
PASTA: Pessimistic Assortment Optimization

Juncheng Dong^{*1} Weibin Mo^{*2} Zhengling Qi³ Cong Shi⁴ Ethan X. Fang⁵ Vahid Tarokh¹

Abstract

We consider a class of assortment optimization problems in an *offline* data-driven setting. A firm does not know the underlying customer choice model but has access to an offline dataset consisting of the historically offered assortment set, customer choice, and revenue. The objective is to use the offline dataset to find an optimal assortment. Due to the combinatorial nature of assortment optimization, the problem of insufficient data coverage is likely to occur in the offline dataset. Therefore, designing a provably efficient offline learning algorithm becomes a significant challenge. To this end, we propose an algorithm referred to as *Pessimistic Assortment Optimization* (PASTA for short) designed based on the principle of pessimism, that can correctly identify the optimal assortment by only requiring the offline data to cover the optimal assortment under general settings. In particular, we establish a regret bound for the offline assortment optimization problem under the celebrated multinomial logit model. We also propose an efficient computational procedure to solve our pessimistic assortment optimization problem. Numerical studies demonstrate the superiority of the proposed method over the existing baseline method.

1. Introduction

One of the most critical problems faced by a seller is to select products for presentation to potential buyers. Often faced with limited display spaces and storage costs in both

brick-and-mortar and online retailing, the seller needs to carefully choose a set of products from the vast collection of all available products for displaying to its customers. In this line, the problem of selecting an *assortment*, i.e., a collection of products from all available products, in order to maximize the seller’s revenue is the *assortment optimization problem*. Obviously, the choice behavior of customers (McFadden, 1981) is of great importance in the problem of assortment optimization. Without loss of generality, we assume the choice of each customer can be described by a preference vector θ^c . This subsumes the seminal multinomial logit (MNL) model (McFadden, 1973) which is arguably the most well-studied and widespread models in assortment optimization literature (Please see Section 6 for more details) (Talluri & van Ryzin, 2004; Caro & Gallien, 2007; Rusmevichientong et al., 2010; Davis et al., 2013; Chen et al., 2021a; Aouad et al., 2022).

In practice, θ^c is often unknown and needs to be estimated. Assuming no historical data of customers, *dynamic assortment optimization* adaptively learns θ^c in a trial-and-error fashion by updating the assortment and observing the subsequent choices of customers sequentially (Caro & Gallien, 2007; Chen et al., 2020; Rusmevichientong et al., 2020; Chen et al., 2021b; Li et al., 2022). Meanwhile, in our era of Big Data, companies often collect abundant customer data. Therefore, it is often in companies’ best interest to learn from the existing (potentially massive) offline datasets rather than starting from scratch. Moreover, offline learning is beneficial since online exploration can sometimes be expensive or infeasible. Hence, we take the first stab to formally study the following important question faced by every seller.

Research Question: *Given a pre-collected offline dataset of historically offered assortment, customers choices, and revenue, how can we find an efficient and theoretically justified offline algorithm to estimate the optimal assortment set without unrealistic assumptions on the offline dataset?*

When the dataset is not adaptively collected, it is not uncommon to encounter the challenge of insufficient coverage of data. For estimators such as the maximum likelihood estimator (MLE) to approximate θ^c accurately, the offline dataset must include sufficiently many assortments and customer choices. In other words, the data-collecting process needs to sufficiently explore different assortments (by the

^{*}Equal contribution ¹Department of Electrical and Computer Engineering, Duke University, Durham, NC 27705, United States. ²Mitchell E. Daniels, Jr. School of Business, Purdue University, West Lafayette, IN 47907, United States. ³Department of Decision Sciences, George Washington University, Washington, DC 20052, United States. ⁴Herbert Business School, University of Miami, Coral Gables, FL 33146, United States. ⁵Department of Biostatistics and Bioinformatics, Duke University, Durham, NC 27705, United States.. Correspondence to: Zhengling Qi <qizhengling@gwu.edu>.

seller) and different choices (by the customers). This is unlikely to happen for offline datasets because the seller would not choose unreasonable assortments whose expected revenues are obviously suboptimal, and the customers would not choose products against their preferences.

Major Contributions. The main contribution of this work is two-fold. First, based on the principle of pessimism, we propose the *Pessimistic Assortment Optimization* (PASTA for short) framework, which correctly identifies the optimal assortment. In particular, our framework only requires that the offline dataset covers the optimal assortment set instead of all possible (combinatorially many) assortment sets. Second, we derive the first finite-sample regret bound for offline assortment optimization under the multinomial logit (MNL) model (Please see Section 6), one of the most widely used models for modeling customers' choices. We subsequently propose an algorithm, also with the name *PASTA*, that can efficiently solve the pessimistic assortment optimization problems. Experiments on the simulated datasets (so that θ is known) corroborate the efficacy of pessimistic assortment optimization.

Paper Organization. We briefly review the related work in Section 2 and the preliminaries in Section 3. We propose the pessimistic assortment optimization in Section 4. In Section 5, we present the theoretical results. In Section 6, we study pessimistic assortment optimization under the MNL model as a concrete example. In Section 7, we propose an algorithm that can solve the problem efficiently. We provide experimental results in Section 8, after which we conclude.

2. Related Work

Assortment Optimization. The assortment optimization problem under the MNL model without any constraints was first studied in (Talluri & van Ryzin, 2004). Then more complicated assortment optimization problems under various types of constraints, including space requirement (Rusmevichientong et al., 2009) and cardinality (Rusmevichientong et al., 2010), were considered. (Davis et al., 2013) proposed a linear programming (LP) formulation of the assortment optimization problem that includes several previous works as special cases corresponding to different constraints in the formulation of LP. This line of work assumes that the true parameters of the customer models are known (or at least can be accurately estimated) (Gallego & Topaloglu, 2014; Feldman & Topaloglu, 2015; Flores et al., 2019; Désir et al., 2020; Liu et al., 2020; Aouad et al., 2021). Another closely related line of work is *dynamic assortment optimization* (Caro & Gallien, 2007). In the setting of dynamic assortment optimization, the seller without any prior information about the customers, has finite selling horizons in which it observes the choices of customers and, based on the observed behaviors, optimize their assortments in

an adaptive, trial-and-error fashion (Sauré & Zeevi, 2013; Wang et al., 2018; Chen et al., 2021a; Rusmevichientong et al., 2020; Chen et al., 2020; 2021b). In comparison with the online setting used in dynamic assortment optimization, our work departs from the existing literature by focusing on the offline setting where the seller only has collected datasets but not any control on the data-collecting process.

Pessimism in Offline Learning. The principle of pessimism has been successfully used in reinforcement learning (RL) for finding an optimal policy with pre-collected datasets. On the empirical side, it has helped with improving the performance of both the model-based approach and value-based approach in offline setting (e.g., Yu et al., 2020; Kidambi et al., 2020; Kumar et al., 2020). The importance of pessimism has been analyzed and verified theoretically in the setting of RL (Jin et al., 2021; Fu et al., 2022). Our work main contribution is to take a pessimistic approach to assortment optimization problems and demonstrate its empirical and theoretical values. Moreover, our work differs from the above works by focusing on a decision-making problem with exponentially many choices.

3. Preliminary

Let $\mathcal{N} = \{1, 2, \dots, N\}$ denote the set of N distinct items. For each item i , a feature vector $x_i \in \mathbb{R}^d$ is available. Assume that x_i and μ_i are fixed vectors. Denote the collection of all possible assortments under consideration by $\mathcal{S} \subseteq 2^{\mathcal{N}}$. For the offline data, we define a random vector $(\mathbf{p}, \mathbf{A}, \mathbf{R})$ from each customer, where $\mathbf{p} \in \mathcal{N}$ denotes an assortment presented to the customer, $\mathbf{A} \in \mathcal{S}$ denotes the item purchased by the customer for $\mathbf{A} \in \mathcal{S}$ ($\mathbf{A} = \emptyset$ where no purchase is made), and $\mathbf{R}$ denotes the corresponding revenue. The ultimate goal of assortment optimization is to find an optimal set of items $\mathbf{S}^* \in \mathcal{S}$ for all customers to maximize the expected revenue. A specific goal of this work is to study how to leverage the offline data, which consists of i.i.d. samples of the random triplet $(\mathbf{p}, \mathbf{A}, \mathbf{R})$ in order to learn an optimal assortment.

For the assortment optimization with offline data, a fundamental question is to estimate the expected revenue for an unexplored assortment $\mathbf{S} \in \mathcal{S}$. This amounts to addressing the causal relationship between assortment and revenue. Under the celebrated potential outcome framework (Rubin, 1974), let the random variable $R_{\mathbf{p}|\mathbf{S}}$ be the potential revenue under an intervention that the assortment is set to be $\mathbf{S} \in \mathcal{S}$. Our goal is to find an optimal assortment

$$\mathbf{S}^* = \arg \max_{\mathbf{S} \in \mathcal{S}} \mathbb{E} R_{\mathbf{p}|\mathbf{S}}$$

Note that the expected potential revenue $\mathbb{E} R_{\mathbf{p}|\mathbf{S}}$ defined in the counterfactual world may not be identifiable from the observed data without additional assumptions. Throughout this paper, we make the following standard consistency and

un-confoundedness assumptions in causal inference.

Assumption 3.1. [CONSISTENCY] With probability one, the observed revenue coincides with the potential revenue of the observed assortment. That is, $R = R_{\mathbf{p}|\mathbf{S}}$ almost surely.

Assumption 3.2. [UN-CONFOUNDEDNESS] The potential revenues are independent variables of the observed assortment, i.e., $\mathbf{r}_{\mathbf{p}|\mathbf{S}} \perp \mathbf{S}$.

Assumption 3.1 ensures that the observed revenue is consistent with the potential revenue of purchasing item $\mathbf{A}$ or no purchase if $\mathbf{A} = 0$, under the observed assortment $\mathbf{S}$. Assumption 3.2 rules out possible unobserved factors that could confound the causal effect of assortment on revenue¹.

Denote $\pi_{\mathbf{S}|\mathbf{p}} = \Pr(\mathbf{S} = \mathbf{s} | \mathbf{p})$ as the probability of observing assortment $\mathbf{S}$ in the offline data. To non-parametrically identify $\mathbf{r}_{\mathbf{p}|\mathbf{S}}$ for every $\mathbf{S} \in \mathcal{S}$, we further require the following positivity assumption (Imbens & Rubin, 2015).

Assumption 3.3. [POSITIVITY EVERYWHERE] For all $\mathbf{S} \in \mathcal{S}$, the probability $\pi_{\mathbf{S}|\mathbf{p}}$ of observing assortment $\mathbf{S}$ is positive (i.e. $\pi_{\mathbf{S}|\mathbf{p}} > 0$).

Assumption 3.3 requires that every assortment can be observed with a positive chance in the offline data. *This is a strong assumption that will be later relaxed in Assumption 5.1 (I), i.e., requiring positivity only at optimum.* With Assumptions 3.1–3.3, we can identify the effect of an assortment set via inverse propensity score weighting (Rosenbaum & Rubin, 1983): for any $\mathbf{S} \in \mathcal{S}$,

$$\mathbf{r}_{\mathbf{p}|\mathbf{S}} = \mathbb{E} \left[\frac{\mathbf{r}_{\mathbf{p}|\mathbf{S}}}{\pi_{\mathbf{S}|\mathbf{p}}} \right], \quad (1)$$

where the expectation in the right-hand-side is taken with respect to the data distribution of $\mathbf{p}, \mathbf{A}, \mathbf{R}$. However, when the number of possible assortments $|\mathcal{S}|$ grows exponentially in N , Assumption 3.3 rarely holds for all $\mathbf{S} \in \mathcal{S}$ in practice, given potentially limited offline data particularly when N is large. Moreover, when an assortment $\mathbf{S}$ corresponds to an inferior expected revenue $\mathbf{r}_{\mathbf{p}|\mathbf{S}}$ it may not be considered by the seller at all. As a consequence, the probability of observing such an assortment $\pi_{\mathbf{S}|\mathbf{p}}$ is zero. These may prevent us from estimating (1) for every assortment $\mathbf{S} \in \mathcal{S}$.

