
Sample Complexity of Posted Pricing for a Single Item

Billy Jin*

Thomas Kesselheim†

Will Ma‡

Sahil Singla§

Abstract

Selling a single item to n self-interested buyers is a fundamental problem in economics, where the two objectives typically considered are welfare maximization and revenue maximization. Since the optimal mechanisms are often impractical and do not work for sequential buyers, posted pricing mechanisms, where fixed prices are set for the item for different buyers, have emerged as a practical and effective alternative. This paper investigates how many samples are needed from buyers' value distributions to find near-optimal posted prices, considering both independent and correlated buyer distributions, and welfare versus revenue maximization. We obtain matching upper and lower bounds (up to logarithmic factors) on the sample complexity for all these settings.

1 Introduction

A fundamental problem in Economics is how to sell a single indivisible item to a set of n self-interested buyers. This problem is challenging because the value associated to the item is private knowledge of each buyer and thus can be strategically misreported. The mechanism designer usually has one of two objectives: (a) *welfare* maximization where we want to maximize the value of the winning buyer, and (b) *revenue* maximization where we want to maximize the price paid by the winning buyer. The welfare maximization problem was resolved by Vickrey in 1961 [Vic61] and the revenue maximization problem was resolved (for independent distributions) by Myerson in 1981 [Mye81], both of whom were awarded the Nobel prize in Economics for their contributions. In either case the optimal mechanism is a truthful auction, in which buyers report their values and have no incentive to misreport.

Although we know optimal mechanisms for selling a single item, these auctions are often difficult/impossible to implement in practice. For instance, the buyer does not know exactly how much they will pay if they win the item and these auctions do not work for online settings where the buyers arrive one-by-one [AM⁺06, HR09]. Thus, the last two decades has seen a lot of progress on understanding the power of simple but near-optimal mechanisms. In this regard, *posted pricing* mechanisms have been identified to be very successful (see books and surveys [Rou16, Har13, CFH⁺19, Luc17]). Here, the mechanism designer puts a (carefully chosen) price π_i on the item and then the i -th buyer takes the item if the item is not already sold and if they value it above π_i . Posted pricing mechanisms have several advantages: they are truthful since π_i does not depend on the i -th buyer's value; the buyer knows exactly how much they pay on winning an item; and they work even for online settings so the buyers do not have to be simultaneously present. Additionally, in the setting of Myerson's 1981 paper of independent buyer valuations, they are known to give a 2-approximation to both the optimal welfare (which is usually called Prophet Inequality) and the optimal revenue

*Cornell University, School of Operations Research and Information Engineering, bzj3@cornell.edu.

†University of Bonn, Institute of Computer Science and Lamarr Institute for Machine Learning and Artificial Intelligence, thomas.kesselheim@uni-bonn.de

‡Columbia University, Graduate School of Business and Data Science Institute, wm2428@gsb.columbia.edu.

§Georgia Tech, School of Computer Science, ssingla@gatech.edu. Supported in part by NSF award CCF-2327010.

[KS77, KS78, CHK07, CFPV19], as long as the posted prices are set correctly. This raises the main question of the current paper:

How many samples are necessary from the value distributions of the buyers to find near-optimal posted prices for single item? Is there a difference between independent vs. correlated distributions or between welfare vs. revenue maximization?

Despite being a natural question, the tight sample complexity bounds for single item posted pricing are unknown, both for independent/correlated and welfare/revenue settings.

1.1 Model

Before stating our results, we formally describe our model.

Posted pricing. There is single copy of an item. Buyers $i = 1, \dots, n$ arrive in order, each with a private value (willingness to pay) V_i for that item. Price π_i is offered to buyer i , and the first buyer (if any) for whom $V_i \geq \pi_i$ end up buying the item, “winning” the auction. The objective is to maximize either *welfare*, which is the value of the item to the winner, or *revenue*, which is the price paid by the winner.

Distributions and policies. We will assume that the vector of values $\mathbf{V} = (V_1, \dots, V_n)$ is drawn from an unknown but fixed distribution $\mathbf{D}$ over $[0, 1]^n$. If each V_i is drawn *independently* from some marginal distribution D_i , then $\mathbf{D}$ is called a *product distribution*, written as $\mathbf{D} = D_1 \times \dots \times D_n$. Meanwhile, *posted pricing policies* are defined by a vector of prices $\pi = (\pi_1, \dots, \pi_n)$ that must be fixed before the first buyer arrives, and possibly restricted to some subclass Π . The objective of policy π under buyer values $\mathbf{V}$ is denoted by $\pi(\mathbf{V})$, which equals $V_{\arg\min\{i:\pi_i \leq V_i\}}$ and $\pi_{\arg\min\{i:\pi_i \leq V_i\}}$ for the welfare and revenue objectives respectively, understood to be 0 if the item is not sold. We (abusively) let $\pi(\mathbf{D}) := \mathbb{E}_{\mathbf{V} \sim \mathbf{D}}[\pi(\mathbf{V})]$ denote the expected objective of π , and let $\Pi(\mathbf{D}) := \sup_{\pi \in \Pi} \pi(\mathbf{D})$ denote the best expected objective of a policy in Π knowing $\mathbf{D}$.

Note that if we are given the product distribution on buyer values, the optimal policy (for either objective) can be easily computed using dynamic programming [CFPV19]. Consequently, for product distributions we consider the full policy class $\Pi = [0, 1]^n$, and omit the dependence on Π .

Learning problem. A *learning algorithm* takes as input samples $\mathbf{V}$ that are drawn independently and identically distributed (IID) from the unknown distribution $\mathbf{D}$ and an error parameter $\varepsilon \in (0, 1)$, and seeks to return a policy $\pi \in \Pi$ that satisfies $\pi(\mathbf{D}) \geq \Pi(\mathbf{D}) - \varepsilon$, which is called an ε -*approximation*. Given a failure probability $\delta \in (0, 1)$, the *sample complexity* is the minimum number of samples required for there to exist a learning algorithm that, under any distribution $\mathbf{D}$, returns an ε -additive approximation with probability at least $1 - \delta$, noting that the randomness is over both the samples and any random bits in the algorithm. The sample complexity and learning algorithm will depend on the problem variant, defined by the objective (welfare or revenue), any parameters of the policy class Π , and whether $\mathbf{D}$ is restricted to be a product distribution or not. The sample complexity will also depend on the parameters $\varepsilon, \delta > 0$.

1.2 Our Contributions to Posted Pricing for Product Distributions

The study of sample complexity for posted pricing goes back to at least the seminal work of Kleinberg and Leighton [KL03] who study revenue maximization for selling a single item to a single buyer (welfare maximization is trivial for single buyer since we just allocate the item by setting 0 price). For the general posted pricing problem with n buyers, the best known sample complexity bounds were due to Guo et al. [GHTZ21], who showed that $\tilde{O}(n/\varepsilon^2)$ samples suffice from each buyer’s distribution for both welfare and revenue maximization objectives. Although there is a simple $\Omega(1/\varepsilon^2)$ lower bound for both welfare and revenue maximization settings, prior to our work it was unclear if polynomial dependency on the number of buyers n is necessary.

Our first result for product distribution shows that for welfare maximization using posted pricing (a.k.a. prophet inequality), the sample complexity is independent of the number of buyers n .

Theorem 1 (proved in Subsection 2.1). *For product distributions, the sample complexity of welfare maximization is $O(1/\varepsilon^2 \cdot \log^2(1/\delta))$.*

Proof sketch. We start by constructing the product empirical distribution using the samples, where for each i we take the uniform distribution on the $\tilde{O}(1/\varepsilon^2)$ samples and then take the product distribution for different i . Our main result is that the optimal policy (dynamic programming solution) on this product empirical distribution is an ε -approximation with high probability. Although one could use standard concentration bounds to bound the gap between the learned and optimal thresholds, such arguments lose a $\text{poly}(n)$ factor. This is because the difference between successive thresholds is bounded by 1, so naively the total variance of n thresholds could be $\Omega(n)$. Our main idea is to instead study a martingale that adds up the errors made in the dynamic programming solution. This way, too high values for one buyer and too low values for another buyer balance each other. On this martingale, we use Freedman's inequality (which is a martingale variant of Bernstein's inequality), which allows us to bound the total variance in terms of conditional variances. A careful application of another concentration bound allows us to bound these conditional variances in terms of the change in the optimal value, whose sum is always at most 1. This latter concentration bound is what causes the additional factor of $\log(1/\delta)$ in the sample complexity.

Our second result for product distribution shows a separation between welfare and revenue maximization using posted pricing, by proving that the $\tilde{O}(n/\varepsilon^2)$ sample-complexity result of [GHTZ21] for revenue maximization is tight up to log factors.

Theorem 2 (proved in Subsection 2.2). *For revenue maximization on product distributions, any learning algorithm requires $\Omega(\frac{n}{\varepsilon^2})$ samples to return an ε -additive approximation with probability greater than $6/7$.*

Proof sketch. We construct for each buyer i two possible value distributions that add the same amount (roughly $\frac{1}{n}$) to the value-to-go of the optimal dynamic program. However, these distributions have different optimal prices, and making a mistake (choosing an incorrect price) in isolation loses roughly $\frac{\varepsilon}{n}$ value. Although these mistakes accumulate in a non-linear fashion, we show that making M mistakes must lose in total $\Omega(\frac{\varepsilon M^2}{n^2})$. Finally, these value distributions have probabilities on the scale of $\frac{1 \pm \varepsilon}{n}$ with the same supports (essentially, only $1/n$ of samples provide information), which means that $\Omega(\frac{n}{\varepsilon^2})$ samples are needed to avoid making a constant fraction of mistakes.

