
Learning from Stochastically Revealed Preference

Anonymous Author(s)

Affiliation

Address

email

Abstract

1 We study the learning problem of revealed preference in a stochastic setting: a
2 learner observes the utility-maximizing actions of a set of agents whose utility
3 follows some unknown distribution, and the learner aims to infer the distribution
4 through the observations of actions. The problem can be viewed as a single-
5 constraint special case of the inverse linear optimization problem. Existing works
6 all assume that all the agents share one common utility which can easily be violated
7 under practical contexts. In this paper, we consider two settings for the underlying
8 utility distribution: a Gaussian setting where the customer utility follows the von
9 Mises-Fisher distribution, and a δ -corruption setting where the customer utility
10 distribution concentrates on one fixed vector with high probability and is arbitrarily
11 corrupted otherwise. We devise Bayesian approaches for parameter estimation and
12 develop theoretical guarantees for the recovery of the true parameter. We illustrate
13 the algorithm performance through numerical experiments.

14 1 Introduction

15 The problem of learning from revealed preference refers to the learning of a common utility function
16 for a set of agents based on the observations of the utility-maximizing actions from the agents. The
17 revealed preference problem has a long history in economics [Sam48, Afr67] (See [Var06] for a
18 review). A line of works [BV06, ZR12, BDM⁺14, ACD⁺15] formulate the problem as a learning
19 problem with two objectives: (i) rationalizing a set of observations, i.e., to find a utility function which
20 explains a set of past observations; (ii) predicting the future behavior of a utility-maximization agent.
21 Mathematically, the action of the agents is modeled by an optimization problem that maximizes a
22 linear (or concave) utility function subject to one linear budget constraint. The learner (decision
23 maker) aims to learn the unknown utility function through a set of observations of the constraints
24 and the optimal solutions. The problem can be viewed as a single-constraint special case of the
25 inverse optimization problem [AO01] which covers a wider range of applications: geoscience [BT92],
26 finance [BGP12], market analysis [BHP17], energy [ASS18], etc.

27 In this paper, we study the problem under a stochastic setting where the agents have a linear utility
28 function randomly distributed according to some unknown distribution. Such a stochastic setting
29 is well-motivated by some application context where the agents are customers and the constraint
30 models the prices and the customer's budget. The optimal solution encodes the customer's purchase
31 behavior and the stochastic utility (objective function) captures the heterogeneity of the customer
32 preference for the products. The goal of learning in this stochastic setting thus becomes to learn the
33 utility distribution through observations of the actions. To the best of our knowledge, we provide the
34 first result of learning a stochastic utility for the revealed preference problem and even in the more
35 general literature of the inverse optimization problem.

36 **Related Literature:** The existing approaches to the problem can be roughly divided into two
37 categories.

Query-based: In a query-based model, the learner aims to learn the utility function by querying an oracle for the agent’s optimal actions, and the goal is to derive the sample complexity guarantee for a sufficiently accurate estimation of the utility function. [BV06] initiates this line of research and studies a statistical setup where the input data is a set of observations and the learner’s performance is evaluated by sample complexity bounds. [ZR12] studies the case of a linear or linearly separable concave utility function, and [BDM⁺14] generalizes the setting and devises learning algorithms for several classes of utility functions. Other than a statistical setup where the observations are sampled from some distribution, both of these two works study an “active” learning setting where the learner has the power to choose the linear constraint (set the prices of the products). Some subsequent works along this line study the associated revenue management problem [ACD⁺15] and a game-theoretic setting [DRS⁺18] where the agents act strategically to hide the true actions.

Optimization-based: The optimization-based approach is usually adopted in the literature of inverse optimization, and some algorithms developed therein can be applied to the special case of the revealed preference problem. [ZL96] and [AO01] study the inverse optimization with one single observation and develop linear programming formulations to solve the problem. Later, [KWB11] and [ASS18] study the statistical (or data-driven) setting where the observations are sampled from some distribution. Specifically, [ASS18] considers a setting where the optimal actions of the agents are contaminated with some independent noises, but all the agents still follow a common utility parameter vector. [MESAHK18] studies the distributional robust version of the problem and [BFL21] considers a contextual formulation. A recent line of works [BMPS18, DCZ18, DZ20, CKK20] cast the inverse optimization problem in an online context and develop (online) gradient-based algorithms.

As we understand, all the existing algorithms and analyses under this topic rely on the assumption that all the agents share one common utility function (or a common utility parameter vector), and thus can fail in the stochastic setting. In this paper, we formulate the problem in Section 2 and focus on the statistical data input where the budget constraint is sampled from some unknown distribution. We consider two stochastic settings: a Gaussian setting in Section 3 and a δ -corruption setting in Section 4. We conclude with numerical experiments and discuss (i) how the results can be generalized to the inverse optimization problem and (ii) the implications for the query-based model where the learner has the power to choose the budget constraints.

2 Model Setup

Consider a customer who purchases a bundle of products subject to some budget constraint. The customer’s utility-maximizing action can be modeled by the following linear program:

$$\begin{aligned} \text{LP}(\mathbf{u}, \mathbf{a}, b) := \max_{\mathbf{x}} \quad & \sum_{i=1}^n u_i x_i \\ \text{s.t.} \quad & \sum_{i=1}^n a_i x_i \leq b, \quad 0 \leq x_i \leq 1, i = 1, \dots, n, \end{aligned}$$

where $\mathbf{u} = (u_1, \dots, u_n) \in \mathbb{R}^n$, $\mathbf{a} = (a_1, \dots, a_n) \in \mathbb{R}_+^n$, and $b \in \mathbb{R}_+$ are the inputs of the LP. Here the decision variables $\mathbf{x}$ encode the purchase decisions where a partial purchase is allowed, u_i denotes the customer’s utility for the i -th product, and a_i denotes the price or cost of purchasing the i -th product. The right-hand-side b denotes the budget of the customer.

Throughout this paper, we make the following assumption.

Assumption 1. *We assume:*

- The utility vector $\mathbf{u}$ follows some unknown distribution $\mathcal{P}_u$.
- The LP’s input $(\mathbf{a}, b)$ follows some unknown distribution $\mathcal{P}_{\mathbf{a}, b}$ independent of $\mathcal{P}_u$.
- There exists $\underline{a} > 0$ such that $a_i \in [\underline{a}, 1]$ almost surely for $i = 1, \dots, n$.