We may tempt to use the following identification strategy:

$$\begin{aligned} \mathbf{r}_{\mathbf{p}|\mathbf{S}} &= \mathbb{E}[\mathbf{r}_{\mathbf{p}|\mathbf{S}} | \mathbf{S} = \mathbf{s}] \stackrel{\text{by Assumptions 3.1-3.2}}{=} \\ &= \mathbb{E}[\mathbf{r}_{\mathbf{p}|\mathbf{S}} | \mathbf{S} = \mathbf{s}, \mathbf{A} = \mathbf{a}] \stackrel{\text{by Assumptions 3.1-3.2}}{=} \mathbb{E}[\pi_{\mathbf{A}|\mathbf{S}}(\mathbf{s}; \mathbf{x}) \mathbf{r}_{\mathbf{S},i} | \mathbf{S} = \mathbf{s}], \end{aligned} \quad (2)$$

¹With the observed features $\mathbf{x}_j \in \mathbb{R}^N$, it can be possible to relax Assumption 3.2 to a more plausible condition: the independence holds conditional on the observed features. However, for notation simplicity, without loss of generality, we consider Assumption 3.2.

where $\mathbf{x} = (\mathbf{x}_j)_{j \in [N]}$ are the features across items, $\pi_{\mathbf{A}|\mathbf{S}}(\mathbf{s}; \mathbf{x}) = \Pr(\mathbf{A} = \mathbf{i} | \mathbf{S} = \mathbf{s}, \mathbf{x})$ is the customer's choice probability (McFadden, 1973) of purchasing the i -th item given an assortment $\mathbf{S}$, $\mathbf{r}_{\mathbf{S},i} = \mathbb{E}[\mathbf{r}_{\mathbf{p}|\mathbf{S}} | \mathbf{S} = \mathbf{s}, \mathbf{A} = \mathbf{i}]$ is the conditional expected revenue given the assortment $\mathbf{S}$ with the i -th item being purchased. For ease of notation, we omit the features $\mathbf{x}$ in $\pi_{\mathbf{A}|\mathbf{S}}$ when there is no confusion. Identifying $\mathbf{r}_{\mathbf{p}|\mathbf{S}}$ as above requires the knowledge of $\pi_{\mathbf{A}|\mathbf{S}}(\mathbf{s}; \mathbf{x})$ and $\mathbf{r}_{\mathbf{S},i}$, which can be learned from data. Although such an identification approach does not explicitly depend on $\pi_{\mathbf{S}|\mathbf{p}}$, full identification of $\pi_{\mathbf{A}|\mathbf{S}}(\mathbf{s}; \mathbf{x})$ and $\mathbf{r}_{\mathbf{S},i}$ requires the positivity of $\pi_{\mathbf{S}|\mathbf{p}}$ for every $\mathbf{S} \in \mathcal{S}$ as assumed above.

Despite the aforementioned challenge of insufficient coverage over assortments, we argue that finding an optimal assortment $\mathbf{S}^*$ may not necessarily require $\pi_{\mathbf{S}|\mathbf{p}} > 0$ everywhere but only at the optimal assortment $\mathbf{S}^*$. In particular, based on (2), when computing

$$\mathbf{S}^* = \arg \max_{\mathbf{S} \in \mathcal{S}} \sum_{i \in \mathcal{I}} \pi_{\mathbf{A}|\mathbf{S}}(\mathbf{s}; \mathbf{x}) \mathbf{r}_{\mathbf{S},i}, \quad (3)$$

we may not necessarily need to estimate $\pi_{\mathbf{A}|\mathbf{S}}(\mathbf{s}; \mathbf{x})$ and $\mathbf{r}_{\mathbf{S},i}$ well for $\mathbf{S} \neq \mathbf{S}^*$, as long as *sub-optimal assortments can be safely ruled out during the optimization*. Our insight is that the estimation of $\pi_{\mathbf{A}|\mathbf{S}}(\mathbf{s}; \mathbf{x})$ and $\mathbf{r}_{\mathbf{S},i}$ for the less seen assortment $\mathbf{S}$ in the data often incurs large errors. Deploying pessimism by taking the estimation error into consideration can rule out those assortments (Jin et al., 2021), while standard predict-then-optimize (Bertsimas & Kallus, 2020) or empirical maximization approaches (Zhao et al., 2012) may suffer from an overestimation of $\mathbf{r}_{\mathbf{p}|\mathbf{S}}$. Hence, in our proposed pessimistic assortment optimization framework, we only require the positivity at optimum $\pi_{\mathbf{S}^*|\mathbf{p}} > 0$, which is a much weaker assumption than that of Assumption 3.3.

In this paper, we focus on handling the estimation error from $\pi_{\mathbf{A}|\mathbf{S}}(\mathbf{s}; \mathbf{x})$ while assuming that $\mathbf{r}_{\mathbf{S},i}$ is known. This is a typical assumption in the literature of assortment optimization (Talluri & van Ryzin, 2004; Davis et al., 2013; Flores et al., 2019; Aouad et al., 2021). Our framework can be naturally extended to the scenario where we need to estimate $\mathbf{r}_{\mathbf{S},i}$'s. For optimization tractability, we further assume that $\mathbf{r}_{\mathbf{S},i} = \mathbf{r}_i$ that the expected revenue depends only on the purchased item but not on the underlying assortment. This assumption is reasonable in many applications where the revenue is a deterministic consequence of a purchased item. This can also be easily extended under our pessimism framework but could result in a more complicated assortment optimization problem.

Below, without loss of generality, we assume that $\mathbf{r}_i \geq 0$ for $i \in [N]$ while $\mathbf{r}_0 = 0$ (no purchase incurs zero revenue). For any vector $\mathbf{x}$, let $\mathbf{x}^T$ and $\|\mathbf{x}\|_2$ respectively denote the transpose and ℓ_2 -norm of $\mathbf{x}$. For any set $\mathcal{A}$, let $|\mathcal{A}|$ denote the cardinal number of $\mathcal{A}$. For any two sequences $\mathbf{a}, \mathbf{b} \in \mathbb{R}^N$

and $\mathbb{P}(\hat{\theta}_n \in \mathcal{B}) \geq 1 - \alpha$, we write $\hat{\theta}_n \in \mathcal{B}$ (resp. $\hat{\theta}_n \notin \mathcal{B}$) whenever there exist constants $c_1 \geq 0$ (respectively $c_2 \geq 0$) such that $\hat{\theta}_n \in \mathcal{B}$ (resp. $\hat{\theta}_n \notin \mathcal{B}$) with probability at least $1 - c_1 \exp(-c_2 n)$. Moreover, we write $\hat{\theta}_n = \theta^*$ whenever $\hat{\theta}_n \in \mathcal{B}$ and $\hat{\theta}_n \notin \mathcal{B}$.

4. Pessimistic Offline Assortment Optimization

In this section, we introduce our pessimistic offline assortment optimization framework. To this end, based on Eq. (2), we first estimate the choice probability $\pi_A(p_j | S)$ from offline data. Subsequently we calculate optimizing values in optimization problem (3) using a plug-in estimator of $\pi_A(p_j | S)$. Consider a generic form of model $\pi_A(p_j | S; \theta)$ with the unknown true parameter θ^* . Again, for ease of notation, we omit the features $\mathbf{x}$ in π_A when there is no confusion. We remark that θ^* could be either finite-dimensional or infinite-dimensional. Given an offline dataset $D = \{S_i, A_i, R_i\}_{i=1}^n$, where n is the sample size, one can estimate the model parameter θ^* via maximum likelihood estimator (MLE). Specifically, define the likelihood-based loss function $\ell_n(\theta)$ as

$$\ell_n(\theta) = -\frac{1}{n} \sum_{i=1}^n \log \pi_A(p_{A_i} | S_i; \theta)$$

Then the MLE of the unknown parameter θ^* is $\hat{\theta}_{ML,n} = \arg \min_{\theta \in \Theta} \ell_n(\theta)$ where Θ is a pre-specified parameter space. Let

$$V_{PS; \hat{\theta}_{ML,n}} = \sum_{i \in P} \pi_A(p_i | S; \hat{\theta}_{ML,n})$$

Here, we define $V_{PS; \theta}$ as the *value function* of S with the customer choice model for π_A depending on the parameter θ . The plug-in estimator of the optimal assortment based on (3) is

$$\hat{S}_{ML,n} = \arg \max_{S \in \mathcal{S}} V_{PS; \hat{\theta}_{ML,n}}(S)$$

The MLE-based approach first plugs in the MLE of θ^* , and then directly optimizes the corresponding estimated value function.

As discussed before, a disadvantage of the above estimate-then-optimize approach is that the estimation error of $\hat{\theta}_{ML,n}$ caused by insufficient data coverage may result in the overestimation of $V_{PS; \hat{\theta}_{ML,n}}$ which will propagate to downstream optimization. Alternatively, we can quantify the estimation uncertainty by considering the following likelihood-ratio-test-based confidence region (Owen, 1990):

$$\Omega_n(\alpha_n) = \{ \theta \in \Theta : \ell_n(\theta) - \ell_n(\hat{\theta}_{ML,n}) \leq \alpha_n \}$$

where $\alpha_n \geq 0$ is pre-specified. Later we analyze the MNL model as a special case (Please see Section 6). With α_n chosen as $O(\frac{1}{\sqrt{n}})$ we establish in Theorem 6.1 that $\hat{\theta}_{ML,n} \in \Omega_n(\alpha_n)$

with high probability. Such a guarantee does not require any data coverage assumption on assortments.

For now on, for simplicity, we drop α_n and write Ω_n for $\Omega_n(\alpha_n)$ when there is no ambiguity.

In order to robustify assortment optimization against plug-in estimation errors, we consider a pessimistic version of (3) by taking the estimation uncertainty from Ω_n into account. Specifically, we propose the **Pessimistic Assortment Optimization (PASTA)** by solving

$$\hat{S}_{PASTA,n} = \arg \max_{S \in \mathcal{S}} \min_{\theta \in \Omega_n} V_{PS; \theta}(S) \quad (4)$$

Here, for a fixed assortment $S \in \mathcal{S}$, the inner layer of minimization computes the worst-case value among all possible model parameters θ within the confidence set Ω_n . In particular, if the estimated value $V_{PS; \hat{\theta}_{ML,n}}(S)$ for S is highly uncertain due to insufficient data coverage, the worst-case value $\min_{\theta \in \Omega_n} V_{PS; \theta}(S)$ is likely much smaller than $V_{PS; \hat{\theta}_{ML,n}}(S)$. In that case, the outer layer of (4) may prefer another assortment with a relatively higher worst-case value. In this way, the inner layer of (4) rules out those assortments with less frequency in the offline data. Hence, one essential advantage of such a strategy is that it avoids an overestimation of the value function. In other words, by the plug-in approach, with a non-negligible chance, the estimated value $V_{PS; \hat{\theta}_{ML,n}}(S)$ can be much larger than the truth $V_{PS; \theta^*}(S)$ which further leads to a possibly sub-optimal assortment but optimized by the MLE-based approach. In contrast, PASTA is aware of insufficient data coverage, and hence more pessimistic about those highly uncertain value estimates. In the next section, we theoretically analyze the advantage of the PASTA approach.

5. Theoretical Results

In this section, we show that the PASTA method (4) enjoys a generic regret guarantee under a weak assumption of *positivity at optimum* that is $\pi_{SP^*}(p_i) > 0$. Specifically, given $\hat{S}_{PASTA,n}$ in (4), we adopt the following regret as the performance metric to evaluate the PASTA's performance

$$R_{PS; \hat{S}_{PASTA,n}} = V_{PS; \theta^*}(\hat{S}_{PASTA,n}) - V_{PS; \hat{\theta}_{ML,n}}(\hat{S}_{PASTA,n})$$

We aim to derive a regret bound for $R_{PS; \hat{S}_{PASTA,n}}$ under generic conditions. Denote $L_{PS}(S) = \mathbb{E}[\log \pi_A(p_{A_i} | S; \theta^*)]$ as the population loss function. All detailed proofs can be found in the Appendix 9.

We first show that whenever $\theta^* \in \Omega_n$ that the confidence region covers the true parameter, the regret of the PASTA method can be calibrated by the worst-case estimation error among $\theta \in \Omega_n$ of the value function at the optimal assortment S^* .

