

Robust Data-Driven Auction Design

Qilong Lin
Shanghai Jiao Tong University
Shanghai, China
edersnow@sjtu.edu.cn

Zhenzhe Zheng*
Shanghai Jiao Tong University
Shanghai, China
zhengzhenzhe@sjtu.edu.cn

Fan Wu
Shanghai Jiao Tong University
Shanghai, China
fwu@cs.sjtu.edu.cn

Yangsuo Liu
Alibaba Group
Beijing, China
liuyangsuo.lys@taobao.com

Jian Xu
Alibaba Group
Beijing, China
xiyu.xj@taobao.com

Guihai Chen
Shanghai Jiao Tong University
Shanghai, China
gchen@cs.sjtu.edu.cn

Dagui Chen
Alibaba Group
Beijing, China
dagui.cd@taobao.com

Bo Zheng
Alibaba Group
Beijing, China
bozheng@alibaba-inc.com

Abstract

In the field of auction design, leveraging deep learning to solve optimal auctions from sampled data has become a promising direction. However, real-world contexts often involve uncertain data, which would severely affect the auction performance, but it is lacking consideration in existing works. To address this challenge, we incorporate these uncertainties into auction design metrics, and frame this challenge as a robust data-driven auction design problem. To solve this problem, we first propose the GAT method, where we introduce the process of problem relaxation and transformation to address the non-differentiable variable presented in the original problem, and further propose an adversarial training algorithm to solve the mini-max problem after transformation. Moreover, to obtain moderately robust auctions, we propose two methods to select the robust coefficient, which provides guidance and insights for selecting robust auctions based on generalization and performance metrics. Finally, with the insights from the GAT method, we further propose the SAT method, where we employ a strict and unified IC constraint that extends from the GAT method, which provides strong IC guarantees and stable revenue in uncertain environments. Experiments on both constructed and real-world datasets show that our robust methods effectively improve the performance of auctions in terms of revenue and IC guarantees.

CCS Concepts

• **Applied computing** → Online auctions; • **Computing methodologies** → Adversarial learning; • **Theory of computation** → Sample complexity and generalization bounds.

*Zhenzhe Zheng is the corresponding author.

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than the author(s) must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.
KDD '25, Toronto, ON, Canada.

© 2025 Copyright held by the owner/author(s). Publication rights licensed to ACM.
ACM ISBN 979-8-4007-1454-2/25/08
<https://doi.org/10.1145/3711896.3737111>

Keywords

Robust Auction Design, Data-Driven Optimization, Adversarial Learning, Online Advertising

ACM Reference Format:

Qilong Lin, Yangsu Liu, Dagui Chen, Zhenzhe Zheng, Jian Xu, Bo Zheng, Fan Wu, and Guihai Chen. 2025. Robust Data-Driven Auction Design. In *Proceedings of the 31st ACM SIGKDD Conference on Knowledge Discovery and Data Mining V.2 (KDD '25)*, August 3–7, 2025, Toronto, ON, Canada. ACM, New York, NY, USA, 12 pages. <https://doi.org/10.1145/3711896.3737111>

1 Introduction

Auction theory has generated significant revenue for many commercial activities, such as digital advertising [20], spectrum auctions [12], and electricity markets [22]. In this field, the goal is to design the decision rules for resource allocation that depend on the values of bidders [46], where the core feature of the problem is that the value information is private to bidders, and they may misreport the values to maximize their own utilities. This introduces additional considerations into the decision-making process, such as the Incentive Compatibility (IC) constraint [53] that characterizes the bidders' willingness to truthfully report their private values.

Our work focuses on the optimal auction design problem, where the goal is to maximize revenue. This objective aligns with the interests of auctioneers, making it a topic of great attention in real-world applications. Under this goal, the most influential work is from Roger B. Myerson [45], which explored the optimal rules for the single-item auctions. Since then, many efforts have focused on theoretically finding optimal auctions in the generalized settings, such as the single-bidder auctions [14, 25, 43], the auctions with symmetric value distribution [15], the combinatorial auctions [56], etc. However, these results are often limited to specific settings or cannot guarantee optimality, which hinders their wide applications.

To detour these theoretical difficulties, Dütting et al. [17] proposed a data-driven solution for the optimal auction design problem. Their approach leverages the concept of regret [19] to transform the auction design problem into a data-driven optimization problem, thereby enabling the use of deep learning to derive the optimal auctions. Currently, this data-driven solution has been validated and

expanded by existing works [11, 27, 50], and has been effectively deployed in various fields, including online advertising [61], energy markets [42] and data markets [52].

However, existing data-driven solutions lack consideration of uncertainty in the collected data samples, which would lead to suboptimal and even quite poor performance. On the one hand, in practice, we often derive the optimal auctions based on a limited amount of sampled data, and thus the empirical distribution may differ from the underlying true distribution. On the other hand, even if we can obtain or approximate the underlying true distribution, the true distribution itself is usually changing over time. These two factors together contribute to the uncertainty of the sampling data, which is widely present in practice. However, existing solutions did not consider this uncertainty in both the design and analysis procedures. As we observed in the experimental results, this can lead to a decline in auction performance in uncertain environments, such as the significant violation of IC under few sampled data and the declines of revenue under distribution-shift over time. Please refer to Section 5 for more details and specific examples.

To address the challenge of uncertainty, we model the problem of data-driven auction design in uncertain environments as a robust optimization problem. Specifically, our modeling is inspired by the distributionally robustness idea [38, 49]: we use the concept of *ambiguity set* to capture the true distribution in uncertain environments, and model our objective as optimizing the worst-case performance within the ambiguity set. This formalizes our goal of optimizing the lower bound of performance under true distribution.

To solve the robust optimization problem, we first propose the General Adversarial Training (GAT) method. Specifically, the main challenge in solving the distributionally robust problem lies in its non-differentiable optimization variable, which is used to capture the distribution uncertainty. To address this, we first introduce a problem relaxation and transformation process, which is supported by rigorous strong-duality theory [24]. Based on this, we further propose an adversarial training algorithm to solve the transformed mini-max optimization problem. These ideas together provide a general approach for solving robust data-driven auctions.

Beyond deriving the optimization framework to solve the problem, the robustness degree is also a key parameter that affects auction performance, but is lacking discussion in existing works [1, 4]. To address this, we propose two methods to select the robust coefficient. Among these, the generalization-guided method has a solid theoretical foundation, while the performance-guided method is more applicable to practical situations. Specifically, we introduce the Repeated Holdout algorithm to select the robust coefficient in online environment as an illustrated example, in order to cope with the continuously changing data distribution.

Finally, we conducted an in-depth analysis of how the robust IC constraint in GAT works in the uncertain data-driven environments, which inspired us to further propose the Stable Adversarial Training (SAT) method. Compared to the robust IC constraint, the SAT method employs a stricter and unified IC constraint, thereby providing stronger IC guarantees for the auctions in the uncertain data-driven environments. It is worth noting that, despite lacking explicit robust considerations, the SAT method still achieves excellent and stable revenue performance in experimental results, which further highlights the advantages of the SAT method.

Our main contributions can be summarized as follows:

- We propose the GAT method to solve the robust data-driven auction design problem, where we address the challenge of non-differentiable variable through a problem relaxation and transformation process supported by duality theory. We further propose an adversarial training algorithm to solve the transformed problem.
- We propose two methods for the selection of the robustness degree based on theoretical foundation and practical requirements, respectively. This is a key parameter that affects the performance of robust auctions, but lacks enough discussion in the literature.
- Inspired by the analysis of the GAT method, we further propose the SAT method with a strict and unified IC constraint. This provides stronger IC guarantees for the uncertain environments, while maintaining excellent and stable revenue performance.
- We apply our methods in two common uncertain scenarios, and analyze their performance with respect to the data amount and the robust coefficient. The results on the constructed and real-world datasets show the efficiency of the above proposed methods.

2 Related Works

Data-Driven Auction Design. Since Dütting et al. [17] successfully proposed RegretNet to solve optimal auctions from sampled data, this idea of data-driven auction design has been adopted and expanded in many subsequent works [11, 27, 50]. Specifically, we focus on the line of work with approximate IC, which inherit the core concept of regret from RegretNet, and achieve approximate IC by minimizing regret. In this context, Feng et al. [23], Kuo et al. [35], and Sakurai et al. [55] extended RegretNet framework to scenarios involving budget constraints, fairness, and false-name-proof requirements, respectively. Moreover, Rahme et al. [50, 51], Ivanov et al. [28] and Duan et al. [16] further improved RegretNet from the perspective of loss functions and model structures. However, these works did not consider the impact of uncertainties in real-world environments on the performance of such data-driven auctions.

Robust Mechanism Design. The goal of robust mechanism design is to relax the assumptions of the Myerson mechanism [45] to develop more practical revenue-maximizing mechanisms. On this topic, a series of works [5, 8, 41] tried to identify the optimal mechanisms when relaxing the common knowledge assumption. Recently, another series of works [1, 4, 9, 31, 57] derived the optimal mechanisms in the presence of ambiguity in value distribution, which is more related to our topic. However, the settings in these works are not universal, and the ambiguity misses the statistical relation between the empirical and true distribution, making them unsuitable for our general data-driven auction design setting.

3 Preliminaries

3.1 Auction Model

We consider a general auction setting where there are N bidders $\mathcal{N} = \{1, \dots, N\}$ competing for M items $\mathcal{M} = \{1, \dots, M\}$. Here the allocation results of items are finite, that is, all subsets $\{S_1, S_2, \dots, S_{2^M}\}$ of $\mathcal{M}$. Each bidder i has a private value vector v_i , where v_{ij} denotes her value for subset S_j . In subsequent discussions, we denote $\mathbf{v} = (v_1, v_2, \dots, v_n)$ as the value vectors of all bidders, and $\mathbf{v}_{-i}$ as all the vectors without v_i . To model randomness, we assume that $\mathbf{v}$ follows a true distribution P_0 over a compact value space $\mathcal{V}$.