1.3 Our Contributions to Posted Pricing for Correlated Distributions

Independence among buyer valuations can be a strong modeling assumption for many applications. Although for arbitrary correlated distributions the optimal policy is not learnable, one could hope to learn the best policy in the class of all posted pricing policies. A recent work of Balcan et al. [BDD⁺21] can be applied to this setting to show that $\tilde{O}(n/\varepsilon^2)$ samples are sufficient to learn ε -optimal posted pricing. As discussed in Theorem 2, this linear in n dependency is necessary for revenue maximization, even for product distributions. Our first observation is that for correlated distributions, we need to lose this factor even for welfare maximization.

Theorem 3 (corollary of Theorem 5). *For welfare maximization with correlated buyers, any learning algorithm requires $\Omega(\frac{n}{\varepsilon^2})$ samples to return an ε -additive approximation with constant probability.*

Given this lower bound, a natural next question is to consider the subclass of posted pricing policies where the algorithm is only allowed to change its threshold at a small number of given locations. Can we now remove the linear in n dependency from sample complexity? The motivation to study this class comes from posted pricing applications where it is not possible for the algorithm designer to update the prices at each time step, e.g., due to business constraints.

Formally, for $\mathcal{S} \subseteq \{2, \dots, n\}$, we say that posted pricing policy π respects change-points $\mathcal{S}$ if π_i can differ from π_{i-1} only when $i \in \mathcal{S}$. We let $\Pi_{\mathcal{S}}$ denote the class of all policies that respect change-points $\mathcal{S}$, noting that policies in $\Pi(\emptyset)$ post a static price for all buyers, and $\Pi(\{2, \dots, n\}) = [0, 1]^n$. Our following result shows that one can obtain sample complexity that is independent of n , depending only on the size of $\mathcal{S}$.

Theorem 4 (proved in Subsection 3.1). *For correlated distributions, the sample complexity of welfare or revenue maximization is $O\left(\frac{(1+|\mathcal{S}|) \log(1+|\mathcal{S}|) + \log(1/\delta)}{\varepsilon^2}\right)$ when the policy is restricted to $\Pi_{\mathcal{S}}$, for any $\mathcal{S} \subseteq \{2, \dots, n\}$.*

Proof sketch. By existing results in learning theory, it suffices to bound the pseudo-dimension of $\Pi_{\mathcal{S}}$. This will boil down to understanding the structure of “good sets”, which are sets of the form $\{\pi \in \Pi_{\mathcal{S}} : \pi(\mathbf{v}) \geq z\}$, for some input $\mathbf{v}$ and some target z . We will show that for a natural parameterization of $\Pi_{\mathcal{S}}$, any good set can be expressed as the union and intersection of $O(|\mathcal{S}| + 1)$ halfspaces, which implies the pseudo-dimension of $\Pi_{\mathcal{S}}$ is $O((|\mathcal{S}| + 1) \log(|\mathcal{S}| + 1))$ by a result of [BDD⁺21]. This bound on the pseudo-dimension translates to the above sample complexity bound.

The learning algorithm which achieves the sample complexity bound in Theorem 4 is simply sample average approximation (SAA), which returns the policy in $\Pi_{\mathcal{S}}$ with the highest objective value averaged over the samples. SAA can be computed in time $O(Tn(Tn)^{1+|\mathcal{S}|})$. This is because there are $1 + |\mathcal{S}|$ prices to decide, each of which can take Tn possible values (one for each realized value in the T samples of length n), and evaluating each combination of prices over each of the T samples takes runtime linear in n . We leave as an open question whether there is a more efficient algorithm.

Finally, we complement our sample complexity upper bounds for correlated distributions by giving matching lower bounds (up to polylogs).

Theorem 5 (proved in Appendix A.5). *For welfare or revenue maximization on correlated distributions, a learning algorithm requires $\Omega(\frac{1+|\mathcal{S}|}{\varepsilon^2})$ samples to return an ε -additive approximation with constant probability, when restricted to the policy class $\Pi_{\mathcal{S}}$ for any $\mathcal{S} \subseteq \{2, \dots, n\}$.*

Proof sketch. In the construction for Theorem 5, on each trajectory exactly one of the $1 + |\mathcal{S}|$ decision points (randomly selected) will be relevant, essentially diluting the samples by a factor of $1 + |\mathcal{S}|$ and leading to a lower bound of $\Omega(\frac{1+|\mathcal{S}|}{\varepsilon^2})$.

1.4 Further Related Work

In 2007, Hajiaghayi et al. [HKS07] discovered connections between auction design and posted pricing via prophet inequalities. Since then, there is a long line of work on understanding the power of posted pricing for selling multiple items to combinatorial buyers [CMS10, FGL15, DKL20, AKS21, CC23]. For single item revenue maximization with known distributions, Correa et al. [CFPV19] showed the equivalence of welfare and revenue maximization objectives for single item posted pricing.

For background on sample complexity, we suggest Wainwright’s excellent textbook [Wai19]. Sample complexity of auction design has been greatly studied; e.g., see [KL03, CR14, MR16]. We refer the readers to Guo et al. [GHZ19], who recently resolved single item revenue maximization in the offline setting, for an overview of the literature. In [GHTZ21], Guo et al. generalized their techniques for revenue maximization over product distributions to all “strongly monotone” problems, which includes posted pricing for welfare and revenue maximization.

Recently, there is also a lot interest in learning auctions in limited feedback models like bandit and pricing queries [GKSW24, SW24, LSTW23]. We should note that sample complexity of optimal stopping (equivalent to our welfare maximization problem) has been previously studied in [GM22], who analyze linear stopping rules in a contextual setting. Our application of techniques from [BDD⁺21] to online algorithms over correlated sequences is also similar in spirit to some results from [XMX24], who study a different application of inventory optimization.

Another related but tangential line of work focuses on prophet inequalities with samples [AKW14, RWW, CDFS22, CDF⁺22, DKL⁺24]. The key distinction in these works is that their benchmark is the expected hindsight optimum, rather than the optimal online policy. Notably, any online algorithm incurs at least a factor of 2 loss compared to the hindsight optimum, even in the case of single-item prophet inequalities. As a result, this line of research aims to achieve $O(1)$ -approximation guarantees, rather than the sublinear regret guarantees pursued in the current paper. Furthermore, their techniques differ significantly, as they often assume unbounded distributions. Finally, there is also work that explores the (non-)robustness of algorithms for the prophet inequality problem to inaccuracies in the distributions [DK19] and to dependencies in distributions [ISW20, LPS24].

2 Product Distributions

In this section we first prove our improved upper bound on the sample complexity of welfare for product distributions, and then prove a new lower bounds on the sample complexity of revenue for product distributions.

2.1 Positive Result for Welfare: Proof of Theorem 1

The reward of the optimal policy is given by the following backward induction: $r_{n+1}^* = 0$ and $r_i^* = \mathbb{E}[\max\{r_{i+1}^*, V_i\}] = \mathbb{E}[(V_i - r_{i+1}^*)^+] + r_{i+1}^*$. It sets $\pi_i = r_{i+1}^*$.

When we do not know the distributions but only have T samples from each of them, we can consider the optimal policy on the product empirical distribution, which corresponds to replacing the expectations in the above definitions by the empirical average. That is, we will analyze the policy that sets $\pi_i = \hat{r}_{i+1}$, where $\hat{r}_i$ is defined recursively by $\hat{r}_{n+1} = 0$, $\hat{r}_i = \frac{1}{T} \sum_{t=1}^T (V_i^{(t)} - \hat{r}_{i+1})^+ + \hat{r}_{i+1}$.

Let r_i be the expected reward of this policy when starting with the i -th arrival. In order to prove the theorem, it is sufficient to show that with probability at least $1 - \delta$, we have $r_1 \geq r_1^* - \epsilon$ if $T \geq (5 \ln(2e/\delta)/\epsilon)^2$ for any choices of $\epsilon, \delta \in (0, 1)$.

Let us define $\eta_i = \hat{r}_i - \mathbb{E}[\max\{\hat{r}_{i+1}, V_i\}]$ or equivalently as

$$\eta_i = \frac{1}{T} \sum_{t=1}^T (V_i^{(t)} - \hat{r}_{i+1})^+ - \mathbb{E}[(V_i - \hat{r}_{i+1})^+] .$$

That is, η_i denotes the error introduced by using the empirical distribution for buyer i instead of the actual one. Note that η_i can both be positive and negative.

The two key steps of our proof are as follows. We first show that

$$r_1 \geq r_1^* - 2 \max_{j \geq 1} \left| \sum_{i=1}^j \eta_i \right| . \quad (1)$$

This inequality holds point-wise, that is for any samples drawn. Then, we show that for $T \geq (5 \ln(2e/\delta)/\epsilon)^2$ samples, we have $\max_{j \geq 1} \left| \sum_{i=1}^j \eta_i \right| \leq \epsilon$ with probability at least $1 - \delta$.

In order to show (1), we first lower bound r_1 in terms of $\hat{r}_1$.

Lemma 1. $r_1 \geq \hat{r}_1 - \max_{j \in \{0, \dots, n\}} \sum_{i=1}^j \eta_i$.