Lemma 5.1. Let $\hat{S}_{PASTA,n}$ be the solution by the PASTA

method defined in (4). If $\theta \in \mathcal{P}(\Omega_n)$, then

$$R_{\mathcal{P}(\Omega_n)}^{\text{PASTA},n}(\mathbf{q}) \leq \max_{\theta \in \mathcal{P}(\Omega_n)} \left(V_{\mathcal{P}}^{\theta}(\mathbf{q}) - V_{\mathcal{P}}^{\theta^*}(\mathbf{q}) \right).$$

Proof of Lemma 5.1.

$$\begin{aligned} & V_{\mathcal{P}}^{\theta}(\mathbf{q}) - V_{\mathcal{P}}^{\theta^*}(\mathbf{q}) \\ & \leq V_{\mathcal{P}}^{\theta}(\mathbf{q}) - \min_{\theta \in \mathcal{P}(\Omega_n)} V_{\mathcal{P}}^{\theta}(\mathbf{q}) \stackrel{\text{by } \theta \in \mathcal{P}(\Omega_n)}{\leq} V_{\mathcal{P}}^{\theta}(\mathbf{q}) - V_{\mathcal{P}}^{\theta^*}(\mathbf{q}) \\ & \leq V_{\mathcal{P}}^{\theta}(\mathbf{q}) - \min_{\theta \in \mathcal{P}(\Omega_n)} V_{\mathcal{P}}^{\theta}(\mathbf{q}) \stackrel{\text{by } \theta \in \mathcal{P}(\Omega_n)}{\leq} V_{\mathcal{P}}^{\theta}(\mathbf{q}) - V_{\mathcal{P}}^{\theta^*}(\mathbf{q}) \\ & \leq \max_{\theta \in \mathcal{P}(\Omega_n)} \left(V_{\mathcal{P}}^{\theta}(\mathbf{q}) - V_{\mathcal{P}}^{\theta^*}(\mathbf{q}) \right). \quad \square \end{aligned}$$

Lemma 5.1 highlights the benefit of pessimistic optimization. In particular, guarantees for the regret of estimate-then-optimize approaches require the control of estimation error uniformly over all $\mathbf{s} \in \mathcal{S}$, that is, $\sup_{\mathbf{s} \in \mathcal{S}} V_{\mathcal{P}}^{\theta}(\mathbf{s}) - V_{\mathcal{P}}^{\theta^*}(\mathbf{s})$, in order to control error caused by the greedy policy, i.e., $V_{\mathcal{P}}^{\theta}(\mathbf{s}) - V_{\mathcal{P}}^{\theta^*}(\mathbf{s})$. This may typically entail that the value function is estimated uniformly well across $\mathbf{s} \in \mathcal{S}$. It further explains the reason why an estimate-then-optimize approach would require the positivity Assumption 3.3 for all $\mathbf{s} \in \mathcal{S}$. In contrast, our Lemma 5.1 suggests that controlling the estimation error at $\mathbf{s}^*$ can be enough. Therefore, it is sufficient for our PASTA method to require the positivity only at optimum.

Next, we impose the following assumptions to obtain the regret guarantee of our algorithm.

Assumption 5.1.

- (I) [POSITIVITY AT OPTIMUM] The probability of observing the optimal assortment is positive, that is, $\pi_{\mathcal{P}}^{\theta^*}(\mathbf{q}) > 0$.
 (II) [LIKELIHOOD-BASED CONCENTRATION] For any $0 \leq \delta \leq 1$, with probability at least $1 - \delta$, we have: (1) $\theta \in \mathcal{P}(\Omega_n)$, and (2)

$$\sup_{\theta \in \mathcal{P}(\Omega_n)} \left(V_{\mathcal{P}}^{\theta}(\mathbf{q}) - V_{\mathcal{P}}^{\theta^*}(\mathbf{q}) \right) \leq \alpha_n.$$

We emphasize that *PASTA only requires the positivity at optimum*. Compared to the positivity at all assortments in Assumption 3.3, our Assumption 5.1 (I) is much weaker and hence more plausible to be satisfied. Assumption 5.1 (II) is a generic condition for likelihood-based concentration. We later justify that (II) above indeed holds under the general MNL model in Theorem 6.2. In particular, Statement (1) of Part (II) requires the validity of the likelihood-ratio-test-based confidence region Ω_n while Statement (2) of Part (II) requires the concentration of the likelihood-based localized empirical process (van der Vaart & Wellner, 1996).

The positivity at optimum is associated with a finite constant $C_{\mathbf{s}^*} = 1/\pi_{\mathcal{P}}^{\theta^*}(\mathbf{q})$ related to the learning performance. We also denote $r_{\mathbf{s}^*} = \max_{j \in \mathcal{P}} r_j$ as the largest possible revenue among all items in $\mathcal{S}$. Notice that both constants $C_{\mathbf{s}^*}$

and $r_{\mathbf{s}^*}$ depend on the optimal assortment $\mathbf{s}^*$ only. In the following lemma, we establish the estimation error bound at the optimal assortment $\mathbf{s}^*$.

Lemma 5.2. *Under Assumption 5.1, for any $0 \leq \delta \leq 1$, with probability at least $1 - \delta$, we have for any $\theta \in \mathcal{P}(\Omega_n)$,*

$$V_{\mathcal{P}}^{\theta}(\mathbf{q}) - V_{\mathcal{P}}^{\theta^*}(\mathbf{q}) \leq r_{\mathbf{s}^*} C_{\mathbf{s}^*} \alpha_n.$$

Combining Lemmas 5.1 and 5.2, we summarize the regret bound for PASTA in the following theorem.

Theorem 5.3. *Under Assumption 5.1, for any $0 \leq \delta \leq 1$, with probability $1 - \delta$, we have*

$$R_{\mathcal{P}(\Omega_n)}^{\text{PASTA},n}(\mathbf{q}) \leq r_{\mathbf{s}^*} C_{\mathbf{s}^*} \alpha_n.$$

6. Application: Multinomial Logit Model

In this section, we consider the Multinomial Logit Model (MNL) for customer choices $\pi_{\mathcal{A}}(\mathbf{s}|\mathbf{q})$. This is one of the most widely used models in assortment optimization literature (Feng et al., 2022). Under the MNL model, we will verify Assumption 5.1 (II) and establish the regret bound for PASTA in this case.

Given the item-specific features $\mathbf{x}_i \in \mathbb{R}^N$, MNL assumes that customer's preference for the i -th item is proportional to $\exp(\mathbf{x}_i^T \boldsymbol{\theta})$ where $\boldsymbol{\theta} \in \Theta$ is the underlying unknown parameter. Here, we assume that the parameter space $\Theta \subseteq \mathbb{R}^d$ is compact with $\theta_{\max} = \sup_{\boldsymbol{\theta} \in \Theta} \|\boldsymbol{\theta}\|_2 \leq 8$. Given an assortment $\mathcal{S}$, the customer choice probability under MNL is given by

$$\pi_{\mathcal{A}}(\mathbf{j}|\mathbf{s}; \boldsymbol{\theta}) = \frac{\exp(\mathbf{x}_j^T \boldsymbol{\theta})}{1 + \sum_{j \in \mathcal{P}} \exp(\mathbf{x}_j^T \boldsymbol{\theta})}, \quad \forall \mathbf{s} \in \mathcal{S}. \quad (5)$$

Moreover, the probability of no-purchase is normalized to $\pi_{\mathcal{A}}(\mathbf{0}|\mathbf{s}; \boldsymbol{\theta}) = 1/(1 + \sum_{j \in \mathcal{P}} \exp(\mathbf{x}_j^T \boldsymbol{\theta}))$. Based on (3) and the MNL model (5), the objective function for assortment optimization can be written as

$$V_{\mathcal{P}}(\mathbf{s}; \boldsymbol{\theta}) = \frac{\sum_{i \in \mathcal{P}} r_i \exp(\mathbf{x}_i^T \boldsymbol{\theta})}{1 + \sum_{i \in \mathcal{P}} \exp(\mathbf{x}_i^T \boldsymbol{\theta})}.$$

We first justify Statement (1) of Assumption 5.1 under the MNL model. To this end, given the compactness of Θ , there exists a finite constant $C_{\mathbf{A}} > 0$ such that for all $\boldsymbol{\theta} \in \Theta$, $\mathbf{s} \in \mathcal{S}$ and $i \in \mathcal{P}$, we have $1/\pi_{\mathcal{A}}(\mathbf{j}|\mathbf{s}; \boldsymbol{\theta}) \leq C_{\mathbf{A}}$.

Lemma 6.1. *Consider the MNL model (5) with a compact set Θ . Assume that $\boldsymbol{\theta} \in \Theta$. For any $0 \leq \delta \leq 1$, with probability at least $1 - \delta$, we have*

$$\pi_{\mathcal{A}}(\mathbf{0}|\mathbf{s}; \boldsymbol{\theta}) \leq \pi_{\mathcal{A}}(\mathbf{0}|\mathbf{s}; \boldsymbol{\theta}_{\max}) \leq \frac{C_{\mathbf{A}}}{n} \log \frac{\theta_{\max}}{\delta}.$$

Lemma 6.1 suggests that, with α_n chosen as $\frac{C_A d}{n} \log \frac{\theta_{\max}}{\delta}$, we can guarantee that $\theta^* \in \Omega_n$ with high probability, which justifies Statement (1) of Assumption 5.1 (II). In particular, the order of α_n is $O(d \log n)$. Notice that Lemma 6.1 does not depend on the distribution of S , which implies that no data coverage assumption on the observed assortments is required. The assumption that $\theta^* \in \Theta$ for Θ a compact set requires that given any assortment, every product has a chance of being selected by the customer in the data. This is a mild requirement as θ^* is always finite.

Next, we justify Statement (2) of Assumption 5.1 (II) in the following theorem.

Lemma 6.2. *Consider the MNL model (5). Suppose conditions in Lemma 6.1 hold, and L_{θ^*} and L_n are uniformly and strongly convex. Let $\alpha_n \approx \frac{C_A d}{n} \log \frac{\theta_{\max}}{\delta}$. For $0 < \delta < 1$, with probability $1 - \delta$ we have*

$$\sup_{\theta \in \Omega_n} \tilde{L}_{\theta^*} - L_{\theta^*} - \tilde{L}_n - L_n \leq \alpha_n.$$

Finally, with the Assumption of positivity only at optimum (Assumption 5.1 (I)), we can apply Theorem 5.3 to establish the regret bound for PASTA in MNL.

Theorem 6.3. *Consider the MNL model (given in Equation (5)). Assume that the conditions in Lemma 6.2 hold and that $\pi_{S^*} > 0$. Fix a $\delta \in (0, 1]$. Suppose $\hat{S}_{\text{PASTA}, n}$ is output of PASTA with $\alpha_n \approx \frac{C_A d}{n} \log \frac{\theta_{\max}}{\delta}$. Then with probability at least $1 - \delta$, we have*

$$R_{\hat{S}_{\text{PASTA}, n}} \leq C_S \frac{C_A d}{n} \log \frac{\theta_{\max}}{\delta}.$$

We remark that under the MNL model (given in Equation (5)), the order of regret is $O(d \log n)$. This is due to the concentration rate of MNL's empirical likelihood ratio in Lemma 6.1. Such a rate of regret bound matches those in the literature under parametric model assumptions (Qian & Murphy, 2011; Mo & Liu, 2022). However, existing literature requires the positivity $\pi_{S^*} > 0$ at every $S \in \mathcal{S}$. In contrast, Theorem 6.3 only requires positivity $\pi_{S^*} > 0$ at the optimal assortment S^* . Furthermore, we can show that $\min_{i \in P_N} \pi_{A_i} | S^* = r_N \alpha_n d^{-1} N$ for any $\theta \in \Theta$, which implies that $C_A \leq N$. Therefore, our regret is of order at least $\sqrt{N}$, where N is the total number of available items. It is an interesting problem to establish the minimax lower bound of offline assortment optimization in terms of N , n , d and the cardinal number of $\mathcal{S}$. This will be investigated in a subsequent work.

7. PASTA Algorithm

In this section, we propose an efficient algorithm for solving the max-min problem given in Optimization Problem (4) for

the MNL model. Specifically, let

$$V_{\mathbf{p}; \theta} = \sum_{i \in \mathbf{p}} \frac{r_i \exp x_i^T \theta}{\sum_{j \in \mathbf{p}} \exp x_j^T \theta}$$

and given the confidence set Ω_n , we wish to solve

$$\max_{S \in \mathcal{S}} \min_{\theta \in \Omega_n} V_{\mathbf{p}; \theta}$$

The proposed iterative algorithm is executed for a maximum of T iterations. At the t -th iteration, given S_t and θ_t from the previous iteration, we consecutively execute the following two steps:

- Step 1: Compute the optimal assortment S_{t+1} given θ_t (see Section 7.1).
- Step 2: Compute the optimal θ_{t+1} using S_{t+1} (see Section 7.2).

The corresponding pseudo-code is presented in Algorithm 1 below.

Algorithm 1 PASTA

Input: offline dataset $\{p_i, A_i, R_i\}_{i=1}^n$; α_n ; $\{r_i\}_{i=1}^N$; $\{x_i\}_{i=1}^N$; maximum number of iterations T

Output: the solution to pessimistic assortment optimization $\hat{S}$

$\hat{\mathbf{p}}_n \leftarrow \frac{1}{n} \sum_{i=1}^n \log \pi_{A_i} | S_i; \theta_i$

$\theta_{\text{ML}, n} \leftarrow \arg \min_{\theta \in \Theta} \hat{\mathbf{p}}_n$

$\Omega_n \leftarrow \{\theta \in \Theta : \hat{\mathbf{p}}_n - \hat{\mathbf{p}}_{\text{ML}, n} \leq d \alpha_n\}$

$t \leftarrow 0; \theta \leftarrow \theta_{\text{ML}, n}$ /* Initialize θ_0 as $\theta_{\text{ML}, n}$ */

for $t \leftarrow 1$ **to** T **do**

$S_t \leftarrow \text{SolveLP}(\theta_{t-1}, \{r_i\}_{i=1}^N, \{x_i\}_{i=1}^N)$

/* Section 7.1 */

$\theta_t \leftarrow \text{SolveGD}(\mathbf{p}_{S_t}, \Omega_n, \{r_i\}_{i=1}^N, \{x_i\}_{i=1}^N)$

/* Section 7.2 */

end for

$\hat{S} \leftarrow S_T$

7.1. Optimal Assortment Computation

Given the MNL model parameter θ_t , computing the assortment S_{t+1} that maximizes the expected revenue can be formulated as a linear programming (LP) problem.