Our goal is to infer the distribution through observations of customers’ optimal actions. Mathematically, we aim to estimate the distribution of $\mathcal{P}_u$ through the dataset

$$\mathcal{D}_T = \{(\mathbf{x}_t^*, \mathbf{a}_t, b_t)\}_{t=1}^T.$$

79 Here the t -th sample corresponds to an unobservable $\mathbf{u}_t$ generated from $\mathcal{P}_u$ and $\mathbf{x}_t^*$ is one optimal
80 solution of $\text{LP}(\mathbf{u}_t, \mathbf{a}_t, b_t)$. Due to the scale invariance of the utility vector, we restrict the distribution
81 $\mathcal{P}_u$ to the unit sphere $\mathcal{S}^{n-1} = \{\mathbf{u} : \|\mathbf{u}\|_2 = 1\}$. In the following two sections, we consider
82 two settings: (i) Gaussian: $\mathcal{P}_u$ follows the von Mises–Fisher distribution, i.e., the restriction of a
83 multivariate Gaussian distribution to the unit sphere; (ii) δ -corruption: $\mathcal{P}_u$ concentrates on one point
84 $\mathbf{u}^*$ with probability $1 - \delta$ and follows an arbitrarily corrupted distribution with probability δ .

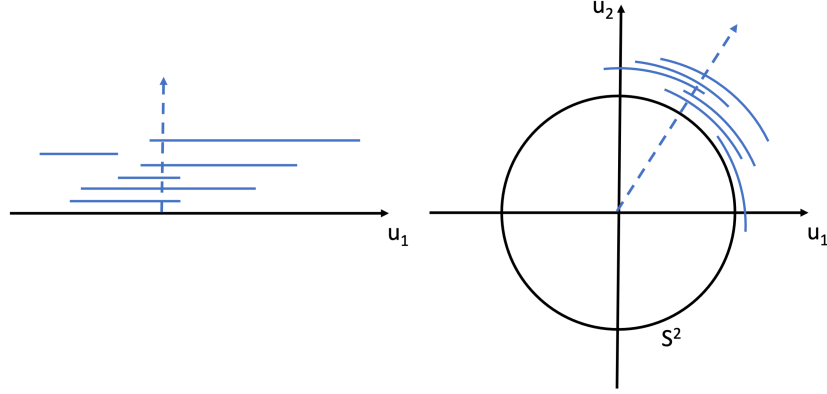


Figure 1: Visualizing the challenge of the problem in 1-D and 2-D.

The challenge of the problem. Each observation $(\mathbf{x}_t^*, \mathbf{a}_t, b_t)$ prescribes a region $\mathcal{U}_t \subset \mathcal{S}^{n-1}$,

$$\mathcal{U}_t := \{\mathbf{u} \in \mathcal{S}^{n-1} : \mathbf{x}_t^* \text{ is an optimal solution of } \text{LP}(\mathbf{u}, \mathbf{a}_t, b_t)\}.$$

85 The set $\mathcal{U}_t$ captures all the possible values of $\mathbf{u}_t$ that is consistent with the t -th observation. The
86 following lemma states that the set $\mathcal{U}_t$ can be expressed by a group of linear constraints.

Lemma 1. For each $\mathcal{U}_t$, there exists a matrix $\mathbf{V}_t$ and a vector $\mathbf{w}_t$ such that

$$\mathcal{U}_t = \{\mathbf{u} \in \mathcal{S}^{n-1} : \mathbf{V}_t \mathbf{u} \leq \mathbf{w}_t\}.$$

87 In the deterministic setting of the revealed preference problem, all the $\mathbf{u}_t$'s are identical and the
88 learning problem is thus reduced to finding one feasible $\mathbf{u}$ in the set of $\cap_{t=1}^T \mathcal{U}_t$. But in a stochastic
89 setting, it may happen that the set of $\cap_{t=1}^T \mathcal{U}_t$ is empty. Figure 1 provides a conceptual visualization
90 of this challenge of “empty intersection”. Each blue solid segment denotes one such $\mathcal{U}_t$ and the blue
91 dashed line represents a value of $\mathbf{u}$ that appears most frequently in these $\mathcal{U}_t$'s. We remark that the
92 figure is just for illustrative purpose as the problem may not be well-defined in the 1-dimensional
93 case.

94 From an estimation viewpoint, the goal is to estimate the distribution of $\mathcal{P}_u$ without the knowledge
95 of the realized samples $\mathbf{u}_t$'s, but merely with the knowledge of $\mathcal{U}_t$ to which $\mathbf{u}_t$ belongs. The
96 sample efficiency of the estimation procedure is naturally contingent on the dispersion of $\mathcal{U}_t$ which is
97 essentially determined by the generation of $(\mathbf{a}_t, b_t)$. For example, if all the $\mathcal{U}_t$'s coincide with each
98 other, then one can hardly learn much about the underlying $\mathcal{P}_u$. In this paper, we aim to pinpoint
99 conditions for $\mathcal{P}_{\mathbf{a},b}$ such that the learning of $\mathcal{P}_u$ is possible. Also, an alternative way to measure
100 the estimation accuracy is to evaluate the predictive performance of the estimated model on new
101 observations generated from $\mathcal{P}_{\mathbf{a},b}$, and such performance bounds generally bear less dependency on
102 the distribution of $\mathcal{P}_{\mathbf{a},b}$. We also provide theoretical guarantees in this sense.

103 3 Gaussian Setting

In this section, we consider a setting where the distribution $\mathcal{P}_u$ follows the von Mises-Fisher distribu-
tion parameterized by $\boldsymbol{\theta} = (\boldsymbol{\mu}, \kappa)$ with the density function

$$f(\mathbf{u}; \boldsymbol{\theta}) := \frac{\exp(-\kappa \boldsymbol{\mu}^\top \mathbf{u})}{\int_{\mathbf{u} \in \mathcal{S}^{n-1}} \exp(-\kappa \boldsymbol{\mu}^\top \mathbf{u}) d\mathbf{u}} \propto \exp(-\kappa \boldsymbol{\mu}^\top \mathbf{u}).$$

Here the vector $\boldsymbol{\mu} \in \mathbb{R}^n$ represents the mean direction and the parameter $\kappa > 0$ controls the concentration of the distribution. The deterministic setting of the revealed preference problem can be viewed as the case when $\kappa = \infty$ and then the distribution degenerates to a point-mass distribution on the unit sphere. Denote the true parameters of the distribution $\mathcal{P}_u$ by $\boldsymbol{\theta}^* = (\boldsymbol{\mu}^*, \kappa^*)$. Then the likelihood of the dataset $\mathcal{D}_t$ under a parameter $\boldsymbol{\theta}$ is

$$\mathbb{P}(\mathcal{D}_t | \boldsymbol{\theta}) := \prod_{t=1}^T \mathbb{P}((\mathbf{x}_t^*, \mathbf{a}_t, b_t) | \boldsymbol{\theta}) = \prod_{t=1}^T \int_{\mathbf{u} \in \mathcal{U}_t} f(\mathbf{u}; \boldsymbol{\theta}) d\mathbf{u}.$$

We remark that the maximum likelihood approach cannot be applied here for two reasons. First, the integration of $f(\mathbf{u}; \boldsymbol{\theta})$ over the region $\mathcal{U}_t$ is not closed-form. The first point is not only pertaining to the Gaussian parameterization of $\mathcal{P}_u$. The scale-invariant property of the utility vector restricts the distribution $\mathcal{P}_u$ to a unit sphere or a simplex, and consequently, the likelihood function inevitably involves the non-closed-form integration. This issue can be partially resolved by using the Monte Carlo method to approximate the integration, and a good thing is that the same integrand is shared across all the observations. Second, the likelihood function is not analytical in $\boldsymbol{\theta}$. Thus this prevents the usage of gradient-based algorithms to solve the problem and also makes it difficult to derive theoretical guarantees for the maximum likelihood estimator.