Proof. Let $j \in \{0, 1, \dots, n\}$ be the smallest index for which $r_{j+1} \geq \hat{r}_{j+1}$, which exists because $0 = r_{n+1} \geq \hat{r}_{n+1} = 0$. Note that we have $r_i < \hat{r}_i$ for all $1 \leq i \leq j$. We rewrite r_1 as

$$r_1 = \mathbf{Pr}[V_1 < \hat{r}_2] r_2 + \mathbf{Pr}[V_1 \geq \hat{r}_2] \mathbb{E}[V_1 \mid V_1 \geq \hat{r}_2] = r_2 + \mathbb{E}[(V_1 - r_2) \mathbb{1}_{V_1 \geq \hat{r}_2}] .$$

Repeating this argument,

$$r_1 = \sum_{i=1}^{j-1} \mathbb{E}[(V_i - r_{i+1}) \mathbb{1}_{V_i \geq \hat{r}_{i+1}}] + \mathbb{E}[V_j \mathbb{1}_{V_j \geq \hat{r}_{j+1}} + r_{j+1} \mathbb{1}_{V_j < \hat{r}_{j+1}}] .$$

Inductively, we can also establish that

$$\hat{r}_1 = \sum_{i=1}^{j-1} \frac{1}{T} \sum_t (V_i^{(t)} - \hat{r}_{i+1})^+ + \frac{1}{T} \sum_t \max\{V_j^{(t)}, \hat{r}_{j+1}\} .$$

Combining these two equalities,

$$\begin{aligned} \hat{r}_1 - r_1 &= \sum_{i=1}^{j-1} \left(\frac{1}{T} \sum_{t=1}^T (V_i^{(t)} - \hat{r}_{i+1})^+ - \mathbb{E}[(V_i - \underbrace{r_{i+1}}_{< \hat{r}_{i+1}}) \mathbb{1}_{V_i \geq \hat{r}_{i+1}}] \right) \\ &\quad + \frac{1}{T} \sum_{t=1}^T \max\{V_j^{(t)}, \hat{r}_{j+1}\} - \mathbb{E}[V_j \mathbb{1}_{V_j \geq \hat{r}_{j+1}} + \underbrace{r_{j+1}}_{\geq \hat{r}_{j+1}} \mathbb{1}_{V_j < \hat{r}_{j+1}}] \\ &\leq \sum_{i=1}^{j-1} \left(\frac{1}{T} \sum_{t=1}^T (V_i^{(t)} - \hat{r}_{i+1})^+ - \mathbb{E}[(V_i - \hat{r}_{i+1})^+] \right) + \frac{1}{T} \sum_{t=1}^T \max\{V_j^{(t)}, \hat{r}_{j+1}\} - \mathbb{E}[\max\{V_j, \hat{r}_{j+1}\}] \\ &= \sum_{i=1}^j \eta_i . \end{aligned}$$

This implies $\hat{r}_1 - r_1 \leq \max_{j \in \{0, \dots, n\}} \sum_{i=1}^j \eta_i$, completing the proof. $\square$

A similar proof allows us to prove the following lemma.

Lemma 2. $\hat{r}_1 \geq r_1^* - \max_{j \in \{0, \dots, n\}} (-\sum_{i=1}^j \eta_i)$.

Proof. Let $j \in \{0, 1, \dots, n\}$ be the smallest index for which $\hat{r}_{j+1} \geq r_{j+1}^*$, which exists because $0 = \hat{r}_{n+1} \geq r_{n+1}^* = 0$. Note that we have $\hat{r}_i < r_i^*$ for all $1 \leq i \leq j$, allowing us to derive

$$\begin{aligned} r_1^* - \hat{r}_1 &= \sum_{i=1}^{j-1} \left(\mathbb{E}[(V_i - \underbrace{r_{i+1}^*}_{> \hat{r}_{i+1}})^+] - \frac{1}{T} \sum_{t=1}^T (V_i^{(t)} - \hat{r}_{i+1})^+ \right) \\ &\quad + \mathbb{E}[\max\{V_j, \underbrace{r_{j+1}^*}_{\leq \hat{r}_{j+1}}\}] - \frac{1}{T} \sum_{t=1}^T \max\{V_j^{(t)}, \hat{r}_{j+1}\} \\ &\leq \sum_{i=1}^{j-1} \left(\mathbb{E}[(V_i - \hat{r}_{i+1})^+] - \frac{1}{T} \sum_{t=1}^T (V_i^{(t)} - \hat{r}_{i+1})^+ \right) \\ &\quad + \mathbb{E}[\max\{V_j, \hat{r}_{j+1}\}] - \frac{1}{T} \sum_{t=1}^T \max\{V_j^{(t)}, \hat{r}_{j+1}\} \\ &= -\sum_{i=1}^j \eta_i. \end{aligned}$$

This implies $r_1^* - \hat{r}_1 \leq \max_{j \in \{0, \dots, n\}} (-\sum_{i=1}^j \eta_i)$, completing the proof. $\square$

The last two lemmas imply (1). Thus, we need to bound $\max_j |\sum_{i=1}^j \eta_i|$ to complete the proof.

Lemma 3. For every $\epsilon, \delta > 0$, with probability at least $1 - \delta$, we have $\max_j |\sum_{i=1}^j \eta_i| \leq \epsilon$ if $T \geq (5 \ln(2e/\delta)/\epsilon)^2$.

Proof. Observe that

$$\left| \sum_{i=1}^j \eta_i \right| = \left| \sum_{i=1}^n \eta_i - \sum_{i=j+1}^n \eta_i \right| \leq \left| \sum_{i=1}^n \eta_i \right| + \left| \sum_{i=j+1}^n \eta_i \right| \leq 2 \max_{\tau} \left| \sum_{i=\tau}^n \eta_i \right|,$$

so it suffices to show that $\max_{\tau} |\sum_{i=\tau}^n \eta_i| \geq \epsilon/2$ with probability at most δ . Reversing the quantity of interest to $\max_{\tau} |\sum_{i=\tau}^n \eta_i|$ allows us to define a martingale, and use Freedman's inequality, which is a martingale version of Bernstein's inequality.

Lemma 4 (Freedman, Theorem 1.6 in [Fre75]). Consider a real-valued sequence $\{X_t\}_{t \geq 0}$ such that $X_0 = 0$ and $\mathbb{E}[X_{t+1} \mid X_t, X_{t-1}, \dots, X_0] = 0$ for all t . Assume that the sequence is uniformly bounded, i.e., $|X_t| \leq M$ almost surely for all t . Now define the predictable quadratic variation process of the martingale to be $W_t = \sum_{j=0}^t \mathbb{E}[X_j^2 \mid X_{j-1}, \dots, X_0]$ for all $t \geq 1$. Then for all $\ell \geq 0$ and $\sigma^2 \geq 0$, and any stopping time τ , we have

$$\Pr \left[\left| \sum_{j=0}^{\tau} X_j \right| \geq \ell \text{ and } W_{\tau} \leq \sigma^2 \right] \leq 2 \exp \left(-\frac{\ell^2/2}{\sigma^2 + M\ell/3} \right).$$

A corollary of Freedman's inequality is that $\Pr \left[\left| \sum_{j=0}^{\tau} X_j \right| \geq \ell \right] \leq 2 \exp \left(-\frac{\ell^2/2}{\sigma^2 + M\ell/3} \right) + \Pr[W_{\tau} \leq \sigma^2]$. In order to use Freedman's inequality, we consider nT random variables $X_1, \dots, X_{nT}$, where $X_{iT+t} = \frac{1}{T} \left((V_{n-i}^{(t)} - \hat{r}_{n-i-1})^+ - \mathbb{E}[(V_{n-i}^{(t)} - \hat{r}_{n-i-1})^+] \right)$ for $i \in \{0, \dots, n-1\}$ and $t \in \{1, \dots, T\}$. By this definition,

$$\eta_i = \sum_{t=1}^T X_{(n-i)T+t}$$

because $V_{n-i}^{(t)}$ and V_{n-i} are identically distributed and independent of $\hat{r}_{n-i-1}$. Moreover, for all j ,

$$\mathbb{E}[X_j \mid X_1, \dots, X_{j-1}] = 0 \text{ and } |X_j| \leq \frac{1}{T}.$$

Let $\mathcal{E}$ be the event that $\sum_{j=1}^{nT} \mathbb{E}[X_j^2 \mid X_1, \dots, X_{j-1}] > \sigma^2 := \frac{e}{e+1} \frac{\ln(2e/\delta)}{T}$. By Freedman's inequality, using $T \geq (5 \ln(2e/\delta)/\epsilon)^2$, we have

$$\begin{aligned} \Pr \left[\max_{\tau} \left| \sum_{i=\tau}^n \eta_i \right| \geq \frac{\epsilon}{2} \right] &\leq \Pr \left[\max_{\tau} \left| \sum_{j=1}^{\tau} X_j \right| \geq \frac{\epsilon}{2} \right] \\ &\leq 2 \exp \left(-\frac{\epsilon^2/8}{\sigma^2 + \frac{\epsilon}{6T}} \right) + \Pr[\mathcal{E}] \\ &\leq 2 \exp \left(-\frac{\left(\frac{5 \ln(2e/\delta)}{\sqrt{T}} \right)^2 / 8}{\frac{e}{e+1} \frac{\ln(2e/\delta)}{T} + \frac{5 \ln(2e/\delta)}{6T\sqrt{T}}} \right) + \Pr[\mathcal{E}] \\ &= 2 \exp \left(-\frac{25}{\frac{8e}{e+1} + \frac{40}{6\sqrt{T}}} \ln \left(\frac{2e}{\delta} \right) \right) + \Pr[\mathcal{E}] \\ &\leq 2 \exp \left(-(1) \ln \left(\frac{4}{\delta} \right) \right) + \Pr[\mathcal{E}] = \frac{\delta}{2} + \Pr[\mathcal{E}]. \end{aligned}$$