In each auction, every bidder can submit a value vector $\tilde{v}_i$ as her bid to the auctioneer. We consider the data-driven approach [17], where the auction's allocation function $a^{\mathbf{w}}$ and payment function $p^{\mathbf{w}}$ are determined by the parameters $\mathbf{w}$ of the neural network model. We use $a_{ij}^{\mathbf{w}}(\tilde{v})$ and $p_i^{\mathbf{w}}(\tilde{v})$ to represent the probability that S_j is allocated and the payment charged for bidder i , respectively, according to the submitted values $\tilde{v}$. We would like to emphasize that the submitted values $\tilde{v}$ could be different from the true values due to the strategic behaviors of bidders. Therefore, in an auction, the utility of bidder i can be represented as:

$$u_i^{\mathbf{w}}(v_i; \tilde{v}) = \sum_{j=1}^{2^M} v_{ij} a_{ij}^{\mathbf{w}}(\tilde{v}) - p_i^{\mathbf{w}}(\tilde{v}). \quad (1)$$

To model the auction design problem as an optimization problem, we consider the following two metrics for auctions. Firstly, the revenue of the auction when all bidders submit truthful values:

$$\text{rev}(\mathbf{v}, \mathbf{w}) = \sum_{i \in N} p_i^{\mathbf{w}}(v_i, v_{-i}), \quad (2)$$

Another metric is the regret of each bidder i in the auction:

$$\text{rgt}_i(\mathbf{v}, \mathbf{w}) = \max_{v'_i} \{u_i^{\mathbf{w}}(v_i; (v'_i, v_{-i})) - u_i^{\mathbf{w}}(v_i; (v_i, v_{-i}))\}, \quad (3)$$

which reflects the additional utility that bidder i can gain by misreporting when all the other bidders submit their values truthfully. This concept converts IC into a quantifiable metric, thereby being widely adopted by existing works [16, 17, 23, 28] to transform auction design problems into data-driven optimization problems.

It should be noted that in classical auction design, Individual Rationality (IR) is also an important constraint [53]. However, for our model-determined payment rules $p^{\mathbf{w}}$, we can easily satisfy the IR constraint through the design of the model structure, as many existing works [17, 23, 28]. Therefore, in the remaining part of our work, we confidently focus on the revenue and regret metrics.

3.2 Uncertainty Environments

In many practical situations, we cannot directly know the true distribution P_0 that the bidders' value vectors $\mathbf{v}$ follows, but can only infer the possible range for it through collected samples, such as the online advertising environment [37, 40]. We will start from discussing these uncertain environments, and then develop a universal and automated method for designing robust data-driven auctions.

Formally, we denote all collected sampled data as a dataset $\mathcal{D}$, and denote the empirical distribution on it as $\hat{P}(\mathcal{D})$. For brevity, we use $\hat{P}$ to represent $\hat{P}(\mathcal{D})$. To model the uncertain range of the true distribution P_0 , we leverage the idea of distributionally robustness [49], assuming the distance between the true distribution P_0 and the empirical distribution $\hat{P}$ does not exceed a certain limit ρ . That is, we consider P_0 to be any distribution within an ambiguity set:

$$\mathcal{P}(\hat{P}, \rho) = \{P | d_c(\hat{P}, P) \leq \rho\}, \quad (4)$$

where d_c is Wasserstein distance [58] using norm function c . The main reason for using Wasserstein distance lies in its advanced theoretical results and its ability to capture the relation between the empirical and the true distribution, while some other well-known measurements, such as the KL divergence [34], fail to capture such a relation. This will be further discussed in Section 4.2.1.

3.3 Problem Formulation

In previous uncertain modeling, since we have captured the true distribution P_0 through the ambiguity set $\mathcal{P}(\hat{P}, \rho)$, a natural and robust strategy is to optimize the worst-case performance within the ambiguity set $\mathcal{P}(\hat{P}, \rho)$, thereby optimizing the lower bound of the auction performance under the true distribution P_0 . Following this idea, we consider the robust versions of the two auction metrics above. The worst-case expected revenue can be written as:

$$\text{rrev}(\mathbf{w}, \hat{P}, \rho) = \min_{P \in \mathcal{P}(\hat{P}, \rho)} \mathbb{E}_{\mathbf{v} \sim P} [\text{rev}(\mathbf{v}, \mathbf{w})], \quad (5)$$

where we consider a perturbation $P \in \mathcal{P}(\hat{P}, \rho)$ that $\mathbf{v}$ would follow, and we regard the expected revenue in the worst-case scenario as the optimization objective. Similarly, the expected regret in the worst-case scenario for any bidder i can be written as:

$$\text{rrgt}_i(\mathbf{w}, \hat{P}, \rho) = \max_{P \in \mathcal{P}(\hat{P}, \rho)} \mathbb{E}_{\mathbf{v} \sim P} [\text{rgt}_i(\mathbf{v}, \mathbf{w})], \quad (6)$$

where its corresponding constraint $\text{rrgt}_i(\mathbf{w}, \hat{P}, \rho) = 0$ that describes the IC objective for the auction is similar to the definition of the distributionally robust incentive compatible (DRIC) constraint in the robust mechanism design literature [31], which requires the expected regret of the auction to be 0 for any distribution P within the ambiguity set $\mathcal{P}(\hat{P}, \rho)$. Therefore, we refer to this corresponding constraint of (6) as a robust IC constraint. Then, the robust auction design problem can be formulated as optimizing the worst-case expected revenue under the robust IC constraint:

$$\begin{aligned} \max_{\mathbf{w}} \quad & \text{rrev}(\mathbf{w}, \hat{P}, \rho) \\ \text{s.t.} \quad & \text{rrgt}_i(\mathbf{w}, \hat{P}, \rho) = 0, \forall i \in N. \end{aligned} \quad (7)$$

Here, since the parameter ρ determines the range of the worst-case scenario $\mathcal{P}$, a larger ρ leads to a larger range, resulting in more conservative solution for problem (7). That is, it is directly related to the robustness degree of the auctions, and we refer to it as the robust coefficient in the remaining discussion.

The characteristics of the robust problem (7) mainly lie in its general settings and its special ambiguity sets tailored to data-driven scenarios. In contrast, existing works [1, 4, 9, 31, 57] did not include combinatorial auctions, where Suzdaltsev [57] only considered a single bidder and Albert et al. [1] only considered discrete value setting. Furthermore, the ambiguity sets in these works are merely described by simple information such as mean and variance, making it difficult to capture the relation between the empirical and the true distribution. Therefore, we refer to problem (7) as the robust data-driven auction design problem, to highlight the distinctive features and challenges of our problem.

Furthermore, we hope to clarify the practical significance of robustness in our method. Although we optimize for the worst-case distribution within the uncertain range, we do not consider the full range of potential distributions, but rather regard the range as a parameter to select. This enables us to select moderately robust auctions based on generalization and performance metrics, allowing us to achieve better performance in uncertain environments. We leave the detailed selection methods for Section 4.2.

4 Design

4.1 General Adversarial Training

We first focus on solving problem (7) given an empirical distribution $\hat{P}$ and a robust coefficient ρ . We present the process of problem relaxation and transformation, which addresses the core challenge of non-differentiable variable in the original problem (7), making it solvable by the proposed adversarial training algorithm. We refer to our solution as the **GAT** (General Adversarial Training) method, and its details are described as follows.

We first derive the Lagrangian function of problem (7) to approximate its solution:

$$\min_{\mathbf{w}} \left\{ -rrev(\mathbf{w}, \hat{P}, \rho) + \sum_{i \in \mathcal{N}} \lambda_i \cdot rrgt_i(\mathbf{w}, \hat{P}, \rho) \right\}, \quad (8)$$

where the Lagrange multipliers $\lambda_i > 0$ for any i .

We next relax the problem based on our practical consideration. Specifically, the original problem (7) is a standard robust problem formulation, where the objective (5) and the constraint (6) consider two separate worst-case scenarios, independently. However, from our practical perspective, the worst-case scenarios are actually correlated. This is because, although the true distribution P is unknown, it uniquely exists in real-world scenarios. Therefore, the perturbation terms P in both objective (5) and constraint (6) of our robust problem should be consistent and be regarded as the same variable. Based on this insight, we relax the original problem (8) by introducing a unified perturbation as follows.

For convenience, we first define the Lagrangian objective in the absence of perturbation. Given the model parameters $\mathbf{w}$ and values $\mathbf{v}$, the value of the Lagrangian function can be written as:

$$lag(\mathbf{w}, \mathbf{v}, \boldsymbol{\lambda}) = -rev(\mathbf{v}, \mathbf{w}) + \sum_{i \in \mathcal{N}} \lambda_i \cdot rgt_i(\mathbf{v}, \mathbf{w}), \quad (9)$$

where we use $\boldsymbol{\lambda}$ to denote vectors $(\lambda_1, \lambda_2, \dots, \lambda_N)$. Then our relaxed optimization objective can be written as:

$$\min_{\mathbf{w}} \max_{P \in \mathcal{P}(\hat{P}, \rho)} \mathbb{E}_{\mathbf{v} \sim P} [lag(\mathbf{v}, \mathbf{w}, \boldsymbol{\lambda})], \quad (10)$$

where a unified perturbation within the ambiguity set $\mathcal{P}$ is used to replace the origin two separated ones. Here, we use “relaxed” to indicate that the worst-case considerations of the new objective (10) are weaker than those of the original objective (8).