We propose a Bayesian perspective for the problem: instead of identifying the parameter that maximizes the likelihood function, we directly draw samples from the posterior distribution. We will see shortly that the approach can be justified through a concentration property of the posterior distribution. Suppose we have a prior distribution $\mathbb{P}_0(\boldsymbol{\theta})$ and then we can define the posterior distribution by

$$\begin{aligned} \mathbb{P}_T(\boldsymbol{\theta}) &:= \frac{\mathbb{P}_0(\boldsymbol{\theta}) \cdot \mathbb{P}(\mathcal{D}_t | \boldsymbol{\theta})}{\mathbb{P}(\mathcal{D}_t)} \\ &\propto \mathbb{P}_0(\boldsymbol{\theta}) \cdot \prod_{t=1}^T \int_{\mathbf{u} \in \mathcal{U}_t} f(\mathbf{u}; \boldsymbol{\theta}) d\mathbf{u}. \end{aligned}$$

With slight abuse of notation, we use $\mathbb{P}_T(\cdot)$ (or $\mathbb{P}_0(\cdot)$) to refer to both the density function and the probability measure of the posterior (or prior) distribution. We make the following assumption on the prior distribution.

Assumption 2. We assume the concentration parameter $\kappa^* \in (\underline{\kappa}, \bar{\kappa})$ where $\underline{\kappa}, \bar{\kappa}$ are two known positive constants. The prior distribution $\mathbb{P}_0(\boldsymbol{\theta})$ is a uniform distribution on $\mathcal{S}^{n-1} \times (\underline{\kappa}, \bar{\kappa})$.

Theorem 1. Let

$$\Theta_T := \left\{ \boldsymbol{\theta} \in \mathcal{S}^{n-1} \times (\underline{\kappa}, \bar{\kappa}) : \mathcal{W}(\mathbb{P}((\mathbf{x}_t^*, \mathbf{a}_t, b) | \boldsymbol{\theta}), \mathbb{P}((\mathbf{x}_t^*, \mathbf{a}_t, b) | \boldsymbol{\theta}^*)) \leq \max(9, 9\bar{\kappa}) \frac{n \cdot \log T}{T^{1/2-\alpha}} \right\}$$

where $\mathcal{W}(\cdot, \cdot)$ is the Wasserstein distance between two distributions supported on $\mathcal{X} \times \mathbb{R}_+^n \times \mathbb{R}_+$ equipped with Euclidean metric. Then, under Assumptions 1-2,

$$1 - \mathbb{P}_T(\Theta_T) \rightarrow 0 \text{ in probability as } T \rightarrow \infty,$$

for any $\alpha \in [0, 1/2]$. Specifically, the following inequality holds

$$\mathbb{E}[\mathbb{P}_T(\Theta_T)] \geq 1 - \frac{2}{T} - \frac{2}{16T^{2\alpha} \log^2 T}.$$

where the expectation is taken with respect to the random distribution $\mathbb{P}_T(\cdot)$ (essentially, with respect to the dataset $\mathcal{D}_T$.)

Theorem 1 justifies the approach of posterior sampling. We first remark that the Bayesian sampling approach is just proposed to estimate the parameters, but all the theoretical results are stated in frequentist language. The proof of Theorem 1 follows the standard analysis of the convergence of the posterior distribution [GGVDV00, CDBW21]. While similar results should also hold for other underlying distribution of $\mathcal{P}_u$, the von Mises-Fisher distribution provides much analytical convenience in deriving the bound. Each $\boldsymbol{\theta}$, together with the distribution of $\mathcal{P}_{\mathbf{a}, b}$, defines a distribution over the space of $(\mathbf{x}_t^*, \mathbf{a}_t, b)$. As we use observations $(\mathbf{x}_t^*, \mathbf{a}_t, b)$'s to identify the true $\boldsymbol{\theta}^*$, the set Θ_T defines a set of indistinguishable $\boldsymbol{\theta}$'s based on the Wasserstein distance between distributions of $(\mathbf{x}_t^*, \mathbf{a}_t, b)$.

133 The set Θ_T shrinks as $T \rightarrow \infty$. The posterior sampling approach samples from the distribution
 134 $\mathbb{P}_T(\cdot)$, and Theorem 1 states that the samples will be concentrated in set Θ_T with high probability.
 135 The posterior distribution $\mathbb{P}_T(\cdot)$ is dependent on the dataset $\mathcal{D}_T$, so it is a random distribution itself
 136 and the results in Theorem 1 are stated in either convergence in probability or expectation. As a side
 137 note, the Wasserstein distance in the theorem is not critical and it can be replaced with other distances
 138 such as the total variation distance and the Hellinger distance.

139 Intuitively, Theorem 1 says that for some θ such that the likelihood distribution $\mathbb{P}((\mathbf{x}_t^*, \mathbf{a}_t, b)|\theta)$
 140 differs from $\mathbb{P}((\mathbf{x}_t^*, \mathbf{a}_t, b)|\theta^*)$ to a certain extent, the posterior $\mathbb{P}_T(\cdot)$ is unlikely to generate such
 141 θ . In other words, the posterior distribution identifies the true θ^* up to some “equivalence” in the
 142 likelihood distribution space. The following corollary formalizes this intuition that if there is an
 143 equivalence between the likelihood distribution space and the underlying parameter space, then the
 144 posterior distribution is capable of identifying the true parameter.