It remains to show that $\Pr[\mathcal{E}] \leq \frac{\delta}{2}$. To this end, we observe that for any $i \in \{0, \dots, n-1\}$ and $t \in \{1, \dots, T\}$,

$$\begin{aligned} &\mathbb{E}[X_{iT+t}^2 \mid X_1, \dots, X_{iT+t-1}] \\ &= \frac{1}{T^2} \mathbb{E} \left[\left((V_{n-i}^{(t)} - \hat{r}_{n-i-1})^+ - \mathbb{E}[(V_{n-i}^{(t)} - \hat{r}_{n-i-1})^+] \right)^2 \mid X_1, \dots, X_{iT+t-1} \right] \\ &\leq \frac{1}{T^2} \mathbb{E} \left[(V_{n-i}^{(t)} - \hat{r}_{n-i-1})^+ \mid X_1, \dots, X_{iT+t-1} \right] \\ &= \frac{1}{T^3} \sum_{t'=1}^T \mathbb{E} \left[(V_{n-i}^{(t')} - \hat{r}_{n-i-1})^+ \mid X_1, \dots, X_{iT} \right] \\ &= \frac{1}{T^2} \mathbb{E}[\hat{r}_i - \hat{r}_{i+1} \mid \hat{r}_{i+1}, \dots, \hat{r}_n], \end{aligned}$$

where the inequality uses $\mathbb{E}[Y] \geq \mathbb{E}[Y^2] \geq \mathbb{E}[Y^2] - (\mathbb{E}[Y])^2 = \mathbb{E}[(Y - \mathbb{E}[Y])^2]$ for any random variable Y with $0 \leq Y \leq 1$ almost surely. In total, we have

$$\sum_{j=1}^{nT} \mathbb{E}[X_j^2 \mid X_1, \dots, X_{j-1}] \leq \frac{1}{T} \sum_{i=1}^n \mathbb{E}[\hat{r}_i - \hat{r}_{i+1} \mid \hat{r}_{i+1}, \dots, \hat{r}_n].$$

As $\sum_{i=1}^n (\hat{r}_i - \hat{r}_{i+1}) = \hat{r}_1 \leq 1$ almost surely, we have by Lemma 5 (see below) that

$$\sum_{i=1}^n \mathbb{E}[\hat{r}_i - \hat{r}_{i+1} \mid \hat{r}_{i+1}, \dots, \hat{r}_n] \geq \frac{e}{e-1} \ln \left(\frac{2e}{\delta} \right)$$

with probability at most $\frac{\delta}{2}$. Therefore, we have

$$\Pr[\mathcal{E}] \leq \Pr \left[\sum_{i=1}^n \mathbb{E}[\hat{r}_i - \hat{r}_{i+1} \mid \hat{r}_{i+1}, \dots, \hat{r}_n] \geq \frac{e}{e-1} \ln \left(\frac{2e}{\delta} \right) \right] \leq \frac{\delta}{2}. \quad \square$$

Lemma 5. Let $Y_1, Y_2, \dots, Y_n$ be a sequence of (not necessarily independent) random variables in $[0, 1]$ such that $\sum_{i=1}^n Y_i \leq 1$ almost surely. Then, for any $\delta > 0$, with probability at most δ , we have

$$\sum_{i=1}^n \mathbb{E}[Y_i \mid Y_1, \dots, Y_{i-1}] \geq \frac{e}{e-1} \ln \left(\frac{e}{\delta} \right).$$

We defer the proof of Lemma 5 to the appendix.

2.2 Negative Result for Revenue: Proof of Theorem 2

Each buyer $i = 1, \dots, n$ has a marginal value distribution that could be D_i^H ("High") or D_i^L ("Low"):

$$D_i^H = \begin{cases} \frac{1}{2} + \frac{n-i}{4n} & \text{w.p. } \frac{1}{2n} \\ \frac{1}{4} + \frac{n-i}{4n} & \text{w.p. } \frac{1}{2n} - 16\frac{\epsilon}{n} \\ 0 & \text{w.p. } 1 - \frac{1}{n} + 16\frac{\epsilon}{n} \end{cases} \quad D_i^L = \begin{cases} \frac{1}{2} + \frac{n-i}{4n} & \text{w.p. } \frac{1}{2n} - 8\frac{\epsilon}{n} \\ \frac{1}{4} + \frac{n-i}{4n} & \text{w.p. } \frac{1}{2n} + 8\frac{\epsilon}{n} \\ 0 & \text{w.p. } 1 - \frac{1}{n} \end{cases}$$

which are valid distributions as long as $\epsilon \leq 1/32$ and $n \geq 2$. All 2^n configurations of whether each buyer has the High or Low distribution are possible.

Optimal policy Fix any configuration of whether each buyer has the High or Low distribution. Let r_i^* denote the expected revenue to be earned under the optimal dynamic program, if buyer i is about to arrive and the item is not yet sold. We show inductively that $r_i^* = \frac{n+1-i}{4n}$. By definition $r_{n+1}^* = 0$, establishing $r_i^* = \frac{n+1-i}{4n}$ for $i = n+1$. Now consider $i = n, \dots, 1$, and assume $r_{i+1}^* = \frac{n-i}{4n}$. Note that it is better to offer one of the prices $\frac{1}{2} + \frac{n-i}{4n}$ or $\frac{1}{4} + \frac{n-i}{4n}$ than to reject the buyer by offering a price of 1, because both of these prices are greater than r_i^* (by the induction hypothesis). Thus, if buyer i has distribution D_i^H , then

$$\begin{aligned} r_i^* &= \max \left\{ \left(\frac{1}{2} + \frac{n-i}{4n} \right) \frac{1}{2n} + r_{i+1}^* \left(1 - \frac{1}{2n} \right), \left(\frac{1}{4} + \frac{n-i}{4n} \right) \left(\frac{1}{n} - 16 \frac{\varepsilon}{n} \right) + r_{i+1}^* \left(1 - \frac{1}{n} + 16 \frac{\varepsilon}{n} \right) \right\} \\ &= \max \left\{ \frac{1}{2} \cdot \frac{1}{2n}, \frac{1}{4} \left(\frac{1}{n} - 16 \frac{\varepsilon}{n} \right) \right\} + \frac{n-i}{4n} = \frac{n+1-i}{4n}. \end{aligned}$$

Similarly, if buyer i instead has D_i^L , then

$$r_i^* = \max \left\{ \frac{1}{2} \left(\frac{1}{2n} - 8 \frac{\varepsilon}{n} \right), \frac{1}{4} \cdot \frac{1}{n} \right\} + \frac{n-i}{4n} = \frac{n+1-i}{4n}.$$

In either case, we have $r_i^* = \frac{n+1-i}{4n}$, completing the induction. Note that $r_1^* = 1/4$.

Bounding a policy's objective by its number of mistakes Now, consider an arbitrary policy $\pi = (\pi_1, \dots, \pi_n)$ decided by the learning algorithm. Let r_i denote its expected revenue earned if buyer i is about to arrive and the item is not yet sold (under the true configuration of buyer distributions). We can without loss assume π to lie in $\{\frac{1}{2} + \frac{n-i}{4n}, \frac{1}{4} + \frac{n-i}{4n}\}^n$, because both prices are higher than r_{i+1} , the expected revenue from rejecting (both prices are in fact higher than r_{i+1}^* , an upper bound on r_{i+1}). We say that the policy *makes a mistake* for buyer i if either $\pi_i = \frac{1}{4} + \frac{n-i}{4n}$ when i has the High distribution, or $\pi_i = \frac{1}{2} + \frac{n-i}{4n}$ when i has the Low distribution. If the policy makes a mistake for i , then we have

$$\begin{aligned} r_i - r_{i+1} &\leq \max \left\{ \left(\frac{1}{4} + \frac{n-i}{4n} - r_{i+1} \right) \left(\frac{1}{n} - 16 \frac{\varepsilon}{n} \right), \left(\frac{1}{2} + \frac{n-i}{4n} - r_{i+1} \right) \left(\frac{1}{2n} - 8 \frac{\varepsilon}{n} \right) \right\} \\ &= \frac{1}{4n} - 4 \frac{\varepsilon}{n} + \left(\frac{n-i}{4n} - r_{i+1} \right) \left(\frac{1}{n} - 16 \frac{\varepsilon}{n} \right); \end{aligned}$$

on the other hand, if the policy does not make a mistake for i , then we have

$$r_i - r_{i+1} \leq \max \left\{ \left(\frac{1}{2} + \frac{n-i}{4n} - r_{i+1} \right) \frac{1}{2n}, \left(\frac{1}{4} + \frac{n-i}{4n} - r_{i+1} \right) \frac{1}{n} \right\} = \frac{1}{4n} + \left(\frac{n-i}{4n} - r_{i+1} \right) \frac{1}{n}.$$

Hence, if the policy makes M mistakes for some $M \in \{1, \dots, n\}$, then

$$r_1 = \sum_{i=1}^n (r_i - r_{i+1}) \leq \frac{1}{4} - M \frac{4\varepsilon}{n} + \sum_{i=1}^n \left(\frac{n-i}{4n} - r_{i+1} \right) \frac{1}{n}.$$

Now, observe that $\frac{n-i}{4n} - r_{i+1} = r_{i+1}^* - r_{i+1} \leq \frac{4\varepsilon}{n} \min\{n-i, M\}$, because the loss from making a mistake is at most $4 \frac{\varepsilon}{n}$, and starting from buyer $i+1$, the number of mistakes can be at most $n - (i+1) + 1 = n-i$ and also at most M . Therefore,

$$\begin{aligned} r_1 &\leq \frac{1}{4} - M \frac{4\varepsilon}{n} + \left((n-M) \frac{4\varepsilon}{n} M + \frac{4\varepsilon}{n} (M-1) + \dots + \frac{4\varepsilon}{n} \right) \frac{1}{n} \\ &= \frac{1}{4} - M \frac{4\varepsilon}{n} \left(1 - \frac{n-M}{n} - \frac{M-1}{2n} \right) = \frac{1}{4} - 2\varepsilon \frac{M}{n} \left(\frac{M+1}{n} \right). \end{aligned}$$

Recalling that $r_1^* = 1/4$, this shows the additive error is $\Omega(\varepsilon)$ as long as the fraction of mistakes M/n is a constant.