Although the relaxation process simplifies the problem, the core challenge, namely the non-differentiable optimization variable P in the problem (10), still remains. Since such a distribution-type variable cannot be differentiated, it is impossible for us to directly solve the problem (10) through classic methods such as gradient descent [2]. To address this challenge, we consider transforming the problem into an equivalent but differentiable one. Specifically, we denote the expected robust forms of the Lagrangian objective lag in (9) with penalty as follows:

$$\phi(\mathbf{v}, \mathbf{w}, \boldsymbol{\lambda}, \gamma) = \max_{\mathbf{v}'} \{ lag(\mathbf{v}', \mathbf{w}, \boldsymbol{\lambda}) - \gamma c(\mathbf{v} - \mathbf{v}') \}, \quad (11)$$

where c is the norm function as defined in the ambiguity set (4). Equation (11) measures the worst-case performance of objectives lag , where perturbations are introduced as adversarial variables $\mathbf{v}'$ to replace origin values $\mathbf{v}$, and the perturbation range is limited by a penalty weighted by γ . We arrive at the following result:

Algorithm 1: General Adversarial Training Approach

Input: Mini-batches $\{\mathcal{D}_1, \mathcal{D}_2, \dots, \mathcal{D}_L\}$
Hyper-parameters: $T, \epsilon, \eta_{\mathbf{w}}, \eta_{\boldsymbol{\lambda}}, \eta_{\gamma}$

```

1 Initialize  $\mathbf{w}, \boldsymbol{\lambda}, \gamma$ ;
2 while stopping criterion not met do
3   for  $l = 1$  to  $L$  do
4     for  $i = 1$  to  $T$  do
5       foreach  $\hat{\mathbf{v}} \in \mathcal{D}_l$  do
6          $\gamma \leftarrow \gamma - \eta_{\gamma} \nabla_{\gamma} h_l(\mathbf{w}, \boldsymbol{\lambda}, \gamma)$ ;
7        $\mathbf{w} \leftarrow \mathbf{w} - \eta_{\mathbf{w}} \nabla_{\mathbf{w}} h_l(\mathbf{w}, \boldsymbol{\lambda}, \gamma)$ ;
8     if updating  $\boldsymbol{\lambda}$  criterion met then
9        $\boldsymbol{\lambda} \leftarrow \boldsymbol{\lambda} + \eta_{\boldsymbol{\lambda}} \nabla_{\boldsymbol{\lambda}} h_l(\mathbf{w}, \boldsymbol{\lambda}, \gamma)$ ;

```

THEOREM 1. *Problem (10) is equivalent to the following problem:*

$$\min_{\mathbf{w}} \min_{\gamma \geq 0} \{ \gamma \rho + \mathbb{E}_{\mathbf{v} \sim \hat{P}} [\phi(\mathbf{v}, \mathbf{w}, \boldsymbol{\lambda}, \gamma)] \}. \quad (12)$$

This equivalence conclusion is obtained by the strong duality property of Wasserstein distributionally robust optimization problem [24], where the original optimization over the worst-case distribution in problem (10) is equivalently formulated as a mini-max problem (12) with definition (11) over the empirical distribution $\hat{P}$. A detailed proof for this is provided in Appendix A.1. Theorem 1 provides an equivalent formulation that does not include the non-differentiable distribution-type optimization variable P , and thus addressing the core challenge in the relaxed problem (10).

After relaxing and transforming the problem, we can formally present the adversarial training method for the transformed problem (12). We adopt the mini-batch gradient descent method [54], assuming that the dataset $\mathcal{D}$ can be divided into L sub-datasets $\{\mathcal{D}_1, \mathcal{D}_2, \dots, \mathcal{D}_L\}$, where the size of each dataset is B . In each batch, we only use one subset $\mathcal{D}_l$ to update the variables. Note that for any $\mathcal{D}_l$, we can expand the expectation terms over the distribution $\hat{P}(\mathcal{D}_l)$ as $\mathbb{E}_{\mathbf{v} \sim \hat{P}(\mathcal{D}_l)} [\dots] = \sum_{\hat{\mathbf{v}} \in \mathcal{D}_l} (\dots) / B$, then our objective in problem (12) with any subset $\mathcal{D}_l$ can be rewritten as:

$$h_l(\mathbf{w}, \boldsymbol{\lambda}, \gamma) = \gamma \rho + \frac{1}{B} \sum_{\hat{\mathbf{v}} \in \mathcal{D}_l} \phi(\hat{\mathbf{v}}, \mathbf{w}, \boldsymbol{\lambda}, \gamma). \quad (13)$$

Then our adversarial training process can be summarized in Algorithm 1. The outer loop of this process is similar to the standard training method [2]. In each iteration, we go through each batch and optimize the model parameters $\mathbf{w}$ using the objective function in (13). The only difference is that we need to gradually increase $\boldsymbol{\lambda}$ during training so that the solution to the objective problem (10) can approximate the original problem. Specifically, we update $\boldsymbol{\lambda}$ using gradient ascent method, and control the increasing rate through the learning rate $\eta_{\boldsymbol{\lambda}}$ and updating criterion. These correspond to the loop in Lines 2-3 and the criterion in Lines 9-10 in Algorithm 1.

Next is the core part of our adversarial training method, that is, optimizing model parameters $\mathbf{w}$ based on the objective h_l . Note that the origin problem in (12) contains multiple min or max operations, which, from outer to inner, correspond to the optimization variables

$\mathbf{w}$, γ , and the adversarial variable $\mathbf{v}'$ in ϕ as well as each $\mathbf{v}'_i$ in rgt_i . Correspondingly, our process updates these variables from the inside out. First, before updating the model parameters $\mathbf{w}$, we perform T updates of the inner variables, and then keep them fixed while updating $\mathbf{w}$ to replace the inner min or max operations. Similarly, before each update of γ , we also need to update the inner adversarial variables. Unlike the fixed number of updates above, we continuously update the adversarial variables until the change in ϕ is less than a certain threshold ϵ . We refer to the adversarial variables $\mathbf{v}'$ in ϕ and each $\mathbf{v}'_i$ in rgt_i at this point as ϵ -approximate maximizers. In particular, we updated all adversarial variables simultaneously during the training, and practical results have shown that this approach is feasible. The corresponding pseudo-code can be found in Lines 4-8 of Algorithm 1, where $\eta_{\mathbf{w}}$ and η_{γ} are learning rates for updating $\mathbf{w}$ and γ , respectively. Thus, we have comprehensively introduced the process of our Algorithm 1.

4.2 Robust Coefficient Selection

In the previous discussion, we have introduced how to solve the robust optimization problem, given the robust coefficient ρ . Next, we will introduce the methods for selecting this robust coefficient for the generalization and performance objectives, respectively.

4.2.1 Generalization-Guided Selection. To illustrate the logic of selecting the robust coefficient through generalization, we first intuitively explain the relation between the robust optimization problem and generalization performance. Intuitively, given the dataset $\mathcal{D}$ and the robust coefficient ρ (which essentially defines an ambiguity set $\mathcal{P}$), we can compute the probability that the true distribution belongs to $\mathcal{P}$ based on the optimal transport theory [29, 59]. Consequently, under this probability, the worst-case performance we optimize in (12) can serve as a lower bound of generalization performance. Conversely, if we hope to provide a degree of guarantee for generalization performance, we can deduce an appropriate robust coefficient for a given dataset $\mathcal{D}$. This is essentially the core idea of applying robust optimization to data-driven scenarios [6, 44].

We formally present this selection method. We first provide the definition of the generalization performance metric. Given the model parameters $\mathbf{w}$, our metric can be formally defined as the expected Lagrangian function value under the true distribution P_0 . Fixed the Lagrange multipliers λ , we denote this metric as:

$$J(\mathbf{w}, \lambda, P_0) = \mathbb{E}_{\mathbf{v} \sim P_0} [\text{lag}(\mathbf{v}, \mathbf{w}, \lambda)]. \quad (14)$$

Then the generalization goal can be described as for some given λ and probability $\beta \in [0, 1]$, we aim to obtain the model parameters $\hat{\mathbf{w}}$ based on the dataset, as well as guarantee $\hat{J}$, such that:

- **Generalization Bound:** We provide guarantee for generalization performance $J(\hat{\mathbf{w}}, \lambda, P_0) \leq \hat{J}$ with probability β .
- **Asymptotic Consistency:** As the size of the dataset $|\mathcal{D}| \rightarrow \infty$, our guarantee approaches optimality $\hat{J} \rightarrow \min_{\mathbf{w}} J(\mathbf{w}, \lambda, P_0)$.

These two metrics are widely adopted to measure generalization performance [26]. We can now formally introduce the selecting method that meets the above generalization requirements:

THEOREM 2. *Let c be l_1 -norm function, then for any compact value space $\mathcal{V}$, for any size of dataset $|\mathcal{D}|$ and β , we can construct an appropriate robust coefficient $\hat{\rho}$, such that the optimal solution and the corresponding optimization variables of problem (10) can be used*

as the guarantee $\hat{J}$ and the model parameters $\hat{\mathbf{w}}$. Such chosen $\hat{\mathbf{w}}$ and $\hat{J}$ satisfy both Generalization Bound and Asymptotic Consistency.

The idea of satisfying the *Generalization Bound* requirement here is consistent with the previous intuitive explanation. Moreover, when the amount of sampled data $|\mathcal{D}|$ is sufficiently large and the empirical distribution approaches the true distribution P_0 , we can choose a small ρ to reduce the robustness and approximate the optimal result, thus satisfying *Asymptotic Consistency* requirement. A detailed proof can be found in Appendix A.2.

However, although this selection method has a strong theoretical explanation for generalization, it is not convenient to apply in practice. The main issue lies in the guarantees provided in the form of probabilities. On one hand, the value of $\hat{\rho}$ obtained from Theorem 2 is overly conservative, resulting in relatively poor auction performance in practice. This conservative theoretical outcome has also been corroborated by the work of [44], a study that applies data-driven distributionally robust optimization to convex optimization problems. In fact, for a given probability, [44] employed an alternative bootstrap method [21] to obtain an empirically suitable $\hat{\rho}$. On the other hand, we find that the adversarial objectives we compute during training are underestimated, which is similar to the underestimate of regret in existing works [3, 13]. This affects the solution obtained for the problem (10), making it difficult to provide a precise guarantee $\hat{J}$ in terms of *Generalization Bound*.

It is worth noting that, despite the above limitations of this generalization-guided selection, its underlying concept still holds significant reference value. Firstly, it reveals that our robust method is not only focused on optimizing worst-case performance but also possesses a certain degree of generalization capability. This provides a solid theoretical foundation for the effectiveness of our method in uncertain environments. In addition, it guides us to employ the Wasserstein distance. Note that the generalization property of our method in Theorem 2 is established under the assumption of using the Wasserstein distance. In contrast, some other well-known distance functions, such as the KL divergence [34], fail to derive similar results. This factor, along with the widespread use of the Wasserstein distance in distributionally robust optimization [24, 33, 44], and the feasibility of solving corresponding problems through adversarial training, led us to choose the Wasserstein distance as the distance function for the ambiguity set.