Suppose that an assortment S can be represented by an N -dimensional binary vector $\mathbf{y} \in \{0, 1\}^N$ where $y_j = 1$ if and only if $j \in S$. Suppose that $S \in \mathcal{S}$ corresponds to the following feasible set for $\mathbf{y}$ with M linear inequality constraints:

$$\mathbf{y} \in \{ \mathbf{y} \in \{0, 1\}^N : \sum_{j \in P_N} a_{ij} y_j \leq b_i \text{ for } i \in M_S \},$$

where the matrix of constraint coefficients $\mathbf{a}_{ij} \in \mathbb{R}^{M_S \times P_N}$ is a totally unimodular matrix (Pang, 2017). In other words,

based on the one-to-one correspondence between S and Y , we have $S \in \mathcal{P}S$ if and only if $Y \in \mathcal{P}Y$.

Next, we denote $v_i = \exp x_i^j \theta_i q$ as the preference score for the i -th item. The customer choice probability under the MNL model (5) becomes $\pi_{A|j|S} q = \frac{v_i}{1 + \sum_{i \in S} v_i}$. The optimization for S_{t-1} can be formulated as

$$\max_{Y \in \mathcal{P}Y} \frac{\sum_{i \in S} r_i v_i Y_i}{1 + \sum_{i \in S} v_i Y_i}, \quad (6)$$

which is equivalent to the following linear programming problem (Davis et al., 2013):

$$\begin{aligned} \max_{w_j : j \in \mathcal{P}N_S} \quad & \sum_{j \in \mathcal{P}N_S} r_j w_j \\ \text{subject to} \quad & \sum_{j \in \mathcal{P}N_S} w_j = w_0 = 1 \\ & \sum_{j \in \mathcal{P}N_S} a_{ij} \frac{w_j}{v_j} \leq d_i b_i w_0 \quad \forall i \in \mathcal{P}M_S \\ & 0 \leq \frac{w_j}{v_j} \leq d_i w_0 \quad \forall i \in \mathcal{P}N_S \end{aligned} \quad (7)$$

In particular, we can recover the optimal solution to Problem (6), denoted as Y^* , using the optimal solution to Problem (7), denoted by W^* , via the following formula:

$$Y_j^* = \frac{w_j^*}{v_j w_0^*} \quad \forall j \in \mathcal{P}N_S \quad (8)$$

To conclude, at the t -th iteration, in order to compute an optimal assortment S_{t-1} for a given θ_t , we first solve an LP problem in (7) for W^* . Then we recover Y^* via (8). Finally, the updated assortment S_{t-1} is obtained by the correspondence $i \in \mathcal{P}S_{t-1}$ if and only if $Y_j^* = 1$.

7.2. Model Parameter Computation

For a given optimized assortment S_{t-1} from Section 7.1, we aim to search for the worst-case MNL parameter θ_{t-1} from the confidence set Ω_n that minimizes the expected revenue. In particular, we employ a gradient descent with line search (GDLS) method to compute θ_{t-1} by solving the following problem

$$\min_{\theta \in \Omega_n} V_{\mathcal{P}S_{t-1}}; \theta q \quad (9)$$

Here, we remark that $V_{\mathcal{P}S_{t-1}}; \theta q = \frac{\sum_{i \in S_{t-1}} r_i \exp x_i^j \theta q}{1 + \sum_{i \in S_{t-1}} \exp x_i^j \theta q}$ is a locally Lipschitz function in θ . Given a feasible initial parameter $\theta^{0q} \in \Omega_n$, we run at most L gradient descent steps. Suppose β_ℓ is the step size for gradient descent in the ℓ -th step. At each step $\ell = 1, 2, \dots, L$, we do a line search to maintain the feasibility. In particular, given $\theta^{\ell-1q} \in \Omega_n$, we first evaluate the gradient as $\xi_\ell = \nabla_{\theta} V_{\mathcal{P}S_{t-1}}; \theta^{\ell-1q} q$. Then we initiate β_ℓ with a pre-specified step size $\beta_\ell = \beta$, and check whether $\theta^{\ell q} = \theta^{\ell-1q} + \beta_\ell \xi_\ell$ is feasible, i.e.

$\theta^{\ell q} \in \Omega_n$. If not, we set $\beta_\ell = c\beta_\ell$ for some pre-specified $c \in (0, 1)$ and recompute $\theta^{\ell q} = \theta^{\ell-1q} + \beta_\ell \xi_\ell$. Such a search is repeated until $\theta^{\ell q}$ is feasible. We provide the pseudocode in Algorithm 2 for the overall process. Note that L, β, c are all hyper-parameters. In all of our numerical studies, we set $L = 2, \beta = 0.01$ and $c = \frac{1}{2}$, which performs well empirically.

Algorithm 2 Gradient Descent with Line Search (GDLS)

Input: assortment S_{t-1} ; feasible set Ω_n ; initial parameter θ^{0q} ; initial step size β ; step shrinkage constant c ; number of descent steps L

Output: the updated parameter θ_{t-1}

```

 $\ell \leftarrow 0$ 
for  $\ell = 1$  to  $L$  do
     $\xi_\ell \leftarrow \nabla_{\theta} V_{\mathcal{P}S_{t-1}}; \theta^{\ell-1q} q$  /* compute the gradient */
     $\beta_\ell \leftarrow \beta$ 
     $\theta^{\ell q} \leftarrow \theta^{\ell-1q} + \beta_\ell \xi_\ell$ 
    while  $\theta^{\ell q} \notin \Omega_n$  do
         $\beta_\ell \leftarrow c\beta_\ell$  /* decrease the step size */
         $\theta^{\ell q} \leftarrow \theta^{\ell-1q} + \beta_\ell \xi_\ell$ 
    end while
end for
 $\theta_{t-1} \leftarrow \theta^{\ell q}$ 
    
```

8. Experiments

We compare the PASTA method with assortment optimization without pessimism (referred to as the *baseline* method in the sequel). Our method and the baseline method are evaluated on synthetic data for which the optimal assortment S^* and true parameter θ^* are known so that the true regrets can be computed. We describe the data generation process and the baseline method in details below.

8.1. Data Generation

We consider the assortment optimization scenarios described by N, K, d, n and p , where N is the total number of available products; K is the cardinality constraint of the assortments, i.e., $S = \{i : |S| \leq K\}$; d is the dimension of θ^* and $\mathbf{x}_i \in \mathbb{R}^d$; n is the sample size of the offline dataset; p is the probability for sampling the optimal assortment S^* . Similar to (Chen et al., 2020), we first generate the true preference vector θ^* as a uniformly random unit d -dim vector. For $i \in \{1, \dots, N\}$, we generate r_i (the reward of product i) uniformly from the range $[0.5, 0.8]$ and generate $\mathbf{x}_i$ (the feature of product i) as uniformly random unit d -dim vector such that $\exp x_i^j \theta^* q \leq \exp 0.6q$ to avoid degenerate cases, where the optimal assortments include too few items. Given such information, the true optimal assortment S^* can be computed. Then, we generate an offline dataset $D = \{S_i, A_i, R_i\}_{i=1}^n$ with n samples. For $i \in \{1, \dots, n\}$, we generate S_i following the distribution π_{S^*}

such that $\pi_{\mathcal{S}}^{\mathcal{S}^*} \leq \frac{1}{|\mathcal{S}|} \frac{p}{1-p}$ and $\pi_{\mathcal{S}}^{\mathcal{S}^*} \leq \frac{1}{|\mathcal{S}|} \frac{p}{1-p}$, where $0 \leq p \leq 1$ is the probability of observing the optimal assortment $\mathcal{S}^*$. After the assortment $\mathcal{S}_i$ is sampled, the customer choice (action) A_i is sampled according to the probability computed by MNL as in Eq. (5) with the true parameter θ^* .

8.2. Baseline

In our experiments, we use the gradient descent method to find $\theta_{ML,n}$ that minimizes the empirical negative log-likelihood function. Then given $\theta_{ML,n}$, the baseline method solves the assortment optimization problem by solving the linear programming problem in (7).

8.3. Performance Comparison

For a given N, K, d, n, p we repeat the data generation process in Section 8.1 to randomly generate 50 offline datasets. The solutions of PASTA and the baseline method are recorded in these experiments. For hyper-parameters, we set $\alpha_n = 2\beta_{ML}$ where $\beta_{ML} = \beta_n \theta_{ML,n}$ and the maximum of iteration $T = 30$. We measure the performance with two metrics: (1) the *average regret* of the solutions which indicates how far the performance of the solutions is to that of the *optimal* performance (i.e., revenue of $\mathcal{S}^*$); (2) the *assortment accuracy* of the solutions (with respect to the optimal assortment $\mathcal{S}^*$). The assortment accuracy of an assortment $\mathcal{S}$ is defined as the ratio of the number of correctly chosen products to the number of products in $\mathcal{S}^*$. The key results are summarized below.

Effect of Sample Size. We set $N = 40, K = 8, d = 16$ and $p = 0.9$. We then gradually increase the number of samples n . The result is presented in Figure 1 indicating that PASTA significantly outperforms the baseline method. While the performance of the baseline method improves with increasing number of samples, the PASTA method maintains a regret that is less than 25% of that of the baseline method. The same experiment repeated with an increased number of products ($N = 60, K = 15$) demonstrates that the gain of the PASTA method is stable, as presented in Figure 1.

Effect of Probability of Sampling Optimal Assortment in Offline Data. We set $N = 40, K = 8, d = 16, n = 150$ and let $p \in \{0.1, 0.3, 0.5, 0.7, 0.9\}$. We also study the effect of p in scenarios with an increased total number of products ($N = 60, K = 15$). As can be seen in Figure 2, the gain of pessimistic assortment optimization is consistent and robust for varying values of p .

Effect of Dimension of Features. We set $N = 20, K = 5, p = 0.9, n = 150$, and let $d \in \{8, 20, 32, 64, 128\}$. In order to characterize the effect of dimension d , we generate d elements of θ^* independently from Uniform $[-1, 1]$. The results are presented in Figure 3. We observe that while both the regret of the baseline method and that of the pes-

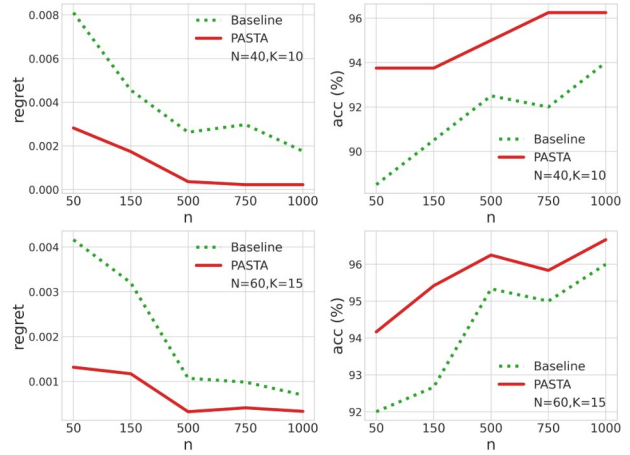


Figure 1. Performance comparison between PASTA and the baseline method with varying number of samples (n). On the left is the average regret (the lower the better) while the assortment accuracy (the higher the better) is on the right.

simistic assortment optimization increase with increasing dimensions of features, the PASTA method maintains its performance gain as the dimension d varies.

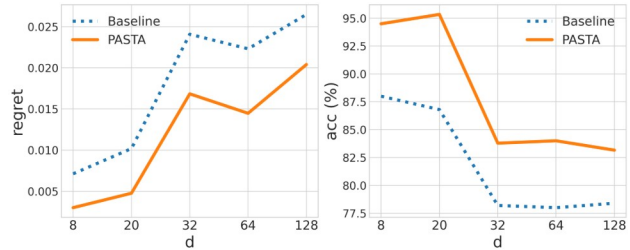


Figure 3. Comparison between PASTA and the baseline method with increasing dimensions of product features (d).

9. Conclusion

This work addresses the issue of insufficient data coverage in offline assortment optimization problems. This becomes more challenging as the number of choices grows quickly as a function of the number of items N . We presented a framework of pessimistic assortment optimization and provided theoretical justifications for our approach. We then performed an in-depth study of the Multinomial Logit Model (MNL), and derived a finite-sample regret bound of pessimistic assortment optimization for this popular model. We presented an efficient algorithm to solve the pessimistic assortment optimization problem for MNL, and demonstrated significant improvements of our approach over the baseline method by extensive numerical studies.