Computing the Hellinger distance We first analyze the Hellinger distance $H(D_i^H, D_i^L)$ between (a single observation of) D_i^H vs. D_i^L , which we note does not depend on the buyer i . The squared Hellinger distance can be bounded using $1 - H^2(D_i^H, D_i^L)$

$$\begin{aligned} &= \sqrt{\frac{1}{2n} \left(\frac{1}{2n} - 8 \frac{\varepsilon}{n} \right)} + \sqrt{\left(\frac{1}{2n} + 8 \frac{\varepsilon}{n} - 24 \frac{\varepsilon}{n} \right) \left(\frac{1}{2n} + 8 \frac{\varepsilon}{n} \right)} + \sqrt{\left(1 - \frac{1}{n} + 16 \frac{\varepsilon}{n} \right) \left(1 - \frac{1}{n} \right)} \\ &\geq \left(\frac{1}{2n} - \frac{4\varepsilon}{n} - \frac{(8 \frac{\varepsilon}{n})^2}{\frac{1}{2n}} \right) + \left(\frac{1}{2n} - \frac{4\varepsilon}{n} - \frac{(24 \frac{\varepsilon}{n})^2}{\frac{1}{2n} + 8 \frac{\varepsilon}{n}} \right) + \left(1 - \frac{1}{n} + \frac{8\varepsilon}{n} - \frac{(16 \frac{\varepsilon}{n})^2}{1 - \frac{1}{n}} \right) = 1 - O\left(\frac{\varepsilon^2}{n}\right), \end{aligned}$$

where the inequality applies Lemma 6 below to each square root. This shows that $H^2(D_i^H, D_i^L) = O(\frac{\varepsilon^2}{n})$, and the squared Hellinger distance is additive across independent samples. Using the fact that the Total Variation distance is upper-bounded by $\sqrt{2}$ times the Hellinger distance [GS02], we see that the Total Variation distance between T independent samples of D_i^H vs. T independent samples of D_i^L is $O(\sqrt{\frac{T}{n}} \varepsilon)$. The proof of Lemma 6 is deferred because it is elementary.

Lemma 6. Suppose $C \in (0, 1)$ and $x \in [-\frac{3}{4}C, \frac{3}{4}C]$. Then $\sqrt{C(C+x)} \geq C + \frac{x}{2} - \frac{x^2}{C}$.

Completing the proof of Theorem 2 Suppose the distribution of each buyer is equally likely to be High or Low, independently across buyers. Fix any learning algorithm and the constant probability $1/2$. By the computation of Hellinger distance above, if the number of samples T is less than $C \frac{n}{\epsilon^2}$ for some constant C , then for any buyer i , the Total Variation distance between the samples observed under D_i^H vs. D_i^L is at most $1/2$. (C depends on the choice of constant $1/2$.) This means that w.p. at least $1 - 1/2$, the price π_i decided by the learning algorithm cannot depend on whether buyer i had distribution D_i^H vs. D_i^L , which means that there exists an adversarial choice of D_i^H or D_i^L for each buyer i under which the probability of making a mistake for buyer i is at least $(1 - 1/2)/2 = 1/4$.

Let M denote the (random) number of mistakes, under this adversarial configuration of whether each buyer i has distribution D_i^H or D_i^L . We have that $\mathbb{E}[\frac{M}{n}] \geq 1/4$. Although whether the algorithm makes a mistake could be arbitrarily correlated across buyers i , applying $\frac{M}{n} \leq 1$, we can employ Markov's inequality on the random variable $1 - \frac{M}{n}$ to see that

$$\mathbb{P}[(1 - \frac{M}{n}) \geq 7/8] \leq \frac{3/4}{7/8} = \frac{6}{7}.$$

The LHS equals $\mathbb{P}[\frac{M}{n} \leq \frac{1}{8}] = 1 - \mathbb{P}[\frac{M}{n} > \frac{1}{8}]$, and hence $\frac{1}{7} \leq \mathbb{P}[\frac{M}{n} > \frac{1}{8}]$. We have shown that unless $T = \Omega(\frac{n}{\epsilon^2})$, there is probability at least $1/7$ of making a constant fraction of mistakes, in which case we showed above that the additive error would be $\Omega(\epsilon)$. This completes the proof of Theorem 2.

3 Correlated Distributions

In this section we prove our upper bound on the sample complexity of welfare/revenue maximization for correlated buyer distributions. We show nearly matching lower bounds in Appendix A.5.

3.1 Positive Result for Welfare and Revenue: Proof of Theorem 4

To bound the sample complexity of posted pricing for correlated distributions, it suffices to bound the pseudo-dimension of the policy class $\Pi_{\mathcal{S}}$. We use the same approach for both the welfare and revenue objectives. By standard learning theory results [BBL03], bounds on the pseudo-dimension translate to bounds on the sample complexity as follows.

Theorem 6. Let $\text{PDim}(\Pi_{\mathcal{S}})$ denote the pseudo-dimension of $\Pi_{\mathcal{S}}$. For any $\epsilon > 0$, any $\delta \in (0, 1)$ and any distribution $\mathbf{D}$ over $[0, 1]^n$, $T = O(\frac{1}{\epsilon^2}(\text{PDim}(\Pi_{\mathcal{S}}) + \log \frac{1}{\delta}))$ samples are sufficient to ensure that with probability at least $1 - \delta$ over the draw of samples $(\mathbf{v}_1, \dots, \mathbf{v}_T) \sim \mathbf{D}^T$, for all $\pi \in \Pi_{\mathcal{S}}$,

$$\left| \frac{1}{T} \sum_{t=1}^T \pi(\mathbf{v}_t) - \pi(\mathbf{D}) \right| \leq \epsilon.$$

We will show that $\text{PDim}(\Pi_{\mathcal{S}}) = O(k \log k)$, where $k = |\mathcal{S}| + 1$. This together with the above theorem immediately implies that Theorem 4 holds for the sample average approximation algorithm.

Let $\mathcal{S} = \{i_1, i_2, i_3, \dots, i_{k-1}\}$, where $1 < i_1 < i_2 < \dots < i_{k-1}$. For each $j = 1, 2, \dots, k$, let $I_j = \{i_{j-1}, \dots, i_j - 1\}$, with the convention that $i_0 = 1$ and $i_k = n + 1$. In other words, $I_1, I_2, \dots, I_k$ are consecutive intervals that partition $[n]$, and $\Pi_{\mathcal{S}}$ is the class of policies that offers every customer in I_j the same price. Note that every policy $\pi \in \Pi_{\mathcal{S}}$ can be parameterized by k prices $\rho = (\rho_1, \rho_2, \dots, \rho_k)$, where ρ_j is the price offered to customers in I_j . In the rest of this proof, we will use π_{ρ} to denote the policy in $\Pi_{\mathcal{S}}$ parameterized by ρ .

By the definition of pseudo-dimension,

$$\text{PDim}(\Pi_{\mathcal{S}}) = \text{VCdim}(\tilde{\Pi}_{\mathcal{S}}), \quad (2)$$

where $\tilde{\Pi}_{\mathcal{S}} := \{(\mathbf{v}, z) \mapsto \mathbb{1}\{\pi_{\rho}(\mathbf{v}) \geq z\} : \rho \in \mathbb{R}^k\}$. Let $\tilde{\Pi}_{\mathcal{S}}^*$ denote the dual class of $\tilde{\Pi}_{\mathcal{S}}$, so

$$\tilde{\Pi}_{\mathcal{S}}^* := \{\rho \mapsto \mathbb{1}\{\pi_{\rho}(\mathbf{v}) \geq z\} : \mathbf{v} \in [0, 1]^n, z \in \mathbb{R}\}.$$

We will use a result from [BDD⁺21] to bound the VC dimension of $\tilde{\Pi}_{\mathcal{S}}$. We state this result below in Definition 1 and Theorem 7. This result essentially says that the pseudo-dimension of the primal class is bounded if the dual class is well-structured. Here, “well-structured” essentially means that the domain can be partitioned into pieces defined by a small number of boundary functions, and the function is simple on each piece.

Definition 1 (Definition 3.2 in [BDD⁺21]). A function class $\mathcal{H} \subseteq \mathbb{R}^{\mathcal{Y}}$ that maps a domain $\mathcal{Y}$ to $\mathbb{R}$ is $(\mathcal{F}, \mathcal{G}, l)$ -piecewise decomposable for a class $\mathcal{G} \subseteq \{0, 1\}^{\mathcal{Y}}$ of boundary functions and a class $\mathcal{F} \subseteq \mathbb{R}^{\mathcal{Y}}$ of piece functions if the following holds: for every $h \in \mathcal{H}$, there are l boundary functions $g^{(1)}, \dots, g^{(l)} \in \mathcal{G}$ and a piece function $f_b \in \mathcal{F}$ for each bit vector $b \in \{0, 1\}^l$ such that for all $y \in \mathcal{Y}$, $h(y) = f_{b_y}(y)$ where $b_y = (g^{(1)}(y), \dots, g^{(l)}(y)) \in \{0, 1\}^l$.

The main theorem in [BDD⁺21] states that if the dual class is $(\mathcal{F}, \mathcal{G}, l)$ -piecewise decomposable, then the pseudo-dimension of the primal class is bounded.