Corollary 1. *Suppose*

$$\mathcal{W}(\mathbb{P}((\mathbf{x}_t^*, \mathbf{a}_t, b)|\theta), \mathbb{P}((\mathbf{x}_t^*, \mathbf{a}_t, b)|\theta^*)) > 0$$

145 *for all $\theta \neq \theta^* \in \mathcal{S}^{n-1} \times [\underline{\kappa}, \bar{\kappa}]$. Then the posterior distribution $\mathbb{P}_T(\cdot)$ will converge to the point-mass*
 146 *distribution on θ^* almost surely as $T \rightarrow \infty$. Moreover, suppose there exists a constant $L > 0$*
 147 *satisfying*

$$\mathcal{W}(\mathbb{P}((\mathbf{x}_t^*, \mathbf{a}_t, b)|\theta), \mathbb{P}((\mathbf{x}_t^*, \mathbf{a}_t, b)|\theta^*)) \geq L \cdot \|\theta - \theta^*\|_2, \quad (1)$$

148 *for all $\theta \neq \theta^* \in \mathcal{S}^{n-1} \times [\underline{\kappa}, \bar{\kappa}]$.*

149 *Under Assumptions 1-2, the following inequality holds with probability no less than $1 - \frac{1}{8T^{2\alpha} \log^2 T}$,*

$$\mathbb{E}_T[\|\theta_T - \theta^*\|_2] \leq \max(9, 9\bar{\kappa}) \cdot \frac{n \cdot \log T}{L \cdot T^{1/2-\alpha}} + \frac{2\sqrt{n}}{T}$$

150 *where θ_T is sampled from the posterior distribution $\mathbb{P}_T(\cdot)$.*

151 The corollary states that when there is some equivalence between the likelihood distribution space
 152 and the parameter space as (1), the true parameter is identifiable. The first part of the corollary states
 153 a consistency result that as long as all the $\theta \neq \theta^*$ are distinguishable from θ^* through the likelihood
 154 function, then the posterior sampling will eventually identify the true θ^* . The second part relates to
 155 the convergence rate with an equivalence parameter L .

156 In Assumption 1, we assume the constraint input $(\mathbf{a}_t, b_t)$ is generated from some distribution $\mathcal{P}_{\mathbf{a},b}$.
 157 We note that Theorem 1 and Corollary 1 hold without any additional assumption on $\mathcal{P}_{\mathbf{a},b}$, but the
 158 space topology of the likelihood distribution is highly dependent on $\mathcal{P}_{\mathbf{a},b}$. Specifically, a different
 159 distribution of $(\mathbf{a}_t, b_t)$ determines the separateness of the parameter space through affecting the value
 160 of L in (1) or even its existence. The value L of a specific distribution of $\mathcal{P}_{\mathbf{a},b}$ can be examined
 161 through simulation. So if the learner has some flexibility in choosing the distribution of $\mathcal{P}_{\mathbf{a},b}$, the
 162 optimal choice would be the one that corresponds to a larger value of L . If the constraint input $(\mathbf{a}_t, b_t)$
 163 is not randomly generated but can be actively chosen as the query-based preference learning problem,
 164 the results in Theorem 1 and Corollary 1 still hold by conditioning on all the $(\mathbf{a}_t, b_t)$ ’s. Unlike the
 165 deterministic case where the utility vector $\mathbf{u}$ is fixed for all the observations, the stochastic nature of
 166 the problem setup here makes it generally very complicated to fully extract the benefit of designing
 167 $(\mathbf{a}_t, b_t)$ ’s by the learner. We leave it as a future open question.

Corollary 2. *Suppose $\mathcal{P}_{\mathbf{a},b}$ is a discrete distribution with a finite support. Let $(\mathbf{a}, b)$ be a new
 sample from $\mathcal{P}_{\mathbf{a},b}$, i.e., independent from the dataset $\mathcal{D}_T$, and let $\theta_T = (\boldsymbol{\mu}_T, \kappa_T)$ be a sample
 from the posterior distribution $\mathbb{P}_T(\cdot)$. Denote $\tilde{\mathbf{x}}^*$ and $\mathbf{x}^*$ as the optimal solutions of $LP(\boldsymbol{\mu}_T, \mathbf{a}, b)$
 and $LP(\boldsymbol{\mu}^*, \mathbf{a}, b)$, respectively. Then, under Assumptions 1-2, the following inequality holds with
 probability no less than $1 - \frac{1}{8T^{2\alpha} \log^2 T}$,*

$$\mathbb{E}[\|\tilde{\mathbf{x}}^* - \mathbf{x}^*\|_2] \leq \max(9, 9\bar{\kappa}) \frac{n \cdot \log T}{T^{1/2-\alpha}} + \frac{2\sqrt{n}}{T},$$

168 *where the expectation is taken with respect to both the posterior distribution $\mathbb{P}_T(\cdot)$ and $(\mathbf{a}, b)$.*

169 Corollary 2 provides an upper bound on the predictive performance of the posterior distribution.
 170 Specifically, we want to predict the optimal solution of a linear program specified by $\boldsymbol{\mu}^*$ (proportion-
 171 ally to $\mathbb{E}[\mathbf{u}]$) and a new sample of the constraint $(\mathbf{a}, b)$, and the prediction $\tilde{\mathbf{x}}^*$ is based on a posterior

sample. We know from Theorem 1 that the posterior distribution concentrates on those θ 's that are indistinguishable from the true θ^* in terms of the likelihood. Speaking of the predictive performance, we only concern the distribution of the optimal solution (equivalently, the likelihood), but do not require the identification of exact true θ^* , so Corollary 2 does not require the condition (1) to hold. Intuitively, the prediction of the optimal solution on a new observation (a, b) can be viewed as a condition distribution of the optimal solution given (a, b) . While the definition of Θ_T in Theorem 1 concerns the joint distribution of the optimal solution and (a, b) , the finite-support condition on $\mathcal{P}_{a,b}$ in Corollary 2 transforms the result on the joint distribution to the conditional distribution.

4 δ -Corruption Case

In this section, we consider a setting where the utility vector is specified by

$$\mathbf{u}_t = \begin{cases} \mathbf{u}^*, & \text{w.p. } 1 - \delta, \\ \mathcal{P}'_u, & \text{w.p. } \delta, \end{cases} \quad (2)$$

where $\mathbf{u}^* \in \mathcal{S}^{n-1}$ is a fixed vector, $\delta \in [0, 1]$, and $\mathcal{P}'_u$ is an arbitrary distribution that corrupts the inference of $\mathbf{u}^*$. The deterministic setting of the revealed preference problem in literature can be viewed as the case of $\delta = 0$, and the Gaussian setting in the previous section can be viewed as the case of $\delta = 1$ and $\mathcal{P}'_u$ being the von-Mises Fisher distribution. In this setting, we do not aim to learn the distribution of $\mathcal{P}'_u$, but rather our goal is to identify the vector $\mathbf{u}^*$ using the dataset $\mathcal{D}_T$.