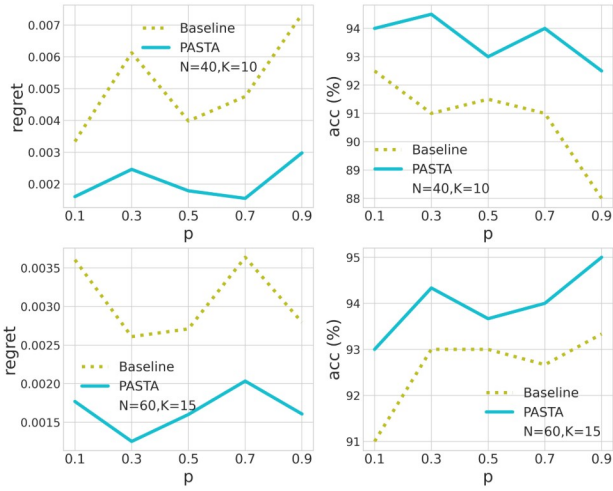


Figure 2. Comparison between PASTA and the baseline method with varying probability of the optimal assortment (p). Top row: $N = 40$; bottom row: $N = 60$.

Acknowledgements

Ethan X. Fang is partially supported by NSF DMS-2230795, NSF DMS-2230797. Cong Shi is partially supported by an Amazon Research Award.

References

- Aouad, A., Farias, V., and Levi, R. Assortment optimization under consider-then-choose choice models. *Management Science*, 67(6):3368–3386, 2021.
- Aouad, A., Feldman, J., and Segev, D. The exponential choice model for assortment optimization: an alternative to the MNL model? *Management Science*, 2022.
- Bertsimas, D. and Kallus, N. From predictive to prescriptive analytics. *Management Science*, 66(3):1025–1044, 2020.
- Caro, F. and Gallien, J. Dynamic assortment with demand learning for seasonal consumer goods. *Management Science*, 53(2):276–292, 2007.
- Chen, X., Wang, Y., and Zhou, Y. Dynamic assortment optimization with changing contextual information. *Journal of Machine Learning Research*, 2020.
- Chen, X., Shi, C., Wang, Y., and Zhou, Y. Dynamic assortment planning under nested logit models. *Production and Operations Management*, 30(1):85–102, 2021a.
- Chen, X., Wang, Y., and Zhou, Y. Optimal policy for dynamic assortment planning under multinomial logit models. *Mathematics of Operations Research*, 46(4):1639–1657, 2021b.

- Davis, J. M., Gallego, G., and Topaloglu, H. Assortment planning under the multinomial logit model with totally unimodular constraint structures, 2013. URL <https://people.orie.cornell.edu/jmd388/publications/MNLConstr.pdf>. Working Paper, Cornell University, Ithaca, NY.
- Désir, A., Goyal, V., Segev, D., and Ye, C. Constrained assortment optimization under the Markov chain-based choice model. *Management Science*, 66(2):698–721, 2020.
- Diaconis, P. and Saloff-Coste, L. Logarithmic sobolev inequalities for finite markov chains. *The Annals of Applied Probability*, 6(3):695–750, 1996.
- Feldman, J. and Topaloglu, H. Bounding optimal expected revenues for assortment optimization under mixtures of multinomial logits. *Production and Operations Management*, 24(10):1598–1620, 2015.
- Feng, Q., Shanthikumar, J. G., and Xue, M. Consumer choice models and estimation: A review and extension. *Production and Operations Management*, 31(2):847–867, 2022.
- Flores, A., Berbeglia, G., and Van Hentenryck, P. Assortment optimization under the sequential multinomial logit model. *European Journal of Operational Research*, 273(3):1052–1064, 2019.
- Fu, Z., Qi, Z., Wang, Z., Yang, Z., Xu, Y., and Kosorok, M. R. Offline reinforcement learning with instrumental variables in confounded markov decision processes, 2022. URL <https://arxiv.org/abs/2209.08666>. Working Paper, Northwestern University, Evanston, IL.
- Gallego, G. and Topaloglu, H. Constrained assortment optimization for the nested logit model. *Management Science*, 60(10):2583–2601, 2014.
- Harsha, P. Communication complexity, 2011. URL <https://www.tifr.res.in/~prahladh/teaching/2011-12/comm/lectures/l12.pdf>. Lecture Notes, Tata Institute of Fundamental Research (TIFR), Mumbai, India.
- Imbens, G. W. and Rubin, D. B. *Causal inference in statistics, social, and biomedical sciences*. Cambridge University Press, Cambridge, UK, 2015.
- Jin, Y., Yang, Z., and Wang, Z. Is pessimism provably efficient for offline RL? In Meila, M. and Zhang, T. (eds.), *Proceedings of the 38th International Conference on Machine Learning*, volume 139 of *Proceedings of Machine Learning Research*, pp. 5084–5096. PMLR, 18–24 Jul 2021.

- Kidambi, R., Rajeswaran, A., Netrapalli, P., and Joachims, T. Morel: Model-based offline reinforcement learning. In Larochelle, H., Ranzato, M., Hadsell, R., Balcan, M., and Lin, H. (eds.), *Advances in Neural Information Processing Systems*, volume 33, pp. 21810–21823. Curran Associates, Inc., 2020.
- Kumar, A., Zhou, A., Tucker, G., and Levine, S. Conservative q-learning for offline reinforcement learning. In *Proceedings of the 34th International Conference on Neural Information Processing Systems*, NIPS’20, Red Hook, NY, USA, 2020. Curran Associates Inc.
- Li, S., Luo, Q., Huang, Z., and Shi, C. Online learning for constrained assortment optimization under Markov chain choice model, 2022. URL https://papers.ssrn.com/sol3/papers.cfm?abstract_id=4079753. Working Paper, University of Michigan, Ann Arbor, MI.
- Liu, N., Ma, Y., and Topaloglu, H. Assortment optimization under the multinomial logit model with sequential offerings. *INFORMS Journal on Computing*, 32(3):835–853, 2020.
- McFadden, D. *Conditional logit analysis of qualitative choice behavior*. Institute of Urban and Regional Development, Berkeley, CA, 1973.
- McFadden, D. Econometric models of probabilistic choice. *Structural analysis of discrete data with econometric applications*, 198272, 1981.
- Mo, W. and Liu, Y. Efficient learning of optimal individualized treatment rules for heteroscedastic or misspecified treatment-free effect models. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 84(2): 440–472, 2022.
- Owen, A. Empirical likelihood ratio confidence regions. *The Annals of Statistics*, 18(1):90–120, 1990.
- Pang, R. 1 totally unimodular matrices - stanford university, 2017. URL <https://theory.stanford.edu/~jvondrak/MATH233B-2017/lec3.pdf>.
- Qian, M. and Murphy, S. A. Performance guarantees for individualized treatment rules. *The Annals of Statistics*, 39(2):1180–1210, 2011.
- Rosenbaum, P. R. and Rubin, D. B. The central role of the propensity score in observational studies for causal effects. *Biometrika*, 70(1):41–55, 1983.
- Rubin, D. B. Estimating causal effects of treatments in randomized and nonrandomized studies. *Journal of educational Psychology*, 66(5):688, 1974.
- Rusmevichientong, P., Shen, M., and Shmoys, D. A PTAS for capacitated sum-of-ratios optimization. *Operations Research Letters*, 37:230–238, 07 2009.
- Rusmevichientong, P., Shen, Z.-J. M., and Shmoys, D. B. Dynamic assortment optimization with a multinomial logit choice model and capacity constraint. *Operations Research*, 58(6):1666–1680, 2010.
- Rusmevichientong, P., Sumida, M., and Topaloglu, H. Dynamic assortment optimization for reusable products with random usage durations. *Management Science*, 66(7): 2820–2844, 2020.
- Sauré, D. and Zeevi, A. Optimal dynamic assortment planning with demand learning. *Manufacturing & Service Operations Management*, 15(3):387–404, 2013.
- Sen, B. A gentle introduction to empirical process theory and applications, 2018. URL <http://www.stat.columbia.edu/~bodhi/Talks/Emp-Proc-Lecture-Notes.pdf>. Lecture Notes, Columbia University, New York, NY.
- Talluri, K. and van Ryzin, G. Revenue management under a general discrete choice model of consumer behavior. *Management Science*, 50(1):15–33, 2004.
- van der Vaart, A. W. and Wellner, J. A. *Weak Convergence and Empirical Processes: with Applications to Statistics*. Springer, New York, NY, 1996.
- Wang, Y., Chen, X., and Zhou, Y. Near-optimal policies for dynamic multinomial logit assortment selection models. In Bengio, S., Wallach, H., Larochelle, H., Grauman, K., Cesa-Bianchi, N., and Garnett, R. (eds.), *Advances in Neural Information Processing Systems*, volume 31. Curran Associates, Inc., 2018.
- Yu, T., Thomas, G., Yu, L., Ermon, S., Zou, J. Y., Levine, S., Finn, C., and Ma, T. Mopo: Model-based offline policy optimization. In Larochelle, H., Ranzato, M., Hadsell, R., Balcan, M., and Lin, H. (eds.), *Advances in Neural Information Processing Systems*, volume 33, pp. 14129–14142. Curran Associates, Inc., 2020.
- Zhao, Y., Zeng, D., Rush, A. J., and Kosorok, M. R. Estimating individualized treatment rules using outcome weighted learning. *Journal of the American Statistical Association*, 107(499):1106–1118, 2012.

A. Proof of Theoretical Results

Throughout the proofs, we use θ^* to denote the true parameter and n to denote the number of samples. We use $\hat{p}_n$ to denote the empirical negative log-likelihood function, i.e., $\hat{p}_n(\theta) = \frac{1}{n} \sum_{i=1}^n \log \pi_A(p_i | S_i; \theta)$ where $\{(S_i, A_i, R_i)\}_{i=1}^n$ is the offline dataset. We use L and $\theta_{ML,n}$ to respectively denote the negative log-likelihood function, i.e., $L(\theta) = \mathbb{E} \log \pi_A(p | S; \theta)$ and the MLE of θ^* , i.e., $\theta_{ML,n} = \arg \min_{\theta \in \Theta} \hat{p}_n(\theta)$. The confidence region Ω_n is defined as $\Omega_n = \{\theta \in \Theta : \hat{p}_n(\theta) - \hat{p}_n(\theta_{ML,n}) \leq d' \alpha_n\}$.

For a general pair of random variables X, Y , assume that the conditional probability density function of Y given X is parametrically modeled by $p(y|x; \theta)$ for parameter θ . For technical reasons, we will consider the following distances.

Definition A.1 (Squared Hellinger Distance).

$$h^2(p_1; \theta_1, p_2; \theta_2) = \frac{1}{2} \int \left(\sqrt{p_1(y|x; \theta_1)} - \sqrt{p_2(y|x; \theta_2)} \right)^2 dy. \quad (10)$$

Definition A.2 (Hellinger Distance).

$$h(p_1; \theta_1, p_2; \theta_2) = \sqrt{h^2(p_1; \theta_1, p_2; \theta_2)} \quad (11)$$

Definition A.3 (Generalized Squared Hellinger Distance).

$$H^2(p_1, p_2) = \mathbb{E}_X [h^2(p_1; \theta_1, p_2; \theta_2)]. \quad (12)$$

Definition A.4 (Generalized Hellinger Distance).

$$H(p_1, p_2) = \sqrt{\mathbb{E}_X [h^2(p_1; \theta_1, p_2; \theta_2)]}. \quad (13)$$

In our theoretical results, we particularly consider $p(y|x; \theta) = \pi_A(p | S; \theta)$ as the conditional density of A given S (hereafter denoted as $A|S$).

A.1. Proof of Lemma 5.2

Under Assumption 5.1, for any $0 < \delta \leq 1$, with probability at least $1 - \delta$, we have for any $\theta \in \Omega_n$,

$$V(p^*; \theta) \geq V(p^*; \theta_{ML,n}) - \frac{r_{S^*} C_{S^*}}{\alpha_n}, \quad (14)$$

where $r_{S^*} = \max_{j \in p^*} r_j$ is the largest possible revenue among all items in the optimal assortment.

Proof of Lemma 5.2. For any θ such that $\hat{p}_n(\theta) - \hat{p}_n(\theta_{ML,n}) \leq d' \alpha_n$, i.e., $\theta \in \Omega_n$, we have

$$\begin{aligned} & V(p^*; \theta) - V(p^*; \theta_{ML,n}) \\ & \leq \mathbb{E} [V(p^*; \theta) - V(p^*; \theta_{ML,n})] \\ & \leq \mathbb{E} [V(p^*; \theta) - V(p^*; \theta_{ML,n})] \\ & \leq \mathbb{E} [V(p^*; \theta) - V(p^*; \theta_{ML,n})] \\ & \leq \mathbb{E} [V(p^*; \theta) - V(p^*; \theta_{ML,n})] \\ & \leq 2 \max_{\theta \in \Omega_n} V(p^*; \theta) - V(p^*; \theta_{ML,n}) \quad (\text{By Assumption 5.1 (II)} \theta \in \Omega_n). \end{aligned}$$

With Lemma B.2, we have that for any $\theta \in \Theta$,

$$V(p^*; \theta) - V(p^*; \theta_{ML,n}) \leq \frac{r_{S^*} C_{S^*}}{\alpha_n} \mathbb{E} [|\pi_A(p^* | S; \theta) - \pi_A(p^* | S; \theta_{ML,n})|]$$

where $\|\cdot\|_1$ is the ℓ_1 -norm, $r_S = \max_{j \in S} r_j$ is the largest possible revenue among all items and $C_S = 1/\pi_S p^\circ q$.