4.2.2 Performance-Guided Selection. If we disregard the generalization and focus on performance only, we can simply regard the robust coefficient ρ as a hyper-parameter and use model selection methods to choose it. In this work, we mainly consider the cross-validation method [32], which is a classic approach for model selection. Its core idea is to split the dataset into parts, using one part for training and another for performance evaluation. For our problem, consistent with the definition in (14), given model parameters $\mathbf{w}$, validation set $\mathcal{D}^{va}$ and Lagrangian multipliers λ , the model performance can be written as $J(\mathbf{w}, \lambda, \hat{P}(\mathcal{D}^{va}))$. Therefore, we can select the best-performing ρ according to the cross-validation method.

This performance-guided selection method is more feasible in practice and is also more scalable. Specifically, we consider the online environment as an example, where the true distribution P_0 would change over time, which is a widely prevalent auction environment in practice. In this context, the selection of ρ has to

Algorithm 2: Repeated Holdout for Robust Coefficient

Input: Datasets $\{\mathcal{D}_t^{tr}\}, \{\mathcal{D}_t^{va}\}$, time points $\mathcal{T}$, candidates C
Parameters: Lagrangian multipliers λ

```

1 Initialize average performances  $\bar{J} \leftarrow \{\}$ ;
2 foreach  $\rho \in C$  do
3    $\bar{J}[\rho] \leftarrow 0$ ;
4   foreach  $t \in \mathcal{T}$  do
5     Solve (7) on  $\mathcal{D}_t^{tr}$  and get model parameters  $\mathbf{w}_t^*$ ;
6      $\bar{J}[\rho] \leftarrow \bar{J}[\rho] + J(\mathbf{w}_t^*, \lambda, \hat{P}(\mathcal{D}_t^{va}))$ ;
7    $\bar{J}[\rho] = \bar{J}[\rho] / |\mathcal{T}|$ ;
8 Return  $\arg_{\rho} \max \bar{J}[\rho]$ ;
```

consider not only the difference between the empirical distribution $\hat{P}$ and the true distribution P_0 , but also the variations of the true distribution P_0 itself, therefore previous generalization-guided method might not be able to yield a good robust coefficient.

To select ρ for performance in the online environment, we follow the idea of Repeated Holdout [10, 48, 60] strategy, which has been shown to achieve relatively excellent results in the unstable online environment [10]. We consider a common strategy of online model updating, where we update the model with some recent data at regular intervals, and then evaluate the model over the following period. Formally, we select a set of split time points $\mathcal{T}$ to divide the dataset into training sets and test sets. For any given split time point $t \in \mathcal{T}$, we take some most recent data from a period before t as the training set $\mathcal{D}_t^{tr}$, and take some following data from a period after t as the validation set $\mathcal{D}_t^{va}$. Then for each given ρ , we iterate over $t \in \mathcal{T}$, training on each training set $\mathcal{D}_t^{tr}$ to obtain the model parameters $\mathbf{w}_t^*$, and then evaluate the performance on each validation set $\mathcal{D}_t^{va}$. After iterations, the average performance over all the split time point in $\mathcal{T}$ is taken as the evaluation performance $\bar{J}_\rho$ for each candidate ρ . Finally, we select the robust coefficient $\hat{\rho}$ with the largest performance $\bar{J}_{\hat{\rho}}$ as the best robust coefficient. The whole process above can be summarized as Algorithm 2.

4.3 Stable Adversarial Training

The previous GAT method and the robust coefficient selection methods have provided a comprehensive solution to the robust auction design problem (7). However, the IC performance of this solution still lacks in-depth investigation. For example, why does the robust IC constraint in (7) yield better IC performance for uncertain data-driven contexts? As the robust coefficient within this constraint is adjustable, can we provide a stricter and unified constraint for a stronger IC guarantee? With these ideas, we propose the **SAT** (Stable Adversarial Training) method, which is detailed below.

We first explain why the robust IC constraint works in uncertain data-driven environments. Note that the regret constraint $rgt_i(\mathbf{v}, \mathbf{w}) = 0$ essentially provides a “single-point” Dominant Strategy Incentive Compatibility (DSIC) guarantee [53] for each given data point $\mathbf{v}$. However, in uncertain environments, the sampled data would be too sparse or biased, making it difficult for the DSIC guarantee to cover the entire value space $\mathcal{V}$. The results in Section 5.2, where existing methods significantly violate the IC guarantee,

also illustrate this point. In contrast, the robust IC constraint considers the distributions near the empirical distribution through the conception of the ambiguity set $\mathcal{P}$, thereby better extending the DSIC guarantee to cover the entire value space $\mathcal{V}$.

Following this line of idea, recall that a larger ρ implies that the uncertainty set $\mathcal{P}$ includes more data distributions. Therefore, when ρ is sufficiently large, does the robust IC constraint imply DSIC over the entire value space? The answer is positive:

THEOREM 3. *For any compact value space $\mathcal{V}$ and any data distribution $\hat{P}$, there exist some $\hat{\rho}$ such that the robust IC guarantee with ambiguity set $\mathcal{P}(\hat{P}, \rho)$ implies DSIC guarantees on value space $\mathcal{V}$.*

The key to this conclusion lies in the fact that, over a “finite” space, the Wasserstein distance from any unimodal distribution to the empirical distribution is always finite, allowing it to be captured by the ambiguity set. In this context, robust IC under any unimodal distribution implies DSIC, which leads to our Theorem 3. We defer the rigorous proof to Appendix A.3. This result provides us with a stricter IC constraint. However, calculating $\hat{\rho}$ for each $\mathcal{V}$ is still cumbersome. Therefore, we hope to find a unified constraint to avoid this calculation. We have the following result:

THEOREM 4. *For any bidder i , any model parameters $\mathbf{w}$, any distribution $\hat{P}$ and any robust coefficient $\hat{\rho}$, we have:*

$$rrgt_i(\mathbf{w}, \hat{P}, \hat{\rho}) \leq \mathbb{E}_{\mathbf{v} \sim \hat{P}}[\phi_i(\mathbf{v}, \mathbf{w})], \quad (15)$$

where for each i , ϕ_i is the robust version of rgt_i without penalty:

$$\phi_i(\mathbf{v}, \mathbf{w}) = \max_{\mathbf{v}'} \{rgt_i(\mathbf{v}', \mathbf{w})\}. \quad (16)$$

This conclusion similarly applies the robust-objective transformation idea in Theorem 1; however, here we remove the robust coefficient $\hat{\rho}$ and further obtain a unified upper bound. The detailed proof is provided in Appendix A.4. Note that the right-hand side of inequality (15) provides a strict and unified IC constraint without the specification of $\mathcal{V}$ and $\hat{\rho}$, which is exactly what we expect.

Based on this insight, we propose the SAT method as follows. We consider a variation of the robust auction design problem (7):

$$\begin{aligned} \max_{\mathbf{w}} \quad & \mathbb{E}_{\mathbf{v} \sim \hat{P}}[rev(\mathbf{v}, \mathbf{w})] \\ \text{s.t.} \quad & \mathbb{E}_{\mathbf{v} \sim \hat{P}}[\phi_i(\mathbf{v}, \mathbf{w})] = 0, \forall i \in \mathcal{N}, \end{aligned} \quad (17)$$

where we take the revenue on the empirical distribution as the objective, while introducing the strict IC constraint obtained in this section. Then similar to GAT, we can optimize its Lagrangian function to approximate this original objective:

$$\min_{\mathbf{w}} \left\{ -\mathbb{E}_{\mathbf{v} \sim \hat{P}}[rev(\mathbf{v}, \mathbf{w})] + \sum_{i \in \mathcal{N}} \lambda_i \mathbb{E}_{\mathbf{v} \sim \hat{P}}[\phi_i(\mathbf{v}, \mathbf{w})] \right\}. \quad (18)$$

Then we can solve problem (18) though a way similar to Algorithm 1, where objective function for any batch $\mathcal{D}_l$ can be written as:

$$h'_l(\mathbf{w}, \lambda) = \frac{1}{B} \sum_{\hat{\mathbf{v}} \in \mathcal{D}_l} rev(\hat{\mathbf{v}}, \mathbf{w}) + \sum_{i \in \mathcal{N}} \lambda_i \cdot \frac{1}{B} \sum_{\hat{\mathbf{v}} \in \mathcal{D}_l} \phi_i(\hat{\mathbf{v}}, \mathbf{w}). \quad (19)$$

Thus, we can use adversarial training to solve problem (18), as detailed in Algorithm 3, which is a simple version of Algorithm 1.

Note that for each i , ϕ_i in (16) is a straightforward DSIC requirement and is theoretically independent of the origin data point $\mathbf{v}$.

Algorithm 3: Stable Adversarial Training Approach**Input:** Mini-batches $\{\mathcal{D}_1, \mathcal{D}_2, \dots, \mathcal{D}_L\}$ **Parameters:** $\epsilon, \eta_w, \eta_\lambda$

```

1 Initialize  $w, \lambda$ ;
2 while stopping criterion not met do
3   for  $l = 1$  to  $L$  do
4     foreach  $\hat{v} \in \mathcal{D}_l$  do
5       for  $i = 1$  to  $N$  do
6          $\phi_i \leftarrow \epsilon$ -approximate maximizers for  $\phi_i$ ;
7        $w \leftarrow w - \eta_w \nabla_w h'_i(w, \lambda)$ ;
8   if updating  $\lambda$  criterion met then
9      $\lambda \leftarrow \lambda + \eta_\lambda \nabla_\lambda h'_i(w, \lambda)$ ;

```

In practice, we initialize the adversarial variable v' with v , thus allowing the dataset $\mathcal{D}$ to influence our adversarial training results.

It can be seen that SAT is more concise compared to GAT, while improving the IC performance. Moreover, although there is no theoretical guarantee, the auctions obtained from SAT exhibit excellent and stable performance. Therefore we use “stable” to describe this method to distinguish it from GAT. We leave the related experiment results and more discussion for this in Section 5.2.

5 Experiment Results

In this section, we test the previously proposed GAT method, robust coefficient selection methods, and SAT method. We first introduce our experimental setup, followed by the experiment results conducted under the limited data and distribution shift scenarios.