A natural idea to estimate $\mathbf{u}^*$ is by solving the following optimization problem:

$$\text{OPT}_\delta := \max_{\mathbf{u} \in \mathcal{S}^{n-1}} \sum_{t=1}^T I_{\mathcal{U}_t}(\mathbf{u})$$

where the indicator function $I_{\mathcal{E}}(e) = 1$ if $e \in \mathcal{E}$ and $I_{\mathcal{E}}(e) = 0$ otherwise. The rationale for the optimization problem is that for the t -th observation, a vector $\mathbf{u}$ is consistent with the observation, i.e., $\mathbf{x}_t^*$ is the optimal solution of $\text{LP}(\mathbf{u}, \mathbf{a}_t, b_t)$, if and only if $I_{\mathcal{U}_t}(\mathbf{u}) = 1$. Thus the optimization problem finds a vector $\mathbf{u}$ that is consistent with the maximal number of observations. The objective function is discontinuous in $\mathbf{u}$, so we propose the simulated annealing algorithm – Algorithm 2 to solve for its optimal solution.

We first build some connection between the optimization problem OPT_δ and that of the deterministic setting with $\delta = 0$. Let $\bar{\mathbf{x}}_t^*$ be the optimal solution of $\text{LP}(\mathbf{u}^*, \mathbf{a}_t, b_t)$ and define

$$\bar{\mathcal{U}}_t := \{\mathbf{u} \in \mathcal{S}^{n-1} : \bar{\mathbf{x}}_t^* \text{ is an optimal solution of } \text{LP}(\mathbf{u}, \mathbf{a}_t, b_t)\}.$$

Then the deterministic setting of the revealed preference problem solves

$$\bar{\text{OPT}}_\delta := \max_{\mathbf{u} \in \mathcal{S}^{n-1}} \sum_{t=1}^T I_{\bar{\mathcal{U}}_t}(\mathbf{u}).$$

By the setup of the problem, $\mathbf{u}^*$ is an optimizer of $\bar{\text{OPT}}_\delta$ and the optimal objective value is T . The following proposition establishes that the optimization problem OPT_δ is a contaminated version of $\bar{\text{OPT}}_\delta$ and the effect that the contamination has on the objective function can be bounded using δ .

Proposition 1. *Under Assumption 1, the following inequality holds*

$$\mathbb{P} \left(\max_{\mathbf{u} \in \mathcal{S}^{n-1}} \left| \frac{1}{T} \sum_{t=1}^T I_{\mathcal{U}_t}(\mathbf{u}) - \frac{1}{T} \sum_{t=1}^T I_{\bar{\mathcal{U}}_t}(\mathbf{u}) \right| \leq \delta + \frac{\log T}{\sqrt{T}} \right) \leq \frac{1}{T}.$$

When the constraints $(\mathbf{a}_t, b_t)$'s are generated from some distribution $\mathcal{P}_{a,b}$, it can happen that there exist some vectors $\mathbf{u}'$ that are indistinguishable from $\mathbf{u}^*$ based on the observations $\mathcal{D}_t$ as in the previous Gaussian case. So, we do not hope for an exact recovery of $\mathbf{u}^*$, but alternatively, we aim to derive a generalization bound for our estimator $\hat{\mathbf{u}}$. Specifically, we define and analyze the accuracy

$$\text{Acc}(\hat{\mathbf{u}}) := \mathbb{E} [I_{\mathcal{U}}(\hat{\mathbf{u}})] \text{ with } \mathcal{U} := \{\mathbf{u} \in \mathcal{S}^{n-1} : \mathbf{x}^* \text{ is an optimal solution of } \text{LP}(\mathbf{u}, \mathbf{a}, b)\}$$

where $\hat{\mathbf{u}}$ is our estimator of $\mathbf{u}^*$, $(\mathbf{a}, b)$ is a new sample from the distribution $\mathcal{P}_{a,b}$, $\mathbf{u}$ is a new sample following the law of (2), and $\mathbf{x}^*$ is the optimal solution of $\text{LP}(\mathbf{u}, \mathbf{a}, b)$. In other words, the quantity

captures the probability that $\hat{\mathbf{u}}$ is consistent with a new (unseen) observation, and we know that for the true parameter, $\text{Acc}(\mathbf{u}^*) \geq 1 - \delta$, which serves as a performance benchmark.

The challenge for deriving a bound on $\text{Acc}(\hat{\mathbf{u}})$ arises from the discontinuity of the objective function OPT_δ . The existing methods for deriving generalization bound largely rely on the continuity and the Lipschitzness of the loss function. To make it worse, from Lemma 1, we know that $\mathcal{U}_t$ is specified by $(\mathbf{V}_t, \mathbf{w}_t)$ and the $\mathbf{V}_t$'s are of different dimensions for different t 's. To overcome these challenges, we devise the following γ -margin objective function. Specifically, we first define a parameterized version of $\mathcal{U}_t$ by

$$\mathcal{U}_t(\gamma) := \{\mathbf{u} \in \mathcal{S}^{n-1} : \mathbf{V}_t \mathbf{u} \leq \mathbf{w}_t - \gamma \mathbf{e}\}$$

where γ is a positive constant and $\mathbf{e}$ is an all-one vector. It is obvious that $\mathcal{U}_t(\gamma) \subset \mathcal{U}_t$. Accordingly, we define the γ -margin optimization problem by

$$\text{OPT}_\delta(\gamma) := \max_{\mathbf{u} \in \mathcal{S}^{n-1}} \sum_{t=1}^T I_{\mathcal{U}_t(\gamma)}(\mathbf{u}).$$

Proposition 2. *Under Assumption 1, the following inequality holds with probability no less than $1 - \epsilon$,*

$$\max_{\mathbf{u} \in \mathcal{S}^{n-1}} \text{Acc}(\mathbf{u}) - \frac{1}{T} \sum_{t=1}^T I_{\mathcal{U}_t(\gamma)}(\mathbf{u}) \leq 9\sqrt{\frac{\log(T)}{\underline{a}^2 \gamma^2 T}} + 15\sqrt{\frac{\log(T/\epsilon)}{T}}$$

for $\epsilon \in (0, 1)$.

Proposition 2 relates the generalization accuracy of any arbitrary $\mathbf{u}$ with the corresponding objective value of the γ -margin optimization problem. As γ increases, the objective function will decrease, so the right-hand-side becomes tighter. Importantly, the accuracy is defined by the original indicator function (or equivalently, $\mathcal{U}_t$), while the objective value is defined by the γ -margin indicator function (or equivalently, $\mathcal{U}_t(\gamma)$). The implication is that when we optimize the γ -margin objective, we can still obtain a bound on the original accuracy $\text{Acc}(\mathbf{u})$ for sufficiently large γ .

