

Safe Multi-Objective Linear Bandits with Hierarchical Preferences

Bo Xue^{1,2}, Mengxia He³, Yilu Liu^{1,2}, Ji Cheng^{1,2}, Zhe Zhao¹ and Qingfu Zhang^{1,2,*}

¹Department of Computer Science, City University of Hong Kong, Hong Kong, China

²The City University of Hong Kong Shenzhen Research Institute, Shenzhen, China

³Department of Electrical and Computer Engineering, The University of Hong Kong, Hong Kong, China
 {boxue4-c, yilu.liu, J.Cheng}@my.cityu.edu.hk, mxhe@eee.hku.hk, zz4543@mail.ustc.edu.cn, qingfu.zhang@cityu.edu.hk

Abstract

Multi-objective bandits with hierarchical preferences and safety constraints is central to many real-world decision-making tasks such as health-care treatment planning and safe autonomous control, where multiple objectives must be optimized according to their priorities while ensuring safety requirements are satisfied. In this paper, we study a multi-objective stochastic linear bandit framework that incorporates *hierarchical preferences* together with *safety constraints*, requiring the learner to remain competitive with respect to a known baseline policy. We consider two practically motivated safety models: (i) *cumulative constraints*, which require the cumulative performance to exceed the baseline, and (ii) *stage-wise constraints*, which impose this requirement at each time step. We propose two algorithms, LexUCB-C and LexTS-S, designed for the cumulative and stage-wise settings, respectively. We establish regret bounds showing that both algorithms achieve performance comparable to existing single-objective safe linear bandit methods, while simultaneously optimizing multiple objectives. In addition to theoretical guarantees, we develop an experimental framework that captures the interaction between hierarchical preferences and safety constraints. Experiments on synthetic and real-world datasets demonstrate the effectiveness of the proposed methods.

1 Introduction

Multi-armed bandits (MAB) provide a foundational framework for modeling sequential decision-making under uncertainty [Lai and Robbins, 1985; Auer *et al.*, 2002; Abbasi-yadkori *et al.*, 2011; Lattimore and Szepesvári, 2020]. In a standard MAB problem, a learner interacts with an environment over T rounds. At each round $t \in [T]$, the learner selects an arm from a finite or structured decision set and receives a stochastic reward corresponding to the chosen arm. The goal of the learner is to maximize the cumulative expected reward. A key challenge in this problem is the

exploration-exploitation trade-off: the learner must balance exploring less-known arms to gather information and exploiting currently known arms that believed to yield high rewards.

To address this, two widely used strategies have been developed: Upper Confidence Bound (UCB) and Thompson Sampling (TS). UCB algorithms construct confidence intervals for each arm’s expected reward and select the arm with the highest upper bound [Auer, 2002; Dani *et al.*, 2008; Magureanu *et al.*, 2014; Xue *et al.*, 2020; Xue *et al.*, 2023]. This strategy favors arms with either high empirical rewards or high uncertainty, thereby balancing exploration and exploitation. TS adopts a Bayesian perspective by maintaining a posterior distribution over the arm parameters and selecting arms according to their probability of being optimal, thus encouraging exploration through randomization [Thompson, 1933; Chapelle and Li, 2011; Xu *et al.*, 2023].

While the classical MAB focuses on a single objective [Bubeck and Cesa-Bianchi, 2012], many real-world problems involve multiple, often conflicting objectives. For example, recommender systems [Rodriguez *et al.*, 2012] must balance user satisfaction, content diversity, and long-term engagement, while autonomous driving [Zhang *et al.*, 2023] requires trade-offs among safety, passenger comfort, and energy efficiency. To capture such multi-criteria decision-making scenarios, the multi-objective multi-armed bandit (MOMAB) framework has been proposed [Drugan and Nowe, 2013], in which the learner receives a vector-valued reward at each round, with each component representing a distinct objective.

However, the MOMAB framework treats each arm independently and thus fails to exploit contextual or structural information often available in real-world applications, such as product categories, prices, or user profiles in recommender systems [Li *et al.*, 2010]. To address this limitation, the multi-objective stochastic linear bandit (MOSLB) model was later introduced [Lu *et al.*, 2019; Xue *et al.*, 2025b], where each arm is associated with a feature vector, and the expected reward for each objective is assumed to be a linear function of this vector. This assumption enables the algorithm to generalize across arms: by leveraging the shared linear structure, the learner can use feedback from one arm to estimate the rewards of other arms with similar features.

