
Multinomial Logit Bandit with Low Switching Cost

Kefan Dong¹ Yingkai Li² Qin Zhang³ Yuan Zhou⁴

Abstract

We study multinomial logit bandit with limited adaptivity, where the algorithms change their exploration actions as infrequently as possible when achieving almost optimal minimax regret. We propose two measures of adaptivity: the assortment switching cost and the more fine-grained item switching cost. We present an anytime algorithm (AT-DUCB) with $O(N \log T)$ assortment switches, almost matching the lower bound $\Omega(\frac{N \log T}{\log \log T})$. In the fixed-horizon setting, our algorithm FH-DUCB incurs $O(N \log \log T)$ assortment switches, matching the asymptotic lower bound. We also present the ESUCB algorithm with item switching cost $O(N \log^2 T)$.

1. Introduction

The dynamic assortment selection problem with the multinomial logit (MNL) choice model, also called MNL-bandit, is a fundamental problem in online learning and operations research. In this problem we have N distinct items, each of which is associated with a known reward r_i and an *unknown* preference parameter v_i . In the MNL choice model, given a subset $S \subseteq [N] \stackrel{\text{def}}{=} \{1, 2, 3, \dots, N\}$, the probability that a user chooses $i \in S$ is given by

$$p_i(S) = \begin{cases} \frac{v_i}{v_0 + \sum_{j \in S} v_j} & \text{if } i \in S \cup \{0\} \\ 0 & \text{otherwise} \end{cases}, \quad (1)$$

where “0” stands for the case that the user does not choose any item, and v_0 is the associated preference parameter. As a convention (see, e.g. Agrawal et al., 2019), we assume

Author names are listed in alphabetical order. ¹Institute for Interdisciplinary Information Sciences, Tsinghua University, Beijing, China. ²Department of Computer Science, Northwestern University, Evanston, Illinois, USA. ³Computer Science Department, Indiana University, Bloomington, Indiana, USA. ⁴Department of ISE, University of Illinois at Urbana-Champaign, Urbana, Illinois, USA. Correspondence to: Yuan Zhou <yuanz@illinois.edu>.

that no-purchase is the most frequent choice, which is very natural in retailing. W.l.o.g., we assume $v_0 = 1$, and $v_i \leq 1$ for all $i \in [N]$. The *expected reward* of the set S under the preference vector $\mathbf{v} = \{v_0, v_1, \dots, v_N\}$ is defined to be

$$R(S, \mathbf{v}) = \sum_{i \in S} r_i p_i(S) = \sum_{i \in S} \frac{r_i v_i}{1 + \sum_{j \in S} v_j}. \quad (2)$$

For any online policy that selects a subset $S_t \subseteq [N]$ ($|S_t| \leq K$, where K is a predefined capacity parameter) at each time step t , observes the user’s choice a_t to gradually learn the preference parameters $\{v_i\}$, and runs for a horizon of T time steps, we define the *regret* of the policy to be

$$\text{Reg}_T \stackrel{\text{def}}{=} \sum_{t=1}^T (R(S^*, \mathbf{v}) - R(S_t, \mathbf{v})), \quad (3)$$

where $S^* = \arg \max_{S \subseteq [N], |S| \leq K} R(S, \mathbf{v})$ is the optimal assortment in hindsight. The goal is to find a policy to minimize the expected regret $\mathbb{E}[\text{Reg}_T]$ for all MNL-bandit instances.

To motivate the definition of the MNL-bandit problem, let us consider a fast fashion retailer such as Zara or Mango. Each of its product corresponds to an item in $[N]$, and by selling the i -th item the retailer takes a profit of r_i . At each specific time in each of its shops, the retailer can only present a certain number of items (say, at most K) on the shelf due to the space constraints. As a consequence, customers who visit the store can only pick items from the presented assortment (or, just buy nothing which corresponds to item 0), following a choice model. There has been a number of choice models being proposed in the literature (see, e.g., (Train, 2009; Luce, 2012) for overviews), and the MNL model is arguably the most popular one. The retailer certainly wants to maximize its profit by identifying the best assortment S^* to present. However, it does not know in advance customers’ preferences to items in $[N]$ (i.e., the preference vector $\mathbf{v}$), to get which it has to learn from customers’ actual choices. More precisely, the retailer needs to develop a policy to choose at each time step t an assortment $S_t \subseteq [N]$ ($|S_t| \leq K$) based on the previous presented assortments $S_1, \dots, S_{t-1}$ and customers’ choices in the past $(t-1)$ time steps. The retailer’s expected reward in a time horizon T can be expressed by $\sum_{t=1}^T R(S_t, \mathbf{v})$, which is typically reformulated as the regret compared with the best policy in the form of (3).

The MNL-bandit problem has attracted quite some attention in the past decade (Rusmevichientong et al., 2010; Sauré & Zeevi, 2013; Agrawal et al., 2016; 2017; Chen & Wang, 2018). However, all these works do not consider an important practical issue for regret minimization: in reality it is often impossible to *frequently* change the assortment display. For example, in retail stores it may not be possible to change the display in the middle of the day, not mentioning doing it after each purchase. We thus hope to minimize the number of assortment switches in the selling time horizon *without* increasing the regret by much. Another advantage of achieving a small number of assortment switches is that such algorithms are easier to parallelize, which enables us to learn users’ preferences much faster. This feature is particularly useful in applications such as online advertising where it is easy to show the same assortment (i.e., a set of ads) in a large amount of end users’ displays simultaneously.

We are interested in two kinds of switching costs under a time horizon T . The first is the *assortment switching cost*, defined as

$$\Psi_T^{(\text{asst})} \stackrel{\text{def}}{=} \sum_{t=1}^T \mathbb{I}[S_t \neq S_{t+1}].$$

The second is the *item switching cost*, defined as

$$\Psi_T^{(\text{item})} \stackrel{\text{def}}{=} \sum_{t=1}^T |S_t \oplus S_{t+1}|,$$

where binary operator $\oplus$ computes the symmetric difference of the two sets. In comparison, the item switching cost is more fine-grained and put less penalty if two neighboring assortments are “almost the same”. As a straightforward observation, we always have that

$$\Psi_T^{(\text{asst})} \leq \Psi_T^{(\text{item})} \leq \min\{2K, N\} \cdot \Psi_T^{(\text{asst})}. \quad (4)$$

Our results. In this paper we obtain the following results for MNL-bandit with low switching cost. By default all log’s are of base 2.

We first introduce an algorithm, AT-DUCB, that achieves almost optimal regret (up to a logarithmic factor) and incurs an assortment switching cost of $O(N \log T)$; this algorithm is *anytime*, i.e., it does *not* need to know the time horizon T in advance. We then show that the AT-DUCB algorithm achieves almost optimal assortment switching cost. In particular, we prove that every anytime algorithm that achieves almost optimal regret must incur an assortment switching cost of at least $\Omega(N \log T / \log \log(NT))$. These results are presented in Section 2.

When the time horizon is known beforehand, we obtain an algorithm, FH-DUCB, that achieves almost optimal regret (up to a logarithmic factor) and incurs an assortment switching cost of $O(N \log \log T)$. We also prove the optimality of

this switching cost by establishing a matching lower bound. See Section 3.

For item switches, while the trivial application of (4) leads to $O(N^2 \log T)$ and $O(N^2 \log \log T)$ item switching cost bounds for AT-DUCB and FH-DUCB respectively, in Section 4, we design a new algorithm, ESUCB, to achieve an item switching cost of $O(N \log^2 T)$. In Appendix F, we show that a more careful modification to the algorithm further improves the item switching cost to $O(N \log T)$.

We make two interesting observations from the results above: (1) there is a *separation* between the assortment switching complexities when knowing the time horizon T and when not; in other words, the time horizon T is useful for achieving a smaller assortment switching cost; (2) the item switching cost is only at most a logarithmic factor higher than the assortment switching cost.

Technical contributions. We combine the epoch-based offering algorithm for MNL-bandits (Agrawal et al., 2019) and a natural delayed update policy in the design of AT-DUCB. Although a similar delayed update rule has been recently analyzed for multi-armed bandits and Q-learning (Bai et al., 2019), and such a result does not seem surprising, we present it in the paper as a warm-up to help the readers get familiar with a few algorithmic techniques commonly used for the MNL-bandit problem.

Our first main technical contribution comes from the design of FH-DUCB algorithm, where we invent a novel delayed update policy that uses the horizon information to improve the switching cost from $O(N \log T)$ to $O(N \log \log T)$. We note that for the ordinary multi-armed bandit problem, recent works (Gao et al., 2019) and (Simchi-Levi & Xu, 2019) managed to show a similar $O(N \log \log T)$ switching cost with known horizon. However, their update rules do not have to utilize the learned parameters for the arms, and a straightforward conversion of such update rules to the MNL-bandit problem does not produce the desired guarantees. In contrast, our update rule, formally described in (6), carefully exploits the structure of the MNL-bandits and uses the information of the partially learned preference parameters (more specifically, $\hat{v}_{i, \tau_i}$ in (6)) to adaptively decide when to switch to a different assortment.

Our second main technical contribution is the ESUCB algorithm for the low item switching cost. The technical challenge here stems from the fact that the low item switching cost is a much stronger requirement than the low assortment switching cost, and simple lazy updates with the doubling trick and the straightforward analysis will show that the item switching cost is at most N times the assortment switching cost (see (4)), leading to a total item switching cost of $O(N^2 \log T)$. To reducing the extra factor N , we propose the idea of decoupling the learning for the optimal revenue

and the assortment, so that the offering of the assortment is decided via optimizing a new objective function based on the (usually) fixed revenue estimate. Since the revenue estimates are fixed, the offered assortments enjoy improved stability, and the item switching cost can be upper bounded by careful analysis.

We remark that the item switching cost is a particularly interesting goal that arises in online learning problems when the actions are sets of elements, which is very different from traditional MAB and linear bandits. Thanks to our novel technical ingredients, we are able to bring the item switching cost down to almost the same order as the assortment switching cost. We hope our results will inspire future study of the switching costs in both settings for other online learning problems with set actions.

Related work. MNL-bandit was first studied in (Rusmevichientong et al., 2010) and (Sauré & Zeevi, 2013), where the authors took the “explore-then-commit” approach, and proposed algorithms with regret $O(N^2 \log^2 T)$ and $O(N \log T)$ respectively under the assumption that the gap between the best and second-to-the-best assortments is known. (Agrawal et al., 2016) removed this assumption using a UCB-type algorithm, which achieves a regret of $O(\sqrt{NT \log T})$. An almost tight regret lower bound of $\Omega(\sqrt{NT})$ was later given by (Chen & Wang, 2018). (Agrawal et al., 2017) proposed an algorithm using Thompson Sampling, which achieves comparable regret bound to the UCB-type algorithms while demonstrates a better numerical performance.

Learning with low policy switches (also called learning in the *batched model* or *limited adaptivity*) has recently been studied in reinforcement learning for several other problems, including stochastic multi-armed bandits (Perchet et al., 2015; Jun et al., 2016; Agarwal et al., 2017; Gao et al., 2019; Esfandiari et al., 2019; Simchi-Levi & Xu, 2019), Q-learning (Bai et al., 2019), and online-learning (Cesa-Bianchi et al., 2013). This research direction is motivated by the fact that in many practical settings, the change of learning policy is very costly. For example, in clinical trials, every treatment policy switch would trigger a separate approval process. In crowdsourcing, it takes time for the crowd to answer questions, and thus a small number of rounds of interactions with the crowd is desirable. The performance of the learning would be much better if the data is processed in batches and during each batch the learning policy is fixed.

2. Warm-up: An anytime algorithm with $O(N \log T)$ assortment switches

As a warm-up, we begin with a simple anytime algorithm using at most $O(N \log T)$ assortment switches. Our algorithm combines the epoch-based offering framework introduce in

(Agrawal et al., 2016) and a deferred update policy. We will first briefly explain the epoch-based offering procedure, and then present and analyze our algorithm.

The epoch-based offering. In the epoch-based offering framework, whenever we are to offer an assortment S , instead of offering it for only one time period, we keep offering S until a no-purchase decision (item 0) is observed, and refer to all the consecutive time periods involved in this procedure as an *epoch*. The detailed offering procedure is described in Algorithm 1, where t is the global counter for the time period, and $\{\Delta_i\}$ records the number of purchases made for each item i in the epoch.

Algorithm 1: EXPLORATION(S)

```

1 Initialize:  $\Delta_i \leftarrow 0$  for all  $i \in [N]$ ;
2 while TRUE do
3    $t \leftarrow t + 1$ ;
4   Offer assortment  $S$ , and observe purchase decision  $a_t$ ;
5   If  $a_t = 0$  then return  $\{\Delta_i\}$ ;
6    $\Delta_{a_t} \leftarrow \Delta_{a_t} + 1$ ;
    
```

The following key observation for EXPLORATION(S) states that $\{\Delta_i\}$ forms an unbiased estimate for the utility parameters of all items in S .

Observation 1. *Let $\{\Delta_i\}$ be returned by EXPLORATION(S). For each $i \in S$, Δ_i is an independent geometric random variable with mean v_i . Moreover, one can verify that $\mathbb{E}[\Delta_i] = v_i$ and*

$$\Pr[\Delta_i = k] = \left(\frac{v_i}{1 + v_i} \right)^k \left(\frac{1}{1 + v_i} \right), \forall k \in \mathbb{N}.$$

At any time of the algorithm when an epoch has ended, for each item $i \in [N]$, we let $\bar{v}_i = n_i/T_i$ where T_i is the number of the past epochs in which i is included in the offered assortment, and n_i is the total number of purchases for item i during all past epochs. By Observation 1, we know that $\bar{v}_i$ is also an unbiased estimate of v_i . In (Agrawal et al., 2016), the following upper confidence bound (UCB) is constructed for each $i \in [N]$,

$$\hat{v}_i = \bar{v}_i + \sqrt{\frac{48 \bar{v}_i \ln(\sqrt{N} \ell + 1)}{T_i}} + \frac{48 \ln(\sqrt{N} \ell + 1)}{T_i}. \quad (5)$$

We will compute the assortment for the next epoch based on the vector of UCB values $\hat{\mathbf{v}} = (\hat{v}_1, \hat{v}_2, \dots, \hat{v}_n)$.

We describe our algorithm in Algorithm 2, which can be seen as an adaptation of the one in (Agrawal et al., 2016). The main difference from (Agrawal et al., 2016) is that the UCB values (and hence the assortment) is updated only when T_i reaches an integer power of 2 for any item $i \in [N]$.

Algorithm 2: Anytime Deferred Update UCB (AT-DUCB)

```

1 Initialize:  $\hat{v}_i \leftarrow 1, T_i \leftarrow 0$  for all  $i \in [N], t \leftarrow 0$ ;
2 for  $\ell \leftarrow 1, 2, 3, \dots$ , do
3   Compute  $S_\ell = \arg \max_{S \subseteq [N]: |S| \leq K} R(S, \hat{\mathbf{v}})$ ;
4    $\{\Delta_i\} \leftarrow \text{EXPLORATION}(S)$ ;
5   for  $i \in S$  do
6      $n_i \leftarrow n_i + \Delta_i$  and  $T_i \leftarrow T_i + 1$ ;
7     if  $T_i = 2^k$  for some  $k \in \mathbb{Z}$  then
8        $\bar{v}_i \leftarrow n_i/T_i; \hat{v}_i \leftarrow \min \{ \hat{v}_i, \bar{v}_i + \sqrt{\frac{48\bar{v}_i \ln(\sqrt{N}\ell+1)}{T_i}} + \frac{48 \ln(\sqrt{N}\ell+1)}{T_i} \}$ ;

```

This deferred update strategy is implemented in Line 7. Also note that instead of directly evaluating (5), the update in Line 8 makes sure that $\hat{v}_i$ is non-increasing as the algorithm proceeds. We comment that the optimization task in Line 3 can be done efficiently, as studied in, for example, (Rusmevichientong et al., 2010).