In Lemma B.3, we establish that

$$E_S \|\pi_A p^\circ | S; \theta_q - \pi_A p^\circ | S; \theta_{ML,n} q\|_1 \leq \frac{1}{d} \frac{1}{2} \frac{b}{2} \frac{1}{H^2 p^\circ, \theta_{ML,n} q}$$

where H^2 is the generalized squared Hellinger distance defined in (12) with $p|X; \theta_q = \pi_A p^\circ | S; \theta_q$ as the conditional density of $A|S$.

Combining the above two inequalities, we have that for any $\theta \in \Theta$,

$$\|\check{V}_S^\circ; \theta_q - V_S^\circ; \theta_{ML,n} q\|_1 \leq r_S C_S \frac{b}{H^2 p^\circ, \theta_{ML,n} q} \quad (15)$$

In the following, we use the fact that $\log x \leq \frac{1}{2} \frac{1}{x} - 1$ for any $x \in (0, 1]$ to show that for any $S \subseteq \mathcal{S}$ and any θ :

$$\begin{aligned} & \int \pi_A p^\circ | S; \theta q \log \frac{\pi_A p^\circ | S; \theta_q}{\pi_A p^\circ | S; \theta q} da \\ & \leq \int \frac{\pi_A p^\circ | S; \theta_q}{\pi_A p^\circ | S; \theta q} \frac{\pi_A p^\circ | S; \theta_q}{\pi_A p^\circ | S; \theta q} - 1 da \\ & = \int \frac{\pi_A p^\circ | S; \theta_q}{\pi_A p^\circ | S; \theta q} \frac{\pi_A p^\circ | S; \theta_q}{\pi_A p^\circ | S; \theta q} - \frac{\pi_A p^\circ | S; \theta_q}{\pi_A p^\circ | S; \theta q} da \\ & = \int \frac{\pi_A p^\circ | S; \theta_q}{\pi_A p^\circ | S; \theta q} \frac{\pi_A p^\circ | S; \theta_q}{\pi_A p^\circ | S; \theta q} - \frac{\pi_A p^\circ | S; \theta_q}{\pi_A p^\circ | S; \theta q} da \\ & = \int \frac{\pi_A p^\circ | S; \theta_q}{\pi_A p^\circ | S; \theta q} \frac{\pi_A p^\circ | S; \theta_q}{\pi_A p^\circ | S; \theta q} - \frac{\pi_A p^\circ | S; \theta_q}{\pi_A p^\circ | S; \theta q} da, \end{aligned}$$

which implies that

$$L p^\circ q - L p^\circ q \leq 2 H^2 p^\circ, \theta q \quad (16)$$

By Lemma B.4, we have that for any $\theta \in \Omega_n$,

$$H^2 p^\circ, \theta_{ML,n} q \leq 2 H^2 p^\circ, \theta_q \leq 2 H^2 p^\circ, \theta_{ML,n} q \leq L p^\circ_{ML,n} q - L p^\circ q \leq L p^\circ q - L p^\circ q \quad (17)$$

From Assumption 5.1, we have that with probability at least $1 - \delta$, for any $\theta \in \Omega_n$,

$$\|\check{L} p^\circ q - L p^\circ q\|_1 \leq \alpha_n \|\check{p}_n p^\circ q - p_n p^\circ q\|_1 \leq \alpha_n.$$

In other words, under Assumption 5.1, with probability at least $1 - \delta$ for any $\theta \in \Omega_n$,

$$L p^\circ q - L p^\circ q \leq \alpha_n \|\check{p}_n p^\circ q - p_n p^\circ q\|_1 \leq 2 \alpha_n. \quad (18)$$

Plugging Eq. (18) into Eq. (17), we have that with probability at least $1 - \delta$, for any $\theta \in \Omega_n$,

$$H^2 p^\circ, \theta_{ML,n} q \leq 4 \alpha_n. \quad (19)$$

Combining the above inequality and Eq. (15), we have that, with probability at least $1 - \delta$, $\|\check{V}_S^\circ; \theta_q - V_S^\circ; \theta_{ML,n} q\|_1 \leq r_S C_S \frac{b}{H^2 p^\circ, \theta_{ML,n} q} \leq \alpha_n$ for all $\theta \in \Omega_n$. This concludes the proof. $\square$

A.2. Proof of Lemma 6.1

Consider the MNL model (5) with a compact set Θ . Assume that $\theta \in \Theta$. For $0 \leq \delta \leq 1$, with probability at least $1 - \delta$, we have

$$\|\check{p}_n p^\circ q - p_n p^\circ_{ML,n} q\|_1 \leq \frac{C_A d}{n} \log \frac{\theta_{\max}}{\delta}. \quad (20)$$

Proof of Lemma 6.1. Fix $0 \leq \delta \leq 1$. Suppose $\alpha_n \approx \frac{C_A d}{n} \log \frac{2\theta_{\max}}{\delta}$. Define an oracle confidence set as

$$\Omega_n \triangleq \{\theta \in \Theta : L_{\theta} q \leq L_{\theta^*} q + d\alpha_n\}.$$

In particular, $\theta^* \in \Omega_n$. By Lemmas B.1 and A.2, we also have with probability at least $1 - \delta/2$,

$$L_{\theta_{ML,n}} q \leq L_{\theta^*} q + d2C_A H^2 \theta_{ML,n}, \quad \theta_{ML,n} \in \Omega_n,$$

that is, $\theta_{ML,n} \in \Omega_n$.

Define $F_n \triangleq \left\{ \log \frac{\pi_A p_A |S; \theta^* q}{\pi_A p_A |S; \theta q} : \theta \in \Omega_n \right\}$. In particular, for any $f \in F_n$, we have $\|f\|_{\infty} \leq d \log C_A$. Let $\{P_n \leq P\}_{F_n}$ be the envelope. By Talagrand's inequality, with probability at least $1 - \delta/2$, we have for any $f \in F_n$,

$$|P_n \leq P|_{F_n} \leq \frac{1}{n} \log C_A + \sup_{f \in F_n} E f^2 \leq \frac{1}{n} \log C_A + \log \frac{2}{\delta} \cdot \frac{\log C_A}{n} \log \frac{2}{\delta}. \quad (21)$$

For the variance term, we have

$$\begin{aligned} \sigma_{F_n}^2 &\triangleq \sup_{f \in F_n} E f^2 \leq \sup_{\theta \in \Omega_n} E \left(\log \frac{\pi_A p_A |S; \theta^* q}{\pi_A p_A |S; \theta q} \right)^2 \\ &\leq \sup_{\theta \in \Omega_n} E \left(2C_A H^2 \theta_{ML,n} \right)^2 \quad \text{by Lemma B.5} \\ &\leq 2C_A \sup_{\theta \in \Omega_n} H^2 \theta_{ML,n}^2 \\ &\leq 2C_A \sup_{\theta \in \Omega_n} L_{\theta^*} q \leq L_{\theta^*} q \quad \text{by (17)} \\ &\leq 2C_A \alpha_n. \quad \text{by definition of } \Omega_n \end{aligned}$$

For the expected envelope, our goal below is to apply Sen (2018, Theorem 7.13) (stated in Theorem B.6). Consider the covering number $N_{\mathcal{F}, F_n, L^2(\mathbb{Q})}$ for any given $\epsilon \geq 0$ and finitely supported probability measure $\mathbb{Q}$. By Lemma A.1, based on the MNL model (5), for some $L \leq 8$, F_n is a class of L -Lipschitz functions with respect to the index space Θ . Then in terms of the bracketing number N_{rs} and covering number N , for any $\epsilon \geq 0$ and probability measure $\mathbb{Q}$, we have

$$N_{\mathcal{F}, F_n, L^2(\mathbb{Q})} \leq N_{\text{rs}}(\mathcal{F}, F_n, L^2(\mathbb{Q}), \epsilon) \leq N_{\mathcal{F}, \Theta, L^2(\mathbb{Q})} \leq \frac{1}{\epsilon^d}.$$

By $\Theta \subseteq \mathbb{R}^d$ and Θ is compact, we further have $N_{\mathcal{F}, \Theta, L^2(\mathbb{Q})} \leq \frac{1}{\epsilon^d}$. Therefore,

$$N_{\mathcal{F}, F_n, L^2(\mathbb{Q})} \leq \frac{1}{\epsilon^d}.$$

By Theorem B.6, we further have

$$E|P_n \leq P|_{F_n} \leq \frac{d}{n} \sigma_{F_n}^2 \log \frac{L}{\sigma_{F_n}} + \frac{d}{n} \log C_A \log \frac{L}{\sigma_{F_n}} \leq \frac{C_A d}{n} \alpha_n \approx \frac{C_A d}{n} \log \frac{2\theta_{\max}}{\delta} \leq \alpha_n.$$

The Talagrand's inequality (21) becomes

$$|P_n \leq P|_{F_n} \leq \alpha_n.$$

In particular, in the case of $\theta_{ML,n} \in \Omega_n$ corresponding to $f_{\theta_{ML,n}} \in F_n$, we have

$$|P_n \leq P|_{F_n} \leq |P_n \leq P|_{F_n} \leq \alpha_n \leq \frac{L_{\theta^*} q}{n} \leq \frac{L_{\theta_{ML,n}} q}{n} \leq \alpha_n.$$

This complete the proof. $\square$

Lemma A.1. Consider the MNL model:

$$\pi_A(j|s; \theta) = \frac{\exp(x_j^T \theta)}{1 + \sum_{j \in S} \exp(x_j^T \theta)}; \quad \pi_A(0|s; \theta) = \frac{1}{1 + \sum_{j \in S} \exp(x_j^T \theta)}; \quad \forall s \in \mathcal{S}, \theta \in \Theta.$$

Let $\theta_{\max} = \max_{\theta \in \Theta} \|\theta\|_2$, $x_{\max} = \max_{j \in [N]} \|x_j\|_2$. If Θ is compact, that is, $\theta_{\max} < \infty$, then the log-likelihood ratio $\log \frac{\pi_A(p^A|s; \theta^*)}{\pi_A(p^A|s; \theta)}$ is a uniformly Lipschitz function in $\theta \in \Theta$.

Proof of Lemma A.1.

$$\begin{aligned} \left| \frac{1}{B} \log \frac{\pi_A(p^A|s; \theta^*)}{\pi_A(p^A|s; \theta)} \right| &\leq \frac{1}{\pi_A(p^A|s; \theta)} \left| \frac{1}{B} \sum_{j \in S} \log \frac{\exp(x_j^T \theta^*)}{\exp(x_j^T \theta)} \right| \\ &\leq \frac{1}{\pi_A(p^A|s; \theta)} \sum_{j \in S} \frac{|x_j^T (\theta^* - \theta)|}{\exp(x_j^T \theta)} \\ &\leq \frac{1}{\pi_A(p^A|s; \theta)} \sum_{j \in S} x_{\max} \exp(x_j^T \theta) \|\theta^* - \theta\|_2 \\ &\leq \frac{1}{\pi_A(p^A|s; \theta)} \sum_{j \in S} \exp(x_{\max} \theta_{\max}) x_{\max} \|\theta^* - \theta\|_2 \\ &\leq \frac{1}{\pi_A(p^A|s; \theta)} \exp(x_{\max} \theta_{\max}) x_{\max} \|\theta^* - \theta\|_2. \end{aligned} \quad (22)$$

That is, $\log \frac{\pi_A(p^A|s; \theta^*)}{\pi_A(p^A|s; \theta)}$ is L -Lipschitz. $\square$

Lemma A.2 (Concentration of Parametric MLE in Hellinger Distance). Consider the MNL model (5). For $0 < \delta < 1$, with probability at least $1 - \delta$, we have

$$H^2(p_{\text{ML},n}, p^*) \leq \frac{d}{n} \log \frac{\theta_{\max}}{\delta}.$$

Proof of Lemma A.2. We follow from Fu et al. (2022, Corollary 2) as a special case, where our data are generated i.i.d. instead of being a general Markov chain. $\square$

A.3. Proof of Lemma 6.2

Consider the MNL model (5). Suppose conditions in Lemma 6.1 hold, and $L(p^*)$ and $L_n(p^*)$ are uniformly and strongly convex. Let $\alpha_n = \frac{C_A d}{n} \log \frac{\theta_{\max}}{\delta}$. For $0 < \delta < 1$, with probability at least $1 - \delta$, we have

$$\sup_{\theta \in \Omega_n} |L(p^*) - L_n(p^*)| \leq \alpha_n.$$

Proof of Lemma 6.2. Fix $0 < \delta < 1$. By the strong convexity assumption on $L(p^*)$ and $L_n(p^*)$ there exists a constant $\mu > 0$ such that for any $\theta \in \Theta$,

$$\mu \|\theta - \theta^*\|_2^2 \leq L(p^*) - L_n(p^*) \leq \mu \|\theta - \theta_{\text{ML},n}\|_2^2 \leq L_n(p^*) - L_n(\theta_{\text{ML},n})$$