Most existing methods in the multi-objective bandit literature rely on scalarization [Drugan and Nowe, 2013; Wani-gasekara *et al.*, 2019] or Pareto-based criteria [Auer *et al.*,

*Qingfu Zhang is the corresponding author.

2016; Cheng *et al.*, 2026], which either collapse the multi-objective problem into a single-objective one by imposing fixed preference weights, or seek to identify Pareto-optimal solutions without explicitly modeling priority among objectives. However, these approaches may fall short in applications where objectives are inherently prioritized. For example, in clinical treatment planning, ensuring safety or efficacy is often paramount, while secondary factors such as cost or comfort are considered only after primary goals are satisfied. To accommodate such structured preferences, lexicographic ordering has been proposed [Ehrgott, 2005; Skalse *et al.*, 2022; Abernethy *et al.*, 2024], wherein objectives are ranked by importance, and reward vectors are compared based on the first dimension where they differ. Incorporating lexicographic preferences into the bandit learning framework leads to the lexicographic multi-objective bandit problem [Tekin, 2019; Hüyük and Tekin, 2021], where the higher-priority objectives must be optimized before lower-priority ones.¹

Despite recent progress in lexicographic bandits [Hüyük and Tekin, 2021; Xue *et al.*, 2024; Xue *et al.*, 2025a; Xue *et al.*, 2026], existing methods often neglect safety constraints, which are essential in many real-world applications. For example, in high-stakes environments such as healthcare systems, deploying exploratory policies without performance guarantees may lead to unacceptable outcomes. This motivates the study of conservative (or safe) bandits [Wu *et al.*, 2016; Katariya *et al.*, 2019; Kazerouni *et al.*, 2017; Moradipari *et al.*, 2020; Lin *et al.*, 2022], where the learner must ensure that its performance does not fall significantly below that of a known baseline. ***While conservative constraints have been well-studied in single-objective settings, whether comparable guarantees can be attained under multi-objective formulations remains an open problem.***

In this paper, we introduce the framework of safe MOSLB with hierarchical preferences, which combines lexicographic preferences with two types of safety constraints: cumulative constraints and stage-wise constraints. This is the first work to address multi-objective bandits under safety constraints, and our main contributions are summarized as follows:

- Under cumulative constraints, we present LexUCB-C, a UCB-based algorithm that achieves a regret bound of $\tilde{O}(W^i(w) \cdot d\sqrt{T} + d^2 \cdot \alpha^{-2})$ for each objective $i \in [m]$ where $W^i(w) = 1 + w + \dots + w^{i-1}$, w is a weight parameter reflecting the trade-off among conflicting objectives (introduced in Assumption 4), d is the dimension, T is the time horizon, and α is the safety threshold.
- Under stage-wise constraints, we develop LexTS-S, a Thompson Sampling-based method that enforces per-round safety. LexTS-S achieves a regret bound of $\tilde{O}(W^i(w) \cdot d^{3/2} \sqrt{T} + d \cdot \alpha^{-3})$ for each objective $i \in [m]$.
- Both algorithms achieve regret bounds comparable to those of single-objective safe bandit methods [Kazerouni *et al.*, 2017; Moradipari *et al.*, 2020], in terms of their dependence on d , T and α . Moreover, since $W^1(w) = 1$ for any trade-off value w , the performance

of the highest-priority objective ($i = 1$) remains nearly unaffected when optimizing additional objectives.

- We conduct experiments on both synthetic and real-world datasets using carefully designed settings that reflect hierarchical preferences and safety requirements. The results validate our theoretical findings.

2 Problem Setting

This section formalizes the problem setting. First, we introduce the notation and the standard MOSLB model. Second, we define the lexicographic ordering and extend the conservative bandit framework to the multi-objective setting. Finally, we state the key assumptions used in our analysis.

Notation. We denote by $\mathbf{x} \in \mathbb{R}^d$ a feature (context) vector in a d -dimensional real space. The standard Euclidean norm of $\mathbf{x}$ is written as $\|\mathbf{x}\|$, while the Mahalanobis norm with respect to a positive definite matrix $\mathbf{V} \in \mathbb{R}^{d \times d}$ is defined as $\|\mathbf{x}\|_{\mathbf{V}} = (\mathbf{x}^\top \mathbf{V} \mathbf{x})^{1/2}$. For a positive integer m , we use $[m]$ to represent the index set $\{1, 2, \dots, m\}$. The superscript $i \in [m]$ indicates quantities associated with the i -th objective.