Theorem 2. For any time horizon T , the expect regret incurred by Algorithm 2 is

$$\mathbb{E}[\text{Reg}_T] \lesssim \sqrt{NT \log T},$$

and the expected number of assortment switches $\mathbb{E}[\Psi_T^{(\text{asst})}]$ is $O(N \log T)$.¹

The proof of the regret upper bound in Theorem 2 is similar to that of (Agrawal et al., 2016), except for a more careful analysis about the deferred update rule. For completeness, we prove this part in Appendix A.

Proof of the assortment switch upper bound in Theorem 2.

Let $\mathcal{D}_i^{(\ell)}$ be the event that Line 8 is executed in Algorithm 2 for item i at the ℓ -th epoch. Recall that the assortment S_ℓ is computed by $S_\ell = \arg \max_{S \subseteq [N], |S| \leq K} R(S, \hat{\mathbf{v}})$, and $\hat{\mathbf{v}}$ is updated after epoch ℓ only when $\mathcal{D}_i^{(\ell)}$ happens for some $i \in [N]$. Let L be the total number of epochs at or before time T ; we thus have $\sum_{\ell=1}^L \mathbb{I}[\mathcal{D}_i^{(\ell)}] \leq \log T$. We then have that

$$\begin{aligned} \mathbb{E}[\Psi_T^{(\text{asst})}] &= \mathbb{E} \sum_{t=1}^{T-1} \mathbb{I}[S_t \neq S_{t+1}] \\ &\leq \sum_{\ell=1}^L \sum_{i=1}^N \mathbb{I}[\mathcal{D}_i^{(\ell)}] = \sum_{i=1}^N \sum_{\ell=1}^L \mathbb{I}[\mathcal{D}_i^{(\ell)}] \lesssim N \log T. \end{aligned}$$

□

¹For two sequences $\{a_n\}$ and $\{b_n\}$, we write $a_n = O(b_n)$ or $a_n \lesssim b_n$ if there exists a universal constant $C < \infty$ such that $\limsup_{n \rightarrow \infty} |a_n|/|b_n| \leq C$. Similarly, we write $a_n = \Omega(b_n)$ or $a_n \gtrsim b_n$ if there exists a universal constant $c > 0$ such that $\liminf_{n \rightarrow \infty} |a_n|/|b_n| \geq c$.

The lower bound. We complement our algorithmic result with the following almost matching lower bound. The theorem states that the number of assortment switches has to be $\Omega(N \log T / \log \log(NT))$, if the algorithm is anytime and incurs only $\sqrt{NT} \times \text{poly} \log(NT)$ regret. The proof of Theorem 3 can be found in Appendix E.1.

Theorem 3. There exist universal constants $d_0, d_1 > 0$ such that the following holds. For any constant $C \geq 1$, if an anytime algorithm $\mathcal{A}$ achieves expected regret at most $d_0 \sqrt{NT} (\ln(NT))^C$ for all T and all instances with N items, then for any $N \geq 2$, $T_0 \geq N$ and T_0 greater than a sufficiently large constant that only depends on C , there exists an instance with N items and a time horizon $T \in [T_0, T_0^2]$, such that the expected number of assortment switches before time T is at least $d_1 N \log T / (C \log \log(NT))$.

3. Achieving $O(N \log \log T)$ assortment switch with known horizons

When the time horizon is known to the algorithm, we can exploit this advantage via more carefully designed update policy to achieve only $O(N \log \log T)$ assortment switches. For the convenience of presentation, we first introduce a few notations.

Algorithm 3: UPDATE(i)

```

1  $\tau_i \leftarrow \tau_i + 1; T_i^{(\tau_i)} \leftarrow T_i^{(\tau_i-1)} + |\mathcal{T}(i, \tau_i - 1)|$ ;
2  $n_i^{(\tau_i)} \leftarrow n_i^{(\tau_i-1)} + n_{i, \tau_i-1}; \bar{v}_{i, \tau_i} \leftarrow n_i^{(\tau_i)} / T_i^{(\tau_i)}$ ;
3  $\hat{v}_{i, \tau_i} \leftarrow \min \left\{ \hat{v}_{i, \tau_i-1}, \bar{v}_{i, \tau_i} + \sqrt{\frac{48 \bar{v}_{i, \tau_i} \ln(\sqrt{N} T^2 + 1)}{T_i^{(\tau_i)}}} + \frac{48 \ln(\sqrt{N} T^2 + 1)}{T_i^{(\tau_i)}} \right\}$ ;

```

For each item $i \in [N]$, we divide the time periods into consecutive *stages* where the boundaries between any two neighboring stages are marked by the UCB updates for item i . Note that the division for the stages may be different for different items. For any $\tau \in \{1, 2, 3, \dots\}$, let $\mathcal{T}(i, \tau)$ be the set of epochs to offer item i , in stage τ for the item. Let $T_i^{(\tau)} = \sum_{\tau'=1}^{\tau-1} |\mathcal{T}(i, \tau')|$ be the total number of epochs to offer item i , before stage τ for the item, and let $n_i^{(\tau)}$ be the total number of purchases for item i in the epochs counted by $T_i^{(\tau)}$. We can therefore define $\bar{v}_{i, \tau} \stackrel{\text{def}}{=} n_i^{(\tau)} / T_i^{(\tau)}$ as an unbiased estimate of v_i based on the observations before stage τ . Similarly to (5), we can define $\hat{v}_{i, \tau}$ as a UCB for v_i . The UPDATE(i) procedure (formally described in Algorithm 3) is invoked whenever the main algorithm decides to conclude the current stage for item i and update the UCB for v_i together with the quantities defined above, where τ_i is the counter for the number of stages for item i , and $n_{i, \tau}$ is the number of purchases observed in stage τ for

item i .

The key to the design of our main algorithm for the fixed time horizon setting is a new trigger for updating the UCB values. Let $\tau_0 = \lceil \log \log(T/N) + 1 \rceil$, for each item $i \in [N]$, we will conclude the current stage τ_i and invoke $\text{UPDATE}(i)$ whenever the following condition $\mathcal{P}(i, \tau_i)$ is satisfied. Note that $\mathcal{P}(i, \tau_i)$ is adaptive to the estimated parameters $\hat{v}_{i, \tau_i}$ to customize the number of epochs between assortment switches for each item. More specifically, the smaller $\hat{v}_{i, \tau_i}$ is, the less regret may be incurred by offering item i , and therefore the longer we can offer item i without switching and incurring too large regret, and this is reflected in the design of $\mathcal{P}$.

$$\mathcal{P}(i, \tau_i) \stackrel{\text{def}}{=} \begin{cases} |\mathcal{T}(i, \tau_i)| \geq 1 + \sqrt{\frac{T \cdot T_i^{(\tau_i)}}{N}} & \text{if } \tau_i < \tau_0 \\ |\mathcal{T}(i, \tau_i)| \geq 1 + \sqrt{\frac{T \cdot T_i^{(\tau_i)}}{N \cdot \hat{v}_{i, \tau_i}}} \\ \text{and } \hat{v}_{i, \tau_0} > 1/\sqrt{NT} & \text{if } \tau_i \geq \tau_0 \end{cases}. \quad (6)$$

For each epoch ℓ , we use $\tau_i(\ell)$ to denote the stage (in terms of item i) where epoch ℓ belongs to. We present the details of our main algorithm in Algorithm 4. The algorithm is terminated whenever the time step t reaches the horizon T .

Theorem 4. *For any given time horizon $T \geq N^4$, we have the following upper bound for the expected regret:*

$$\mathbb{E}[\text{Reg}_T] \lesssim \sqrt{NT \ln(\sqrt{NT^2} + 1)} \cdot \log \log T,$$

and the following upper bound for the expected number of assortment switches:

$$\mathbb{E}[\Psi_T^{(\text{asst})}] \lesssim N \log \log T.$$

To prove Theorem 4, we first define the desired events. Let

$$\mathcal{E}_{i, \tau}^{(1)} \stackrel{\text{def}}{=} \left\{ \hat{v}_{i, \tau} \geq v_i \text{ and } \hat{v}_{i, \tau} \leq v_i + \sqrt{\frac{144v_i \ln(\sqrt{NT^2} + 1)}{T_i^{(\tau)}} + \frac{144 \ln(\sqrt{NT^2} + 1)}{T_i^{(\tau)}}} \right\},$$

and

$$\mathcal{E}^{(1)} \stackrel{\text{def}}{=} \cap_{i, \tau} \mathcal{E}_{i, \tau}^{(1)}.$$

We also let

$$\mathcal{E}_{i, \tau}^{(2)} \stackrel{\text{def}}{=} \left\{ n_{i, \tau} \geq \frac{1}{2} v_i |\mathcal{T}(i, \tau)|, \right. \\ \left. \text{if } v_i \geq \frac{1}{2} \sqrt{\frac{1}{NT}} \text{ and } |\mathcal{T}(i, \tau)| \geq \frac{T}{4N \cdot v_i} \right\},$$

and

$$\mathcal{E}^{(2)} \stackrel{\text{def}}{=} \cap_{i, \tau} \mathcal{E}_{i, \tau}^{(2)}.$$

Finally, let $\mathcal{E} = \mathcal{E}^{(1)} \cap \mathcal{E}^{(2)}$. In Appendix B.1, we prove the following lemma.

Algorithm 4: Deferred Update UCB for Fixed Time Horizon (FH-DUCB)

Input : The time horizon T .

```

1 Initialize:  $\tau_i \leftarrow 1, \hat{v}_{i, \tau_i} \leftarrow 1, n_{i, \tau_i} \leftarrow 0, \mathcal{T}(i, \tau_i) \leftarrow \emptyset, T_i^{(1)} \leftarrow 0, n_i^{(1)} \leftarrow 0$  for all  $i \in [N]$ ;
2  $t \leftarrow 0, S_0 \leftarrow [N]$ ;
3 for  $\ell \leftarrow 1, 2, 3, \dots$ , do
4    $S_\ell \leftarrow S_{\ell-1}$ ;
5   if  $\exists i : \mathcal{P}(i, \tau_i)$  holds then
6      $\text{UPDATE}(i)$  for all  $i$  such that  $\mathcal{P}(i, \tau_i)$  holds;
7     Compute  $S_\ell \leftarrow \arg \max_{S \subseteq [N]: |S| \leq K} R(S, \hat{v}_\ell)$ 
      where  $\hat{v}_\ell = (\hat{v}_{i, \tau_i(\ell)})_{i \in [N]}$ ;
8    $\{\Delta_i\} \leftarrow \text{EXPLORATION}(S_\ell)$ ;
9   for  $i \in S$  do
10     $n_{i, \tau_i} \leftarrow n_{i, \tau_i} + \Delta_i$ ; Add  $\ell$  to  $\mathcal{T}(i, \tau_i)$ ;
```

Lemma 5. *If $T \geq N^4$ and T is greater than a large enough universal constant, then $\Pr[\mathcal{E}] \geq 1 - \frac{14}{T}$.*

Bounds for the stage lengths. When $\mathcal{E}$ happens, we can infer the following useful lower bound for the lengths of the stages after τ_0 . The lemma is proved in Appendix B.2.

Lemma 6. *Assume that $T \geq N^4$ and T is greater than a sufficiently large universal constant. Conditioned on $\mathcal{E}^{(1)}$, for each $i \in [N]$, if τ_0 is not the last stage for item i , we have that $v_i \geq \frac{1}{2} \sqrt{\frac{1}{NT}}$. Additionally, if $\hat{v}_{i, \tau_0} > 1/\sqrt{NT}$, then for all $\tau > \tau_0$ such that τ is not the last stage for i , we have that $|\mathcal{T}(i, \tau)| \geq (T/(2Nv_i))^{1-2^{-\tau+\tau_0+1}}$.*

Upper bounding the number of assortment switches.

Suppose that there are L epochs before the algorithm terminates. We only need to upper bound $\mathbb{E} \sum_{i=1}^N \tau_i(L)$ which upper bounds the number of assortment switches $\mathbb{E}[\Psi_T^{(\text{asst})}]$. For each $i \in [N]$, if $\tau_i(L) \geq \tau_0$ and $\hat{v}_{i, \tau_0} \leq 1/\sqrt{NT}$, we easily deduce that $\tau_i(L) \leq \tau_0 + 1$ because of the condition $\mathcal{P}(i, \tau_0)$. Otherwise, assuming that $\hat{v}_{i, \tau_0} > 1/\sqrt{NT}$, by Lemma 6, conditioned on $\mathcal{E}^{(1)}$, we have that $v_i \geq \frac{1}{2} \sqrt{\frac{1}{NT}}$ and $|\mathcal{T}(i, \tau)| \geq \frac{T}{4Nv_i}$ for all $\tau \in [\tau_0 + \log \log \frac{T}{2Nv_i} + 1, \tau_i(L) - 1]$. Because of $\mathcal{E}^{(2)}$, we have $n_{i, \tau} \geq \frac{v_i}{2} \cdot |\mathcal{T}(i, \tau)| \geq \frac{T}{8N}$ for all $\tau \in [\tau_0 + \log \log \frac{T}{2Nv_i} + 1, \tau_i(L) - 1]$. Therefore, we know that there are no more than $8N$ pairs of (i, τ) satisfying $\tau \in [\tau_0 + \log \log \frac{T}{2Nv_i} + 1, \tau_i(L) - 1]$. In total, conditioned on $\mathcal{E}$, we have that

$$\mathbb{E} \sum_{i=1}^N \tau_i(L)$$

$$\begin{aligned}
 &\lesssim N\tau_0 + \sum_{i=1}^N \mathbb{I}[\hat{v}_{i,\tau_0} > 1/\sqrt{NT}] \log \log \frac{T}{2Nv_i} \\
 &\quad + \mathbb{E} \sum_{i=1}^N \max\{\tau_i(L) - \tau_0 - \log \log \frac{T}{2Nv_i}, 0\} \\
 &\lesssim N \log \log T + \sum_{i=1}^N \log \log \frac{T^{3/2}}{N^{1/2}} \lesssim N \log \log T, \quad (7)
 \end{aligned}$$

where the second inequality is because of Lemma 6. Finally, since the contribution to the expected number of assortment switches when $\mathcal{E}$ fails is at most $\Pr[\mathcal{E}] \cdot T \leq O(1)$ (because of Lemma 5), we prove the upper bound for the number of assortment switches in Theorem 4.

Upper bounding the expected regret. Let $E^{(\ell)}$ be the length of epoch ℓ , i.e., the number of time steps taken in epoch ℓ . Note that $E^{(\ell)}$ is a geometric random variable with mean value $(1 + \sum_{i \in S_\ell} v_i)$. Also recall that there are L epochs in total. Letting S^* be the optimal assortment, conditioned on event $\mathcal{E}^{(1)}$, we have that

$$\begin{aligned}
 \mathbb{E}[\text{Reg}_T] &= \mathbb{E} \sum_{\ell=1}^L E^{(\ell)} (R(S^*, v) - R(S_\ell, v)) \\
 &= \mathbb{E} \sum_{\ell=1}^L \left(1 + \sum_{i \in S_\ell} v_i \right) (R(S^*, v) - R(S_\ell, v)) \\
 &\leq \mathbb{E} \sum_{\ell=1}^L \sum_{i \in S_\ell} (\hat{v}_{i,\tau_i(\ell)} - v_i) \\
 &= \mathbb{E} \sum_{i=1}^N \sum_{\ell: i \in S_\ell} (\hat{v}_{i,\tau_i(\ell)} - v_i) \\
 &= \mathbb{E} \sum_{i=1}^N \sum_{\tau=1}^{\tau_i(L)} \sum_{\ell \in \mathcal{T}(i,\tau)} (\hat{v}_{i,\tau} - v_i), \quad (8)
 \end{aligned}$$

where the inequality is due to Lemma 17. In the next lemma, we upper bound the contribution from each item i and stage τ to the upper bound in (8). The lemma is proved in Appendix B.3.