By Lemma 6.1, with probability at least $1 - \delta/2$, we have $\theta_{\text{ML},n} \in \Omega_n$. Then for any $\theta \in \Omega_n$, we have

$$\begin{aligned} \mu \|\theta - \theta^*\|_2^2 &\leq L(p^*) - L_n(p^*) \leq L_n(p^*) - L_n(\theta_{\text{ML},n}) \\ &\leq \frac{1}{\mu} |L_n(p^*) - L_n(\theta_{\text{ML},n})| \leq \frac{1}{\mu} |L(p^*) - L_n(p^*)| + |L_n(p^*) - L_n(\theta_{\text{ML},n})| \\ &\leq \frac{1}{\mu} \alpha_n + |L_n(p^*) - L_n(\theta_{\text{ML},n})| \end{aligned} \quad \begin{aligned} &\text{by triangular inequality} \\ &\text{by strong convexity} \\ &\text{by } \theta \in \Omega_n \text{ and } \theta_{\text{ML},n} \in \Omega_n \text{ respectively} \end{aligned}$$

The above implies that with probability at least $1 - \delta/2$, we have $\Omega_n \subseteq \mathcal{B}(\theta^*, \bar{\alpha}_n)$, where $\mathcal{B}(\theta^*, \bar{\alpha}_n)$ is a ball centered around θ^* with radius $\bar{\alpha}_n$:

$$\mathcal{B}(\theta^*, \bar{\alpha}_n) = \{\theta \in \Theta : \|\theta - \theta^*\|_2 \leq \bar{\alpha}_n\}.$$

Define $\mathcal{F}_n \leftarrow \left\{ \log \frac{\pi_A p^A |S; \theta^\circ q}{\pi_A p^A |S; \theta q} : \theta \in \Theta_n \right\}$. Let $\{P_n \in \mathcal{P}(\mathcal{F}_n) : \sup_{f \in \mathcal{F}_n} |P_n - P q f q| \leq \alpha_n\}$ be the envelop. By Talagrand's inequality, with probability at least $1 - \delta/2$, we have for any $f \in \mathcal{F}_n$,

$$|P_n - P q f q| \leq \mathbb{E} |P_n - P q f q| + \sqrt{\frac{1}{n} \log \frac{C_A}{\delta} \mathbb{E} |P_n - P q f q|} + \sup_{f \in \mathcal{F}_n} \mathbb{E} f^2 - \mathbb{E} f q^2 \log \frac{2}{\delta} + \frac{\log C_A}{n} \log \frac{2}{\delta}. \quad (23)$$

For the variance term, we have

$$\begin{aligned} \sigma_{\mathcal{F}_n}^2 &\leq \sup_{f \in \mathcal{F}_n} \mathbb{E} f^2 - \mathbb{E} f q^2 \leq \sup_{\theta \in \Theta_n} \mathbb{E} \log^2 \frac{\pi_A p^A |S; \theta^\circ q}{\pi_A p^A |S; \theta q} \\ &\leq \sup_{\theta \in \Theta_n} |\theta^\circ - \theta|_2^2 \quad \text{by Lipschitzness in Lemma A.1} \\ &\leq \alpha_n \quad \text{by definition of } \Theta_n \end{aligned}$$

For the expected envelope, by Theorem B.6, we further have

$$\mathbb{E} |P_n - P q f q| \leq \frac{d}{n} \sigma_{\mathcal{F}_n}^2 \log \frac{L}{\sigma_{\mathcal{F}_n}} + \frac{d}{n} \wedge 2 \log C_A q \log \frac{L}{\sigma_{\mathcal{F}_n}} \leq \frac{d}{n} \alpha_n \leq \alpha_n.$$

Therefore, the Talagrand's inequality (23) gives that for any $f \in \mathcal{F}_n$

$$|P_n - P q f q| \leq \alpha_n.$$

In other words, with probability at least $1 - \delta/2$, Θ_n corresponding to $f \in \mathcal{F}_n$, we have

$$\sup_{\theta \in \Theta_n} |L p^\theta q - L p^\theta q| \leq \sup_{\theta \in \Theta_n} |L p^\theta q - L p^\theta q| \leq \alpha_n.$$

□

B. Technical Lemmas

Lemma B.1. Suppose $C_A \geq 2$ and $C_A \leq 1/(\pi_A p^\circ |S; \theta q)$ for all $\theta \in \Theta$, $S \in \mathcal{S}$ and $i \in \mathcal{S}$. Then for any $\theta \in \Theta$:

$$|L p^\theta q - L p^\theta q| \leq 2C_A H^2 p^\theta, \theta q \quad (24)$$

Proof of Lemma B.1. By definition,

$$|L p^\theta q - L p^\theta q| \leq \mathbb{E} \left| \log \frac{\pi_A p^A |S; \theta^\circ q}{\pi_A p^A |S; \theta q} \right|.$$

In particular, for a fixed $S \in \mathcal{S}$, we have

$$\begin{aligned} \mathbb{E} \left| \log \frac{\pi_A p^A |S; \theta^\circ q}{\pi_A p^A |S; \theta q} \right| &\leq \mathbb{E} \left| \log \frac{\pi_A p^\circ |S; \theta^\circ q}{\pi_A p^\circ |S; \theta q} \right| \\ &\leq \frac{\log C_A + 1}{1 - 2/C_A} \leq 1 + h^2 p^\circ |S; \theta q| \pi_A p^\circ |S; \theta q| \\ &\quad \text{by log-Sobolev inequality (Diaconis & Saloff-Coste, 1996, Theorem A.1)} \\ &\leq \frac{C_A \log C_A + 2}{C_A - 2} \frac{1}{h^2} \leq h^2 p^\circ |S; \theta q| \pi_A p^\circ |S; \theta q| \\ &\leq 2C_A h^2 p^\circ |S; \theta q| \pi_A p^\circ |S; \theta q| \end{aligned}$$

Therefore,

$$|L p^\theta q - L p^\theta q| \leq 2C_A h^2 p^\circ |S; \theta q| \pi_A p^\circ |S; \theta q| \leq 2C_A H^2 p^\theta, \theta q$$

□

Lemma B.2. Let $C_S \leftarrow \frac{1}{\pi_S p^S q}$ and $r_S \leftarrow \max_{j \in \mathcal{P}^S} r_j$, then the following inequality holds for any $\theta_1, \theta_2 \in \Theta$:

$$\mathbb{E}_S \left[\left| \pi_A p^S; \theta_1 q - \pi_A p^S; \theta_2 q \right| \right] \leq \mathbb{E}_S \left[\left| \pi_A p^S; \theta_1 q - \pi_A p^S; \theta_2 q \right| \right]$$

where $\|\cdot\|_1$ denotes the L^1 norm.

Proof of Lemma B.2.

$$\mathbb{E}_S \left[\left| \pi_A p^S; \theta_1 q - \pi_A p^S; \theta_2 q \right| \right] \leq \mathbb{E}_S \left[\left| \pi_A p^S; \theta_1 q - \pi_A p^S; \theta_2 q \right| \right] \quad (25)$$

$$\mathbb{E}_S \left[\left| \pi_A p^S; \theta_1 q - \pi_A p^S; \theta_2 q \right| \right] \leq \mathbb{E}_S \left[\left| \pi_A p^S; \theta_1 q - \pi_A p^S; \theta_2 q \right| \right] \quad (26)$$

$$\mathbb{E}_S \left[\left| \pi_A p^S; \theta_1 q - \pi_A p^S; \theta_2 q \right| \right] \leq \mathbb{E}_S \left[\left| \pi_A p^S; \theta_1 q - \pi_A p^S; \theta_2 q \right| \right] \quad (27)$$

$$\mathbb{E}_S \left[\left| \pi_A p^S; \theta_1 q - \pi_A p^S; \theta_2 q \right| \right] \leq \mathbb{E}_S \left[\left| \pi_A p^S; \theta_1 q - \pi_A p^S; \theta_2 q \right| \right] \quad (28)$$

$$\mathbb{E}_S \left[\left| \pi_A p^S; \theta_1 q - \pi_A p^S; \theta_2 q \right| \right] \leq \mathbb{E}_S \left[\left| \pi_A p^S; \theta_1 q - \pi_A p^S; \theta_2 q \right| \right] \quad (29)$$

$$\mathbb{E}_S \left[\left| \pi_A p^S; \theta_1 q - \pi_A p^S; \theta_2 q \right| \right] \leq \mathbb{E}_S \left[\left| \pi_A p^S; \theta_1 q - \pi_A p^S; \theta_2 q \right| \right] \quad (30)$$

where Eq. (25) comes from the sample-based estimation of $\mathbb{E} r_{p^S q}$ (Eq. (1)), Eq. (27) comes from the Hölder's inequality, Eq. (28) comes from the fact that $\frac{\pi_A p^S; \theta_1 q}{\pi_S p^S q} \leftarrow C_S$ because $\frac{\pi_A p^S; \theta_1 q}{\pi_S p^S q}$ has the value zero everywhere except at $S \leftarrow S^*$. The last equality follows from the definition of L^1 norm. $\square$

Lemma B.3. For any $\theta_1, \theta_2 \in \Theta$,

$$\mathbb{E}_S \left[\left| \pi_A p^S; \theta_1 q - \pi_A p^S; \theta_2 q \right| \right] \leq \frac{1}{2} \sqrt{H^2 p_{\theta_1, \theta_2}^S}$$

Proof of Lemma B.3. We first use the facts that (1) $L^1 \leftarrow \frac{1}{2} \text{TV}$ where TV is the total variation distance and (2) $\text{TV} \leftarrow \frac{1}{2} h$ (Harsha, 2011) where h is the Hellinger distance to have that for any $S \in \mathcal{S}$,

$$\left| \pi_A p^S; \theta_1 q - \pi_A p^S; \theta_2 q \right| \leq \frac{1}{2} h \sqrt{\pi_A p^S; \theta_1 q \pi_A p^S; \theta_2 q} \quad (31)$$

From (31), we have

$$\begin{aligned} \left| \pi_A p^S; \theta_1 q - \pi_A p^S; \theta_2 q \right| &\leq \frac{1}{2} h \sqrt{\pi_A p^S; \theta_1 q \pi_A p^S; \theta_2 q}, \\ \left| \pi_A p^S; \theta_1 q - \pi_A p^S; \theta_2 q \right|^2 &\leq \frac{1}{4} h^2 \pi_A p^S; \theta_1 q \pi_A p^S; \theta_2 q. \end{aligned}$$

Taking expectation with respect to S on both sides, we have

$$\begin{aligned} \mathbb{E}_S \left[\left| \pi_A p^S; \theta_1 q - \pi_A p^S; \theta_2 q \right|^2 \right] &\leq \frac{1}{4} h^2 \mathbb{E}_S \left[\pi_A p^S; \theta_1 q \pi_A p^S; \theta_2 q \right], \\ \mathbb{E}_S \left[\left| \pi_A p^S; \theta_1 q - \pi_A p^S; \theta_2 q \right|^2 \right] &\leq \frac{1}{4} h^2 \mathbb{E}_S \left[\pi_A p^S; \theta_1 q \pi_A p^S; \theta_2 q \right] \end{aligned}$$

By the Jensen's inequality, we have

$$\mathbb{E}_S \left[\left| \pi_A p^S; \theta_1 q - \pi_A p^S; \theta_2 q \right| \right] \leq \sqrt{\mathbb{E}_S \left[\left| \pi_A p^S; \theta_1 q - \pi_A p^S; \theta_2 q \right|^2 \right]}.$$

This implies that

$$E_S \left[\left| \pi_A(p|S; \theta_1 q) - \pi_A(p|S; \theta_2 q) \right| \right] \leq \frac{1}{2} \sqrt{\frac{1}{2} H^2(p_{\theta_1}, \theta_2 q)}$$

□

Lemma B.4 (Properties of H and H^2). *For any $\theta_1, \theta_2, \theta_3 \in \Theta$, the following inequalities hold:*

$$\begin{aligned} H(p_{\theta_1}, \theta_2 q) &\leq H(p_{\theta_1}, \theta_3 q) + H(p_{\theta_2}, \theta_3 q) \\ H(p_{\theta_1}, \theta_2 q)^2 &\leq H^2(p_{\theta_1}, \theta_2 q) + H^2(p_{\theta_1}, \theta_2 q) \\ H^2(p_{\theta_1}, \theta_2 q) &\leq 2H^2(p_{\theta_1}, \theta_3 q) + 2H^2(p_{\theta_2}, \theta_3 q) \end{aligned}$$

Proof of Lemma B.4. For ease of notation, for $i = 1, 2, 3$ we use p_i to denote π_A parametrized by θ_i , i.e., $p_i(p|S; \theta_i q)$.