Model. MOSLB is a T -round sequential decision-making problem. At each round $t \in [T]$, the learner observes a set of feasible arms $\mathcal{X}_t \subseteq \mathbb{R}^d$, and selects an arm $\mathbf{x}_t \in \mathcal{X}_t$. Upon selection, the learner receives a stochastic reward vector $\mathbf{y}_t = [y_t^1, y_t^2, \dots, y_t^m] \in \mathbb{R}^m$, where m is the number of objectives. The expected reward for each objective $i \in [m]$ satisfies

$$\mathbb{E}[y_t^i] = \mu^i(\mathbf{x}_t) = \langle \boldsymbol{\theta}_*^i, \mathbf{x}_t \rangle,$$

where $\boldsymbol{\theta}_*^i \in \mathbb{R}^d$ is an unknown vector for the i -th objective.

Next, we adopt the lexicographic order [Hüyük and Tekin, 2021] to compare reward vectors across multiple objectives.

Definition 1 (Lexicographic Order). Suppose $\mathbf{u}, \mathbf{v} \in \mathbb{R}^m$. $\mathbf{u}$ lexicographically dominates $\mathbf{v}$, if there exists an index $i^* \in [m]$ such that $u^j = v^j$ for all $j < i^*$, and $u^{i^*} > v^{i^*}$.

For example, $[1, 3, 2]$ dominates $[1, 0, 8]$ with $i^* = 2$. Based on this, the optimal arm is defined as follows.

Definition 2 (Lexicographic Optimal Arm). An arm $\mathbf{x}_t^* \in \mathcal{X}_t$ is lexicographic optimal if its expected reward vector is not lexicographically dominated by that of any arm in $\mathcal{X}_t$.

Given the optimal arm $\mathbf{x}_t^*$, we can naturally define the regret in multi-objective setting, which is the cumulative gap between the optimal arm $\mathbf{x}_t^*$ and the chosen arm $\mathbf{x}_t$, i.e.,

$$R^i(T) = \sum_{t=1}^T \langle \boldsymbol{\theta}_*^i, \mathbf{x}_t^* - \mathbf{x}_t \rangle, i \in [m].$$

In conservative bandits, the learner is required to maintain performance that does not fall significantly below that of a given baseline policy [Kazerouni *et al.*, 2017; Moradipari *et al.*, 2020]. We extend this principle to the multi-objective setting. Specifically, at each round t , a baseline arm $\mathbf{x}_{b_t}$ is provided. We consider two types of safety constraints:

Definition 3 (Cumulative Constraint). For each objective $i \in [m]$, the learner’s cumulative expected reward must exceed a fraction $(1 - \alpha)$ of the baseline’s cumulative reward:

$$\forall t \in [T], \sum_{\tau=1}^t \mu^i(\mathbf{x}_\tau) \geq (1 - \alpha) \sum_{\tau=1}^t \mu^i(\mathbf{x}_{b_\tau}).$$

¹Due to space limitations, a more comprehensive discussion of related work can be found in the Appendix A.

Definition 4 (Stage-Wise Constraint). *For each objective $i \in [m]$, the learner’s expected reward at each round must exceed a fraction $(1 - \alpha)$ of the baseline reward:*

$$\forall t \in [T], \mu^i(\mathbf{x}_t) \geq (1 - \alpha)\mu^i(\mathbf{x}_{b_t}).$$

Finally, we introduce some common assumptions on linear bandits [Abbasi-yadkori *et al.*, 2011; Jun and Kim, 2024], conservative bandits [Kazerouni *et al.*, 2017; Moradipari *et al.*, 2020] and multi-objective bandits [Xue *et al.*, 2024].