Lemma 7. *Conditioned on event $\mathcal{E}^{(1)}$, for any item i and any stage $\tau \leq \tau_i(L)$, we have that*

$$\sum_{\ell \in \mathcal{T}(i,\tau)} (\hat{v}_{i,\tau} - v_i) \lesssim \sqrt{T \ln(\sqrt{NT^2} + 1)/N}.$$

Combining Lemma 5, Lemma 7, inequalities (7) and (8), we have that

$$\begin{aligned}
 \mathbb{E}[\text{Reg}_T] &\leq T \cdot \Pr[\overline{\mathcal{E}^{(1)}}] + \mathbb{E}[\text{Reg}_T \mid \mathcal{E}^{(1)}] \\
 &\lesssim 1 + \mathbb{E} \sum_{i=1}^N \tau_i(L) \times \sqrt{\frac{T \ln(\sqrt{NT^2} + 1)}{N}}
 \end{aligned}$$

$$\lesssim \sqrt{NT \ln(\sqrt{NT^2} + 1)} \cdot \log \log T,$$

proving the expected regret upper bound in Theorem 4.

The lower bound. We prove the following matching lower bound in Appendix E.2.

Theorem 8. *For any constant $C \geq 0$ and time horizon T , if an algorithm $\mathcal{A}$ achieves expected regret $\mathbb{E}[\text{Reg}_T]$ at most $\frac{1}{7525} \cdot \sqrt{NT} (\ln(NT))^C$ for all N -item instances, then there exists an N -item instance such that the expected number of assortment switches is*

$$\mathbb{E}[\Psi_T^{(\text{asst})}] = \Omega(N \log \log T).$$

4. Optimizing the number of item switches

In this section, we study how to minimize the item switch cost while still achieving $\tilde{O}(\sqrt{NT})$ regret.

Algorithm 5: The Exponential Stride UCB algorithm (ESUCB) for MNL-Bandit

```

1 Initialize:  $\hat{\theta} \leftarrow 1, \epsilon_1 \leftarrow 1/3, c_1 \leftarrow 44840$ ;
2 for  $\tau \leftarrow 1, 2, 3, \dots$  do
3    $t_{\max} \leftarrow c_1 N \ln^3(NT/\delta)/\epsilon_\tau^2$ ;
4   if CHECK( $\hat{\theta} - 3\epsilon_\tau, \hat{\theta} - \epsilon_\tau, t_{\max}$ ) then  $\hat{\theta} \leftarrow \hat{\theta} - \epsilon_\tau$ ;
5    $\epsilon_{\tau+1} \leftarrow \frac{2}{3}\epsilon_\tau$ ;
    
```

We now propose a new algorithm, Exponential Stride UCB (ESUCB), to achieve an item switching cost that is linear with N and poly-logarithmic with T . The specific guarantee of the ESUCB algorithm is presented in Theorem 10, the main theorem of this section. The key idea of the algorithm is to decouple the learning of the optimal expected revenue and the optimal assortment, which is made possible by the following lemma.

Lemma 9. *Define $G(\theta) \stackrel{\text{def}}{=} R(S_\theta, v)$, where $S_\theta \stackrel{\text{def}}{=} \arg \max_{S \subseteq [N]: |S| \leq K} (\sum_{i \in S} v_i (r_i - \theta))$. There exists a unique θ^* such that*

$$G(\theta^*) = \theta^* = \max_{|S| \leq K} R(S, v).$$

Moreover,

- (1) for any $\theta < \theta^*$, we have that $G(\theta) > \theta$, and
- (2) for any $\theta > \theta^*$, we have that $G(\theta) < \theta$.

The proof of Lemma 9 is deferred to Appendix D.1. Motivated by the lemma, we present our ESUCB algorithm in Algorithm 5. The algorithm learns the optimal revenue θ^* in the main loop, using a sequence of exponentially decreasing learning step size ϵ_τ . For each estimate $\hat{\theta}$, the CHECK

procedure (Algorithm 6) learns the assortment $S_{\hat{\theta}}$ via the UCB method with deferred updates. (More precisely speaking, the algorithm learns $S_{\hat{\theta}-\epsilon_\tau}$ and $S_{\hat{\theta}-3\epsilon_\tau}$, and at Line 4, chooses one of them based on the UCB estimation $\hat{\rho}$ for the expected revenue of $S_{\hat{\theta}-\epsilon_\tau}$.) In the CHECK procedure, the variable t keeps the count of time steps and is updated in EXPLORATION. We also make the following notes: 1) The ESUCB algorithm needs the horizon T as input, and uses a confidence parameter δ , which is usually set as $1/T$. The whole algorithm terminates whenever the horizon T is reached. 2) At the optimization steps (Lines 6 and 9 of Algorithm 6), we have to adopt a deterministic tie breaking rule, e.g., we let the $\arg \max$ operator to return the S such that $\sum_{i \in S} 2^i$ is minimized among multiple maximizers.

Theorem 10. *Setting $\delta = 1/T$, we have the following upper bound for the expected regret of ESUCB:*

$$\mathbb{E}[\text{Reg}_T] \lesssim \sqrt{NT} \cdot \log^{1.5}(NT),$$

and the item switching cost for ESUCB is

$$\mathbb{E}[\Psi_T^{(\text{item})}] \lesssim N \log^2 T.$$

To prove Theorem 10, we upper bound the item switching cost and the expected regret separately.

Upper bounding the item switch cost. Since the estimate of θ^* is fixed in CHECK, the outcome of $\arg \max_{S: |S| \leq K} \sum_{i \in S} \hat{v}_i(r_i - \theta)$ (corresponding to Lines 6 and 9 of Algorithm 6) becomes more stable compared to that of $\arg \max_{S: |S| \leq K} R(S, \hat{\theta})$ in previous algorithms. Exploiting this advantage, we upper bound the number of item switches incurred by each call of CHECK as follows. The lemma is proved in Appendix D.2.

Lemma 11. *The item switch cost incurred by any invocation $\text{CHECK}(\theta_l, \theta_r, t_{\max})$ is $O(N \log T)$.*

Since the τ loop in Algorithm 5 iterates for only $O(\log T)$ times, Lemma 11 easily implies an $O(N \log^2 T)$ item switching cost upper bound for ESUCB. We also note that this bound can be improved to $O(N \log T)$ via a slight modification to the algorithm which is elaborated in Appendix F.

Upper bounding the expected regret. We first provide the following guarantees for CHECK.

Lemma 12 (Main Lemma for CHECK). *For any invocation $\text{CHECK}(\theta_l, \theta_r, t_{\max})$, with probability at least $(1 - \delta/T)$, the following statements hold.*

(a) *If CHECK returns true, then $G(\theta_r) < \theta_r$.*

(b) *If CHECK returns false, then*

$$\theta^* \geq \theta_r - \frac{2}{t_{\max}} \left(c_2 \sqrt{N t_{\max} \ln^3 \frac{NT}{\delta}} + c_3 N \ln^3 \frac{NT}{\delta} \right).$$

Algorithm 6: $\text{CHECK}(\theta_l, \theta_r, t_{\max})$

```

1 Initialize:  $\hat{v}_i \leftarrow 1, T_i \leftarrow 0, n_i \leftarrow 0$  for all  $i \in [N]$ ,
    $c_2 \leftarrow 688, c_3 \leftarrow 21732$ ;
2  $\rho \leftarrow 0, \hat{\rho} \leftarrow 1, b \leftarrow \text{false}, t \leftarrow 0$ ;
3 for  $\ell \leftarrow 1, 2, 3, \dots$  do
4   if  $\hat{\rho} < \theta_r$  then
5      $b \leftarrow \text{true}$ ;
6      $S_\ell \leftarrow \arg \max_{S \subseteq [N], |S| \leq K} (\sum_{i \in S} \hat{v}_i(r_i - \theta_l))$ ;
7      $\{\Delta_i\} \leftarrow \text{EXPLORATION}(S_\ell)$ ;
8   else
9      $S_\ell \leftarrow \arg \max_{S \subseteq [N], |S| \leq K} (\sum_{i \in S} \hat{v}_i(r_i - \theta_r))$ ;
10     $\{\Delta_i\} \leftarrow \text{EXPLORATION}(S_\ell)$ ;
11     $\rho \leftarrow \rho + \sum_{i \in S_\ell} \Delta_i \cdot r_i; \hat{\rho} \leftarrow \frac{1}{t}(\rho +$ 
       $c_2 \sqrt{N t_{\max} \ln^3(NT/\delta)} + c_3 N \ln^3(NT/\delta))$ ;
12  if  $t \geq t_{\max}$  then return  $b$ ;
13  for  $i \in S_\ell$  do
14     $n_i \leftarrow n_i + \Delta_i, T_i \leftarrow T_i + 1$ ;
15    if  $T_i = 2^k$  for some  $k \in \mathbb{Z}$  then
16       $\hat{v}_i \leftarrow n_i/T_i; \hat{v}_i \leftarrow \min \{ \hat{v}_i, \hat{v}_i +$ 
         $\sqrt{\frac{196 \hat{v}_i \log(NT/\delta+1)}{T_i}} + \frac{292 \log(NT/\delta+1)}{T_i} \} ;$ 

```

(c) *Let $r_{\text{CHECK}}^{(t)}$ be the reward at time step t in this invocation. If $\theta_l \leq \theta^*$, then we have that*

$$t_{\max} \theta_l - \mathbb{E} \left[\sum_{t=1}^{t_{\max}} r_{\text{CHECK}}^{(t)} \right] \lesssim \sqrt{N t_{\max} \ln^3(NT/\delta)} + N \ln^3(NT/\delta).$$

Proof of Lemma 12 is built upon Lemma 9 and deferred to Appendix D.3.

Let $\mathcal{Q}_\tau$ be the event that the statements (a) – (c) hold for the invocation of CHECK at iteration τ of Algorithm 5, and let $\mathcal{Q}$ be the event that $\mathcal{Q}_\tau$ holds every all τ . By Lemma 12 and a union bound, we immediately have that $\Pr[\mathcal{Q}] \geq 1 - \delta$. The next lemma, built upon Lemma 9 and Lemma 12, shows that $\hat{\theta}$ in Algorithm 5 is always an upper confidence bound for the true parameter θ^* , and converges to θ^* with a decent rate.

Lemma 13. *Let $\hat{\theta}^{(\tau)}$ be the value of $\hat{\theta}$ at the beginning of iteration τ of Algorithm 5. Conditioned on event $\mathcal{Q}$, for any iteration $\tau = 1, 2, 3, \dots$, we have that $\hat{\theta}^{(\tau)} - 3\epsilon_\tau \leq \theta^* \leq \hat{\theta}^{(\tau)}$.*

Proof. Recall that for every $\tau = 1, 2, 3, \dots$, we need to prove

$$\hat{\theta}^{(\tau)} - 3\epsilon_\tau \leq \theta^* \leq \hat{\theta}^{(\tau)}. \quad (9)$$

We prove this by induction. For iteration $\tau = 1$, (9) trivially holds since $0 \leq r_i \leq 1$ and therefore $0 \leq \theta^* \leq 1$.

Now suppose (9) holds for iteration τ , we will establish (9) for iteration $(\tau + 1)$. Consider the invocation of $\text{CHECK}(\theta_l, \theta_r, t_{\max})$ at iteration τ , where $\theta_l = \hat{\theta}^{(\tau)} - 3\epsilon_\tau$ and $\theta_r = \hat{\theta}^{(\tau)} - \epsilon_\tau$. We discuss the following two cases.

Case 1. When the CHECK procedure returns **true**, by Lemma 12 we have that $G(\theta_r) < \theta_r$. By Lemma 9, we have that $\theta_r > \theta^*$. Therefore, by Line 4 and the induction hypothesis we have that $\hat{\theta}^{(\tau+1)} = \hat{\theta}^{(\tau)} - \epsilon_\tau = \theta_r > \theta^*$, and $\hat{\theta}^{(\tau+1)} - 3\epsilon_{\tau+1} = \theta_r - 2\epsilon_\tau = \hat{\theta}^{(\tau)} - 3\epsilon_\tau \leq \theta^*$, proving (9).

Case 2. When the CHECK procedure returns **false**, by Lemma 12, we have that

$$\theta^* \geq \theta_r - \frac{1}{t_{\max}} \left((c_2 + 8) \sqrt{N t_{\max} \ln^3 \frac{NT}{\delta}} + c_3 N \ln^3 \frac{NT}{\delta} \right).$$

Recall that at Line 3 we set $t_{\max} = c_1 N \ln^3(NT/\delta)/\epsilon_\tau^2$. For large enough c_1 , this implies that

$$\theta^* \geq \theta_r - \epsilon_\tau = \hat{\theta}^{(\tau)} - 2\epsilon_\tau = \hat{\theta}^{(\tau+1)} - 3\epsilon_{\tau+1}.$$

By Line 4 and the induction hypothesis we have that $\hat{\theta}^{(\tau+1)} = \hat{\theta}^{(\tau)} \geq \theta^*$, finishing the proof of (9). $\square$

Finally we upper bound the expected regret of Algorithm 5.

Lemma 14. *With probability at least $1 - \delta$, the expected regret incurred by Algorithm 5 is $O(\sqrt{NT} \log^{1.5}(NT/\delta))$. Therefore, if we set $\delta = 1/T$, we have that*

$$\mathbb{E}[\text{Reg}_T] \lesssim \sqrt{NT} \log^{1.5}(NT).$$

Proof. Throughout the proof we condition on the event $\mathcal{Q}$, which happens probability at least $(1 - \delta)$. We first prove that at iteration τ of Algorithm 5, the expected regret for this iteration is bounded by $\tilde{O}(N/\epsilon_\tau)$. Consider the invocation $\text{CHECK}(\theta_l, \theta_r, t_{\max})$ at Line 4. Recall that we define $t_{\max} = c_1 N \ln^3(NT/\delta)/\epsilon_\tau^2$. Combining with statement (c) of Lemma 12 and Lemma 13, the expected regret of this invocation is bounded by (where the $O(N)$ term is due to the last epoch that might run over time $t_{\max}$),

$$\begin{aligned} & \mathbb{E} \left[\theta^* \cdot t_{\max} - \sum_{t=1}^{t_{\max}} r_{\text{CHECK}}^{(t)} \right] + O(N) \\ & \lesssim t_{\max}(\theta^* - \theta_l) + \mathbb{E} \left[\theta_l \cdot t_{\max} - \sum_{t=1}^{t_{\max}} r_{\text{CHECK}}^{(t)} \right] + O(N) \\ & \lesssim t_{\max}(\theta^* - \theta_l) + N \ln^3(NT/\delta)/\epsilon_\tau. \end{aligned} \quad (10)$$

By Lemma 13, we have that $\theta^* - \theta_l \lesssim \epsilon_\tau$. Therefore, (10) is upper bounded by $O(N \ln^3(NT/\delta)/\epsilon_\tau)$.

Since $\text{CHECK}(\theta_l, \theta_r, t_{\max})$ runs for at least $t_{\max}$ time steps, the second to the last iteration $(\tau_{\max} - 1)$ satisfies that $c_1 N \ln^3(NT/\delta)/\epsilon_{\tau_{\max}-1}^2 \leq T$, which means that

$$\epsilon_{\tau_{\max}} \gtrsim \sqrt{N \log^3(NT/\delta)/T}.$$

Since ϵ_τ is an exponential sequence, the overall expected regret is bounded by the order of

$$\sum_{\tau=1}^{\tau_{\max}} N \log^3(NT/\delta)/\epsilon_\tau \lesssim \sqrt{NT \log^3(NT/\delta)}.$$

$\square$

Refined and non-trivial item switching cost upper bound for the AT-DUCB algorithm. Since an assortment switch may incur at most $2K$ item switches, Theorem 2 trivially implies that Algorithm 2 (AT-DUCB) incurs at most $O(KN \log T)$ item switches, which is upper bounded by $O(N^2 \log T)$ since $K = O(N)$.

