc. 2017 ACM Internat. Conf. Management Data* (Association for Computing Machinery, New York), 1307–1322.