Assumption 1. *The reward noise is σ -sub-Gaussian, i.e., for all $t \in [T]$ and $i \in [m]$,*

$$\mathbb{E} \left[e^{\beta(y_t^i - \mu^i(\mathbf{x}_t))} \right] \leq \exp \left(\frac{\beta^2 \sigma^2}{2} \right), \forall \beta \in \mathbb{R}.$$

Assumption 2. *The inherent parameter vectors and contextual vectors are bounded, such that there exists some $B > 0$,*

$$\|\boldsymbol{\theta}_*^i\| \leq B, \forall i \in [m], \text{ and } \|\mathbf{x}\| \leq 1, \forall \mathbf{x} \in \mathcal{X}_t.$$

Meanwhile, $\langle \boldsymbol{\theta}_^i, \mathbf{x} \rangle \leq 1$ holds for all $i \in [m]$ and $\mathbf{x} \in \mathcal{X}_t$.*

The sub-Gaussian assumption allows us to derive high-probability confidence bounds via standard concentration inequalities, while the boundedness of both model parameters and context vectors ensures that the expected rewards remain within a well-defined range.

Assumption 3. *Let $\Delta^i(\mathbf{x}_{b_t}) = \mu^i(\mathbf{x}_t^*) - \mu^i(\mathbf{x}_{b_t})$. There exist constants $0 \leq \Delta_\ell \leq \Delta_u$ and $0 \leq \mu_\ell \leq \mu_u$ such that for all $i \in [m]$ and $t \in [T]$,*

$$\Delta_\ell \leq \Delta^i(\mathbf{x}_{b_t}) \leq \Delta_u, \text{ and } \mu_\ell \leq \mu^i(\mathbf{x}_{b_t}) \leq \mu_u.$$

This assumption ensures that the conservative baseline arm maintains a bounded gap from the optimal arm, which is critical for establishing safe exploration guarantees.

Assumption 4. *There exists a constant $w > 0$ such that for any $i \geq 2$ and $\mathbf{x} \in \mathcal{X}_t$, the trade-off between the i -th objective and higher-priority objectives is bounded as*

$$\mu^i(\mathbf{x}) - \mu^i(\mathbf{x}_t^*) \leq w \cdot \max_{j \in [i-1]} \{\mu^j(\mathbf{x}_t^*) - \mu^j(\mathbf{x})\}.$$

This structural condition captures the lexicographic preference among objectives, ensuring that, relative to the optimal arm $\mathbf{x}_t^*$, any improvement in a lower-priority objective cannot compensate for excessive degradation in higher-priority ones.

3 Algorithms

We develop two algorithms for lexicographic linear bandits with safety constraints. LexUCB-C addresses the cumulative constraint using the UCB principle, while LexTS-S handles the stage-wise constraint via Thompson Sampling. Together, they provide complementary approaches for safe lexicographic bandit optimization.

3.1 Lexicographic UCB for Cumulative Safety

LexUCB-C is a UCB-based algorithm for setting with cumulative safety constraints. At each round, it either selects a high-performing arm under lexicographic preferences or defaults to a conservative baseline to ensure safety. The full procedure is in Algorithm 1.

Algorithm 1 LexUCB-C

Input: $T, m, B, \alpha, \delta, \lambda, w$

```

1: Set  $V_1 = \lambda I$ ,  $c_1 = \sqrt{\lambda}B$ ,  $\hat{\boldsymbol{\theta}}_1^i = \mathbf{0}$ ,  $\mathbf{z}_1 = \mathbf{0}$ ,  $\mathcal{T}_1 = \emptyset$ 
2: for  $t = 1$  to  $T$  do
3:    $\mathbf{x}'_t \leftarrow \text{SELECTARM}(\{\hat{\boldsymbol{\theta}}_t^i\}_{i=1}^m, c_t, V_t, \mathcal{X}_t)$  (Alg. 2)
4:   Compute the lower bounds  $L_t^i$  for  $i \in [m]$  by Eq. (1)
5:   if Safety condition (2) holds for all objectives then
6:     Play  $\mathbf{x}_t = \mathbf{x}'_t$  and observe  $[y_t^1, y_t^2, \dots, y_t^m]$ 
7:     Update  $\mathbf{z}_{t+1} = \mathbf{z}_t + \mathbf{x}_t$ ,  $\mathcal{T}_{t+1} = \mathcal{T}_t \cup \{t\}$ 
8:     Update  $V_{t+1} = V_t + \mathbf{x}_t \mathbf{x}_t^\top$ 
9:     Compute the estimators  $\{\hat{\boldsymbol{\theta}}_{t+1}^i\}_{i=1}^m$  by Eq. (3)
10:    Compute the confidence radius  $c_{t+1}$  by Eq. (4)
11:   else
12:     Play baseline  $\mathbf{x}_t = \mathbf{x}_{b_t}$  (Conservative play)
13:     Maintain  $\mathbf{z}_{t+1} = \mathbf{z}_t$ , update  $\mathcal{T}_{t+1}^c = \mathcal{T}_t^c \cup \{t\}$ 
14:     Keep  $V_{t+1} = V_t$ ,  $c_{t+1} = c_t$ ,  $\hat{\boldsymbol{\theta}}_{t+1}^i = \hat{\boldsymbol{\theta}}_t^i, \forall i \in [m]$ 
15:   end if
16: end for