(1) Notice that for any $S \in \mathcal{S}$, $\sqrt{\frac{1}{2} H^2(p_1(p|S; \theta_1 q), p_2(p|S; \theta_2 q))}$ is just the regular Hellinger distance that satisfies the triangular inequality. Hence we have

$$\sqrt{\frac{1}{2} H^2(p_1(p|S; \theta_1 q), p_2(p|S; \theta_2 q))} \leq \sqrt{\frac{1}{2} H^2(p_1(p|S; \theta_1 q), p_3(p|S; \theta_3 q))} + \sqrt{\frac{1}{2} H^2(p_2(p|S; \theta_2 q), p_3(p|S; \theta_3 q))} \quad (32)$$

Take expectation of both side of Eq. (32) with respect to S , we have

$$E_S \left[\sqrt{\frac{1}{2} H^2(p_1(p|S; \theta_1 q), p_2(p|S; \theta_2 q))} \right] \leq E_S \left[\sqrt{\frac{1}{2} H^2(p_1(p|S; \theta_1 q), p_3(p|S; \theta_3 q))} \right] + E_S \left[\sqrt{\frac{1}{2} H^2(p_2(p|S; \theta_2 q), p_3(p|S; \theta_3 q))} \right].$$

By the definition of H , this means that

$$H(p_{\theta_1}, \theta_2 q) \leq H(p_{\theta_1}, \theta_3 q) + H(p_{\theta_2}, \theta_3 q)$$

(2) For the first inequality, by applying the Jensen's inequality, we have

$$E \left[\sqrt{\frac{1}{2} H^2(p_1(p|S; \theta_1 q), p_2(p|S; \theta_2 q))} \right]^2 \leq E \left[\frac{1}{2} H^2(p_1(p|S; \theta_1 q), p_2(p|S; \theta_2 q)) \right]. \quad (33)$$

Then the inequality follows.

For the second inequality, we have

$$H(p_{\theta_1}, \theta_2 q) \leq H^2(p_{\theta_1}, \theta_2 q) + E_S \left[\frac{1}{2} H^2(p_1(p|S; \theta_1 q), p_2(p|S; \theta_2 q)) \right] \quad (34)$$

$$= E_S \left[1 + \frac{1}{2} H^2(p_1(p|S; \theta_1 q), p_2(p|S; \theta_2 q)) \right] \quad (35)$$

Note that for any S , $\frac{1}{2} H^2(p_1(p|S; \theta_1 q), p_2(p|S; \theta_2 q))$ is a non-negative function because the Hellinger distance is no larger than 1, and $\frac{1}{2} H^2(p_1(p|S; \theta_1 q), p_2(p|S; \theta_2 q))$ is also a positive function because it is a metric. We have Eq. (35) as the expectation of non-negative functions, and thus we have

$$E_S \left[1 + \frac{1}{2} H^2(p_1(p|S; \theta_1 q), p_2(p|S; \theta_2 q)) \right] \geq 1.$$

Then the inequality follows.

(3) Notice that for any a, b, c we have $a^2 + b^2 \geq 2ab \geq c^2 - 2ab \geq c^2$. With this fact, we have that for any S ,

$$\frac{1}{2} H^2(p_1(p|S; \theta_1 q), p_2(p|S; \theta_2 q)) \leq \frac{1}{2} \int_0^1 \frac{1}{2} \left(\frac{p_1(p|S; \theta_1 q)}{p_2(p|S; \theta_2 q)} + \frac{p_2(p|S; \theta_2 q)}{p_1(p|S; \theta_1 q)} - 2 \right) da \quad (36)$$

$$\leq \frac{1}{2} \int_0^1 \frac{1}{2} \left(\frac{p_1(p|S; \theta_1 q)}{p_2(p|S; \theta_2 q)} + \frac{p_2(p|S; \theta_2 q)}{p_3(p|S; \theta_3 q)} - 2 \right) da + \frac{1}{2} \int_0^1 \frac{1}{2} \left(\frac{p_2(p|S; \theta_2 q)}{p_3(p|S; \theta_3 q)} + \frac{p_3(p|S; \theta_3 q)}{p_2(p|S; \theta_2 q)} - 2 \right) da \quad (37)$$

$$\leq \frac{1}{2} \int_0^1 \frac{1}{2} \left(\frac{p_1(p|S; \theta_1 q)}{p_3(p|S; \theta_3 q)} + \frac{p_3(p|S; \theta_3 q)}{p_1(p|S; \theta_1 q)} - 2 \right) da + \frac{1}{2} \int_0^1 \frac{1}{2} \left(\frac{p_2(p|S; \theta_2 q)}{p_3(p|S; \theta_3 q)} + \frac{p_3(p|S; \theta_3 q)}{p_2(p|S; \theta_2 q)} - 2 \right) da \quad (38)$$

$$\leq \frac{1}{2} H^2(p_1(p|S; \theta_1 q), p_3(p|S; \theta_3 q)) + \frac{1}{2} H^2(p_2(p|S; \theta_2 q), p_3(p|S; \theta_3 q)) \quad (39)$$

This implies that

$$H^2(p_{\theta_1}, p_{\theta_2}) \leq 2H^2(p_{\theta_1}, p_{\theta_3}) + 2H^2(p_{\theta_2}, p_{\theta_3}) \quad (40)$$

□

Lemma B.5 (Log-Density Ratio Variance Bound). *Suppose X is an $\mathbb{R}$ -valued random variable with probability density function p , and p_1, p_2 are two other probability density functions for X such that p_1 and p_2 are uniformly bounded from below by C^{-1} on the support of p . Then we have*

$$\mathbb{E}_{X \sim p} \left(\log \frac{p_1(X)}{p_2(X)} \right)^2 \leq 2Ch^2(p_1, p_2)$$

where h^2 is the squared Hellinger distance in (10).

Proof of Lemma B.5. By $\log p(X) \leq 2p(X) - 1$ for any $X \in \mathbb{R}$, we have

$$\begin{aligned} \mathbb{E}_{X \sim p} \left(\log \frac{p_1(X)}{p_2(X)} \right)^2 &\leq \mathbb{E}_{X \sim p} \left(\log \frac{p_1(X)}{p_2(X)} \right)^2 p(X) dx \\ &\leq \mathbb{E}_{X \sim p} \left(\frac{p_1(X)}{p_2(X)} - 1 \right)^2 p(X) dx \\ &\leq \mathbb{E}_{X \sim p} \left(\frac{1}{p_2(X)} - \frac{1}{p_1(X)} \right)^2 p(X) dx \\ &\leq \mathbb{E}_{X \sim p} \left(\frac{1}{p_2(X)} + \frac{1}{p_1(X)} \right)^2 p(X) dx \\ &\leq 4C \mathbb{E}_{X \sim p} \left(\frac{1}{p_1(X)} + \frac{1}{p_2(X)} \right)^2 dx \\ &\leq 2Ch^2(p_1, p_2) \end{aligned}$$

□

Theorem B.6 (Sen (2018, Theorem 7.13)). *Let $\mathcal{F}$ be a measurable function class, such that $\sup_{f \in \mathcal{F}} \|f\|_8 \leq f_{\max}$ for some constant $f_{\max} \leq 8$. Assume that for $A \in \mathbb{R}$, $d \in \mathbb{N}$, $d \leq A$, and every finitely supported probability measure Q , we have the covering number (Sen, 2018) as:*

$$N(\mathcal{F}, L^2(Q)) \leq \frac{A^d}{\epsilon^d}. \quad (41)$$

Let $\sigma_F^2 = \sup_{f \in \mathcal{F}} \mathbb{E} f^2 - (\mathbb{E} f)^2$. Then we have

$$\mathbb{E} \|P_n - P\|_F \leq \frac{C}{n} \frac{A^d}{\sigma_F^2 \log \frac{A}{\sigma_F}} + \frac{d}{n} f_{\max} \log \frac{A}{\sigma_F}.$$

Proof of Theorem B.6. In this proof, we denote X as the underlying random variable, $t_{i=1}^n X_i$ are n i.i.d. copies of X , and for any $f \in \mathcal{F}$, $P_n f = \frac{1}{n} \sum_{i=1}^n f(X_i)$, $P f = \mathbb{E} f(X)$. Without loss of generality, assume that $0 \leq f$, and for any $f \in \mathcal{F}$, $P f \leq 1$. Let ϵ_i be i.i.d. Rademacher random variables that are independent of $t_{i=1}^n X_i$. By symmetrization, we have

$$\mathbb{E} \|P_n - P\|_F \leq 2 \mathbb{E} \sup_{f \in \mathcal{F}} \frac{1}{n} \sum_{i=1}^n \epsilon_i f(X_i). \quad (42)$$

Conditional on $t_{i=1}^n X_i$, by Dudley's entropy bound, we have

$$\mathbb{E} \sup_{f \in \mathcal{F}} \frac{1}{n} \sum_{i=1}^n \epsilon_i f(X_i) \leq \frac{1}{n} \sqrt{\log N(\mathcal{F}, L^2(P_n))} \leq \frac{1}{n} \sqrt{\log N(\mathcal{F}, L^2(P))} \leq \frac{1}{n} \sqrt{\log \frac{A^d}{\epsilon^d}}. \quad (43)$$

In particular, $\log^A\{\sigma_{F,n} \mid q \leq \log^A\{f_{\max} \mid q \leq 1\}$ by assumption, which satisfies the condition for Lemma B.7. Combining (42), (43) and (44), we have

where E takes expectation with respect to $\mathbf{t}^X_{i|\mathcal{U}^n_{i-1}}$. Notice that

where we define $F^2 \leftarrow \inf_{f \in F} \|f\|_{L^2(\mathbb{P}_{n,q})}^2$. We aim to apply (42), (43), (44), (45) to F^2 . Notice that $\sup_{f \in F} \|f\|_{L^2(\mathbb{P}_{n,q})}^2 \leq \frac{2f_{\max}^2}{P_n} + \frac{2f_{\max}^2}{P_n} = \frac{4f_{\max}^2}{P_n}$. By $\|f\|_{L^2(\mathbb{P}_{n,q})}^2 \leq \frac{4f_{\max}^2}{P_n}$, we further have

Therefore, applying (42), (43), (44) and (45) to F^2 , we have

Define $B \stackrel{\text{b}}{=} \frac{1}{2} \mathbb{E} \sigma_{F,n}^2 q \log \frac{A^2}{\mathbb{E} \sigma_{F,n}^2 q}$. Then we have

By $u \tilde{N} u \log \frac{A^2}{u}$ is non-decreasing on $u \in [\rho, A^2]$ and non-increasing on $u \in [A^2, 8q]$, we have

In particular, $B \leq d' \frac{A^2}{2e}, \sigma_F^2 \leq d' f_{\max}^2 \leq d' A^2 \{e^2 - A^2\} e$. Then the cap $A^2 \{e$ is inactive as $d(n - \tilde{N}) \rightarrow 0$ asymptotically. Therefore,

In particular, B is bounded by both roots of the corresponding quadratic equation:

19

Combined with (45), we further have

$$E\{P_n\} - P_F \leq \frac{C}{n} B \leq \frac{C}{n} \frac{d}{\sigma_F^2} \log \frac{A}{\sigma_F} = \frac{d}{n} f_{\max} \log \frac{A}{\sigma_F}^*.$$

□

Lemma B.7. Suppose $a, A \geq 0$ such that $\log \frac{A}{a} \leq 1$. Then we have

$$\int_a^A \frac{C}{\log \frac{A}{u}} du \leq 2a \int_a^A \frac{C}{\log \frac{A}{a}}.$$

Proof of Lemma B.7. Define

$$f(a) = \begin{cases} \int_a^A \frac{C}{\log \frac{A}{u}} du, & a \geq 0; \\ 0, & a = 0. \end{cases}$$

Then f is continuous at 0. Moreover, for $a \geq 0$, we have

$$f'(a) = -\frac{C}{\log \frac{A}{a}}, \quad f''(a) = \frac{1}{a \log \frac{A}{a}},$$

which is nonnegative if $\log \frac{A}{a} \leq 1$. As $a \rightarrow 0$, we further have

$$\begin{aligned} \frac{f(a)}{a} &\leq \frac{C}{2 \log \frac{A}{a}} + \frac{1}{a} \int_a^A \frac{C}{\log \frac{A}{u}} du \\ &\leq \frac{C}{\log \frac{A}{a}} + \frac{1}{2a} \int_a^A \frac{1}{\log \frac{A}{u}} du \quad \text{by integration-by-part} \\ &\leq \frac{C}{\log \frac{A}{a}} + \frac{1}{2 \log \frac{A}{a}}; \\ \liminf_{a \rightarrow 0} \frac{f(a)}{a} &\leq \frac{1}{2}. \end{aligned}$$

Therefore, for any $a \geq 0$, we have $f(a) \leq 0$, which concludes the proof.

□