5.1 Experimental Setup

5.1.1 Bidder Settings. We first introduce the bidder settings in our experiment. We mainly consider two kinds of bidder valuations, which reflect the value of the allocated set of items to the bidder:

- **Additive:** For any bidder i , the value v_{ij} of any subset S_j is equal to the sum of the values of all items in the subset.

- **Unit-demand:** For any bidder i , the value v_{ij} of any subset S_j is equal to the max value of a single item within the subset.

Thus, the bidders participating in the auction in our experiments can mainly be divided into the following two categories:

- **$N \times M$ additive:** N bidders competing for M items, where each bidder has additive valuations, and the value of each item for each bidder follows a uniform distribution over $[0, 1]$.

- **1x2 unit-demand:** 1 bidders competing for 2 items, where the bidder has unit-demand valuation, and the value of each item follows a uniform distribution over $[2, 3]$. The optimal auction for this setting is provided in [47].

5.1.2 Metrics. We mainly consider the following metrics for the auction on the data distribution of the test set $\hat{P}_{te}$:

- **Regret:** Sum of regrets $\sum_{i=1}^N \mathbb{E}_{v \sim \hat{P}_{te}} [rt_i(v, w)]$;

- **Revenue:** Total revenue $\mathbb{E}_{v \sim \hat{P}_{te}} [rev(v, w)]$;

- **Loss:** Lagrange function value $\mathbb{E}_{v \sim \hat{P}_{te}} [lag(v, w, \lambda)]$;

where the Lagrange multipliers λ are set to 100 for all evaluation of loss. Note that in our experiment, all test sets contained sufficient

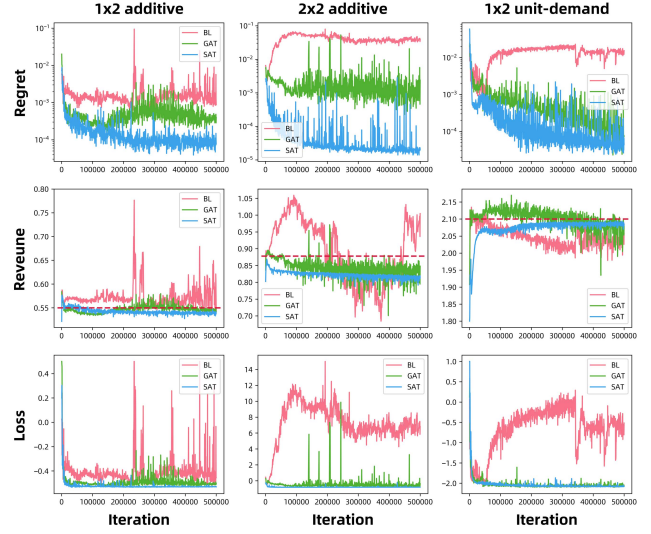


Figure 1: The metrics during training under different settings.

data, enabling the distribution $\hat{P}_{te}$ to closely approximate the true distribution P_0 . This allowed the loss metric to accurately reflect the generalization performance metric $J(w, \lambda, P_0)$.

5.1.3 Methods. We compare three training methods in our work:

- **BL:** The widely adopted baseline method proposed in [17].
- **GAT:** Our general adversarial training method in Section 4.1.
- **SAT:** Our stable adversarial training method in Section 4.3.

It is worth noting that these methods differ only in their training processes and impose no additional requirements on the model structure being trained. Therefore, just as the training method in [17] can be extended to multiple domains [42, 52, 61], our robust training method can also be similarly extended. Moreover, since the RegretNet model structure proposed in [17] has been widely discussed and validated, we adopt this solid model structure in all our experiments to evaluate the above methods.

5.2 Evaluation for Limited Data

We first compare the performance of the three methods under limited data context. Figure 1 shows the results for some settings with 100 sampled data as the training set, where GAT uniformly uses $\rho = 0.05$ as the robust coefficient. We can observe that:

- For the regret metric in the first row, our SAT and GAT methods have both achieved the performance far exceeding that of BL. In particular, the regret of BL is as low as nearly 0.05 for *2x2 additive*, which means that bidders can, on average, gain an additional utility of 0.05 by misreporting. Considering that the values are uniformly distributed in $[0, 1]$, this indicates a significant violation of IC. In contrast, the regrets of SAT and GAT are at the level of 10^{-3} or lower, both providing excellent IC guarantees.

- For the revenue metric in the second row, both GAT and SAT achieve stable results close to optimal (red dashed line). In contrast, the result of BL is quite unstable, indicating a potential overfitting issue. Interestingly, we did not explicitly incorporate a robust revenue objective in SAT; however, SAT still achieved very good

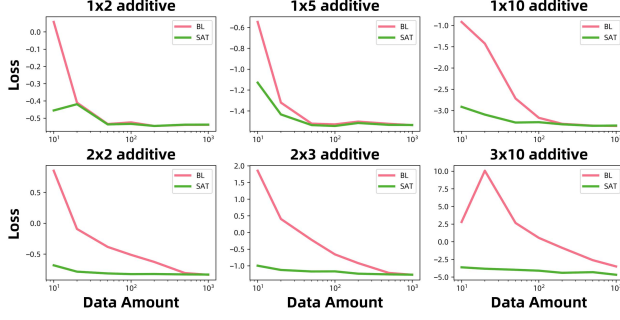


Figure 2: Variation of the loss metrics with data amount under different settings.

generalization results. This may be because the robust regret target of SAT potentially influences the selection of the payment function p^w , leading to better generalization of revenue performance.

- Finally, the losses achieved by GAT and SAT, as displayed in the third row, are significantly lower than that of BL, and their loss on the test set are able to reach a state of convergence, indicating that our auction has better generalization performance. In particular, SAT exhibits very stable generalization performance, which may be attributed to its relatively simple training process.

Overall, our method demonstrates strong generalization performance across all three metrics. In addition, we conducted experiments in more settings, and the results were consistent with previous observations. Meanwhile, we observed that under appropriate hyper-parameter selection, the runtime of the GAT and SAT methods generally does not exceed twice that of the BL method. This indicates that although our training process is more complex, the resulting time overhead is acceptable. Additionally, since our method does not alter the model structure, the memory usage is almost identical to that of BL. We leave more experimental results and hyper-parameter selections in Appendix B.1.

5.3 Impact of Data Amount

Next, we analyze the impact of data amount under the limited data context. For simplicity of the experiment and intuitive comparison of the performance, we only conducted comparisons between SAT and BL. It is important to note that as shown in Figure 1, BL is prone to overfitting when the data amount is limited. For convenience, we overlook this issue and directly use the auction that exhibited the best generalization performance during training as the upper limit of BL. Under these considerations, we test SAT and BL under different data amount for some settings, and the results are shown in Figure 2. We can see that as the number of bidders and items increases, BL requires more and more data to approach the generalization performance of SAT. Specifically, for 3×10 additive, the generalization performance of BL on 1000 sampled data is only able to get close to that of SAT on 10 sampled data. This result well demonstrates the excellent generalization performance of SAT. In addition, we can see that when the data amount is large enough, the performance of SAT can also approach optimal results as BL. This shows that SAT may also exhibit the empirical observation of a "no local optima" phenomenon as BL [18].

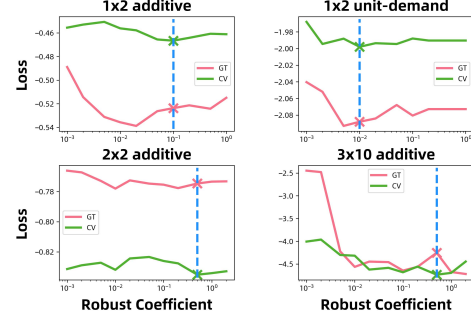


Figure 3: Variation of the loss metrics with robust coefficient under different settings.

5.4 Impact of Robust Coefficient

We next focus on the performance-guided selection for the robust coefficient in the GAT method. We consider a simple case, where the amount of training data is limited to 100 (1000 for 3×10 additive), and we aim to select such robust coefficients so that the auction trained through the GAT method exhibits good generalization performance. To this end, we consider a simple cross-validation method [32], where the dataset is divided into a training set and a validation set in an 8:2 ratio, and we record the optimal loss achieved on the validation set during training with different robust coefficients. The results are shown by the green line in Figure 3. The cross-validation method selects the robust coefficient corresponding to the optimal loss in the validation set, and this choice is indicated by the blue dashed line. For evaluation, we calculate the performance of the auction trained with 100 samples on the test set under different robust coefficients, serving as the ground truth performance for different robust coefficients. The results are shown by the red line in Figure 3. It can be seen that simple cross-validation method generally can help avoid some poor robust coefficient choices.

5.5 Evaluation for Distribution-Shift

We finally consider real-world auction environments, where we have a sufficient amount of sampled data; however, the true value distribution may vary over different times and scenarios:

5.5.1 Distribution-Shift Over Time. We first consider the case where the distribution changes over time. This is a common environment in practice, where a natural solution is to assume that the distribution change will not be too large, and to periodically update the model based on recent data, as described in Section 4.2.2. We aim to explore the performance of GAT and BL methods when periodically updating the model in such an uncertain environment.

We first clarify the uncertain environment and the model update strategy. We selected 10 bidders from the real-world auction dataset and collected their value data between 04:00 and 22:00 in one day¹. Note that such auction scales are consistent with our practice and align with existing reports from the ad platforms [36, 37, 39]. Here, for normalization, we process and scale the values to the range of [0,

¹In practice, since most advertisers rely on the models (such as CTR and CVR models) provided by the platform to assess their value, the platform can retrieve value data after auctions. This background is consistent with existing works [37, 40, 61].

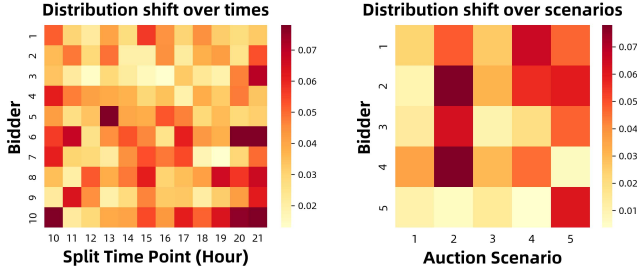


Figure 4: Distribution shift in the real-world auction environments.