```

Initialization. At the beginning of the algorithm, several quantities are initialized. The regularized covariance matrix is set to $V_1 = \lambda I$, where $\lambda > 0$ ensures numerical stability in the estimation process. The confidence radius is initialized as $c_1 = \sqrt{\lambda}B$, where B is the bounded prior in Assumption 2. The initial estimators for each objective $i \in [m]$ are set to zero vectors, i.e., $\hat{\boldsymbol{\theta}}_1^i = \mathbf{0}$. The cumulative decision vector $\mathbf{z}_1$ is initialized to zero and will be used to accumulate the sum of selected arms that are not conservative. The set $\mathcal{T}_1$, which tracks rounds in which optimistic (i.e., non-conservative) arms are played, is initialized to be empty.

Main Loop. At each round t , LexUCB-C first calls the subroutine SELECTARM (Algorithm 2), which returns a candidate arm $\mathbf{x}'_t$ that respects the lexicographic preference structure. The details of this subroutine will be discussed shortly.

Given a selected candidate arm $\mathbf{x}'_t$, LexUCB-C evaluates its safety by computing the worst-case lower bound on the cumulative reward, including $\mathbf{x}'_t$ for each objective $i \in [m]$,

$$L_t^i = \min_{\boldsymbol{\theta} \in \mathcal{C}_t^i} \langle \boldsymbol{\theta}, \mathbf{z}_t + \mathbf{x}'_t \rangle, \quad \mathcal{C}_t^i = \{\boldsymbol{\theta} \mid \|\boldsymbol{\theta} - \hat{\boldsymbol{\theta}}_t^i\|_{V_t} \leq c_t\}. \quad (1)$$

$\mathbf{z}_t$ denotes the cumulative action vector up to round $t - 1$, and $\hat{\boldsymbol{\theta}}_t^i$ is the empirical estimate of the reward vector for objective i . The term L_t^i captures the most pessimistic cumulative reward that could be incurred under current confidence region, assuming $\mathbf{x}'_t$ is selected.

To ensure safety, LexUCB-C verifies whether the total worst-case reward exceeds a certain fraction of the reward obtained by always playing a fixed conservative baseline. Specifically, the arm $\mathbf{x}'_t$ is considered safe if the following condition holds:

$$L_t^i + \sum_{\tau \in \mathcal{T}_t^c} \mu^i(\mathbf{x}_{b_\tau}) \geq (1 - \alpha) \sum_{\tau=1}^t \mu^i(\mathbf{x}_{b_\tau}), \forall i \in [m], \quad (2)$$

where $\mathcal{T}_t^c$ is the set of rounds up to $t - 1$ in which the conservative baseline action $\mathbf{x}_{b_\tau}$ was played. This constraint guarantees that, in each objective, the cumulative reward does not fall below the benchmark established by the baseline policy.

Algorithm 2 SELECTARM

Input: $\{\hat{\theta}_t^i\}_{i=1}^m, c_t, V_t, \mathcal{X}_t$

- 1: Initialize $s = 1, \mathcal{X}_{t,s} = \mathcal{X}_t$
- 2: **repeat**
- 3: **if** $\exists x \in \mathcal{X}_{t,s}$ with $\|x\|_{V_t^{-1}} \leq 1/\sqrt{T}$ **then**
- 4: $\mathcal{X}_{t,T} \leftarrow \text{LEXFILTER}(\{\hat{\theta}_t^i\}_{i=1}^m, c_t, \mathcal{X}_{t,s}, \frac{1}{\sqrt{T}})$
- 5: Return any $x \in \mathcal{X}_{t,T}$
- 6: **else if** $\exists x \in \mathcal{X}_{t,s}$ with $\|x\|_{V_t^{-1}} > 2^{-s}$ **then**
- 7: Return such x
- 8: **else**
- 9: $\mathcal{X}_{t,s+1} \leftarrow \text{LEXFILTER}(\{\hat{\theta}_t^i\}_{i=1}^m, c_t, \mathcal{X}_{t,s}, 2^{-s})$
- 10: $s \leftarrow s + 1$
- 11: **end if**
- 12: **until** an arm is selected