In Appendix C, we present a refined analysis showing that the item switching cost of AT-DUCB is at most $O(N^{1.5} \log T)$. While it is not clear to us whether the dependence on N delivered by this analysis is optimal, we also discuss the relationship between the analysis and an extensively studied (but not yet fully resolved) geometry problem, namely the maximum number of planar K -sets. We hope that further study of this relationship might lead to improvement of both upper and lower bounds of the item switching cost of AT-DUCB. Please refer to Appendix C for more details.

5. Conclusion

In this paper, we present algorithms for MNL-bandits that achieve both almost optimal regret and assortment switching cost, in both anytime and fixed-horizon settings. We also design the ESUCB algorithm that achieves the almost optimal regret and item switching cost $O(N \log^2 T)$. For future directions, it is interesting to study whether it is possible to achieve an item switching cost of $O(N \log T)$ in the anytime setting and $O(N \log \log T)$ in the fixed-horizon setting. Also, as mentioned in Section 4 (and Appendix C), given the simplicity of our AT-DUCB algorithm, it is worthwhile to further refine the bounds for its item switching cost.

Acknowledgement

Part of the work done while Kefan Dong was a visiting student at UIUC. Kefan Dong and Yuan Zhou were supported in part by a Ye Grant and a JPMorgan Chase AI Research Faculty Research Award. Qin Zhang was supported in part by NSF IIS-1633215, CCF-1844234 and CCF-2006591.

References

- Agarwal, A., Agarwal, S., Assadi, S., and Khanna, S. Learning with limited rounds of adaptivity: Coin tossing, multi-armed bandits, and ranking from pairwise comparisons. In *COLT*, pp. 39–75, 2017.
- Agrawal, S., Avadhanula, V., Goyal, V., and Zeevi, A. A near-optimal exploration-exploitation approach for assortment selection. In *EC*, pp. 599–600, 2016.
- Agrawal, S., Avadhanula, V., Goyal, V., and Zeevi, A. Thompson sampling for the MNL-bandit. In *Proceedings of the 30th Conference on Learning Theory, COLT 2017, Amsterdam, The Netherlands, 7-10 July 2017*, pp. 76–78, 2017.
- Agrawal, S., Avadhanula, V., Goyal, V., and Zeevi, A. MNL-bandit: A dynamic learning approach to assortment selection. *Operations Research*, 67(5):1453–1485, 2019.
- Bai, Y., Xie, T., Jiang, N., and Wang, Y.-X. Provably efficient q-learning with low switching cost. In *NeurIPS*, 2019.
- Cesa-Bianchi, N., Dekel, O., and Shamir, O. Online learning with switching costs and other adaptive adversaries. In *NIPS*, pp. 1160–1168, 2013.
- Chen, X. and Wang, Y. A note on a tight lower bound for capacitated MNL-bandit assortment selection models. *Oper. Res. Lett.*, 46(5):534–537, 2018.
- Dey, T. K. Improved bounds for planar k-sets and related problems. *Discrete & Computational Geometry*, 19(3): 373–382, 1998.
- Esfandiari, H., Karbasi, A., Mehrabian, A., and Mirrokni, V. S. Batched multi-armed bandits with optimal regret. *CoRR*, abs/1910.04959, 2019.
- Gao, Z., Han, Y., Ren, Z., and Zhou, Z. Batched multi-armed bandits problem. In *NeurIPS*, 2019.
- Jin, Y., Li, Y., Wang, Y., and Zhou, Y. On asymptotically tight tail bounds for sums of geometric and exponential random variables. *arXiv preprint arXiv:1902.02852*, 2019.
- Jun, K., Jamieson, K. G., Nowak, R. D., and Zhu, X. Top arm identification in multi-armed bandits with batch arm pulls. In *AISTATS*, pp. 139–148, 2016.
- Luce, R. D. *Individual choice behavior: A theoretical analysis*. Courier Corporation, 2012.
- Perchet, V., Rigollet, P., Chassang, S., and Snowberg, E. Batched bandit problems. In *COLT*, pp. 1456, 2015.
- Pinsker, M. S. *Information and information stability of random variables and processes*. Holden-Day, 1964.
- Rusmevichientong, P., Shen, Z. M., and Shmoys, D. B. Dynamic assortment optimization with a multinomial logit choice model and capacity constraint. *Operations Research*, 58(6):1666–1680, 2010.
- Sauré, D. and Zeevi, A. Optimal dynamic assortment planning with demand learning. *Manufacturing & Service Operations Management*, 15(3):387–404, 2013.
- Simchi-Levi, D. and Xu, Y. Phase transitions and cyclic phenomena in bandits with switching constraints. In *NeurIPS*, 2019.
- Tóth, G. Point sets with many k-sets. *Discrete & Computational Geometry*, 26(2):187–194, 2001.
- Train, K. E. *Discrete Choice Methods with Simulation*. Cambridge university press, 2009.

Appendix

A. Proof of the regret upper bound in Theorem 2

In this section we complete the proof of Theorem 2 for completeness. The proof is almost identical to that in (Agrawal et al., 2017) except for the handling of the deferred UCB value updates.

The following lemma proves that $\hat{v}_i$ is indeed an upper confidence bound of true parameter v_i with high probability, and converges to the true value with decent rate.

Lemma 15 (Lemma 4.1 of (Agrawal et al., 2017)). *For any $\ell = 1, 2, 3, \dots$, in Algorithm 2, at Line 7 immediately after the ℓ -th epoch, the following two statements hold,*

1. *With probability at least $1 - \frac{6}{N\ell}$, $\frac{n_i}{T_i} + \sqrt{\frac{48(n_i/T_i) \ln(\sqrt{N}\ell + 1)}{T_i}} + \frac{48 \ln(\sqrt{N}\ell + 1)}{T_i} \geq v_i$ for any $i \in [N]$,*
2. *With probability at least $1 - \frac{7}{N\ell}$, for any $i \in [N]$,*

$$\frac{n_i}{T_i} + \sqrt{\frac{48(n_i/T_i) \ln(\sqrt{N}\ell + 1)}{T_i}} + \frac{48 \ln(\sqrt{N}\ell + 1)}{T_i} - v_i \leq \sqrt{\frac{144v_i \ln(\sqrt{N}\ell + 1)}{T_i}} + \frac{144 \ln(\sqrt{N}\ell + 1)}{T_i}.$$

By the update rule, Lemma 16 can be extended to $\{\hat{v}_i\}$ as follows.

Lemma 16. *For any $\ell = 1, 2, 3, \dots$, the following two statements hold at the end of the ℓ -th iteration of the outer for-loop of Algorithm 2.*

1. *With probability at least $1 - \frac{6}{N\ell}$, $\hat{v}_i \geq v_i$ for any $i \in [N]$,*
2. *With probability at least $1 - \frac{7}{N\ell}$, for any $i \in [N]$,*

$$\hat{v}_i - v_i \lesssim \sqrt{\frac{v_i \log(\sqrt{N}\ell + 1)}{T_i}} + \frac{\log(\sqrt{N}\ell + 1)}{T_i}.$$

Proof. For any epoch ℓ , let T'_i and $\hat{v}'_i$ be the value of T_i and $\hat{v}_i$ at the last update. Then we have, $\hat{v}_i = \hat{v}'_i$ and $T'_i \leq 2T_i$. Inherited from Lemma 15, we have $\hat{v}_i = \hat{v}'_i \geq v_i$. And

$$\hat{v}_i - v_i = \hat{v}'_i - v_i \lesssim \sqrt{\frac{v_i \log(\sqrt{N}\ell + 1)}{T'_i}} + \frac{\log(\sqrt{N}\ell + 1)}{T'_i} \lesssim \sqrt{\frac{v_i \log(\sqrt{N}\ell + 1)}{T_i}} + \frac{\log(\sqrt{N}\ell + 1)}{T_i}.$$

□

Once we establish Lemma 16, the proof of the regret upper bound in Theorem 2 is identical to that in (Agrawal et al., 2017). We include the proof here for completeness.

The next lemma shows that the expect regret for one epoch is bounded by the summation of estimation errors in the assortment.

Lemma 17 (Lemma A.4 of (Agrawal et al., 2017)). *For any epoch ℓ , if $r_i \in [0, 1]$ and $0 \leq v_i \leq \hat{v}_i$ hold for every $i \in [N]$ at the beginning of the ℓ -th iteration of the outer for-loop in Algorithm 2, we have that*

$$\left(1 + \sum_{i \in S_\ell} v_i\right) (R(S_\ell, \hat{\mathbf{v}}) - R(S_\ell, \mathbf{v})) \leq \sum_{i \in S_\ell} (\hat{v}_i - v_i).$$

As a corollary, we have the following lemma, which is an analog to Lemma 4.3 of (Agrawal et al., 2017).

Lemma 18. *Given that $r_i \in [0, 1]$ for every $i \in [N]$, for any epoch $\ell = 1, 2, 3, \dots$, with probability at least $\frac{13}{\ell}$ we have that*

$$\left(1 + \sum_{i \in S_\ell} v_i\right) (R(S_\ell, \hat{\mathbf{v}}) - R(S_\ell, \mathbf{v})) \lesssim \sqrt{\frac{v_i \log(\sqrt{N}\ell + 1)}{T_i}} + \frac{\log(\sqrt{N}\ell + 1)}{T_i}.$$

Proof. Combine Lemma 16 and Lemma 17. □

We will also use the following lemma which is proved in (Agrawal et al., 2017).

Lemma 19 (Lemma A.3 of (Agrawal et al., 2017)). *If $v_i \leq \hat{v}_i$ holds for every $i \in [N]$, then we have that $R(S^*, \hat{\mathbf{v}}) \geq R(S^*, \mathbf{v})$.*

Now we complete the proof of Theorem 2.

Proof of the regret upper bound in Theorem 2. Let $E^{(\ell)}$ be the length of epoch ℓ . That is, the number of time steps taken in epoch ℓ . Note that $E^{(\ell)}$ is a geometric random variable with mean $(1 + \sum_{i \in S_\ell} v_i)$. As a result,

$$\begin{aligned} \mathbb{E}[\text{Reg}_T] &= \mathbb{E} \left[\sum_{\ell=1}^L E^{(\ell)} (R(S^*, \mathbf{v}) - R(S_\ell, \mathbf{v})) \right] \\ &\leq \mathbb{E} \left[\sum_{\ell=1}^L E^{(\ell)} \left(R(S^*, \hat{\mathbf{v}}) - R(S_\ell, \mathbf{v}) + \frac{6}{\ell} \right) \right] \\ &\leq \mathbb{E} \left[\sum_{\ell=1}^L E^{(\ell)} \left(R(S_\ell, \hat{\mathbf{v}}) - R(S_\ell, \mathbf{v}) + \frac{6}{\ell} \right) \right] \\ &= \mathbb{E} \left[\sum_{\ell=1}^L \left(1 + \sum_{i \in S_\ell} v_i\right) \left(R(S^*, \hat{\mathbf{v}}) - R(S_\ell, \hat{\mathbf{v}}) + \frac{6}{\ell} \right) \right], \end{aligned}$$

where the first inequality is due to Lemma 19 and Lemma 16. Let $\Delta R^{(\ell)} \stackrel{\text{def}}{=} (1 + \sum_{i \in S_\ell} v_i) (R(S^*, \hat{\mathbf{v}}) - R(S_\ell, \hat{\mathbf{v}}) + 6/\ell)$ for shorthand. We use $T_i^{(\ell)}$ to denote the value of variable T_i at the beginning of epoch ℓ . By Lemma 18, we have

$$\mathbb{E}[\Delta R^{(\ell)}] \lesssim \frac{1}{\ell} \left(1 + \sum_{i \in S_\ell} v_i\right) + \mathbb{E} \left[\sum_{i \in S_\ell} \left(\sqrt{\frac{v_i \log(\sqrt{N}T + 1)}{T_i^{(\ell)}}} + \frac{\log(\sqrt{N}T + 1)}{T_i^{(\ell)}} \right) \right].$$

As a consequence,

$$\begin{aligned} \mathbb{E}[\text{Reg}_T] &\lesssim \sum_{\ell=1}^L \left(\frac{1}{\ell} \left(1 + \sum_{i \in S_\ell} v_i\right) + \mathbb{E} \left[\sum_{i \in S_\ell} \left(\sqrt{\frac{v_i \log(\sqrt{N}T + 1)}{T_i^{(\ell)}}} + \frac{\log(\sqrt{N}T + 1)}{T_i^{(\ell)}} \right) \right] \right) \\ &\lesssim N \log T + \sum_{\ell=1}^L \mathbb{E} \left[\sum_{i \in S_\ell} \left(\sqrt{\frac{v_i \log(\sqrt{N}T + 1)}{T_i^{(\ell)}}} + \frac{\log(\sqrt{N}T + 1)}{T_i^{(\ell)}} \right) \right] \\ &\lesssim N \log T + \mathbb{E} \left[N \log^2(\sqrt{N}T + 1) + \sum_{i \in [N]} \sqrt{v_i T_i^{(L)}} \log(\sqrt{N}T + 1) \right] \\ &\lesssim N \log^2(\sqrt{N}T + 1) + \sum_{i \in [N]} \sqrt{\mathbb{E}[v_i T_i^{(L)}]} \log(\sqrt{N}T + 1). \end{aligned} \tag{11}$$

Note that $\mathbb{E}[E_\ell] = 1 + \sum_{i \in S_\ell} v_i$. We have

$$\sum_{i \in [N]} v_i T_i^{(L)} = \sum_{\ell=1}^L \sum_{i \in S_\ell} v_i \leq \sum_{\ell=1}^L \mathbb{E}[E_\ell] \leq T.$$

As a result, by Jensen's inequality we get that

$$(11) \lesssim N \log^2(\sqrt{NT} + 1) + \sqrt{NT \log(\sqrt{NT} + 1)},$$

which concludes the proof. $\square$

B. Omitted proofs for the FH-DUCB algorithm in Section 3

B.1. Proof of Lemma 5

By Lemma 15, we have that $\Pr[\neg \mathcal{E}_{i, \tau_i}^{(1)}] \leq \frac{13}{NT^2}$. Via a union bound, we have that

$$\Pr[\neg \mathcal{E}^{(1)}] \leq \sum_{i, \tau_i} \Pr[\neg \mathcal{E}_{i, \tau_i}^{(1)}] \leq \frac{13}{T}.$$

Next we introduce the following concentration inequality for geometric random variables.

Lemma 20 (Theorem 1 and Proposition 1 of (Jin et al., 2019)). *For any m i.i.d. geometric random variables $x_1, \dots, x_m$ with parameter p , i.e., $\Pr[x_i = k] = p(1-p)^k$, we have*

$$\Pr \left[\sum_{i=1}^m x_i < \frac{m(1-p)}{2p} \right] \leq \exp \left(-m \cdot \frac{1-p}{8} \right).$$

Note that n_{i, τ_i} is the sum of $|\mathcal{T}(i, \tau_i)|$ independent geometric random variables with parameter $p = \frac{1}{1+v_i}$ (by Observation 1).

Substituting $v_i \geq \frac{1}{2} \sqrt{\frac{1}{NT}}$ and $m = |\mathcal{T}(i, \tau_i)| \geq \frac{T}{4Nv_i}$, we have $\frac{(1-p)}{2p} = \frac{v_i}{2}$ and

$$\begin{aligned} \Pr \left[n_{i, \tau_i} < \frac{1}{2} v_i \cdot |\mathcal{T}(i, \tau_i)| \right] &\leq \exp \left(-|\mathcal{T}(i, \tau_i)| \cdot \frac{1-p}{8} \right) \\ &\leq \exp \left(-\frac{T}{4Nv_i} \cdot \frac{1 - \frac{1}{1+v_i}}{8} \right) \\ &\leq \exp \left(-\frac{T}{64N} \right) \leq \frac{1}{NT^2}, \end{aligned}$$

where the last inequality holds for T such that $T \geq N^4$ and T greater than a sufficiently large universal constant. By a union bound, we have that

$$\Pr[\neg \mathcal{E}^{(2)}] \leq \frac{1}{T}.$$

Therefore, we have that

$$\Pr[\mathcal{E}] \geq 1 - \Pr[\neg \mathcal{E}^{(1)}] + \Pr[\neg \mathcal{E}^{(2)}] \geq 1 - \frac{14}{T},$$

proving the lemma.