Table 1: Revenue for Different Model Update Times (10^{-1})

Time	16	17	18	19	20	21	Avg.
BL	6.455	6.768	6.572	6.630	6.771	6.625	6.637
GAT	<u>6.575</u>	<u>6.894</u>	<u>6.615</u>	<u>6.678</u>	<u>6.842</u>	<u>6.659</u>	<u>6.711</u>

1]. Thus, we naturally update the model on the hour, selecting the most recent 6 hours of data for training, and then evaluating over the next 1 hour. Following this, we can quantify the uncertainty of the environment by comparing the distance between the two data distributions separated by each hour, as shown in the left heatmap in Figure 4. For example, the color block in the top-left corner represents the distance between the two value distributions for bidder 1, separated at 10:00 (distributions over 04:00-10:00 and 10:00-11:00). We can see that the distance is around 0.02-0.08, indicating a significant distribution change relative to the value range of $[0, 1]$.

Next, we evaluate the performance of GAT and BL. We simulate the 10×1 additive setting. For BL, we can directly evaluate it via the model update strategy above. For GAT, we also need to specify ρ . We run Algorithm 2 on the datasets divided by hours $\mathcal{T} = \{10, 11, \dots, 15\}$, where for each $t \in \mathcal{T}$, the training set $\mathcal{D}_t^{tr}$ contains data in previous 6 hours, and the validation set $\mathcal{D}_t^{va}$ contains data in next 1 hour. The details of the execution are provided in Appendix B.2, and we obtain $\rho = 0.03$ as the parameter for GAT.

Finally, we can perform model updates for the remaining hours 16, 17, ..., 21 and evaluate BL and GAT. For fairness, we select the maximum revenue achievable when the regret is less than 10^{-4} as the metric, and the results are shown in Table 1. We can see that GAT outperformed BL in all update times, while achieving an average improvement of 1.1%. This effectively show the excellent and stable performance of our GAT in this uncertain environment.

5.5.2 Distribution-Shift Over Scenarios. In ad platforms, there are often multiple auction scenarios, such as homepages, shopping cart pages, and order pages in e-commerce apps. For stability, platforms often employ a unified auction model across all scenarios. However, the bidders' value distributions may vary significantly over scenarios, and the unified model needs to be trained on the mixed distribution of these scenarios. In this context, we will better demonstrate the robustness and application value of GAT.

We first illustrate the distribution differences over scenarios. We selected value data from 5 bidders in 5 scenarios for one day, and

Table 2: Revenue for Different Auction Scenarios (10^{-1})

Scenario	1	2	3	4	5	Mixed
BL	<u>4.126</u>	4.943	<u>3.746</u>	4.717	3.620	<u>4.081</u>
GAT	4.094	<u>5.005</u>	3.701	<u>4.744</u>	<u>3.626</u>	4.074

calculated the distances between each bidder's value distribution in each scenario and the mixed distribution of this bidder in all scenarios, as shown in the right heatmap of Figure 4. We can see that the distributions in scenarios 1, 3 are close to the mixed distributions, while there are considerable differences in scenarios 2, 4, 5.

We now can evaluate BL and GAT in the mixed distribution of all bidders, where we take $\rho = 0.05$ as an example in GAT. The result is shown in Table 2. For fairness, we select the auctions that achieve the maximum revenue on the mixed distribution when regret is less than 5×10^{-4} , while also calculating their revenue in each scenario.

As we can see, since both training and evaluation used the mixed distribution, BL performs slightly better than GAT on it. However, GAT achieves better revenue than BL in scenarios 2, 4, 5, where the distribution differences are larger, as shown in Figure 4. This highlights the robustness of GAT and reflects a degree of fairness in practice. Specifically, the mixed distribution is naturally closer to distributions of scenarios with more auctions, while tending to differ from ones with fewer auctions. Our example precisely illustrates this, as scenarios 1, 3, where the distribution differences are small, have more auctions than other scenarios. From this perspective, BL naturally focuses on the more populated scenarios 1, 3. In contrast, GAT demonstrates a more equitable focus on the less populated scenarios 2, 4, 5. This fairness is beneficial for better assessing the revenues of small or new scenarios. In this regard, GAT can offer a more equitable performance for scenarios, including new ones, thus promoting diverse development within the ad platforms.

6 Conclusion

In this work, we study the important optimal auction design problem in uncertain data-driven environments. To derive moderately robust auctions based on uncertain sampled data, we primarily propose three techniques, including the GAT method to solve the robust data-driven auction design problem, the robust coefficient selection methods to select robust auctions, and the extended SAT method to achieve a unified and strict IC guarantee. We evaluate these methods in two types of uncertain environments that are commonly encountered in practice: limited data and distribution shift. The experimental results on both the constructed and real-world datasets demonstrate the effectiveness of the proposed techniques.

Acknowledgments

This work was supported in part by National Key R&D Program of China (No. 2023YFB4502400), in part by China NSF grant No. 62322206, 62025204, 62132018, U2268204, 62272307, 62372296, in part by Alibaba Group through Alibaba Innovative Research Program. The opinions, findings, conclusions, and recommendations expressed in this paper are those of the authors and do not necessarily reflect the views of the funding agencies or the government.