Algorithm 3 LEXFILTER

Input: $\{\hat{\theta}_t^i\}_{i=1}^m, c_t, \mathcal{X}_{t,s}, C$

- 1: Initialize $\mathcal{X}_{t,s}^0 = \mathcal{X}_{t,s}$
- 2: **for** $i = 1$ to m **do**
- 3: $\hat{x}_t^i = \arg \max_{x \in \mathcal{X}_{t,s}^{i-1}} \langle \hat{\theta}_t^i, x \rangle$
- 4: $\mathcal{X}_{t,s}^i = \{x \in \mathcal{X}_{t,s}^{i-1} \mid \langle \hat{\theta}_t^i, \hat{x}_t^i - x \rangle \leq (\sum_{j=0}^{i-1} 4w^j + 2) \cdot c_t \cdot C\}$
- 5: **end for**
- 6: **return** $\mathcal{X}_{t,s}^m$

If the safety condition is satisfied for all objectives, the candidate arm x_t' is considered safe. LexUCB-C plays x_t' and observes the reward vector $[y_t^1, \dots, y_t^m]$. It then updates the cumulative arm vector, $z_{t+1} = z_t + x_t$, adds the current round to the non-conservative set, $\mathcal{T}_{t+1} = \mathcal{T}_t \cup t$, updates the covariance matrix, $V_{t+1} = V_t + x_t x_t^\top$, and recomputes the estimators $\{\hat{\theta}_{t+1}^i\}_{i=1}^m$ via regularized least squares:

$$\hat{\theta}_{t+1}^i = V_{t+1}^{-1} X_{t+1} Y_{t+1}^i, \quad (3)$$

where

$$X_{t+1} = [x_1, x_2, \dots, x_t], Y_{t+1}^i = [y_1^i, y_2^i, \dots, y_t^i].$$

The confidence radius c_{t+1} is then updated accordingly, i.e.,

$$c_{t+1} = \sigma \sqrt{d \log \left(\frac{m(1 + |\mathcal{T}_{t+1}|/\lambda)}{\delta} \right)} + \sqrt{\lambda} B. \quad (4)$$

If the safety test fails for some objective, LexUCB-C plays a predefined conservative arm $x_t = x_{b_t}$. In this case, the cumulative vector z_t , estimators, confidence radius and covariance matrix remain unchanged. The round index is added to the set of conservative rounds: $\mathcal{T}_{t+1}^c = \mathcal{T}_t^c \cup \{t\}$.

Arm Selection. To identify a candidate arm x_t' for the safety test, LexUCB-C relies on a dedicated subroutine, SELECTARM, which aims to find a lexicographically optimal arm under current uncertainty. We now describe this arm selection procedure and its internal mechanism.

SELECTARM begins with the full feasible arm set $\mathcal{X}_{t,1} = \mathcal{X}_t$ and iteratively refines it. At each iteration s , SELECTARM

considers the current set $\mathcal{X}_{t,s}$, and evaluates the uncertainty term $\|x\|_{V_t^{-1}}$ for each arm $x \in \mathcal{X}_{t,s}$. Precisely, **if** there exists an arm with sufficiently small uncertainty (Step 3), SELECTARM invokes the LEXFILTER procedure with a high-confidence threshold $C = 1/\sqrt{T}$, and randomly returns any arm from the resulting filtered set. This ensures the selection of a confident and lexicographically optimal arm. **Otherwise**, if there exists an arm with norm exceeding a coarser threshold 2^{-s} (Step 6), such an arm is directly selected. This branch encourages timely exploration when high-confidence arms are not yet available. **If neither** condition is met (Step 8), the arm set is refined by applying LEXFILTER with the threshold $C = 2^{-s}$. Note that this lexicographic criterion becomes stricter as the iteration progresses, i.e., from stage s to $s + 1$. The updated set $\mathcal{X}_{t,s+1}$ is then used in the next iteration, and such process repeats until an arm is selected.