B.2. Proof of Lemma 6

We first state the following lemma, showing that for any item and before stage τ_0 , the stage lengths quickly grows to T/N .

Lemma 21. *For each $i \in [N]$ and $\tau \leq \tau_0$, if τ is not the last stage for i , it holds that $|\mathcal{T}(i, \tau)| \geq (T/N)^{1-2^{-\tau+1}}$.*

Lemma 21 can be proved by combining the condition $\mathcal{P}(i, \tau)$ for $\tau < \tau_0$ and $\tau = \tau_0$ (also noting that $\hat{v}_{i, \tau} \leq 1$ for all τ) and the following fact (whose proof is via straightforward induction and omitted).

Fact 22. *For $M \geq 0$ and a sequence $a_0, a_1, a_2, \dots$ such that $a_i \geq 1 + \sqrt{M a_{i-1}}$ for all $i \geq 1$, we have that $a_\tau \geq M^{1-2^{-\tau+1}}$ for all $\tau \geq 1$.*

Now we are ready to prove Lemma 6.

Proof of Lemma 6. We have that $|\mathcal{T}(i, \tau_0)| \geq \frac{T}{2N}$ because of Lemma 21. We now prove that $v_i \geq \frac{1}{2}\sqrt{\frac{1}{NT}}$. This is because, suppose the contrary, for T such that $T \geq N^4$ and greater than a sufficiently large universal constant, conditioned on $\mathcal{E}^{(1)}$, we have that

$$\begin{aligned} \hat{v}_{i, \tau_0} &\leq v_i + \sqrt{\frac{144 \ln(\sqrt{NT}^2 + 1)}{T_i^{(\tau_0)}/v_i}} + \frac{144 \ln(\sqrt{NT}^2 + 1)}{T_i^{(\tau_0)}} \\ &\leq \frac{1}{2\sqrt{NT}} + O\left(\sqrt{\frac{\ln(\sqrt{NT}^2 + 1)}{\sqrt{T^3/N}}} + \frac{\ln(\sqrt{NT}^2 + 1)}{T}\right), \end{aligned}$$

which is at most $1/\sqrt{NT}$, contradicting to the condition $\mathcal{P}(i, \tau_0)$ and that τ_0 is not the last stage.

Moreover, for T such that $T \geq N^4$ and greater than a sufficiently large universal constant, when $\tau > \tau_0$, using $T_i^{(\tau)} \geq |\mathcal{T}(i, \tau_0)| \geq \frac{T}{2N}$, we have that

$$\hat{v}_{i, \tau} \leq v_i + \sqrt{\frac{144 v_i \ln(\sqrt{NT}^2 + 1)}{T_i^{(\tau)}}} + \frac{144 \ln(\sqrt{NT}^2 + 1)}{T_i^{(\tau)}} \leq 2v_i.$$

By the condition $\mathcal{P}(i, \tau)$, when $\tau > \tau_0$ and τ is not the last stage, we have that

$$|\mathcal{T}(i, \tau_i)| \geq 1 + \sqrt{\frac{T \cdot T_i^{(\tau_i)}}{N \cdot \hat{v}_{i, \tau_i}}} \geq 1 + \sqrt{\frac{T \cdot |\mathcal{T}(i, \tau_i - 1)|}{2N \cdot v_i}}.$$

Applying Fact 22, we prove the desired inequality of this lemma. $\square$

B.3. Proof of Lemma 7

Proof of Lemma 7. For the first stage, i.e., $\tau = 1$, since the number of epochs in this stage is at most $\sqrt{T/N}$, we have that $\sum_{\ell \in \mathcal{T}(i, 1)} (\hat{v}_{i, 1} - v_i) \leq \sqrt{T/N}$ for any item i . From now on, we only prove the lemma for $\tau \in [2, \tau_i(L)]$.

If $\tau \in [2, \tau_0]$, we have that $|\mathcal{T}(i, \tau)| \leq \sqrt{\frac{T \cdot T_i^{(\tau)}}{N}} + 1$. By $\mathcal{E}^{(1)}$, we upper bound $\sum_{\ell \in \mathcal{T}(i, \tau)} (\hat{v}_{i, \tau} - v_i)$ by the order of

$$\sqrt{\frac{T \cdot T_i^{(\tau)}}{N}} \cdot \left(\sqrt{\frac{v_i \ln(\sqrt{NT}^2 + 1)}{T_i^{(\tau)}}} + \frac{\ln(\sqrt{NT}^2 + 1)}{T_i^{(\tau)}} \right) \lesssim \sqrt{T \ln(\sqrt{NT}^2 + 1)/N},$$

where the inequality holds due to that $v_i \leq 1$ and $T_i^{(\tau)} \geq \sqrt{T/N}$ for any $\tau \in [2, \tau_0]$ (by Lemma 21).

When $\tau > \tau_0$, we prove the lemma by considering the following two cases. The first case is that $\hat{v}_{i, \tau_0} \leq 1/\sqrt{NT}$. In this case, we have that

$$\sum_{\ell \in \mathcal{T}(i, \tau)} (\hat{v}_{i, \tau} - v_i) \leq T \cdot \hat{v}_{i, \tau} \leq \sqrt{T/N}.$$

In the second case where $\hat{v}_{i, \tau_0} > 1/\sqrt{NT}$, by Lemma 6 it holds that $v_i \geq 1/(2\sqrt{NT})$. By $\mathcal{E}^{(1)}$, we have $\hat{v}_{i, \tau} \geq v_i$.

Therefore, $\hat{v}_{i, \tau} \geq 1/(2\sqrt{NT})$. Also note that $T_i^{(\tau)} \geq |\mathcal{T}(i, \tau_0)| \geq \frac{T}{2N}$ by Lemma 21, and $|\mathcal{T}(i, \tau)| \leq 1 + \sqrt{\frac{T \cdot T_i^{(\tau)}}{N \cdot \hat{v}_{i, \tau}}}$. Altogether, we have that $\sum_{\ell \in \mathcal{T}(i, \tau)} (\hat{v}_{i, \tau} - v_i)$ is upper bounded by a universal constant times

$$\sqrt{\frac{T \cdot T_i^{(\tau)}}{N \cdot \hat{v}_{i, \tau}}} \cdot \left(\sqrt{\frac{v_i \ln(\sqrt{NT}^2 + 1)}{T_i^{(\tau)}}} + \frac{\ln(\sqrt{NT}^2 + 1)}{T_i^{(\tau)}} \right) \lesssim \sqrt{\frac{T \ln(\sqrt{NT}^2 + 1)}{N}} + \frac{\sqrt{T} \ln(\sqrt{NT}^2 + 1)}{\sqrt{NT_i^{(\tau)} \hat{v}_{i, \tau}}},$$

which is $O(\sqrt{T \ln(\sqrt{NT}^2 + 1)/N})$ for $T \geq N^4$. $\square$

C. Bounding the number of item switches for Algorithm 2

Since an assortment switch may incur at most $2K$ item switches, Theorem 2 trivially implies that Algorithm 2 (AT-DUCB) incurs at most $O(KN \log T)$ item switches, which is upper bounded by $O(N^2 \log T)$ since $K = O(N)$. In the following theorem, we prove an improved upper bound on item switches for Algorithm 2.

Theorem 23. *For any input instance with N items, before any time T , the number of item switches of Algorithm 2 (AT-DUCB) satisfies that $\Psi_T^{(\text{item})} \lesssim N^{1.5} \log T$.*

The proof of Theorem 23 includes a novel analysis with the careful application of the Cauchy-Schwartz inequality, which will be presented immediately after this paragraph. However, we would like to first add a few remarks on the optimality of the presented analysis. Indeed, we do not know whether the upper bound proved in Theorem 23 can be improved, and leave the possibility of further improvement as an open question. Our preliminary research suggests that the number of the item switches of Algorithm 2 is closely related to the maximal number of planar K -sets (i.e., the number of subsets $P' \subseteq P$ where P is a given set of N points in a 2-dimensional plane, $P' = P \cap H$ for a half-space H). Very roughly, this relation is suggested by Lemma 24, where the optimal assortment $\arg \max_{S \subseteq [N], |S| \leq K} R(S, \mathbf{v})$ can be viewed as a planar K -set whether each item correspond to a 2-dimensional point $(-v_i, v_i r_i)$ and the half plane $H = \{(x, y) : y \geq r^* \cdot x + b\}$ for some parameter b . The continuous change of the the estimated optimal revenue r^* during the UCB algorithm may produce many half planes, and lead to the item change in the K -sets (assortments). Upper bounding the number of the K -sets would result in an upper bound for the number of the item switches. To our best knowledge, the best known upper bound for the number of planar K -sets is $O(NK^{1/3})$ (Dey, 1998), and the best known lower bound is $Ne^{\Omega(\sqrt{\log K})}$ (Tóth, 2001). For future work, it is very interesting to study whether these upper and lower bounds imply the bounds on the number of item switches of our Algorithm 2.

Now we dive into the proof of Theorem 23.

We first analyze the optimization process of $\arg \max_{S \subseteq [N], |S| \leq K} R(S, \mathbf{v})$ for any preference vector $\mathbf{v}$. Define $F(\mathbf{v}) \stackrel{\text{def}}{=} \max_{S \subseteq [N], |S| \leq K} R(S, \mathbf{v})$. The following lemma characterizes the optimal assortment S given the preference vector $\mathbf{v}$. Similar statements can also be found in, e.g., Section 2.1 of (Rusmevichientong et al., 2010).

Lemma 24. *For any preference value vector $\mathbf{v} \geq 0$, let $r^* = F(\mathbf{v})$. Define $g_i = v_i(r_i - r^*)$. Let σ be the minimal permutation of $[N]$ such that $g_{\sigma_i} \geq g_{\sigma_j}$ for all $1 \leq i < j \leq N$. (In other words, σ is the sorted index according to value g , with a deterministic tie-breaking rule). Then the optimal assortment S is given by $S = \{\sigma_i : 1 \leq i \leq K, g_{\sigma_i} > 0\}$.*

Proof. Let $S^* = \arg \max_{S \subseteq [N], |S| \leq K} R(S, \mathbf{v})$. Then we have

$$\frac{\sum_{i \in S^*} r_i v_i}{1 + \sum_{i \in S^*} v_i} = r^*,$$

which implies that

$$\sum_{i \in S^*} v_i(r_i - r^*) = \sum_{i \in S^*} g_i = r^*. \quad (12)$$

Now we prove that $S^* = \arg \max_{S \subseteq [N], |S| \leq K} (\sum_{i \in S} g_i)$. Suppose otherwise that there exists $S' \subseteq [N]$ with $|S'| \leq K$ such that $\sum_{i \in S'} g_i > \sum_{i \in S^*} g_i = r^*$. It follows that $\sum_{i \in S'} v_i(r_i - r^*) > r^*$. Therefore,

$$R(S', \mathbf{v}) = \frac{\sum_{i \in S'} v_i r_i}{1 + \sum_{i \in S'} v_i} > r^*,$$

which contradicts to the definition of S^* .

Now, note that σ is a permutation of $[N]$ such that g_{σ_i} is non-increasing according to i . We have that $\arg \max_{S \subseteq [N], |S| \leq K} (\sum_{i \in S} g_i) = \{\sigma_i : 1 \leq i \leq K, g_{\sigma_i} > 0\}$, which finishes the proof. $\square$

The next lemma shows that $F(\mathbf{v})$ is monotonically decreasing in $\mathbf{v}$.

Lemma 25. *Consider two vectors $\mathbf{v}$ and $\hat{\mathbf{v}}$. If $\hat{v}_i \geq v_i \geq 0$ for all $i \in [N]$, we have $F(\hat{\mathbf{v}}) \geq F(\mathbf{v})$.*

Proof. Let $S^* = \arg \max_{S \subseteq [N], |S| \leq K} R(S, \mathbf{v})$ and $r^* = R(S^*, \mathbf{v})$. Then we have $\sum_{i \in S^*} v_i (r_i - r^*) = r^*$. According to Lemma 24, $r_i - r^* > 0$ for all $i \in S^*$. Combining with the assumption that $\hat{v}_i \geq v_i, \forall i \in [N]$, we get $\sum_{i \in S^*} \hat{v}_i (r_i - r^*) \geq \sum_{i \in S^*} v_i (r_i - r^*) = r^*$. As a result,

$$R(S^*, \hat{\mathbf{v}}) = \frac{\sum_{i \in S^*} r_i \hat{v}_i}{1 + \sum_{i \in S^*} \hat{v}_i} \geq r^*.$$

Therefore, $F(\hat{\mathbf{v}}) = \max_{S \subseteq [N], |S| \leq K} R(S, \hat{\mathbf{v}}) \geq R(S^*, \hat{\mathbf{v}}) \geq r^* = F(\mathbf{v})$. $\square$

Let m be the total number of times that Line 8 of Algorithm 2 is executed, and let $\tau^{(1)} < \tau^{(2)} < \tau^{(3)} < \dots < \tau^{(m)}$ be the time steps that Line 8 of Algorithm 2 is executed. In other words, only in the time steps in $\{\tau^{(p)}\}_{p=0}^m$, the UCB value vector $\hat{\mathbf{v}}$ is updated (where for convenience, we set $\tau^{(0)} = 0$). Let $\hat{\mathbf{v}}^{(p)}$ be the UCB value after the update at time $\tau^{(p)}$, and for convenience we let $\hat{\mathbf{v}}^{(0)} = (1, 1, \dots, 1)$. Define $r^{(p)} = F(\hat{\mathbf{v}}^{(p)})$. Let $\rho_i^{(p)}$ be the rank of item i according to value $g_i^{(p)} \stackrel{\text{def}}{=} \hat{v}_i^{(p)} (r_i - r^{(p)})$ with the tie-breaking rule defined in Lemma 24. We then have the following lemma.

Lemma 26. Let $\delta_{i,j}^{(p)} \stackrel{\text{def}}{=} \mathbb{I}[\rho_i^{(p)} > \rho_j^{(p)}]$. For any two items $i, j \in [N]$, the number of times that the relative order of i, j changes is bounded by $c \log T$ for some universal constant c . That is,

$$\sum_{p=0}^{m-1} \mathbb{I}[\delta_{i,j}^{(p)} \neq \delta_{i,j}^{(p+1)}] \lesssim \log T.$$

As a corollary, we have that

$$\sum_{i,j \in [N]} \sum_{p=0}^{m-1} \mathbb{I}[\delta_{i,j}^{(p)} \neq \delta_{i,j}^{(p+1)}] \lesssim N^2 \log T.$$

Proof. Let $\mathcal{D}_i^{(p)}$ be the event that Line 8 is executed in Algorithm 2 for item i at time $\tau^{(p)}$. In the following we prove that

$$\sum_{p=0}^{m-1} \mathbb{I}[\delta_{i,j}^{(p)} \neq \delta_{i,j}^{(p+1)}] \leq 2 \sum_{p=0}^{m-1} \mathcal{D}_i^{(p)} + 2 \sum_{p=0}^{m-1} \mathcal{D}_j^{(p)}.$$

For a fixed pair of items i, j , let $\{\bar{p}_q\}_{q=1}^Q$ be the time steps that $\mathcal{D}_i^{(\bar{p}_q)}$ or $\mathcal{D}_j^{(\bar{p}_q)}$ occur. We only need to prove that

$$\sum_{p=\bar{p}_q}^{\bar{p}_{q+1}-1} \mathbb{I}[\delta_{i,j}^{(p)} \neq \delta_{i,j}^{(p+1)}] \leq 1$$

for all $q \in [Q]$.