References

- [1] Michael Albert, Vincent Conitzer, and Peter Stone. 2017. Automated Design of Robust Mechanisms. In *Proceedings of the AAAI Conference on Artificial Intelligence*. 298–304.
- [2] Shun-ichi Amari. 1993. Backpropagation and stochastic gradient descent method. *Neurocomputing* 5, 4-5 (1993), 185–196.
- [3] Anish Athalye, Nicholas Carlini, and David Wagner. 2018. Obfuscated gradients give a false sense of security: Circumventing defenses to adversarial examples. In *International Conference on Machine Learning*. 274–283.
- [4] Chaithanya Bandi and Dimitris Bertsimas. 2014. Optimal design for multi-item auctions: A robust optimization approach. *Mathematics of Operations Research* 39, 4 (2014), 1012–1038.
- [5] Dirk Bergemann and Stephen Morris. 2005. Robust mechanism design. *Econometrica* 73, 6 (2005), 1771–1813.
- [6] Dimitris Bertsimas, Vishal Gupta, and Nathan Kallus. 2018. Data-driven robust optimization. *Mathematical Programming* 167 (2018), 235–292.
- [7] Vladimir Igorevich Bogachev and Maria Aparecida Soares Ruas. 2007. *Measure theory*. Vol. 1. Springer.
- [8] Benjamin Brooks and Songzi Du. 2021. Optimal auction design with common values: An informationally robust approach. *Econometrica* 89, 3 (2021), 1313–1360.
- [9] Vinicius Carrasco, Vitor Farinha Luz, Nenad Kos, Matthias Messner, Paulo Monteiro, and Humberto Moreira. 2018. Optimal selling mechanisms under moment conditions. *Journal of Economic Theory* 177 (2018), 245–279.
- [10] Vitor Cerqueira, Luis Torgo, and Igor Mozetič. 2020. Evaluating time series forecasting models: An empirical study on performance estimation methods. *Machine Learning* 109, 11 (2020), 1997–2028.
- [11] Vincent Conitzer, Zhe Feng, David C Parkes, and Eric Sodomka. 2021. Welfare-preserving ϵ -bic to bic transformation with negligible revenue loss. In *International Conference on Web and Internet Economics*. 76–94.
- [12] Peter Cramton. 2013. Spectrum auction design. *Review of industrial organization* 42 (2013), 161–190.
- [13] Michael Curry, Ping-Yeh Chiang, Tom Goldstein, and John Dickerson. 2020. Certifying strategyproof auction networks. In *Advances in Neural Information Processing Systems*. 4987–4998.
- [14] Constantinos Daskalakis, Alan Deckelbaum, and Christos Tzamos. 2015. Strong duality for a multiple-good monopolist. In *Proceedings of the Sixteenth ACM Conference on Economics and Computation*. 449–450.
- [15] Constantinos Daskalakis and Seth Matthew Weinberg. 2012. Symmetries and optimal multi-dimensional mechanism design. In *Proceedings of the 13th ACM Conference on Electronic Commerce*. 370–387.
- [16] Zhijian Duan, Jingwu Tang, Yutong Yin, Zhe Feng, Xiang Yan, Manzil Zaheer, and Xiaotie Deng. 2022. A context-integrated transformer-based neural network for auction design. In *International Conference on Machine Learning*. 5609–5626.
- [17] Paul Dütting, Zhe Feng, Harikrishna Narasimhan, David Parkes, and Sai Srivatsa Ravindranath. 2019. Optimal auctions through deep learning. In *International Conference on Machine Learning*. 1706–1715.
- [18] Paul Dütting, Zhe Feng, Harikrishna Narasimhan, David C Parkes, and Sai Srivatsa Ravindranath. 2024. Optimal auctions through deep learning: Advances in differentiable economics. *J. ACM* 71, 1 (2024), 1–53.
- [19] Paul Dütting, Felix Fischer, Pichayut Jirapinyo, John K Lai, Benjamin Lubin, and David C Parkes. 2012. Payment rules through discriminant-based classifiers. In *Proceedings of the 13th ACM Conference on Electronic Commerce*. 477–494.
- [20] Benjamin Edelman, Michael Ostrovsky, and Michael Schwarz. 2007. Internet advertising and the generalized second-price auction: Selling billions of dollars worth of keywords. *American economic review* 97, 1 (2007), 242–259.
- [21] Bradley Efron and Robert J Tibshirani. 1994. *An introduction to the bootstrap*. Chapman and Hall/CRC.
- [22] Natalia Fabra, Nils-Henrik von der Fehr, and David Harbord. 2006. Designing electricity auctions. *The RAND Journal of Economics* 37, 1 (2006), 23–46.
- [23] Zhe Feng, Harikrishna Narasimhan, and David C. Parkes. 2018. Deep Learning for Revenue-Optimal Auctions with Budgets. In *Proceedings of the 17th International Conference on Autonomous Agents and MultiAgent Systems*. 354–362.
- [24] Rui Gao and Anton Kleywegt. 2023. Distributionally robust stochastic optimization with Wasserstein distance. *Mathematics of Operations Research* 48, 2 (2023), 603–655.
- [25] Yiannis Giannakopoulos and Elias Koutsoupias. 2018. Duality and optimality of auctions for uniform distributions. *SIAM J. Comput.* 47, 1 (2018), 121–165.
- [26] Trevor Hastie, Robert Tibshirani, Jerome H Friedman, and Jerome H Friedman. 2009. *The elements of statistical learning: data mining, inference, and prediction*. Vol. 2. Springer.
- [27] Christoph Hertrich, Yixin Tao, and László A. Végh. 2023. Mode Connectivity in Auction Design. In *Advances in Neural Information Processing Systems*. 52957–52968.
- [28] Dmitry Ivanov, Iskander Safiulin, Igor Filippov, and Ksenia Balabaeva. 2022. Optimal-er auctions through attention. In *Advances in Neural Information Processing Systems*. 34734–34747.
- [29] Ran Ji and Miguel A Lejeune. 2021. Data-driven optimization of reward-risk ratio measures. *INFORMS Journal on Computing* 33, 3 (2021), 1120–1137.
- [30] Alexander S. Kechris. 1995. *Polish Spaces*. Springer New York, 13–17.
- [31] Çağrı Kocyiğit, Garud Iyengar, Daniel Kuhn, and Wolfram Wiesemann. 2020. Distributionally robust mechanism design. *Management Science* 66, 1 (2020), 159–189.
- [32] Ron Kohavi et al. 1995. A study of cross-validation and bootstrap for accuracy estimation and model selection. In *International Joint Conference on Artificial Intelligence*. 1137–1145.
- [33] Daniel Kuhn, Peyman Mohajerin Esfahani, Viet Anh Nguyen, and Soroosh Shafieezadeh-Abadeh. 2024. Wasserstein Distributionally Robust Optimization: Theory and Applications in Machine Learning. arXiv:1908.08729
- [34] Solomon Kullback and Richard A Leibler. 1951. On information and sufficiency. *The annals of mathematical statistics* 22, 1 (1951), 79–86.
- [35] Kevin Kuo, Anthony Ostuni, Elizabeth Horishny, Michael J. Curry, Samuel Doolley, Ping yeh Chiang, Tom Goldstein, and John P. Dickerson. 2020. Proportion-Net: Balancing Fairness and Revenue for Auction Design with Deep Learning. arXiv:2010.06398
- [36] Ningyuan Li, Yunxuan Ma, Yang Zhao, Zhijian Duan, Yurong Chen, Zhilin Zhang, Jian Xu, Bo Zheng, and Xiaotie Deng. 2023. Learning-Based Ad Auction Design with Externalities: The Framework and A Matching-Based Approach. In *Proceedings of the 29th ACM SIGKDD Conference on Knowledge Discovery and Data Mining*. 1291–1302.
- [37] Xuejian Li, Ze Wang, Bingqi Zhu, Fei He, Yongkang Wang, and Xingxing Wang. 2024. Deep Automated Mechanism Design for Integrating Ad Auction and Allocation in Feed. arXiv:2401.01656
- [38] Fengming Lin, Xiaolei Fang, and Zheming Gao. 2022. Distributionally robust optimization: A review on theory and applications. *Numerical Algebra, Control and Optimization* 12, 1 (2022), 159–212.
- [39] Hongtao Liu, Luxi Chen, Yiming Ding, Changcheng Li, Han Li, Peng Jiang, and Weiran Shen. 2025. Balancing Revenue and Privacy with Signaling Schemes in Online Ad Auctions. In *Proceedings of the Eighteenth ACM International Conference on Web Search and Data Mining*. 512–520.
- [40] Xiangyu Liu, Chuan Yu, Zhilin Zhang, Zhenzhe Zheng, Yu Rong, Hongtao Lv, Da Huo, Yiqing Wang, Dagui Chen, Jian Xu, et al. 2021. Neural auction: End-to-end learning of auction mechanisms for e-commerce advertising. In *Proceedings of the 27th ACM SIGKDD Conference on Knowledge Discovery & Data Mining*. 3354–3364.
- [41] Giuseppe Lopomo, Luca Rigotti, and Chris Shannon. 2021. Uncertainty in mechanism design. arXiv:2108.12633
- [42] Nguyen Cong Luong, Le Khac Chau, Nguyen Do Duy Anh, Nguyen Huu Sang, Shaohan Feng, Van-Dinh Nguyen, Dusit Niyato, and Dong In Kim. 2024. Optimal Auction for Effective Energy Management in UAV-Assisted Vehicular Metaverse Synchronization Systems. *IEEE Transactions on Vehicular Technology* 73, 1 (2024), 1207–1222. doi:10.1109/TVT.2023.3302411
- [43] Alejandro M Manelli and Daniel R Vincent. 2006. Bundling as an optimal selling mechanism for a multiple-good monopolist. *Journal of Economic Theory* 127, 1 (2006), 1–35.
- [44] Peyman Mohajerin Esfahani and Daniel Kuhn. 2018. Data-driven distributionally robust optimization using the Wasserstein metric: Performance guarantees and tractable reformulations. *Mathematical Programming* 171, 1 (2018), 115–166.
- [45] Roger B Myerson. 1981. Optimal auction design. *Mathematics of operations research* 6, 1 (1981), 58–73.
- [46] Roger B. Myerson. 1989. *Mechanism Design*. Palgrave Macmillan UK, 191–206.
- [47] Gregory Pavlov. 2011. Optimal Mechanism for Selling Two Goods. *The BE Journal of Theoretical Economics* 11, 1 (2011), 1–35.
- [48] Richard R Picard and R Dennis Cook. 1984. Cross-validation of regression models. *J. Amer. Statist. Assoc.* 79, 387 (1984), 575–583.
- [49] Hamed Rahimian and Sanjay Mehrotra. 2019. Distributionally robust optimization: A review. arXiv:1908.05659
- [50] Jad Rahme, Samy Jelassi, Joan Bruna, and S Matthew Weinberg. 2021. A permutation-equivariant neural network architecture for auction design. In *Proceedings of the AAAI conference on artificial intelligence*. 5664–5672.
- [51] Jad Rahme, Samy Jelassi, and S. Matthew Weinberg. 2021. Auction learning as a two-player game. arXiv:2006.05684
- [52] Sai Srivatsa Ravindranath, Yanchen Jiang, and David C Parkes. 2023. Data Market Design through Deep Learning. In *Advances in Neural Information Processing Systems*. 6662–6689.
- [53] Tim Roughgarden. 2010. Algorithmic game theory. *Commun. ACM* 53, 7 (2010), 78–86.
- [54] Sebastian Ruder. 2017. An overview of gradient descent optimization algorithms. arXiv:1609.04747
- [55] Yuko Sakurai, Satoshi Oyama, Mingyu Guo, and Makoto Yokoo. 2019. Deep false-name-proof auction mechanisms. In *PRIMA 2019: Principles and Practice of Multi-Agent Systems*. 594–601.
- [56] Tuomas Sandholm and Anton Likhodedov. 2015. Automated design of revenue-maximizing combinatorial auctions. *Operations Research* 63, 5 (2015), 1000–1025.
- [57] Alex Suzdaltsev. 2020. An optimal distributionally robust auction. arXiv:2006.05192

- [58] SS Vallender. 1974. Calculation of the Wasserstein distance between probability distributions on the line. *Theory of Probability & Its Applications* 18, 4 (1974), 784–786.
- [59] Cédric Villani et al. 2009. *Optimal transport: old and new*. Vol. 338. Springer.
- [60] Qing-Song Xu and Yi-Zeng Liang. 2001. Monte Carlo cross validation. *Chemometrics and Intelligent Laboratory Systems* 56, 1 (2001), 1–11.
- [61] Ruitao Zhu, Yangsu Liu, Dagui Chen, Zhenjia Ma, Chufeng Shi, Zhenzhe Zheng, Jie Zhang, Jian Xu, Bo Zheng, and Fan Wu. 2024. Contextual Generative Auction with Permutation-level Externalities for Online Advertising. arXiv:2412.11544

A Proofs

A.1 Proof of Theorem 1

Our conclusion is a direct derivation of Gao and Kleywegt’s work [24]. We first repeat the main contributions of their work:

THEOREM 5 (MAIN CONTRIBUTIONS IN [24]). *If $(\mathcal{V}, c)$ forms a Polish space and the lag function (9) is measurable, for fixed $\mathbf{w}, \lambda$:*

$$\max_{P \in \mathcal{P}(\hat{P}, \rho)} \mathbb{E}_{v \sim P} [\text{lag}(v, \mathbf{w}, \lambda)] = \min_{\gamma \geq 0} \{\gamma \rho + \mathbb{E}_{v \sim \hat{P}} [\phi(v, \mathbf{w}, \lambda, \gamma)]\}.$$

We first clarify that $(\mathcal{V}, c)$ constitutes a Polish space. Here, we make a simple assumption about the value space, that is, any v_{ij} will not exceed a certain upper limit $\bar{v}$, thus $\mathcal{V} = [0, \bar{v}]^{N \times 2^M}$. This assumption is reasonable because there are practically no infinitely large values, and it aligns with the setup of our experiment. Then since every finite-dimensional normed vector space is a Banach space, while $\mathcal{V}$ is also a separable space, then by definition $(\mathcal{V}, c)$ constitutes a Polish space [30].

We next demonstrate that the *lag* function is measurable. We consider the general case that the function $a^{\mathbf{w}}$ and $p^{\mathbf{w}}$ are continuous with respect to v . Then after addition, subtraction, and max operations, the resulting *lag* function is also continuous, thereby is a Borel measurable function (Example 2.1.4 in [7]). Moreover, our value space $\mathcal{V}$ is closed, which is a Borel set. Thus, the *lag* function is measurable.

Thus, according to Theorem 5, we can naturally obtain the result of Theorem 1, and the proof is concluded.