Theorem 7 (Theorem 3.3 in [BDD⁺21]). Suppose that the dual function class $\mathcal{U}^*$ is $(\mathcal{F}, \mathcal{G}, l)$ -piecewise decomposable with boundary functions $\mathcal{G} \subseteq \{0, 1\}^{\mathcal{U}}$ and piece functions $\mathcal{F} \subseteq \mathbb{R}^{\mathcal{U}}$. The pseudo-dimension of $\mathcal{U}$ is bounded as follows:

$$\text{PDim}(\mathcal{U}) = O((\text{PDim}(\mathcal{F}^*) + \text{VCdim}(\mathcal{G}^*)) \ln(\text{PDim}(\mathcal{F}^*) + \text{VCdim}(\mathcal{G}^*)) + \text{VCdim}(\mathcal{G}^*) \ln l).$$

We now apply Theorem 7 to bound the VC dimension of $\tilde{\Pi}_{\mathcal{S}}$.⁵ To do so we must show that the dual class $\tilde{\Pi}_{\mathcal{S}}^*$ is $(\mathcal{F}, \mathcal{G}, l)$ -piecewise decomposable for “nice” classes $\mathcal{F}$ and $\mathcal{G}$. We will show that for both the welfare and revenue objectives, we can take $\mathcal{F}$ to be the family of constant functions and $\mathcal{G}$ to be the family of axis-aligned halfspaces. This will follow from the below lemma, whose proof is deferred to the Appendix for space reasons.

Lemma 7. Let $\mathbf{v} \in [0, 1]^n$ and let $z \in \mathbb{R}$. Let $G(\mathbf{v}, z) = \{\rho \in \mathbb{R}^k : \pi_{\rho}(\mathbf{v}) \geq z\}$. For both the welfare and revenue objectives, there are $l_1, u_1, \dots, l_k, u_k \in \mathbb{R} \cup \{\pm\infty\}$ such that

$$G(\mathbf{v}, z) = \bigcup_{j=1}^k (u_1, \infty) \times \dots \times (u_{j-1}, \infty) \times (l_j, u_j] \times \mathbb{R}^{k-j}. \quad (3)$$

Corollary 1. $\tilde{\Pi}_{\mathcal{S}}^*$ is $(\mathcal{F}, \mathcal{G}, l)$ -piecewise decomposable, where $\mathcal{F}$ is the set of constant functions, $\mathcal{G}$ is the set of axis-aligned halfspaces, and $l = 2(|\mathcal{S}| + 1)$.

Proof. Consider a function $h \in \tilde{\Pi}_{\mathcal{S}}^*$. By definition of $\tilde{\Pi}_{\mathcal{S}}^*$, there exist $\mathbf{v} \in [0, 1]^n$ and $z \in \mathbb{R}$ such that $h(\rho) = \mathbb{1}\{\pi_{\rho}(\mathbf{v}) \geq z\}$. By Lemma 7, there are $l_1, u_1, \dots, l_k, u_k$ such that

$$\{\rho \in \mathbb{R}^k : h(\rho) = 1\} = \bigcup_{j=1}^k (u_1, \infty) \times \dots \times (u_{j-1}, \infty) \times (l_j, u_j] \times \mathbb{R}^{k-j}.$$

For each $j \in [k]$, let $L_j = \{\rho \in \mathbb{R}^k : \rho_j > l_j\}$ and $U_j = \{\rho \in \mathbb{R}^k : \rho_j \leq u_j\}$. Note that L_j and U_j are axis-aligned halfspaces, and

$$\{\rho \in \mathbb{R}^k : h(\rho) = 1\} = \bigcup_{j=1}^k \bar{U}_1 \cap \bar{U}_2 \cap \dots \cap \bar{U}_{j-1} \cap L_j \cap U_j.$$

Therefore, in the definition of $(\mathcal{F}, \mathcal{G}, l)$ -piecewise decomposable, we may take the boundary functions to be the $2k$ functions corresponding to the halfspaces $L_1, U_1, \dots, L_k, U_k$. On any given piece defined by these boundary functions, h is a constant function (equal to either 0 or 1). $\square$

For $\mathcal{F}$ the set of constant functions and $\mathcal{G}$ the set of axis-aligned halfspaces in $\mathbb{R}^k$, it is easy to check that $\text{PDim}(\mathcal{F}^*) = 0$ and $\text{VCdim}(\mathcal{G}^*) = k$. Combining Corollary 1 with Theorem 7, we get that

$$\text{VCdim}(\tilde{\Pi}_{\mathcal{S}}) = O(k \ln k + k \ln 2k) = O(k \ln k).$$

This completes the proof of Theorem 4.

Remark. If we directly apply Theorem 7 from [BDD⁺21] to bound $\text{PDim}(\Pi_{\mathcal{S}})$, then we would get

- A $O(k \ln k)$ pseudo-dimension bound for the revenue objective;
- A $O(k \ln(kn))$ pseudo-dimension bound for welfare objective.

In particular, the bound for welfare would grow with n . This dependence on n is in fact unavoidable if one works with $\Pi_{\mathcal{S}}$ directly, because for instances like $\mathbf{v} = (\frac{1}{n}, \frac{2}{n}, \dots, 1)$, the dual function corresponding to $\mathbf{v}$ is piecewise constant with $\Theta(n)$ pieces, even if $k = 1$. This is why our Corollary 1 focuses on the indicator functions $\tilde{\Pi}_{\mathcal{S}}^*$ instead. Regardless, both the bounds for welfare and revenue require our Lemma 7 that analyzes the problem-specific structure of the “good sets”.

⁵Note Theorem 7 is stated to bound the pseudo-dimension, but the pseudo-dimension coincides with the VC dimension for function classes consisting of $\{0, 1\}$ -valued functions.

Acknowledgement This work was done in part while the authors were visiting the Simons Institute for the Theory of Computing for the program on Data-Driven Decision Processes. The authors thank Zhuoxin Chen for identifying typos in an early version.

References

- [AKS21] Sepehr Assadi, Thomas Kesselheim, and Sahil Singla. Improved truthful mechanisms for subadditive combinatorial auctions: Breaking the logarithmic barrier. In *Proceedings of SODA*, pages 653–661, 2021.
- [AKW14] Pablo Daniel Azar, Robert Kleinberg, and S. Matthew Weinberg. Prophet inequalities with limited information. In Chandra Chekuri, editor, *ACM-SIAM Symposium on Discrete Algorithms, SODA*, pages 1358–1377, 2014.
- [AM⁺06] Lawrence M Ausubel, Paul Milgrom, et al. The lovely but lonely vickrey auction. *Combinatorial auctions*, 17(3):22–26, 2006.
- [BBL03] Olivier Bousquet, Stéphane Boucheron, and Gábor Lugosi. Introduction to statistical learning theory. In *Summer School on Machine Learning*, pages 169–207. Springer, 2003.
- [BDD⁺21] Maria-Florina Balcan, Dan DeBlasio, Travis Dick, Carl Kingsford, Tuomas Sandholm, and Ellen Vitercik. How much data is sufficient to learn high-performing algorithms? generalization guarantees for data-driven algorithm design. In *Proceedings of the 53rd Annual ACM SIGACT Symposium on Theory of Computing, STOC*, page 919–932, 2021.
- [CC23] José R. Correa and Andrés Cristi. A constant factor prophet inequality for online combinatorial auctions. In *Proceedings of STOC*, pages 686–697, 2023.
- [CDF⁺22] Constantine Caramanis, Paul Dütting, Matthew Faw, Federico Fusco, Philip Lazos, Stefano Leonardi, Orestis Papadigenopoulos, Emmanouil Pountourakis, and Rebecca Reiffenhäuser. Single-sample prophet inequalities via greedy-ordered selection. In *ACM-SIAM Symposium on Discrete Algorithms, SODA*, pages 1298–1325, 2022.
- [CDFS22] José Correa, Paul Dütting, Felix A. Fischer, and Kevin Schewior. Prophet inequalities for independent and identically distributed random variables from an unknown distribution. *Math. Oper. Res.*, 47(2):1287–1309, 2022.
- [CFH⁺19] Jose Correa, Patricio Foncea, Ruben Hoeksma, Tim Oosterwijk, and Tjark Vredeveld. Recent developments in prophet inequalities. *ACM SIGecom Exchanges*, 17(1):61–70, 2019.
- [CFPV19] José R. Correa, Patricio Foncea, Dana Pizarro, and Victor Verdugo. From pricing to prophets, and back! *Oper. Res. Lett.*, 47(1):25–29, 2019.
- [CHK07] Shuchi Chawla, Jason D. Hartline, and Robert Kleinberg. Algorithmic pricing via virtual valuations. In *Proceedings of the 8th ACM Conference on Electronic Commerce, EC '07*, page 243–251, 2007.
- [CMS10] Shuchi Chawla, David L. Malec, and Balasubramanian Sivan. The power of randomness in bayesian optimal mechanism design. In *Proceedings of the 11th ACM Conference on Electronic Commerce, EC '10*, page 149–158, 2010.
- [CR14] Richard Cole and Tim Roughgarden. The sample complexity of revenue maximization. In *Symposium on Theory of Computing, STOC*, pages 243–252. ACM, 2014.
- [DK19] Paul Dütting and Thomas Kesselheim. Posted pricing and prophet inequalities with inaccurate priors. In *Conference on Economics and Computation, EC*, pages 111–129. ACM, 2019.
- [DKL20] Paul Dütting, Thomas Kesselheim, and Brendan Lucier. An $o(\log \log m)$ prophet inequality for subadditive combinatorial auctions. In *Proceedings of FOCS*, pages 306–317, 2020.