Note that at time interval $[\bar{p}_q, \bar{p}_{q+1} - 1]$, $\bar{v}_i$ and $\bar{v}_j$ does not change. Therefore, $\delta_{i,j}^{(p)} = \mathbb{I}[\bar{v}_i (r_i - r^{(p)}) < \bar{v}_j (r_j - r^{(p)})]$. It is implied by Lemma 25 that $r^{(p)}$ is monotonically decreasing. As a result, $\sum_{p=\bar{p}_q}^{\bar{p}_{q+1}-1} \mathbb{I}[\delta_{i,j}^{(p)} \neq \delta_{i,j}^{(p+1)}] \leq 1$. $\square$

Now we are ready to prove Theorem 23.

Proof of Theorem 23. Let $K^{(p)} = \min \left\{ K, \left| \{i : g_i^{(p)} > 0\} \right| \right\}$. Note that since $r^{(p)}$ is non-increasing, $K^{(p)}$ is non-decreasing. Then we have, $S^{(\tau_p)} = \{i : \rho_i^{(p)} \leq K^{(p)}\}$. Let $\bar{S}^{(\tau_{p+1})} = \{i : \rho_i^{(p+1)} \leq K^{(p)}\}$. Then we have, $\bar{S}^{(\tau_{p+1})} \subseteq S^{(\tau_{p+1})}$ and $|S^{(\tau_{p+1})} \setminus \bar{S}^{(\tau_{p+1})}| = K^{(p+1)} - K^{(p)}$. It follows that

$$|S^{\tau_p} \oplus S^{\tau_{p+1}}| \leq |S^{\tau_p} \oplus \bar{S}^{\tau_{p+1}}| + K^{(p+1)} - K^{(p)}. \quad (13)$$

Let $x^{(p)} = |S^{\tau_p} \oplus \bar{S}^{\tau_{p+1}}|$. In the following we prove that

$$(x^{(p)}/2)^2 \leq \sum_{i,j \in [N]} \mathbb{I}[\delta_{i,j}^{(p)} \neq \delta_{i,j}^{(p+1)}]. \quad (14)$$

Note that $|S^{(\tau_p)}| = |\bar{S}^{(\tau_{p+1})}| = K^{(p)}$. Define $Z = S^{(\tau_p)} \setminus \bar{S}^{(\tau_{p+1})}$ and $Z' = \bar{S}^{(\tau_{p+1})} \setminus S^{(\tau_p)}$. Then we have that $x^{(p)} = 2|Z| = 2|Z'|$. Note that for all $i \in Z$, we have that $\rho_i^{(p)} \leq K^{(p)}$ and $\rho_i^{(p+1)} > K^{(p)}$. And for all $j \in Z'$, we have that $\rho_j^{(p)} > K^{(p)}$ and $\rho_j^{(p+1)} \leq K^{(p)}$. It follows that $\delta_{i,j}^{(p)} = 0, \delta_{i,j}^{(p+1)} = 1$ for all $i \in Z, j \in Z'$. Hence, we have that

$$\sum_{i,j \in [N]} \mathbb{I}[\delta_{i,j}^{(p)} \neq \delta_{i,j}^{(p+1)}] \geq |Z| \times |Z'| = (x^{(p)}/2)^2,$$

which establishes (14).

Combining (14) and Lemma 26, we have that $\sum_{p=1}^{m-1} (x^{(p)}/2)^2 \leq N^2 \log T$. By the deferred update rule in Algorithm 2, we have that $m \leq N(1 + \log T)$. Applying Cauchy-Schwarz inequality, we get that

$$\sum_{p=1}^{m-1} x^{(p)} \lesssim N^{1.5} \log T.$$

Therefore, by (13) we have that

$$\sum_{p=1}^{m-1} |S^{(\tau_p)} \oplus S^{(\tau_{p+1})}| \leq \sum_{p=1}^{m-1} (x^{(p)} + K^{(p+1)} - K^{(p)}) \lesssim N^{1.5} \log T. \quad (15)$$

Note that there is no assortment switch at time steps where $\hat{v}$ is not updated. Therefore (15) directly leads to Theorem 23. $\square$

D. Omitted proofs for the ESUCB algorithm in Section 4

D.1. Proof of Lemma 9

Proof of Lemma 9. We first prove the existence of θ^* . Note that the uniqueness follows directly from statements 1) and 2) in the lemma statement.

Proof of the existence of θ^* . Let $S^* = \arg \max_{S \subseteq [N]: |S| \leq K} R(S, v)$ and $\theta^* = R(S^*, v)$. We only need to prove that $G(\theta^*) = \theta^*$.

On the one hand, since $G(\theta) = R(S_\theta, v)$, we have $G(\theta^*) \leq \theta^*$ be the optimality of S^* . On the other hand, we will prove that $G(\theta^*) \geq \theta^*$. For the sake of contradiction, suppose $G(\theta^*) < \theta^*$. Then we have,

$$\frac{\sum_{i \in S_{\theta^*}} v_i r_i}{1 + \sum_{i \in S_{\theta^*}} v_i} = G(\theta^*) < \theta^*.$$

By algebraic manipulation we get $\sum_{i \in S_{\theta^*}} v_i (r_i - \theta^*) < \theta^*$. By the optimality of S_{θ^*} we have

$$\sum_{i \in S^*} v_i (r_i - \theta^*) \leq \sum_{i \in S_{\theta^*}} v_i (r_i - \theta^*) < \theta^*.$$

As a result, we have $R(S^*, v) = \frac{\sum_{i \in S^*} v_i r_i}{1 + \sum_{i \in S^*} v_i} < \theta^*$, which leads to contradiction.

Proof of statement 1). For the sake of contradiction, suppose $G(\theta) \leq \theta$. Then we have

$$\frac{\sum_{i \in S_\theta} r_i v_i}{1 + \sum_{i \in S_\theta} v_i} \leq \theta,$$

which means that $\sum_{i \in S_\theta} v_i (r_i - \theta) \leq \theta$. Note that $v_i \geq 0$ for all $i \in [N]$. By the optimality of S_θ , we get

$$\sum_{i \in S_{\theta^*}} v_i (r_i - \theta^*) \leq \sum_{i \in S_\theta} v_i (r_i - \theta) \leq \sum_{i \in S_\theta} v_i (r_i - \theta) \leq \theta < \theta^*.$$

By algebraic manipulation, we get $R(S_{\theta^*}, v) < \theta^*$, which leads to contradiction.

Proof of statement 2). By the optimality of S^* , we have $G(\theta) \leq G(\theta^*) = \theta^* < \theta$. $\square$

D.2. Proof of Lemma 11

Proof of Lemma 11. Observe that in the CHECK procedure, when b equals false, S_ℓ is evaluated by Line 9 and with respect to θ_r . When b is set to true, S_ℓ will always be evaluated by Line 6 with respect to θ_l . This switch happens for at most once. Therefore, we only need to show that for fixed any $\theta \in \{\theta_l, \theta_r\}$, and $S'_\ell = \arg \max_{S \subseteq [N], |S| \leq K} (\sum_{i \in S} \hat{v}_i(r_i - \theta))$, it holds that (assuming that there are L epochs)

$$\sum_{\ell=1}^{L-1} |S'_\ell \oplus S'_{\ell+1}| \lesssim N \log T. \quad (16)$$

Suppose that there are n_ℓ items whose UCB values are updated after the ℓ -th epoch. We claim that $|S_\ell \oplus S_{\ell+1}| \leq n_\ell$. This is simply because S_ℓ corresponds to the items $i \in [N]$ such that $\hat{v}_i(r_i - \theta)$ is positive and among the K largest ones (and thanks to the tie breaking rule). Therefore, any update to a single $\hat{v}_i$ will incur at most one item switch to S_ℓ , and n_ℓ updates will incur at most n_ℓ item switches. Now, (16) is established because $\sum_{\ell=1}^{L-1} |S'_\ell \oplus S'_{\ell+1}| \leq \sum_{\ell=1}^{L-1} n_\ell \lesssim N \log T$, where the second inequality is due to the deferred update rule for the UCB values. $\square$

D.3. Proof of Lemma 12

We now prove Lemma 12. For preparation, we first show that the UCB value $\hat{v}_i$ is valid throughout the execution of Algorithm 6.

Lemma 27. *For any invocation of CHECK($\theta_l, \theta_r, t_{\max}$), and for any epoch $\ell = 1, 2, 3, \dots$, during the algorithm, the following two statements hold throughout the execution,*

1. *With probability at least $1 - \frac{\delta}{4NT^2}$, $\hat{v}_i^{(\ell)} \geq v_i$ for any $i \in [N]$,*
2. *With probability at least $1 - \frac{\delta}{4NT^2}$, for any $i \in [N]$,*

$$\hat{v}_i^{(\ell)} - v_i \leq \sqrt{\frac{196v_i \log(NT/\delta)}{T_i^{(\ell)}}} + \frac{292 \log(NT/\delta)}{T_i^{(\ell)}}.$$

Proof. The proof is essentially the same as Lemma 16. $\square$

Let $\mathcal{H}$ be the event that the events described by Lemma 27 holds throughout the execution of Algorithm 6 for any ℓ and $i \in [N]$. We have that $\Pr[\mathcal{H}] \geq 1 - \frac{\delta}{4T}$.

Now we prove the following lemma.

Lemma 28. *For any fixed θ where $G(\theta) \geq \theta$, define $\hat{S}_\theta = \arg \max_{S: S \subseteq [N], |S| \leq K} (\sum_{i \in S} \hat{v}_i(r_i - \theta))$. Suppose $\hat{v}_i \geq v_i$ for all $i \in [N]$. We have that*

$$\left(1 + \sum_{i \in \hat{S}_\theta} v_i\right) (\theta - R(\hat{S}_\theta, \mathbf{v})) \leq \sum_{i \in \hat{S}_\theta} (\hat{v}_i - v_i).$$

Proof. Recall that $S_\theta = \arg \max_{S: S \subseteq [N], |S| \leq K} (\sum_{i \in S} v_i(r_i - \theta))$. We then have that

$$\begin{aligned} & \left(1 + \sum_{i \in \hat{S}_\theta} v_i\right) (\theta - R(\hat{S}_\theta, \mathbf{v})) \\ &= \left(1 + \sum_{i \in \hat{S}_\theta} v_i\right) \left(\theta - \frac{\sum_{i \in \hat{S}_\theta} r_i \hat{v}_i}{1 + \sum_{i \in \hat{S}_\theta} \hat{v}_i} + \frac{\sum_{i \in \hat{S}_\theta} r_i \hat{v}_i}{1 + \sum_{i \in \hat{S}_\theta} \hat{v}_i} - R(\hat{S}_\theta, \mathbf{v})\right) \end{aligned}$$

$$= \left(1 + \sum_{i \in \hat{S}_\theta} v_i\right) \left(\theta - \frac{\sum_{i \in \hat{S}_\theta} r_i \hat{v}_i}{1 + \sum_{i \in \hat{S}_\theta} \hat{v}_i}\right) + \sum_{i \in \hat{S}_\theta} r_i \left(\left(1 + \sum_{i \in \hat{S}_\theta} v_i\right) \frac{\hat{v}_i}{1 + \sum_{i \in \hat{S}_\theta} \hat{v}_i} - v_i\right). \quad (17)$$

Note that by assumption we have $\hat{v}_i \geq v_i$ for all $i \in [N]$. Therefore it holds that $1 + \sum_{i \in \hat{S}_\theta} \hat{v}_i \geq 1 + \sum_{i \in \hat{S}_\theta} v_i$. As a result,

$$\sum_{i \in \hat{S}_\theta} r_i \left(\left(1 + \sum_{i \in \hat{S}_\theta} v_i\right) \frac{\hat{v}_i}{1 + \sum_{i \in \hat{S}_\theta} \hat{v}_i} - v_i\right) \leq \sum_{i \in \hat{S}_\theta} r_i (\hat{v}_i - v_i) \leq \sum_{i \in \hat{S}_\theta} (\hat{v}_i - v_i). \quad (18)$$

On the other hand,

$$\left(1 + \sum_{i \in \hat{S}_\theta} v_i\right) \left(\theta - \frac{\sum_{i \in \hat{S}_\theta} r_i \hat{v}_i}{1 + \sum_{i \in \hat{S}_\theta} \hat{v}_i}\right) = \frac{1 + \sum_{i \in \hat{S}_\theta} v_i}{1 + \sum_{i \in \hat{S}_\theta} \hat{v}_i} \left(\theta - \sum_{i \in \hat{S}_\theta} \hat{v}_i (r_i - \theta)\right). \quad (19)$$

Note that by monotonicity (see Lemma 25) and our assumption (namely, $G(\theta) > \theta$),

$$\frac{\sum_{i \in \hat{S}_\theta} r_i \hat{v}_i}{1 + \sum_{i \in \hat{S}_\theta} \hat{v}_i} = R(\hat{S}_\theta, \hat{\mathbf{v}}) \geq R(S_\theta, \mathbf{v}) = G(\theta) \geq \theta.$$

By algebraic manipulation, we get that

$$\sum_{i \in \hat{S}_\theta} \hat{v}_i (r_i - \theta) \geq \theta. \quad (20)$$

Combining (19) and (20), we get that

$$\left(1 + \sum_{i \in \hat{S}_\theta} v_i\right) \left(\theta - \frac{\sum_{i \in \hat{S}_\theta} r_i \hat{v}_i}{1 + \sum_{i \in \hat{S}_\theta} \hat{v}_i}\right) \leq 0. \quad (21)$$

Plug in (18) and (21) into (17), we have that

$$\left(1 + \sum_{i \in \hat{S}_\theta} v_i\right) \left(\theta - R(\hat{S}_\theta, \mathbf{v})\right) \leq \sum_{i \in \hat{S}_\theta} (\hat{v}_i - v_i).$$

□

We will also need the following Azuma-Hoeffding inequality for martingales.

Theorem 29. Suppose $\{X_k : k = 0, 1, 2, 3, \dots\}$ is a martingale and $|X_k - X_{k-1}| \leq M$ almost surely for all k . Then for all positive integers n and all positive reals ϵ , it holds that

$$\Pr[X_n - X_0 \geq \epsilon] \leq \exp\left(-\frac{\epsilon^2}{2nM^2}\right).$$

Now we are ready to prove Lemma 12.

Proof of Lemma 12. We prove that each of the statements (a)–(c) holds with probability at least $1 - \delta/(4T)$, given that the UCB estimation of value $\mathbf{v}$ is valid (i.e., event $\mathcal{H}$). Then Lemma 12 holds by a union bound.

Proof of statement (a). Note that we only need to prove that if $G(\theta_r) \geq \theta_r$, then with probability at least $1 - \delta/(4T)$, $\text{CHECK}(\theta_l, \theta_r, t_{\max})$ returns false.