A.2 Proof of Theorem 2

We first rigorously present the conclusion from Ji and Lejeune’s work [29], which indicate that given the compact value space $\mathcal{V}$, the sampled data amount $|\mathcal{D}|$, and a specified probability requirement β , it is possible to calculate an appropriate robust coefficient such that the true distribution has a probability greater than β within the ambiguity set. Their conclusions can be stated as:

THEOREM 6 (THEOREM 2 IN [29]). *For compact space $\mathcal{V}$, sampled data amount $|\mathcal{D}|$ and robability requirement β , let $B = \bar{v} \times N \times 2^M$ be the diameter of $\mathcal{V}$, then the robust coefficient can be calculated as:*

$$\hat{\rho} = \left(B + \frac{3}{4}\right) \left(-\frac{1}{|\mathcal{D}|} \log(1 - \beta) + 2\sqrt{-\frac{1}{|\mathcal{D}|} \log(1 - \beta)}\right),$$

such that true distribution $P_0 \in \mathcal{P}(\hat{P}, \hat{\rho})$ with probability at least β .

Here, we continue the assumptions for $\mathcal{V}$ in Section A.1. Then let the guarantee $\hat{J}$ and the model parameters $\hat{\mathbf{w}}$ be:

$$\hat{J} = \max_{P \in \mathcal{P}(\hat{P}, \hat{\rho})} J(\hat{\mathbf{w}}, \lambda, P) = \min_{\mathbf{w}} \max_{P \in \mathcal{P}(\hat{P}, \hat{\rho})} J(\mathbf{w}, \lambda, P),$$

where the right term is exactly problem (10), and we can guarantee that:

$$\hat{J} = \max_{P \in \mathcal{P}(\hat{P}, \hat{\rho})} J(\hat{\mathbf{w}}, \lambda, P) \geq J(\hat{\mathbf{w}}, \lambda, P_0),$$

with probability at least β , thereby satisfying the *Generalization Bound* requirement. At the same time, $\rho \rightarrow 0$ as $|\mathcal{D}| \rightarrow +\infty$, and according to Theorem 6 we have:

$$\max_{P \in \mathcal{P}(\hat{P}, \hat{\rho})} J(\mathbf{w}, \lambda, P) \rightarrow J(\mathbf{w}, \lambda, P_0), \text{ as } \rho \rightarrow 0,$$

therefore as $|\mathcal{D}| \rightarrow +\infty$, we have $\hat{J} \rightarrow \min_{\mathbf{w}} J(\mathbf{w}, \lambda, P_0)$, thereby satisfying the *Asymptotic Consistency* requirement.

Combining the above results, we have proven Theorem 2.

A.3 Proof of Theorem 3

We first clarify that for a given value space $\mathcal{V}$ and empirical distribution $\hat{P}$, there exists some $\hat{\rho}$ such that all unimodal distributions on $\mathcal{V}$ are within the ambiguity set $\mathcal{P}(\hat{P}, \hat{\rho})$. For convenience, we continue the assumptions for $\mathcal{V}$ in Section A.1 and Section A.2, thus for any two value $v_1, v_2 \in \mathcal{V}$, the norm function value is bounded:

$$c(v_1 - v_2) \leq \bar{v} \sum_{i=1}^N \sum_{j=1}^{2^M} c(\mathbf{1}_{ij}),$$

where $\mathbf{1}_{ij}$ is the value vectors that satisfy $v_{ij} = 1$ while all other elements are 0. This conclusion is a simple derivation of the properties of norms. Therefore, by definition [59], the Wasserstein distance between any unimodal distribution $P(\{\hat{v}\})$ and $\hat{P}$ is also bounded:

$$d_c(\hat{P}, P(\{\hat{v}\})) \leq \bar{v} \sum_{i=1}^N \sum_{j=1}^{2^M} c(\mathbf{1}_{ij}),$$

where $P(\{\hat{v}\})$ is the unimodal distribution on $\hat{v}$. Therefore, we choose the value on the right as $\hat{\rho}$, then according to (4), every unimodal distribution $P(\{\hat{v}\}) \in \mathcal{P}(\hat{P}, \hat{\rho})$.

With such chosen $\hat{\rho}$, for any bidder i , the *rrgti* in (6) satisfies:

$$\begin{aligned} \text{rrgt}_i(\mathbf{w}, \hat{P}, \hat{\rho}) &= \max_{P \in \mathcal{P}(\hat{P}, \hat{\rho})} \mathbb{E}_{v \sim P} [\text{rgt}_i(v, \mathbf{w})] \\ &\geq \max_{\hat{v}} \mathbb{E}_{v \sim P(\{\hat{v}\})} [\text{rgt}_i(v, \mathbf{w})] \\ &= \max_{\hat{v}} \text{rgt}_i(\hat{v}, \mathbf{w}) \\ &\geq 0, \end{aligned}$$

hence the robust IC constraint in (7) implies that:

$$\text{rrgt}_i(\mathbf{w}, \hat{P}, \hat{\rho}) = \max_{\hat{v}} \text{rgt}_i(\hat{v}, \mathbf{w}) = 0,$$

which is exactly the DSIC constraint. Thus the proof is concluded.

A.4 Proof of Theorem 4

Similar to Proof A.1, for any bidder i , we replace the origin *lag* by *rgti* function, and the conclusion becomes:

$$\begin{aligned} \text{rrgt}_i(\mathbf{w}, \hat{P}, \hat{\rho}) &= \max_{P \in \mathcal{P}(\hat{P}, \hat{\rho})} \mathbb{E}_{v \sim P} [\text{rgt}_i(v, \mathbf{w})] \\ &= \min_{\gamma \geq 0} \{\gamma \rho + \mathbb{E}_{v \sim \hat{P}} [\phi'_i(v, \mathbf{w}, \gamma)]\}, \end{aligned}$$

where ϕ'_i is the same as ϕ defined in (11):

$$\phi'_i(v, \mathbf{w}, \gamma) = \max_{v'} \{\text{rgt}_i(v', \mathbf{w}) - \gamma c(v - v')\},$$

Table 4: Revenue for Different Robust Coefficients (10^{-1})

R.C.	0	0.001	0.01	<u>0.03</u>	0.1	0.3	1	10
A.R.	6.63	6.68	6.75	<u>6.76</u>	6.69	6.74	6.74	6.72

Table 3: Additional Results on More Valuation Settings

Setting	Method	Regret	Revenue	Loss
2x2SY	BL	0.0488	2.2587	7.5013
	GAT	0.0007	2.7820	-2.6420
	SAT	<u>0.0002</u>	<u>2.8265</u>	<u>-2.7865</u>
2x2AS	BL	0.0309	3.7042	2.4758
	GAT	0.0011	<u>4.1206</u>	-3.9006
	SAT	<u>0.0005</u>	4.1008	<u>-4.0008</u>
2x3AD	BL	0.0915	<u>1.6853</u>	16.6147
	GAT	0.0016	1.2296	-0.9096
	SAT	<u>0.0003</u>	1.2191	<u>-1.1591</u>
3x10AD	BL	0.0793	<u>6.9427</u>	16.8473
	GAT	0.0018	4.8012	-4.2612
	SAT	<u>0.0016</u>	5.0802	<u>-4.6002</u>

and note that $\phi'_i(\mathbf{v}, \mathbf{w}, 0) = \phi_i(\mathbf{v}, \mathbf{w})$ defined in (16). We have:

$$\begin{aligned}
 \text{rrgt}_i(\mathbf{w}, \hat{P}, \hat{\rho}) &= \min_{\gamma \geq 0} \{ \gamma \rho + \mathbb{E}_{\mathbf{v} \sim \hat{P}} [\phi'_i(\mathbf{v}, \mathbf{w}, \gamma)] \} \\
 &\leq \mathbb{E}_{\mathbf{v} \sim \hat{P}} [\phi'_i(\mathbf{v}, \mathbf{w}, 0)] \\
 &= \phi_i(\mathbf{v}, \mathbf{w}),
 \end{aligned}$$

which is exactly the conclusion in (15). Thus the proof is concluded.

B Supplementary Results

B.1 Supplementary Results for Section 5.2

We will first introduce two additional settings that correspond to slightly more complex value distributions:

- **2x2SY**: short for *2x2 symmetric*, which is a variant of *2x2 additive*. Here we denote the subsets as $S_1 = \emptyset$, $S_2 = \{1\}$, $S_3 = \{2\}$ and $S_4 = \{1, 2\}$, then the distribution for this setting can be described as $v_{11} = v_{21} = 0$, $v_{12}, v_{13}, v_{22}, v_{23} \sim U[1, 2]$, $v_{14} = v_{12} + v_{13} + c_1$, $v_{24} = v_{22} + v_{23} + c_2$ and $c_1, c_2 \sim U[-1, 1]$, where $U[l, r]$ for any l and r denotes the uniform distribution over $[l, r]$.

- **2x2AS**: short for *2x2 asymmetric*. The distribution is the same as the **2x2SY** except that $v_{22}, v_{23} \sim U[1, 5]$.

These settings is the same as the settings (IV) and (V) in ([17]). Now we can present the results of our GAT and SAT method on more valuation settings, as shown in Table 3, where **2x3AD** and **3x10AD** are abbreviations for the *2x3 additive* and *3x10 additive* settings, respectively. Since the training process is generally the same as Figure 1, we only consider the metrics after 100,000 iterations as the performance for each method, which are presented in the last three columns of Table 3. It can be observed that the experimental results in the supplementary settings are consistent with the conclusions in Section 5.2.

In addition, we further provide the commonly used values for the hyper-parameters in the GAT and SAT methods. λ is set to 1 at the start of training, $\eta_{\mathbf{w}}$ is set to 10^{-4} , η_{λ} is set to 5 initially and is continuously increased during training. These settings are consistent with the experiments of existing work [17]. Besides, ϵ and η_{γ} are usually set to 0.01 and 0.05, respectively.

B.2 Supplementary Results for Section 5.5

Here, we briefly present part of the average performance corresponding to different robust coefficients obtained during the process of Algorithm 2 on the datasets splitted by $\mathcal{T} = \{10, 11, \dots, 15\}$, as shown in Table 4, where **R.C.** and **A.R.** are abbreviations for robust coefficient and average revenue, respectively. Notably, a robust coefficient of 0 refers to the results obtained using the BL method.

We can observe that the performance of the auctions that introduced the robust coefficient all exceeded that of BL. Therefore, based on these results and the process of Algorithm 2, we adopted $\rho = 0.03$ as the robust coefficient for GAT in the subsequent evaluations, which is consistent with the description in Section 5.5.