Lexicographic Elimination. LEXFILTER performs sequential elimination across objectives following their priorities. Given the current arm set $\mathcal{X}_{t,s}$, it maintains a sequence of filtered subsets $\mathcal{X}_{t,s}^i$ for $i \in [m]$. At each level i , LEXFILTER identifies the empirically optimal arm,

$$\hat{x}_t^i = \arg \max_{x \in \mathcal{X}_{t,s}^{i-1}} \langle \hat{\theta}_t^i, x \rangle.$$

It then eliminates any arm x from $\mathcal{X}_{t,s}^{i-1}$ if its estimated reward falls behind that of $\hat{x}_t^i$ by more than a threshold proportional to $c_t \cdot C$,

$$\mathcal{X}_{t,s}^i = \{x \in \mathcal{X}_{t,s}^{i-1} \mid \langle \hat{\theta}_t^i, \hat{x}_t^i - x \rangle \leq (\sum_{j=0}^{i-1} 4w^j + 2) \cdot c_t \cdot C\}.$$

The threshold is scaled by a factor capturing the cumulative uncertainty propagated from higher-priority objectives. After m rounds, the final surviving set $\mathcal{X}_{t,s}^m$ is returned which consists of arms that remain competitive under all objectives.

In summary, LexUCB-C combines lexicographic optimization with cumulative safety monitoring. It ensures that the chosen arms respect safety constraints over time while optimizing higher-priority objectives first. We provide the following theoretical guarantee for LexUCB-C.

Theorem 1. Suppose Assumptions 1 - 4 hold. Then, with probability at least $1 - \delta$, for any objective $i \in [m]$, the regret of LexUCB-C is bounded as

$$R^i(T) \leq 4W^i(w) \cdot c_T \cdot (\sqrt{T} + 10 \log(T) \sqrt{dT}) + \frac{d^2 K \Delta_u}{\alpha \mu_\ell \cdot (\Delta_\ell \alpha \mu_\ell)},$$

where

$$W^i(w) = 1 + w + \dots + w^{i-1},$$

$$c_T = \sigma \sqrt{d \log \left(\frac{m(1 + |\mathcal{T}_T|/\lambda)}{\delta} \right)} + \sqrt{\lambda} B,$$

$$K = 228(B\sqrt{\lambda} + \sigma)^2 \left[\log \left(\frac{62d\sqrt{m}(B\sqrt{\lambda} + \sigma)}{\sqrt{\delta}(\Delta_\ell + \alpha \mu_\ell)} \right) \right]^2.$$

Remark 1. Theorem 1 shows that LexUCB-C achieves a regret bound of $\tilde{O}(W^i(w) \cdot d\sqrt{T} + d^2 \cdot \alpha^{-2})$ for any objective $i \in [m]$. The first term matches the standard regret rate of single-objective linear bandits [Kazerouni *et al.*, 2017] up to a multiplicative factor $W^i(w)$. This factor reflects the cost of optimizing multiple objectives in a lexicographic manner. Notably, $W^1(w) = 1$ implies that the regret bound on the highest-priority objective remains unaffected even optimizing other objectives simultaneously. Lower-priority objectives incur larger regret so as to maintain prioritization structure.

3.2 Lexicographic TS for Stage-Wise Safety

Under the stage-wise safety setting, we propose LexTS-S, a Thompson Sampling-based algorithm that incorporates lexicographic preferences and enforces safety at every round. The full procedure is summarized in Algorithm 4.

Initialization. At the beginning, LexTS-S set the regularized covariance matrix as $V_1 = \lambda I$, where the parameter $\lambda > 0$ ensures numerical stability. The initial confidence parameter is defined as $c_1 = \beta_1 = \sqrt{\lambda}B$, and the reward estimator for each objective $i \in [m]$ is initialized to zero, i.e., $\hat{\theta}_1^i = 0$. Meanwhile, a mixing parameter $\rho_1 \in (0, \frac{\mu_\ell \alpha}{B + \mu_u})$ is also specified for constructing safe exploration.