- [DKL⁺24] Paul Duetting, Thomas Kesselheim, Brendan Lucier, Rebecca Reiffenhausser, and Sahil Singla. Online combinatorial allocations and auctions with few samples. In *65th IEEE Annual Symposium on Foundations of Computer Science, FOCS*, 2024.
- [FGL15] Michal Feldman, Nick Gravin, and Brendan Lucier. Combinatorial auctions via posted prices. In *Proceedings of SODA*, 2015.
- [Fre75] David A Freedman. On tail probabilities for martingales. *the Annals of Probability*, pages 100–118, 1975.
- [GHTZ21] Chenghao Guo, Zhiyi Huang, Zhihao Gavin Tang, and Xinzhi Zhang. Generalizing complex hypotheses on product distributions: Auctions, prophet inequalities, and pandora’s problem. In *Conference on Learning Theory, COLT*, pages 2248–2288. PMLR, 2021.
- [GHZ19] Chenghao Guo, Zhiyi Huang, and Xinzhi Zhang. Settling the sample complexity of single-parameter revenue maximization. In *Proceedings of the 51st Annual ACM SIGACT Symposium on Theory of Computing, STOC*, pages 662–673. ACM, 2019.
- [GKSW24] Khashayar Gatmiry, Thomas Kesselheim, Sahil Singla, and Yifan Wang. Bandit algorithms for prophet inequality and pandora’s box. In *Proceedings of the thirty-sixth annual ACM-SIAM symposium on Discrete Algorithms, SODA*, 2024.
- [GM22] Xinyi Guan and Velibor V Mišić. Randomized policy optimization for optimal stopping. *arXiv preprint arXiv:2203.13446*, 2022.
- [GS02] Alison L Gibbs and Francis Edward Su. On choosing and bounding probability metrics. *International statistical review*, 70(3):419–435, 2002.
- [Har13] Jason D Hartline. Mechanism design and approximation. *Book draft. October*, 122(1), 2013.
- [HKS07] Mohammad Taghi Hajiaghayi, Robert D. Kleinberg, and Tuomas Sandholm. Automated online mechanism design and prophet inequalities. In *Proceedings of the Twenty-Second AAAI Conference on Artificial Intelligence*, pages 58–65. AAAI Press, 2007.
- [HR09] Jason D. Hartline and Tim Roughgarden. Simple versus optimal mechanisms. In *Proceedings 10th ACM Conference on Electronic Commerce (EC-2009)*, pages 225–234. ACM, 2009.
- [ISW20] Nicole Immorlica, Sahil Singla, and Bo Waggoner. Prophet inequalities with linear correlations and augmentations. In *Conference on Economics and Computation, EC*, pages 159–185, 2020.
- [KL03] Robert D. Kleinberg and Frank Thomson Leighton. The value of knowing a demand curve: Bounds on regret for online posted-price auctions. In *44th Symposium on Foundations of Computer Science, FOCS*, pages 594–605. IEEE Computer Society, 2003.
- [KS77] U. Krengel and L. Sucheston. Semiamarts and finite values. *Bulletin of the American Mathematical Society*, 83:745–747, 1977.
- [KS78] U. Krengel and L. Sucheston. On semiamarts, amarts, and processes with finite value. *Advances in Probability and Related Topics*, 4:197–266, 1978.
- [LPS24] Vasilis Livanos, Kalen Patton, and Sahil Singla. Improved mechanisms and prophet inequalities for graphical dependencies. In *Conference on Economics and Computation, EC*, 2024.
- [LSTW23] Renato Paes Leme, Balasubramanian Sivan, Yifeng Teng, and Pratik Worah. Pricing query complexity of revenue maximization. In *Proceedings of the thirty-fourth annual ACM-SIAM symposium on Discrete Algorithms*, 2023.

- [Luc17] Brendan Lucier. An economic view of prophet inequalities. *ACM SIGecom Exchanges*, 16(1):24–47, 2017.
- [MR16] Jamie Morgenstern and Tim Roughgarden. Learning simple auctions. In *Proceedings of the 29th Conference on Learning Theory, COLT 2016*, volume 49 of *JMLR Workshop and Conference Proceedings*, pages 1298–1318. JMLR.org, 2016.
- [Mye81] Roger B. Myerson. Optimal auction design. *Math. Oper. Res.*, 6(1):58–73, 1981.
- [Rou16] Tim Roughgarden. *Twenty lectures on algorithmic game theory*. Cambridge University Press, 2016.
- [RWW] Aviad Rubinstein, Jack Z. Wang, and S. Matthew Weinberg. Optimal single-choice prophet inequalities from samples. In *11th Innovations in Theoretical Computer Science Conference, ITCS*, volume 151, pages 60:1–60:10.
- [SW24] Sahil Singla and Yifan Wang. Bandit sequential posted pricing via half-concavity. In *Conference on Economics and Computation, EC*. ACM, 2024.
- [Vic61] William Vickrey. Counterspeculation, auctions, and competitive sealed tenders. *The Journal of finance*, 16(1):8–37, 1961.
- [Wai19] Martin J Wainwright. *High-dimensional statistics: A non-asymptotic viewpoint*, volume 48. Cambridge university press, 2019.
- [XMX24] Yaqi Xie, Will Ma, and Linwei Xin. Vc theory for inventory policies. *arXiv preprint arXiv:2404.11509*, 2024.

A Appendix

A.1 Proof of Lemma 2

Proof. Now, consider the smallest j so that $\hat{r}_j \geq r_j^*$. Again, we know such a j exists because $0 = \hat{r}_{n+1} \geq r_{n+1}^* = 0$. We claim that that $\hat{r}_1 \geq r_1^* + \min\{\sum_{i=1}^{j-1} \eta_i, 0\}$.

Observe that there is nothing to be shown if $j = 1$ because then $\hat{r}_1 \geq r_1^*$. So, let's consider the case that $j > 1$. We can now write $\hat{r}_1$ as a telescoping sum

$$\hat{r}_1 = \hat{r}_{j-1} + \sum_{i=1}^{j-2} (\hat{r}_i - \hat{r}_{i+1}) . \quad (4)$$

Since $\hat{r}_j \geq r_j^*$, we have

$$\begin{aligned} \hat{r}_{j-1} &= \hat{r}_j + \frac{1}{T} \sum_{t=1}^T (V_{j-1}^{(t)} - \hat{r}_j)^+ = \hat{r}_j + \mathbb{E}[(V_{j-1} - \hat{r}_j)^+] + \eta_{j-1} \\ &= \mathbb{E}[\max\{V_{j-1}, \hat{r}_j\}] + \eta_{j-1} \geq \mathbb{E}[\max\{V_{j-1}, r_j^*\}] + \eta_{j-1} = r_{j-1}^* + \eta_{j-1} . \end{aligned}$$

Observe that for $i + 1 < j$, we have $\hat{r}_{i+1} < r_{i+1}^*$. Therefore

$$\hat{r}_i - \hat{r}_{i+1} = \frac{1}{T} \sum_{t=1}^T (V_i^{(t)} - \hat{r}_{i+1})^+ = \mathbb{E}[(V_i - \hat{r}_{i+1})^+] + \eta_i \geq \mathbb{E}[(V_i - r_{i+1}^*)^+] + \eta_i = r_i^* - r_{i+1}^* + \eta_i .$$

So, in combination with (4),

$$\hat{r}_1 \geq r_{j-1}^* + \eta_{j-1} + \sum_{i=1}^{j-2} (r_i^* - r_{i+1}^* + \eta_i) = r_1^* + \sum_{i=1}^{j-1} \eta_i . \quad \square$$

A.2 Proof of Lemma 5

Ignoring constants, a slick proof of this lemma is via another application of Freedman's inequality, where we define a martingale $X_i = Y_i - \mathbb{E}[Y_i | Y_1, \dots, Y_{i-1}]$, and cut this martingale (stop) whenever the sum of conditional variances exceeds $\log(1/\delta)$. This bounds the sum of conditional variances by $\log(1/\delta)$, i.e. $W_\tau \leq \sigma^2 := \log(1/\delta)$ w.p. 1. Since the lemma is upper-bounding the probability of the event that $\sum_i X_i > \log(1/\delta)$, we can show that this probability is $O(\delta)$ by substituting $\ell = \log(1/\delta)$ into Freedman's inequality.

Below we give another, elementary proof of this lemma.

Proof. Let $Z = 1$ if $\sum_{i=1}^n \mathbb{E}[Y_i | Y_1, \dots, Y_{i-1}] \geq \frac{e}{e-1} \ln(\frac{e}{\delta})$. By Markov's inequality, we have

$$\Pr \left[\sum_{i=1}^n Y_i \leq 1 \text{ and } Z = 1 \right] = \Pr \left[Z e^{-\sum_{i=1}^n Y_i} \geq e^{-1} \right] \leq \mathbb{E} \left[Z e^{-\sum_{i=1}^n Y_i} \right] \cdot e$$

Now, we have

$$\mathbb{E} \left[Z e^{-\sum_{i=1}^n Y_i} \right] = \mathbb{E} \left[Z \prod_{i=1}^n e^{-Y_i} \right] = \mathbb{E} \left[\mathbb{E}[Z | Y_1, \dots, Y_n] \prod_{i=1}^n \mathbb{E}[e^{-Y_i} | Y_1, \dots, Y_{i-1}] \right]$$

For every single conditional expectation, we obtain by convexity

$$\begin{aligned} \mathbb{E}[e^{-Y_i} | Y_1, \dots, Y_{i-1}] &\leq \mathbb{E}[Y_i e^{-1} + (1 - Y_i) | Y_1, \dots, Y_{i-1}] \\ &= \mathbb{E}[Y_i(e^{-1} - 1) | Y_1, \dots, Y_{i-1}] + 1 \\ &\leq \exp(\mathbb{E}[Y_i | Y_1, \dots, Y_{i-1}](e^{-1} - 1)) . \end{aligned}$$

Now, consider a fixed realization of $Y_1, \dots, Y_n$. This realization fully determines Z . If $Z = 0$, then

$$\mathbb{E}[Z | Y_1, \dots, Y_n] \prod_{i=1}^n \mathbb{E}[e^{-Y_i} | Y_1, \dots, Y_{i-1}] = 0 .$$

Otherwise, with $Z = 1$, we have

$$\mathbb{E}[Z | Y_1, \dots, Y_n] \prod_{i=1}^n \mathbb{E}[e^{-Y_i} | Y_1, \dots, Y_{i-1}] \leq \exp \left(\sum_{i=1}^n \mathbb{E}[Y_i | Y_1, \dots, Y_{i-1}](e^{-1} - 1) \right) \leq \frac{\delta}{e} .$$

So, we have an upper bound of $\frac{\delta}{e}$ regardless of the realization. Taking an expectation, we get

$$\mathbb{E} \left[\mathbb{E} [Z \mid Y_1, \dots, Y_n] \prod_{i=1}^n \mathbb{E} [e^{-Y_i} \mid Y_1, \dots, Y_{i-1}] \right] \leq \frac{\delta}{e},$$

which implies $\Pr [\sum_{i=1}^n Y_i \leq 1 \text{ and } Z = 1] \leq \mathbb{E} [Ze^{-\sum_{i=1}^n Y_i}] \cdot e \leq \frac{\delta}{e} e = \delta$, completing the proof. $\square$

A.3 Proof of Lemma 6.