For simplicity, we use the superscript (ℓ) to denote the value of a variable in Algorithm 6 at the beginning of epoch ℓ . For example, $t^{(\ell)}$ denotes the time steps taken at the beginning of epoch ℓ . Now we prove that for large enough constants c_2 and c_3 , and any fixed L it holds that

$$\Pr \left[\sum_{\tau=1}^{t^{(L)}} \left(R(S_{\theta_r}^{(\tau)}, \mathbf{v}) - \theta_r \right) + (c_2 - 8) \sqrt{N t^{(L)} \log^3(NT/\delta)} \right]$$

$$+ c_3 N \log^3(NT/\delta) \geq 0 \wedge t^{(L)} \leq t_{\max} \Big] \leq 1 - \delta/(8T). \quad (22)$$

Let $\mathcal{J}_\ell$ be the filtration of random variables upto epoch ℓ . Let $S_\theta^{(\ell)} = \arg \max_{S: S \subseteq [N], |S| \leq K} \left(\sum_{i \in S} r_i (\hat{v}_i^{(\ell)} - \theta) \right)$. Then $S_{\theta_r}^{(\ell)}$ is $\mathcal{J}_{\ell-1}$ measurable. For simplicity we define $S_\ell = S_{\theta_r}^{(\ell)}$. As a result,

$$\sum_{\tau=1}^{t^{(L)}} \left(\theta_r^{(\ell)} - R(S_\ell, \mathbf{v}) \right) = \sum_{\ell=1}^L \left(t^{(\ell+1)} - t^{(\ell)} \right) \left(\theta_r^{(\ell)} - R(S_\ell, \mathbf{v}) \right).$$

Note that $(t^{(\ell+1)} - t^{(\ell)})$ follows geometric distribution given $\mathcal{J}_{\ell-1}$ with mean $(1 + \sum_{i \in S_\ell} v_i)$. Therefore with probability at least $1 - \delta/(16T^3)$ we have $t^{(\ell+1)} - t^{(\ell)} \leq 24 \log(T/\delta) (1 + \sum_{i \in S_\ell} v_i)$. Consequently, with probability at least $1 - \delta/(16T^2)$,

$$\sum_{\ell=1}^L \left(t^{(\ell+1)} - t^{(\ell)} \right) \left(\theta_r^{(\ell)} - R(S_\ell, \mathbf{v}) \right) \leq \sum_{\ell=1}^L 24 \log(T/\delta) \left(1 + \sum_{i \in S_\ell} v_i \right) \left(\theta_r^{(\ell)} - R(S_\ell, \mathbf{v}) \right)_+,$$

where the $(x)_+$ notation denotes $\max\{x, 0\}$. Under event $\mathcal{H}$, it follows from Lemma 28 that

$$\begin{aligned} & \sum_{\ell=1}^L 24 \log(T/\delta) \left(1 + \sum_{i \in S_\ell} v_i \right) \left(\theta_r^{(\ell)} - R(S_\ell, \mathbf{v}) \right)_+ \\ & \leq 24 \log(T/\delta) \sum_{\ell=1}^L \sum_{i \in S_\ell} (\hat{v}_i^{(\ell)} - v_i) \\ & \leq 24 \log(T/\delta) \sum_{\ell=1}^L \sum_{i \in S_\ell} \left(\sqrt{\frac{196 v_i \log(NT/\delta)}{T_i^{(\ell)}}} + \frac{292 \log(NT/\delta)}{T_i^{(\ell)}} \right) \\ & \leq 24 \log(T/\delta) \left(\sum_{i \in [N]} \sqrt{392 T_i^{(L)} v_i \log(NT/\delta)} + 876 N \log^2(NT/\delta) \right). \end{aligned}$$

Recall that in Algorithm 6 we define

$$\bar{v}_i^{(L)} = \sum_{\ell=1}^L \Delta_i^{(\ell)} / T_i^{(L)}.$$

Since $\Delta_i^{(\ell)}$ follows geometric distribution, by concentration inequality (namely, Theorem 5 of (Agrawal et al., 2017))

$$\Pr \left[\bar{v}_i^{(L)} < \frac{1}{2} v_i \right] \leq \exp \left(-T_i^{(L)} v_i / 48 \right).$$

Therefore we get with probability at least $1 - \delta/(16T^2)$, for any $i \in [N]$,

$$T_i^{(L)} v_i \leq \max \left\{ 2 \bar{n}_i^{(L)}, 144 \log(NT/\delta) \right\}.$$

Since every time step at most one item can be chosen, we get $\sum_{i \in [N]} \bar{n}_i^{(L)} \leq t^{(L)}$. Consequently,

$$\begin{aligned} & \sum_{i \in [N]} \sqrt{T_i^{(L)} v_i \log(NT/\delta)} \\ & \leq \sum_{i \in [N]} \sqrt{2 \bar{n}_i^{(L)} \log(NT/\delta)} + \sqrt{144 N \log(NT/\delta)} \\ & \leq \sqrt{2 N t^{(L)} \log(NT/\delta)} + \sqrt{144 N \log(NT/\delta)}. \end{aligned}$$

Putting everything together, we prove Eq. (22) with $c_2 = 688$ and $c_3 = 21036$. Note that

$$r_{\text{CHECK}}^{(\tau)} - R(S_{\theta_r}^{(\tau)}, \mathbf{v})$$

is a martingale sequence for $\tau = 0, 1, 2, 3, \dots$. By Theorem 29 (using $M = 2$), with probability $1 - \delta/(8T^2)$, we have that

$$\sum_{\tau=1}^{t^{(L)}} \left(r_{\text{CHECK}}^{(\tau)} - \theta_r \right) \geq \sum_{\tau=1}^{t^{(L)}} \left(R(S_{\theta_r}^{(\tau)}) - \theta_r \right) - 8\sqrt{t_{\max} \log(T/\delta)}.$$

Combining with (22), we get with probability at least $1 - \delta/(4T)$, it holds that

$$\sum_{\tau=1}^{t^{(L)}} \left(r_{\text{CHECK}}^{(\tau)} - \theta_r \right) + c_2 \sqrt{N t_{\max} \log^3(NT/\delta)} + c_3 N \log^3(NT/\delta) \geq 0,$$

in any of the epoch L such that $t^{(L)} \leq t_{\max}$. Consequently, with probability at most $1 - \delta/(4T)$, the event that $\hat{\rho}^{(\ell)} < \theta$ never occur, which means that $\text{CHECK}(\theta_l, \theta_r, t_{\max})$ returns false.

Proof of Statement (b). Note that when the Algorithm returns false, the **if**-condition in Line 4 is always false. By the optimality, we have $\theta^* = G(\theta^*) \geq R(S_{\theta_r}^{(\tau)}, \mathbf{v})$ for any $1 \leq \tau \leq t_{\max}$. Note that $(r_{\text{CHECK}}^{(\tau)} - R(S_{\theta_r}^{(\tau)}, \mathbf{v}))$ is a martingale sequence. Again, invoking Theorem 29, we have that with probability at least $1 - \delta/(8T)$, it holds that

$$\begin{aligned} \theta^* &\geq \frac{1}{t^{(L)}} \sum_{\tau=1}^{t^{(L)}} R(S_{\theta_r}^{(\tau)}, \mathbf{v}) \\ &\geq \frac{1}{t^{(L)}} \sum_{\tau=1}^{t^{(L)}} r_{\text{CHECK}}^{(\tau)} - 8\sqrt{\log(T/\delta)/t^{(L)}} \quad (\text{Martingale concentration}) \\ &\geq \theta_r - \frac{1}{t^{(L)}} \left(c_2 \sqrt{N t^{(L)} \log^3(NT/\delta)} + c_3 N \log^3(NT/\delta) + 8\sqrt{t^{(L)} \log(T/\delta)} \right). \quad (\text{By the if statement in Line 4}) \end{aligned}$$

Note that the time steps taken by the last epoch is bounded by $24(N+1) \log(T/\delta)$ with probability $1 - \delta/(8T)$. As a result, $(c_2 + 8)/t^{(L)} \leq 2/t_{\max}$ and $c_3/t^{(L)} \leq 2/t_{\max}$. Consequently,

$$\begin{aligned} \theta_r - \frac{1}{t^{(L)}} \left(c_2 \sqrt{N t^{(L)} \log^3(NT/\delta)} + c_3 N \log^3(NT/\delta) + 8\sqrt{t^{(L)} \log(T/\delta)} \right) \\ \geq \theta_r - \frac{2}{t_{\max}} \left(c_2 \sqrt{N t_{\max} \log^3(NT/\delta)} + c_3 N \log^3(NT/\delta) \right), \end{aligned}$$

which proves statement (b).

Proof of statement (c). Let $\bar{t}$ be the time step when the **if** condition is first violated (and let $\bar{t} = t_{\max}$ if the condition holds throughout an execution). We first show that

$$\mathbb{E} \left[\sum_{\tau=1}^{\bar{t}} \left(\theta_l - R(S_{\theta_r}^{(\tau)}, \mathbf{v}) \right) \right] \lesssim \sqrt{N t_{\max} \log^3(NT/\delta)} + N \log^3(NT/\delta) \quad (23)$$

holds with high probability. Note that the **if** condition is false for all $t \leq \bar{t}$. Therefore, $\bar{t}\theta_r \leq \sum_{\tau=1}^{\bar{t}} r_{\text{CHECK}}^{(\tau)} + c_2 \sqrt{N t_{\max} \log^3(NT/\delta)} + c_3 N \log^3(NT/\delta)$. Applying Theorem 29, we have that with probability at least $1 - \delta/(8T)$, it holds that $\sum_{\tau=1}^{\bar{t}} r_{\text{CHECK}}^{(\tau)} - \sum_{\tau=1}^{\bar{t}} R(S_{\theta_r}^{(\tau)}, \mathbf{v}) \lesssim \sqrt{t_{\max} \log(T/\delta)}$. Note that $\theta_l \leq \theta_r$, we get (23) with probability at least $1 - \delta/(8T)$.

Then we show that given $\bar{t}$,

$$(t_{\max} - \bar{t})\theta_l - \mathbb{E} \left[\sum_{t=\bar{t}+1}^{t_{\max}} r_{\text{CHECK}}^{(t)} \right] \lesssim \sqrt{N t_{\max} \log^3(NT/\delta)} + N \log^3(NT/\delta), \quad (24)$$

holds with high probability. Note that by assumption we have $\theta_l \leq \theta^*$. It follows from Lemma 9 that $G(\theta_l) \geq \theta_l$. By the same argument in the proof of statement (a), we have with probability $1 - \delta/(8T)$, it holds that

$$\mathbb{E} \left[\sum_{\tau=\bar{\ell}+1}^{t_{\max}} \left(R(S_{\theta_l}^{(\tau)}, \mathbf{v}) - \theta_l \right) \right] + c_2 \sqrt{N t_{\max} \log^3(NT/\delta)} + c_3 N \log^3(NT/\delta) \geq 0,$$

which implies (24).

Combining (23) and (24) with a union bound, we prove statement (c). $\square$

E. Lower bound proofs

E.1. Proof of Theorem 3

To prove Theorem 3, we first introduce the following more general theorem relating the expected regret with the number of assortment switches.

Theorem 30. *For any $N \geq 2$, $T_0 \geq 4$, fix a function $g(T)$ such that $g(T) \in \left[\frac{3}{\log_2 T}, \frac{1}{2}\right]$ and is non-increasing for $T \geq T_0$. For any anytime algorithm, there exists an N -item assortment instance $\mathcal{I}$ with time horizon $T \in [T_0, T_0^2]$ such that either the expected regret of the algorithm for instant $\mathcal{I}$ is*

$$\mathbb{E} [\text{Reg}_T] \geq \frac{1}{7525} \cdot \sqrt{NT}^{\frac{1}{2} + \frac{g(T)}{3}}$$

or the expected assortment switching cost before time T is

$$\mathbb{E} [\Psi_T^{(\text{asst})}] = \mathbb{E} \left[\sum_{t=1}^{T-1} \mathbb{I}[S_t \neq S_{t+1}] \right] \geq \frac{N}{8 \log_2(1 + g(T))}.$$

Before proving Theorem 30, we first prove Theorem 3 using Theorem 30.

Proof of Theorem 3. We set $g(T) = \frac{3C \ln \ln(NT)}{\ln T}$. It is easy to verify that the derivative of $\frac{\ln \ln(NT)}{\ln T}$ is

$$\frac{\ln T - \ln(NT) \cdot \ln \ln(NT)}{T \ln^2 T \ln(NT)} < 0$$

for all $N \geq 2$ and $T \geq 2$. Therefore $g(T)$ is non-increasing for all $N \geq 2$ and $T \geq 2$. Also note that for $T \geq N$ and T greater than a sufficiently large constant that only depends on C , we have that $g(T) \in \left[\frac{3}{\log_2 T}, \frac{1}{2}\right]$.

Now invoke Theorem 30, and we have that there exists an N -item assortment instance $\mathcal{I}$ with time horizon $T \in [T_0, T_0^2]$ such that either $\mathbb{E} [\text{Reg}_T] \geq \frac{1}{7525} \cdot \sqrt{NT}(\ln(NT))^C$ or

$$\mathbb{E} [\Psi_T^{(\text{asst})}] \geq \Omega \left(\frac{N}{g(T)} \right) = \Omega \left(\frac{N \log T}{C \log \log(NT)} \right),$$

proving Theorem 3. $\square$

Proof of Theorem 30. Suppose that the expected number of assortment switches by the given policy for any input instance is at most $\frac{N}{8 \log_2(1+g(T))}$ for any time horizon T , we will prove the theorem by showing that there exists an instance with time horizon $T \in [T_0, T_0^2]$ such that the expected regret is at least $\frac{1}{7525} \cdot T^{\frac{1}{2} + \frac{g(T)}{3}}$.

Consider the assortment instance $\mathcal{I} = (\mathbf{v}, \mathbf{r})$, where $v_i = \frac{1}{2}$ and $r_i = 1$ for any $i \in [N]$. We will let the capacity constraint be $K = 1$ for all assortment instances considered in this proof. By the assumption of the algorithm, the expected number of assortment switches given input instance $\mathcal{I}$ is at most $\frac{N}{8 \log_2(1+g(T_0^2))}$. Thus, there exists T_1 such that $T_1^{1+g(T_0^2)} \in [T_0, T_0^2]$ and the expected number of assortment switches in time interval $[T_1, T_1^{1+g(T_0^2)}]$ is at most $\frac{N}{8}$. Otherwise, there are $\frac{1}{\log_2(1+g(T_0^2))}$

such disjoint intervals in range $[T_0, T_0^2]$ and the expected number of assortment switches is at least $\frac{N}{8 \log_2(1+g(T_0^2))}$, violating the assumption. Let

$$\mathcal{F}_1^{(i)} = \{\text{item } i \text{ is not offered in time interval } [T_1, T_1^{1+g(T_0^2)}] \text{ given instance } \mathcal{I}\}.$$

Note that $\sum_i \Pr_{\mathcal{I}}[\neg \mathcal{F}_1^{(i)}] \leq \frac{N}{8} + 1 \leq \frac{5N}{8}$ for any $N \geq 2$, because the expected number of items get offered in time interval $[T_1, T_1^{1+g(T_0^2)}]$ is at most the expected number of assortment switches plus 1. Therefore, there must exist a set of items $I \subseteq [N]$ such that $|I| \geq \frac{N}{4}$ and for any item $i \in I$, $\Pr_{\mathcal{I}}[\neg \mathcal{F}_1^{(i)}] \leq \frac{5}{6}$. Let

$$\mathcal{F}_2^{(i)} = \{\text{the number of times that item } i \text{ is offered in } [1, T_1] \text{ given instance } \mathcal{I} \text{ is at most } \frac{48T_1}{N}\}.$$

Note that T_1 is at least the expected number of times an item $i \in I$ is chosen between $[1, T_1]$, which implies $T_1 \geq \frac{48T_1}{N} \cdot \sum_{i \in I} \Pr_{\mathcal{I}}[\neg \mathcal{F}_2^{(i)}]$. Thus there exists $k \in I$ such that $\Pr_{\mathcal{I}}[\neg \mathcal{F}_2^{(k)}] \leq \frac{1}{12}$ since $|I| \geq \frac{N}{4}$. Let $\mathcal{F}^{(k)} = \mathcal{F}_1^{(k)} \cap \mathcal{F}_2^{(k)}$, we have

$$\Pr_{\mathcal{I}}[\mathcal{F}^{(k)}] \geq 1 - \Pr_{\mathcal{I}}[\neg \mathcal{F}_1^{(k)}] - \Pr_{\mathcal{I}}[\neg \mathcal{F}_2^{(k)}] \geq \frac{1}{12}. \quad (25)$$

Now we consider the assortment instance $\mathcal{I}^{(k)} = (\mathbf{v}^{(k)}, \mathbf{r})$ where $v_k^{(k)} = \frac{1}{2} + \frac{1}{16} \sqrt{\frac{N}{24T_1}}$ and $v_j^{(k)} = \frac{1}{2}$ for $j \neq k$. We will be interested in the regret of the algorithm at time horizon $T_1^{1+g(T_0^2)}$. First, we show that with high probability, no algorithm can distinguish instance $\mathcal{I}$ and $\mathcal{I}^{(k)}$ at time T_1 with high probability. Formally, we have the following lemma, the proof of which is provided at the end of this section.