Theorem 2. *Suppose $\mathcal{P}_{\mathbf{a},b}$ is a continuous distribution and it has a density function upper bounded by $\bar{p}$. Then, under Assumption 1, the following inequality holds*

$$\mathbb{P}\left(\text{Acc}(\hat{\mathbf{u}}) \geq 1 - \delta - \frac{60n^2 \log(T)}{\underline{a}T^{1/4}}\right) \geq 1 - \frac{3 + 3\bar{p}}{\min_{\{i: u_i^* \neq 0\}} |u_i^*| \cdot T^{1/4}}$$

where $\hat{\mathbf{u}}$ is one optimal solution of $\text{OPT}_\delta(\gamma)$ with $\gamma = \frac{\underline{a}}{4n^2 T^{1/4}}$.

Theorem 2 states a generalization bound on the accuracy of $\hat{\mathbf{u}}$ for continuous distributions of $(\mathbf{a}, b)$. From Proposition 2, a larger γ leads to a smaller gap between the accuracy and the γ -margin objective function. Meanwhile, a smaller γ leads to a smaller gap between the optimal objective value $\text{OPT}_\delta(\gamma)$ and $1 - \delta$. In the extreme case of $\gamma = 0$, $\mathbb{E}[\text{OPT}_\delta(0)] = \text{Acc}(\mathbf{u}^*) \geq 1 - \delta$. Theorem 2 optimizes the value of γ to trade off these two aspects. We note that the continuous distribution and the upper bound on the density function make it possible to bound the gap of the second aspect. Finally, we remark that the design of γ -margin loss function is inspired from the max-margin classifier, but the analysis is entirely different. For the max-margin classifier, the introduction of the margin aims to make the underlying loss function 1-Lipschitz so that a generalization bound using Rademacher complexity can be derived. But for here, our γ -margin objective function is still a discontinuous one.

5 Computational Aspects and Discussions

In the previous section, we developed theoretical results for both Gaussian and δ -corruption settings. Now we discuss computational aspects with respect to the sampling of the posterior $\mathbb{P}_T(\cdot)$ and the optimization of $\text{OPT}_\delta(\gamma)$. As mentioned earlier, the posterior sampling removes the complication of optimizing over θ in the maximum likelihood estimation, but still inevitably needs to deal with the sampling and numeric approximation of the likelihood function. Algorithm 1 describes a standard Metropolis–Hastings algorithm to sample from the posterior distribution $\mathbb{P}_T(\cdot)$. In the numerical experiments, we choose the proposal distribution $\mathbb{Q}$ to be a Gaussian random perturbation, i.e., $\theta' = \text{Proj}(\theta^{(k-1)} + \epsilon)$ where ϵ follows a Gaussian distribution and the projection ensures that θ'

Algorithm 1 Posterior Sampling for the Gaussian Setting

- 1: Input: dataset $\mathcal{D}_T = \{(\mathbf{x}_t^*, \mathbf{a}_t, b_t)\}_{t=1}^T$, number of iterations K
- 2: Initialize $\boldsymbol{\theta}^{(0)}$ by randomly sampling from the prior distribution $\mathbb{P}_0(\boldsymbol{\theta})$
- 3: **for** $k = 1, \dots, K$ **do**
- 4: Draw a random $\boldsymbol{\theta}'$ from a pre-determined proposal distribution $\mathbb{Q}(\boldsymbol{\theta}'|\boldsymbol{\theta}^{(k-1)})$
- 5: Compute the acceptance rate:

$$r = \min \left\{ \frac{\mathbb{P}_T(\boldsymbol{\theta}')}{\mathbb{P}_T(\boldsymbol{\theta}^{(k-1)})}, 1 \right\}$$

- 6: Set

$$\boldsymbol{\theta}^{(k)} = \begin{cases} \boldsymbol{\theta}', & \text{w.p. } r \\ \boldsymbol{\theta}^{(k-1)}, & \text{w.p. } 1 - r \end{cases}$$

- 7:
 - 8: **end for**
 - 9: Output: $\boldsymbol{\theta}^{(K)}$
-

Algorithm 2 Simulated annealing algorithm for δ -corruption

- 1: Input: dataset $\mathcal{D}_T = \{(\mathbf{x}_t^*), \mathbf{a}_t, b_t\}_{t=1}^T$, margin γ , number of iterations K , interval length τ
- 2: Initialize an initial (temperature) $\eta > 0$ and the reduction rate $c \in (0, 1)$
- 3: Randomly generate the first estimate $\mathbf{u}^{(0)}$
- 4: **for** $k = 1, \dots, K$ **do**
- 5: **if** $k \bmod \tau = 0$ **then**
- 6: Update $\eta \leftarrow c \cdot \eta$
- 7: **end if**
- 8: Draw a proposal $\mathbf{u}'$ from a predetermined proposal distribution $\mathbb{Q}(\mathbf{u}'|\mathbf{u}^{(k-1)})$
- 9: Compute the acceptance rate:

$$r = \min \left\{ \exp \left\{ \frac{1}{\eta} \cdot \left(\sum_{t=1}^T I_{\mathcal{U}_t(\gamma)}(\mathbf{u}') - \sum_{t=1}^T I_{\mathcal{U}_t(\gamma)}(\mathbf{u}^{(k-1)}) \right) \right\}, 1 \right\} \quad (3)$$

- 10: Set

$$\mathbf{u}^{(k)} = \begin{cases} \mathbf{u}', & \text{w.p. } r \\ \mathbf{u}^{(k-1)}, & \text{w.p. } 1 - r \end{cases}$$

- 11: **end for**
 - 12: Output: $\mathbf{u}^{(K)}$
-

230 stays on the sphere $\mathcal{S}^{n-1}$. For the acceptance ratio, as the posterior distribution is not in closed form,
231 a Monte Carlo subroutine is needed to estimate the ratio.

232 Algorithm 2 presents a simulated annealing algorithm to solve the optimization problem $\text{OPT}_\delta(\gamma)$
233 in Section 4. It takes a similar MCMC routine as Algorithm 1 and we use the same Gaussian
234 random perturbation for the proposal distribution $\mathbb{Q}$. As the temperature parameter η decreases, the
235 sampling distribution in Algorithm 2 will gradually be more concentrated on the optimal solution set
236 of $\text{OPT}_\delta(\gamma)$. Algorithm 2 can be implemented more efficiently than Algorithm 1 in that the likelihood
237 ratio calculation in (3) is analytical.