Proof. Let $f(x) = \sqrt{C(C+x)} = (C^2 + Cx)^{1/2}$, which is a continuous function over $x \in [-\frac{3}{4}C, \frac{3}{4}C]$. We have $f'(x) = \frac{1}{2}(C^2 + Cx)^{-1/2}C$ and $f''(x) = -\frac{1}{4}(C^2 + Cx)^{-3/2}C^2$, both of which exist over $x \in [-\frac{3}{4}C, \frac{3}{4}C]$. Applying Taylor's theorem around 0, we get

$$\begin{aligned} f(x) &= f(0) + f'(0)x + \frac{f''(y)}{2}x^2 = C + \frac{1}{2}x - \frac{1}{8}(C^2 + Cy)^{-3/2}C^2x^2 \\ &\geq C + \frac{x}{2} - \frac{1}{8}(C^2 + C(-\frac{3}{4}C))^{-3/2}C^2x^2 = C + \frac{x}{2} - \frac{x^2}{C}, \end{aligned}$$

where y lies between 0 and x and the inequality holds because $y \geq -\frac{3}{4}C$, completing the proof. $\square$

A.4 Proof of Lemma 7

If $z \leq 0$ then $G(\mathbf{v}, z) = \mathbb{R}^k$, so we may take any values of $l_1, u_1, \dots, l_k, u_k$ such that $l_1 = l_2 = \dots = l_k - \infty$ and $u_k = \infty$.

For the rest of the proof, assume $z > 0$. First, consider the welfare objective. We show the lemma holds for l_j and u_j as follows:

Define l_j to be $-\infty$ if $\mathbf{v}(I_j)_1 \geq z$, and otherwise to be $\max\{\mathbf{v}(I_j)_1, \dots, \mathbf{v}(I_j)_{m-1}\}$, where m is the smallest index such that $\mathbf{v}(I_j)_m \geq z$. (If $\mathbf{v}(I_j)_m < z$ for all m , set $l_j = \max(\mathbf{v}(I_j))$.)

Define u_j to be $\max(\mathbf{v}(I_j))$.

Let $G = \{\rho \in \mathbb{R}^k : \pi_\rho(\mathbf{v}) \geq z\}$ and $G' = \bigcup_{j=1}^k (u_1, \infty) \times \dots \times (u_{j-1}, \infty) \times (l_j, u_j] \times \mathbb{R}^{k-j}$. To show that $G = G'$, we'll show $G \subseteq G'$ and $G' \subseteq G$.

Case 1 ($G \subseteq G'$). If $G = \emptyset$ we are done, so assume otherwise. Let $\rho \in G$. Let j be the interval such that the algorithm (using thresholds ρ) accepts a value in $\mathbf{v}(I_j)$. Since a value in interval j was accepted, $\rho_j \leq \max(\mathbf{v}(I_j)) = u_j$. Also, since no value in the previous intervals were accepted, $\rho_i > u_i$ for all $i < j$. Finally, we must have $\rho_j > l_j$, since otherwise the accepted value will be less than z . Thus $\rho \in G'$.

Case 2 ($G' \subseteq G$). Let $\rho \in G'$. Then $\rho \in (u_1, \infty) \times \dots \times (u_{j-1}, \infty) \times (l_j, u_j] \times \mathbb{R}^{k-j}$ for some j . Since $\rho_i > u_i$ for all $i < j$, the algorithm using thresholds ρ does not accept any value in the intervals $i < j$. Since $\rho_j \in (l_j, u_j]$, the definitions of l_j and u_j imply that a value in $\mathbf{v}(I_j)$ is accepted, and this value is greater than or equal to z . Thus $\rho \in G$.

For revenue, the only difference is that we define l_j to be $\min(z, \max(\mathbf{v}(I_j)))$. (u_j is still defined to be $\max(\mathbf{v}(I_j))$.) With these definitions of l_j and u_j , a very similar analysis shows that $G = G'$ in the revenue case as well.

A.5 Proof of Theorem 5

We let $\mathcal{S}' \subseteq \{1\} \cup \mathcal{S}$ be a subset of decision points such that price π_i can be freely chosen for all $i \in \mathcal{S}'$. For welfare maximization we require $i+1 \in \{1, \dots, n\} \setminus \mathcal{S}'$ for all $i \in \mathcal{S}'$, which can be achieved while ensuring $|\mathcal{S}'| \geq \lfloor \frac{1+|\mathcal{S}|}{2} \rfloor$. Each decision point $i \in \mathcal{S}'$ has marginal value distribution(s) that could be “High” or “Low”, separately for each decision point (i.e., all $2^{|\mathcal{S}'|}$ combinations are possible). For welfare maximization, they are defined as

- High: $V_i = 1/2$ with probability (w.p.) 1, $V_{i+1} = 1$ w.p. $1/2 + \varepsilon$, $V_{i+1} = 0$ w.p. $1/2 - \varepsilon$;
- Low: $V_i = 1/2$ w.p. 1, $V_{i+1} = 1$ w.p. $1/2 - \varepsilon$, $V_{i+1} = 0$ w.p. $1/2 + \varepsilon$.

For revenue maximization, they are defined as

- High: $V_i = 1$ w.p. $1/2 + \varepsilon$, $V_i = 1/2$ w.p. $1/2 - \varepsilon$;
- Low: $V_i = 1$ w.p. $1/2 - \varepsilon$, $V_i = 1/2$ w.p. $1/2 + \varepsilon$.

The overall distribution over trajectories $\mathbf{V}$ is then correlated as follows. First, a decision point $\tilde{i} \in \mathcal{S}'$ is drawn uniformly at random. Then, $V_{\tilde{i}}$ (as well as $V_{\tilde{i}+1}$ in the case of welfare maximization) is drawn according to the distributions above, depending on whether decision point $\tilde{i}$ is High or Low. All other buyer values are 0 on this trajectory.

Given this, only the prices for decision point $\tilde{i}$ are relevant, which can without loss be restricted to $\{1/2, 1\}$. The policy should optimize π_i for each $i \in \mathcal{S}'$ as if $\tilde{i} = i$. For either welfare or revenue maximization, the constructions above can be checked to satisfy the following properties:

1. Setting π_i (and π_{i+1} in the case of welfare maximization) to 1 is optimal for High and earns objective $1/2 + \varepsilon$ in expectation; setting to $1/2$ is optimal for Low and earns objective $1/2$;
2. Setting the wrong price earns expected objective $1/2$ for High and $1/2 - \varepsilon$ for Low, incurring a loss of ε compared to optimal in both cases.
3. The High and Low distributions have the same support and probabilities that differ by 2ε .

We note that for welfare maximization, it does not matter whether $\pi_{i+1} \in \mathcal{S}$, because the decision is on whether to accept buyer i when $V_i = 1/2$.

Finally, a policy has additive error $\Omega(\varepsilon)$ as long as it has constant probability of setting the wrong price for a decision point. Due to property 3. above, the policy needs $\Omega(\frac{1}{\varepsilon^2})$ relevant observations for a given decision point to avoid setting the wrong price for it with constant probability. However, each sample contains a relevant observation for a given decision point only with probability $1/|\mathcal{S}'|$, so $\Omega(\frac{|\mathcal{S}'|}{\varepsilon^2})$ samples are necessary for the number of relevant observations to be $\Omega(\frac{1}{\varepsilon^2})$ with high probability. Because $|\mathcal{S}'| \geq \lfloor \frac{1+|\mathcal{S}|}{2} \rfloor$, this completes the proof of Theorem 5.

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [\[Yes\]](#)

Justification: The abstract and the introduction clearly reflect our contributions.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: We discuss in the related work section topics that are not in this paper's scope.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory Assumptions and Proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [\[Yes\]](#)

Justification: We clearly state all the assumptions needed for our theorems, and all the proofs are either in the body or the appendix.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental Result Reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [NA]

Justification: This paper does not contain experiments.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [NA]

Justification: This paper does not contain experiments.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental Setting/Details

Question: Does the paper specify all the training and test details (e.g., data splits, hyper-parameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [NA]

Justification: This paper does not contain experiments.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment Statistical Significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [NA]

Justification: This paper does not contain experiments.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer “Yes” if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.

- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments Compute Resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [NA]

Justification: This paper does not contain experiments.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code Of Ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: The research in this paper conforms with the NeurIPS Code of Ethics.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader Impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [NA]

Justification: There is no societal impact of the work performed.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to

generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.

- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: The paper poses no such risks.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [NA]

Justification: The paper does not use existing assets.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. New Assets

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [NA]

Justification: The paper does not release new assets.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. **Crowdsourcing and Research with Human Subjects**

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [NA]

Justification: The paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. **Institutional Review Board (IRB) Approvals or Equivalent for Research with Human Subjects**

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

Justification: The paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.