Lemma 31. *We have that*

$$\left| \Pr_{\mathcal{I}}[\mathcal{F}^{(k)}] - \Pr_{\mathcal{I}^{(k)}}[\mathcal{F}^{(k)}] \right| \leq \frac{1}{24},$$

where $\Pr_{\mathcal{I}}[\cdot]$ uses the probability distribution when running the policy using input instance $\mathcal{I}$.

Combining Lemma 31 with inequality (25), we have

$$\Pr_{\mathcal{I}^{(k)}}[\mathcal{F}^{(k)}] \geq \frac{1}{24}.$$

Now, we lower bound the expected regret of the algorithm for instance $\mathcal{I}^{(k)}$ at time horizon $T_1^{1+g(T_0^2)}$ as

$$\begin{aligned} \mathbb{E}_{\mathcal{I}^{(k)}} \left[\text{Reg}_{T_1^{1+g(T_0^2)}} \right] &\geq \mathbb{E}_{\mathcal{I}^{(k)}} \left[\text{Reg}_{T_1^{1+g(T_0^2)}} \mid \mathcal{F}^{(k)} \right] \cdot \Pr_{\mathcal{I}^{(k)}}[\mathcal{F}^{(k)}] \\ &\geq (T_1^{1+g(T_0^2)} - T_1) \cdot \frac{\frac{1}{16} \sqrt{\frac{N}{24T_1}}}{\frac{3}{2} + \frac{1}{16} \sqrt{\frac{N}{24T_1}}} \cdot \frac{1}{24} \\ &\geq \frac{1}{7525} \cdot \sqrt{N} T_1^{\frac{1}{2}+g(T_0^2)} \geq \frac{1}{7525} \cdot \sqrt{N} T_1^{(1+g(T_0^2))(\frac{1}{2}+\frac{g(T_0^2)}{3})}, \end{aligned}$$

for any $g(T_0^2) \in \left[\frac{3}{\log_2 T_0^2}, \frac{1}{2} \right]$. The third inequality holds because $\frac{3}{2} + \frac{1}{16} \sqrt{\frac{N}{24T_1}} \leq 2$ and $T_1^{1+g(T_0^2)} \geq T_0$, and hence for

$g(T_0^2) \geq \frac{3}{\log_2 T_0^2}$, we have $T_1^{1+g(T_0^2)} \geq T_1 \cdot T_0^{\frac{g(T_0^2)}{1+g(T_0^2)}} \geq 2T_1$. Let $T = T_1^{1+g(T_0^2)} \in [T_0, T_0^2]$. Since by assumption $g(\cdot)$ is a non-increasing function when $T \geq T_0$, we have that $g(T) \geq g(T_0^2)$, therefore

$$\mathbb{E}[\text{Reg}_T] \geq \frac{1}{7525} \cdot T^{\frac{1}{2}+\frac{g(T)}{3}}. \quad \square$$

Finally we need to prove Lemma 31. First we introduce the following theorem on bounding the difference of the probability for a certain event.

Theorem 32 ((Pinsker, 1964)). For any probability distribution P, Q on measurable space (X, Σ) , for any event $\mathcal{F} \in \Sigma$, we have

$$|P(\mathcal{F}) - Q(\mathcal{F})| \leq \sqrt{\frac{1}{2} \text{KL}(P||Q)},$$

where $\text{KL}(P||Q)$ is the KL-divergence between distribution P and Q .

Lemma 33. The KL divergence between two Bernoulli distributions with $p_1 = \frac{1}{3} + \Delta$ and $p_2 = \frac{1}{3}$ is

$$\text{KL}(p_1, p_2) \leq \frac{9\Delta^2}{2}$$

Proof. The KL-divergence between two Bernoulli distributions with parameters p_1, p_2 is

$$\text{KL}(p_1, p_2) = p_1 \ln \frac{p_1}{p_2} + (1 - p_1) \ln \frac{1 - p_1}{1 - p_2}$$

Substituting $p_1 = \frac{1}{3} + \Delta$ and $p_2 = \frac{1}{3}$, we have

$$\text{KL}(p_1, p_2) = \left(\frac{1}{3} + \Delta\right) \ln(1 + 3\Delta) + \left(\frac{2}{3} - \Delta\right) \ln\left(1 - \frac{3\Delta}{2}\right) \leq \frac{9\Delta^2}{2}$$

where the last inequality holds by $\ln(1 + x) \leq x$. $\square$

Proof of Lemma 31. Note that in our construction, the choice distribution at each time t is a Bernoulli distribution. More specifically, under instance $\mathcal{I}$, when item k is offered to the customer, the probability she chooses to purchase item k is $p_2 = \frac{\frac{1}{2}}{1 + \frac{1}{2}} = \frac{1}{3}$, while under instance $\mathcal{I}^{(k)}$, when item k is offered to the customer, the probability she chooses to purchase item k is

$$p_1 = \frac{\frac{1}{2} + \frac{1}{16} \sqrt{\frac{N}{24T_1}}}{\frac{3}{2} + \frac{1}{16} \sqrt{\frac{N}{24T_1}}} = \frac{1}{3} + \frac{\frac{1}{16} \sqrt{\frac{N}{24T_1}}}{\frac{3}{2} + \frac{1}{16} \sqrt{\frac{N}{24T_1}}} \leq \frac{1}{3} + \frac{1}{24} \sqrt{\frac{N}{24T_1}}. \quad (26)$$

In event $\mathcal{F}^{(k)}$, the number of times item k is offered is at most $\frac{48T_1}{N}$. The total information available to the algorithm is the set of choice distributions observed for item k since the choice distributions for other items are the same. Therefore, combining Theorem 32, Lemma 33 and inequality (26), we have

$$\left| \Pr_{\mathcal{I}}[\mathcal{F}^{(k)}] - \Pr_{\mathcal{I}^{(k)}}[\mathcal{F}^{(k)}] \right| \leq \sqrt{\frac{1}{2} \cdot \frac{48T_1}{N} \cdot \text{KL}(p_1, p_2)} \leq \frac{1}{24}. \quad \square$$

E.2. Proof of Theorem 8

The proof of Theorem 8 is similar to that of Theorem 3 except for that we divide the time periods with a different scheme. It suffices to prove the following theorem in order to establish Theorem 8.

Theorem 34. For any $N \geq 2$, $T \geq 4$, and $M \leq \log_2 \log_2 T$, we have that for any algorithm such that the expected number assortment switches before time horizon T is $\mathbb{E}[\Psi_T^{(\text{asst})}] \leq \frac{NM}{8}$, there exists an N -item assortment instance $\mathcal{I}$ such that the expected regret of the algorithm for instance $\mathcal{I}$ at time horizon T is

$$\mathbb{E}[\text{Reg}_T] \geq \frac{1}{7525} \cdot \sqrt{NT^{\frac{1}{2(1-2^{-M})}}}.$$

Before proving Theorem 34, we first prove Theorem 8 using Theorem 34.

Proof of Theorem 8. We set $M = \lfloor \log_2(\frac{\log_2 T}{2C \log_2 \ln(NT)}) \rfloor$. It is easy to verify that M is at most $\log_2 \log_2 T$ for T larger than a universal constant that depends on C . Now invoke Theorem 34, and we have that for any algorithm, there exists an N -item assortment instance $\mathcal{I}$ such that either $\mathbb{E}[\text{Reg}_T] \geq \frac{1}{7525} \cdot \sqrt{NT(\ln(NT))^C}$ or

$$\mathbb{E}[\Psi_T^{(\text{asst})}] = \Omega\left(\frac{NM}{8}\right) = \Omega(N \log \log T),$$

proving Theorem 8. $\square$

Proof of Theorem 34. Suppose that the expected number of assortment switches by the given policy for any input instance is at most $\frac{NM}{8}$ before time horizon T , we will prove the theorem by showing that there exists an instance such that the expected regret incurred by the algorithm is at least $\frac{1}{7525} \cdot \sqrt{NT}^{\frac{1}{2(1-2^{-M})}}$.

Consider the assortment instance $\mathcal{I} = (\mathbf{v}, \mathbf{r})$, where $v_i = \frac{1}{2}$ and $r_i = 1$ for any $i \in [N]$. We will let the capacity constraint be $K = 1$ for all assortment instances considered in this proof. By the assumption of the algorithm, the expected number of assortment switches given input instance $\mathcal{I}$ is at most $\frac{NM}{8}$. For any $j \leq M$, we define

$$T_{(j)} = T^{\frac{1-2^{-j}}{1-2^{-M}}}.$$

By definition, we have that $T_{(M)} = T$. Therefore, there exists j such that $0 \leq j \leq M-1$ and the expected number of assortment switches in time interval $[T_{(j)}, T_{(j+1)}]$ is at most $\frac{N}{8}$ since there are M such disjoint intervals in range $[1, T]$. Let

$$\mathcal{G}_1^{(i)} = \{\text{item } i \text{ is not offered in time interval } [T_{(j)}, T_{(j+1)}] \text{ given instance } \mathcal{I}\}.$$

Note that $\sum_i \Pr_{\mathcal{I}}[\neg \mathcal{G}_1^{(i)}] \leq \frac{N}{8} + 1 \leq \frac{5N}{8}$ for any $N \geq 2$, because the expected number of items get offered during time interval $[T_{(j)}, T_{(j+1)}]$ is at most the expected number of assortment switches plus 1. Therefore, by an averaging argument, we have that there exists a set of items $I \subseteq [N]$ such that $|I| \geq \frac{N}{4}$ and for any item $i \in I$, $\Pr_{\mathcal{I}}[\neg \mathcal{G}_1^{(i)}] \leq \frac{5}{6}$. Define the following event

$$\mathcal{G}_2^{(i)} = \{\text{the number of times that item } i \text{ is offered in } [1, T_{(j)}] \text{ given instance } \mathcal{I} \text{ is at most } \frac{48T_{(j)}}{N}\}.$$

Note that T_1 is at least the expected number of times an item $i \in I$ is chosen between $[1, T_1]$, which implies $T_{(j)} \geq \frac{48T_{(j)}}{N} \cdot \sum_{i \in I} \Pr_{\mathcal{I}}[\neg \mathcal{G}_2^{(i)}]$. Thus there exists $k \in I$ such that $\Pr_{\mathcal{I}}[\neg \mathcal{G}_2^{(k)}] \leq \frac{1}{12}$ since $|I| \geq \frac{N}{4}$. Let $\mathcal{G}^{(k)} = \mathcal{G}_1^{(k)} \cap \mathcal{G}_2^{(k)}$, we have that

$$\Pr_{\mathcal{I}}[\mathcal{G}^{(k)}] \geq 1 - \Pr_{\mathcal{I}}[\neg \mathcal{G}_1^{(k)}] - \Pr_{\mathcal{I}}[\neg \mathcal{G}_2^{(k)}] \geq \frac{1}{12}. \quad (27)$$

Now we consider the assortment instance $\mathcal{I}^{(k)} = (\mathbf{v}^{(k)}, \mathbf{r})$ where $v_k^{(k)} = \frac{1}{2} + \frac{1}{16} \sqrt{\frac{N}{24T_{(j)}}}$ and $v_j^{(k)} = \frac{1}{2}$ for $j \neq k$. Using the same proof of Lemma 31, we have that

$$\left| \Pr_{\mathcal{I}}[\mathcal{G}^{(k)}] - \Pr_{\mathcal{I}^{(k)}}[\mathcal{G}^{(k)}] \right| \leq \frac{1}{24},$$

and combining it with inequality (27), we have that

$$\Pr_{\mathcal{I}^{(k)}}[\mathcal{G}^{(k)}] \geq \frac{1}{24}.$$

Now, we lower bound the expected regret of the algorithm for instance $\mathcal{I}^{(k)}$ as

$$\begin{aligned} \mathbb{E}_{\mathcal{I}^{(k)}}[\text{Reg}_T] &\geq \mathbb{E}_{\mathcal{I}^{(k)}}[\text{Reg}_T \mid \mathcal{G}^{(k)}] \cdot \Pr_{\mathcal{I}^{(k)}}[\mathcal{G}^{(k)}] \\ &\geq (T_{(j+1)} - T_{(j)}) \cdot \frac{\frac{1}{16} \sqrt{\frac{N}{24T_{(j)}}}}{\frac{3}{2} + \frac{1}{16} \sqrt{\frac{N}{24T_{(j)}}}} \cdot \frac{1}{24} \\ &\geq \frac{1}{7525} \cdot T_{(j+1)} \cdot \sqrt{\frac{N}{T_{(j)}}} \geq \frac{1}{7525} \cdot \sqrt{NT}^{\frac{1}{2(1-2^{-M})}}, \end{aligned}$$

The third inequality holds because $\frac{3}{2} + \frac{1}{16} \sqrt{\frac{N}{24T_{(j)}}} \leq 2$ and for $j \leq M-1$, $M \leq \log_2 \log_2 T$, we have that

$$T_{(j+1)} = T^{\frac{1-2^{-j-1}}{1-2^{-M}}} \geq T^{\frac{1-2^{-j}}{1-2^{-M}}} \cdot T^{\frac{2^{-j}-1}{1-2^{-M}}} \geq T^{\frac{1-2^{-j}}{1-2^{-M}}} \cdot T^{\frac{2^{-M}}{1-2^{-M}}} \geq 2T^{\frac{1-2^{-j}}{1-2^{-M}}} = 2T_{(j)}.$$

□

F. $N \log T$ item switch bound for ESUCB

In this section we show that a modification of ESUCB algorithm achieves $O(N \log T)$ item switches.

The modification is to use variables T_i and n_i without initializing in each $\text{CHECK}(\theta_l, \theta_r, t_{\max})$ sub-routine. That is, move the $T_i \leftarrow 0, n_i \leftarrow 0$ statement to the initialize phase of Algorithm 5. Note that n_i/T_i is still an unbiased estimation of v_i , and only concentrates better. As a result, the regret analysis applies directly.

Regarding the number of item switches, since the value of T_i and n_i are not initialized in CHECK procedure, number of updates in value $\hat{v}_i$ is bounded by $\log T$ during the execution of ESUCB algorithm, instead of $\log^2 T$ when initialization is executed in CHECK. Therefore we can give a better upper bound on the item switch of ESUCB algorithm. The following theorem shows the item switch bound of modified ESUCB algorithm.

Theorem 35. *The number of item switches incurred by ESUCB algorithm is bounded by $O(N \log T)$.*

Proof. Recall that S_ℓ is calculated by $S_\ell = \arg \max_{S \in [N], |S| \leq K} (\sum_{i \in S} \hat{v}_i(r_i - \theta))$ for some θ (Line 6 and Line 9 of Algorithm 6). Observe that the value of b in Algorithm 6 can only be switched once in an invocation. Therefore the number of switches in value θ is upper bounded by $O(\log T)$. The item number of item switch introduced by the change of θ is then bounded by $O(N \log T)$. Now, consider an consecutive time steps where θ is unchanged. We only need to show that for fixed any θ , and $S'_\ell = \arg \max_{S \subseteq [N], |S| \leq K} (\sum_{i \in S} \hat{v}_i(r_i - \theta))$, it holds that (assuming that there are L epochs)

$$\sum_{\ell=1}^{L-1} |S'_\ell \oplus S'_{\ell+1}| \lesssim N \log T. \quad (28)$$

Suppose that there are n_ℓ items whose UCB values are updated after the ℓ -th epoch. We claim that $|S_\ell \oplus S_{\ell+1}| \leq n_\ell$. This is simply because S_ℓ corresponds to the items $i \in [N]$ such that $\hat{v}_i(r_i - \theta)$ is positive and among the K largest ones (and thanks to the tie breaking rule). Therefore, any update to a single $\hat{v}_i$ will incur at most one item switch to S_ℓ , and n_ℓ updates will incur at most n_ℓ item switches. Now, (28) is established because $\sum_{\ell=1}^{L-1} |S'_\ell \oplus S'_{\ell+1}| \leq \sum_{\ell=1}^{L-1} n_\ell \lesssim N \log T$, where the second inequality is due to the deferred update rule for the UCB values. $\square$