238 Table 1 reports some numerical results for the two algorithms. For both the Gaussian and δ -corruption
239 settings, we consider three distributions of $\mathcal{P}_{\mathbf{a},b}$: (i) a uniform distribution where $\mathbf{a} \sim \text{Unif}([1, 2]^n)$
240 and $b \sim \text{Unif}([1, n])$; (ii) a discrete distribution where $\text{Unif}(\{1, 2\}^n)$ and $b \sim \text{Unif}(1, \dots, n)$; (iii) a
241 fixed- $\mathbf{a}$ distribution where $\mathbf{a} = (1, \dots, 1)^\top$ and $b \sim \text{Unif}(1, \dots, n)$. For the Gaussian case, the true
242 parameters $(\boldsymbol{\mu}^*, \kappa^*)$ are uniformly generated from $\mathcal{S}^{n-1} \times [1, 10]$, and the accuracy is calculated
243 by $(\boldsymbol{\mu}^* \mathbf{x}^* - \boldsymbol{\mu}^* \tilde{\mathbf{x}}^*) / \boldsymbol{\mu}^* \mathbf{x}^*$ where $\mathbf{x}^*$ and $\tilde{\mathbf{x}}^*$ are defined in Corollary 2. For the Gaussian case,
244 $\mathbf{u}^*$ is uniformly generated from $\mathcal{S}^{n-1}$ and δ is set to be 0.1, and the accuracy is calculated by

245 $\text{Acc}(\hat{\mathbf{u}})/\text{Acc}(\mathbf{u}^*)$ where $\text{Acc}(u)$ is defined in Section 4. The numbers in Table 1 are reported based
 246 on an average of 20 simulation trials, and we run both Algorithm 1 and Algorithm 2 for $K = 1000$
 247 iterations.

248 We make the following observations from the numerical experiments. First, we remark that the
 249 theoretical results in the previous sections provide strong guarantees on the convergence property
 250 of the posterior distribution. So the deterioration of the algorithm performance for the case when
 251 $n = 25$ is solely caused by the inaccuracy of the approximate sampling in either Algorithm 1 or
 252 Algorithm 2. Such inaccuracy can definitely be mitigated to some extent by a more efficient algorithm
 253 implementation such as parallel computing. However, we argue that the performance deterioration as
 254 n grows may point to a curse of dimensionality that is intrinsic to this estimation problem. Essentially,
 255 we aim to estimate a high-dimensional distribution only through partial information, i.e., the sets
 256 $\mathcal{U}_t$'s. On the positive end, the algorithms work well for $n \leq 10$, so if the learner has the power
 257 of choosing $(\mathbf{a}_t, b_t)$, s/he can break up the high-dimensional estimation problem into a number of
 258 low-dimensional estimation problems by focusing on a handful of dimensions each time. Moreover,
 259 we provide a visualization of the condition (1) in Figure 2 for $n = 5$ calculated based on simulation.
 260 The visualization supports the existence of L and thus the identifiability of the true parameters when
 the posterior sampling can be accurately fulfilled. 0

		$n = 3$	$n = 5$	$n = 10$	$n = 25$
Gaussian	(i)	99.9%	99.9%	98.8%	59.9%
	(ii)	99.9%	99.3%	96.9%	56.9%
	(iii)	99.9%	94.8%	92.6%	68.5%
δ -corru.	(i)	99.9%	96.7%	96.0%	55.7%
	(ii)	99.6%	97.9%	97.6%	63.1%
	(iii)	99.9%	98.7%	87.1%	58.7%

Table 1: Predictive accuracies under two settings.

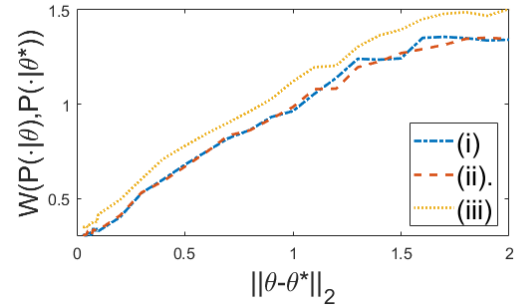


Figure 2: Visualization of (1).

261

262 We conclude our discussion with the following remarks.

263 Query-based model with learner-chosen $(\mathbf{a}_t, b_t)$: In this paper, we have focused on the case where
 264 the constraints $(\mathbf{a}_t, b_t)$'s are stochastically generated. When the concentration parameter κ is known
 265 for the Gaussian case, there is an efficient way of learning $\boldsymbol{\mu}$ through choosing $(\mathbf{a}_t, b_t)$'s (See the
 266 Appendix). In addition, the numerical experiments above also inspire a method that dismantles
 267 the high-dimensional estimation problem into a number of low-dimensional problems. Another
 268 interesting and important question is whether there exist designs of $(\mathbf{a}_t, b_t)$'s such that the posterior
 269 sampling can be more efficiently carried out.

270 Multiple constraints and nonlinear utility: The results in this paper are presented under the setting of
 271 a linear objective and a single constraint. We emphasize that the results can be easily generalized to
 272 the case of multiple constraints and parameterized nonlinear utility. Thus our result can be viewed as
 273 a preliminary effort to address the problem of stochastic inverse optimization. Our conjecture is that
 274 when the set $\mathcal{U}_t$ corresponds to a multiple-constraint problem, it may feature more structure and thus
 275 facilitate the learning of the utility distribution.

276 Choice modeling: The stochastic utility model in our paper also draws an interesting connection with
 277 the literature on choice modeling, which is a pillar for the pricing and assortment problems in revenue
 278 management [TVRVR04, GT⁺19]. For most of the existing choice models, the learning problem
 279 can be viewed as a special case of our study by letting $\mathbf{a} = (1, \dots, 1)^\top$ and $b = 1$. The results in our
 280 paper complement to this line of literature in developing a model where customers can make multiple
 281 purchases.

References

- [ACD⁺15] Kareem Amin, Rachel Cummings, Lili Dworkin, Michael Kearns, and Aaron Roth. Online learning and profit maximization from revealed preferences. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 29, 2015.
- [Afr67] Sydney N Afriat. The construction of utility functions from expenditure data. *International economic review*, 8(1):67–77, 1967.
- [AO01] Ravindra K Ahuja and James B Orlin. Inverse optimization. *Operations Research*, 49(5):771–783, 2001.
- [ASS18] Anil Aswani, Zuo-Jun Shen, and Auyon Siddiq. Inverse optimization with noisy data. *Operations Research*, 66(3):870–892, 2018.
- [BDM⁺14] Maria-Florina Balcan, Amit Daniely, Ruta Mehta, Ruth Urner, and Vijay V Vazirani. Learning economic parameters from revealed preferences. In *International Conference on Web and Internet Economics*, pages 338–353. Springer, 2014.
- [BFL21] Omar Besbes, Yuri Fonseca, and Ilan Lobel. Contextual inverse optimization: Offline and online learning. *Available at SSRN 3863366*, 2021.
- [BGP12] Dimitris Bertsimas, Vishal Gupta, and Ioannis Ch Paschalidis. Inverse optimization: A new perspective on the black-litterman model. *Operations research*, 60(6):1389–1403, 2012.
- [BHP17] John R Birge, Ali Hortaçsu, and J Michael Pavlin. Inverse optimization for the recovery of market structure from market outcomes: An application to the miso electricity market. *Operations Research*, 65(4):837–855, 2017.
- [BMPS18] Andreas Bärmann, Alexander Martin, Sebastian Pokutta, and Oskar Schneider. An online-learning approach to inverse optimization. *arXiv preprint arXiv:1810.12997*, 2018.
- [BT92] Didier Burton and Ph L Toint. On an instance of the inverse shortest paths problem. *Mathematical programming*, 53(1):45–61, 1992.
- [BV06] Eyal Beigman and Rakesh Vohra. Learning from revealed preference. In *Proceedings of the 7th ACM Conference on Electronic Commerce*, pages 36–42, 2006.
- [CDBW21] Minwoo Chae, Pierpaolo De Blasi, and Stephen G Walker. Posterior asymptotics in wasserstein metrics on the real line. *Electronic Journal of Statistics*, 15(2):3635–3677, 2021.
- [CKK20] Violet Xinying Chen and Fatma Kılınç-Karzan. Online convex optimization perspective for learning from dynamically revealed preferences. *arXiv preprint arXiv:2008.10460*, 2020.
- [DCZ18] Chaosheng Dong, Yiran Chen, and Bo Zeng. Generalized inverse optimization through online learning. *Advances in Neural Information Processing Systems*, 31, 2018.
- [DRS⁺18] Jinshuo Dong, Aaron Roth, Zachary Schutzman, Bo Waggoner, and Zhiwei Steven Wu. Strategic classification from revealed preferences. In *Proceedings of the 2018 ACM Conference on Economics and Computation*, pages 55–70, 2018.
- [DZ20] Chaosheng Dong and Bo Zeng. Expert learning through generalized inverse multiobjective optimization: Models, insights, and algorithms. In *International Conference on Machine Learning*, pages 2648–2657. PMLR, 2020.
- [GGVDV00] Subhashis Ghosal, Jayanta K Ghosh, and Aad W Van Der Vaart. Convergence rates of posterior distributions. *Annals of Statistics*, pages 500–531, 2000.
- [GT⁺19] Guillermo Gallego, Huseyin Topaloglu, et al. *Revenue management and pricing analytics*, volume 209. Springer, 2019.

- [KWB11] Arezou Keshavarz, Yang Wang, and Stephen Boyd. Imputing a convex objective function. In *2011 IEEE international symposium on intelligent control*, pages 613–619. IEEE, 2011.
- [MESAHK18] Peyman Mohajerin Esfahani, Soroosh Shafieezadeh-Abadeh, Grani A Hanasusanto, and Daniel Kuhn. Data-driven inverse optimization with imperfect information. *Mathematical Programming*, 167(1):191–234, 2018.
- [Sam48] Paul A Samuelson. Consumption theory in terms of revealed preference. *Economica*, 15(60):243–253, 1948.
- [TVRVR04] Kalyan T Talluri, Garrett Van Ryzin, and Garrett Van Ryzin. *The theory and practice of revenue management*, volume 1. Springer, 2004.
- [Var06] Hal R Varian. Revealed preference. *Samuelsonian economics and the twenty-first century*, pages 99–115, 2006.
- [ZL96] Jianzhong Zhang and Zhenhong Liu. Calculating some inverse linear programming problems. *Journal of Computational and Applied Mathematics*, 72(2):261–273, 1996.
- [ZR12] Morteza Zadimoghaddam and Aaron Roth. Efficiently learning from revealed preference. In *International Workshop on Internet and Network Economics*, pages 114–127. Springer, 2012.

Checklist

The checklist follows the references. Please read the checklist guidelines carefully for information on how to answer these questions. For each question, change the default **[TODO]** to **[Yes]**, **[No]**, or **[N/A]**. You are strongly encouraged to include a **justification to your answer**, either by referencing the appropriate section of your paper or providing a brief inline description. For example:

- Did you include the license to the code and datasets? **[Yes]**
- Did you include the license to the code and datasets? **[N/A]**
- Did you include the license to the code and datasets? **[N/A]**

Please do not modify the questions and only use the provided macros for your answers. Note that the Checklist section does not count towards the page limit. In your paper, please delete this instructions block and only keep the Checklist section heading above along with the questions/answers below.

1. For all authors...
 - (a) Do the main claims made in the abstract and introduction accurately reflect the paper’s contributions and scope? **[Yes]**
 - (b) Did you describe the limitations of your work? **[Yes]**
 - (c) Did you discuss any potential negative societal impacts of your work? **[N/A]**
 - (d) Have you read the ethics review guidelines and ensured that your paper conforms to them? **[Yes]**
2. If you are including theoretical results...
 - (a) Did you state the full set of assumptions of all theoretical results? **[Yes]**
 - (b) Did you include complete proofs of all theoretical results? **[Yes]** In the Appendix.
3. If you ran experiments...
 - (a) Did you include the code, data, and instructions needed to reproduce the main experimental results (either in the supplemental material or as a URL)? **[Yes]**
 - (b) Did you specify all the training details (e.g., data splits, hyperparameters, how they were chosen)? **[Yes]**
 - (c) Did you report error bars (e.g., with respect to the random seed after running experiments multiple times)? **[No]**

- 375 (d) Did you include the total amount of compute and the type of resources used (e.g., type
376 of GPUs, internal cluster, or cloud provider)? [Yes]
- 377 4. If you are using existing assets (e.g., code, data, models) or curating/releasing new assets...
- 378 (a) If your work uses existing assets, did you cite the creators? [N/A]
- 379 (b) Did you mention the license of the assets? [N/A]
- 380 (c) Did you include any new assets either in the supplemental material or as a URL? [N/A]
- 381
- 382 (d) Did you discuss whether and how consent was obtained from people whose data you're
383 using/curating? [N/A]
- 384 (e) Did you discuss whether the data you are using/curating contains personally identifiable
385 information or offensive content? [N/A]
- 386 5. If you used crowdsourcing or conducted research with human subjects...
- 387 (a) Did you include the full text of instructions given to participants and screenshots, if
388 applicable? [N/A]
- 389 (b) Did you describe any potential participant risks, with links to Institutional Review
390 Board (IRB) approvals, if applicable? [N/A]
- 391 (c) Did you include the estimated hourly wage paid to participants and the total amount
392 spent on participant compensation? [N/A]