Main Loop. At each round t , LexTS-S first generates a randomized parameter vector to estimate reward for each objective. Specifically, for each $i \in [m]$, a sample is drawn from a Gaussian posterior,

$$\tilde{\theta}_t^i \sim \mathcal{N}(\hat{\theta}_t^i, c_t^2 \cdot V_t^{-1}),$$

where the posterior mean $\hat{\theta}_t^i$ and covariance $c_t^2 \cdot V_t^{-1}$ are constructed using the same regularized least-squares estimator and covariance matrix as in LexUCB-C. This TS-based sampling introduces randomized exploration, in contrast to the optimistic UCB strategy used in LexUCB-C.

Next, LexTS-S constructs a stage-wise safety region S_t by intersecting the per-objective safe sets:

$$S_t = \bigcap_{i=1}^m S_t^i,$$

where

$$S_t^i = \left\{ x \in \mathcal{X}_t \mid \min_{\theta \in \mathcal{C}_t^i} \langle \theta, x \rangle \geq (1 - \alpha)\mu^i(x_{b_t}) \right\}.$$

Here, $\mathcal{C}_t^i = \{\theta \mid \|\theta - \hat{\theta}_t^i\|_{V_t} \leq c_t\}$ is the confidence region for objective i . Intuitively, S_t^i includes arms whose worst-case reward is guaranteed to exceed a fraction $(1 - \alpha)$ of the baseline’s reward for objective i . By intersecting these sets, S_t contains only arms that are simultaneously safe across all objectives at round t . This differs from LexUCB-C, where safety is verified cumulatively using a global reward budget, and thus the candidate arm could be risky in the short term as long as cumulative safety is maintained.

LexTS-S then checks whether the safety region S_t is nonempty and whether the design matrix is sufficiently well-conditioned, i.e.,

$$\lambda_{\min}(V_t) \geq \left(\frac{2c_t}{\Delta_\ell + \alpha\mu_\ell} \right)^2,$$

Algorithm 4 LexTS-S

Input: $T, m, B, \alpha, \delta, \lambda, w, \Delta_\ell, \mu_\ell$

- 1: Set $V_1 = \lambda I$, $c_1 = \beta_1 = \sqrt{\lambda}B$, $\hat{\theta}_1^i = 0$
- 2: Choose a value $\rho_1 \in (0, \frac{\mu_\ell \alpha}{B + \mu_u})$
- 3: **for** $t = 1$ to T **do**
- 4: Sample $\tilde{\theta}_t^i \sim \mathcal{N}(\hat{\theta}_t^i, c_t^2 \cdot V_t^{-1})$
- 5: Compute the safety region: $S_t = \bigcap_{i=1}^m S_t^i$, where

$$S_t^i = \left\{ x \in \mathcal{X}_t \mid \min_{\theta \in \mathcal{C}_t^i} \langle \theta, x \rangle \geq (1 - \alpha)\mu^i(x_{b_t}) \right\}$$

- 6: **if** $S_t \neq \emptyset$ and $\lambda_{\min}(V_t) \geq \left(\frac{2c_t}{\Delta_\ell + \alpha\mu_\ell} \right)^2$ **then**
- 7: $x_t \leftarrow \text{SELECTARM}(\tilde{\theta}_t^i, c_t + \beta_t, V_t, S_t)$ (Alg. 2)
- 8: **else**
- 9: Sample ζ_t uniformly from the unit sphere in $\mathbb{R}^d$
- 10: $x_t \leftarrow (1 - \rho_1)x_{b_t} + \rho_1\zeta_t$ (Conservative play)
- 11: **end if**
- 12: Play x_t and observe reward vector $[y_t^1, y_t^2, \dots, y_t^m]$
- 13: Update $V_{t+1} = V_t + x_t x_t^\top$
- 14: Compute the estimators $\{\hat{\theta}_{t+1}^i\}_{i=1}^m$ by Eq. (3)
- 15: Compute c_{t+1} and β_{t+1} by Eq. (5)
- 16: **end for**

where Δ_ℓ and μ_ℓ are problem-dependent constants defined in Assumption 3. If these conditions are satisfied, LexTS-S calls the SELECTARM subroutine (Algorithm 2) to select an arm x_t , which passes the sampled parameters $\{\tilde{\theta}_t^i\}_{i=1}^m$, the confidence parameter $c_t + \beta_t$, and the feasible set S_t . This subroutine follows the same iterative filtering mechanism as in LexUCB-C, using LEXFILTER to preserve lexicographic optimality among the arms in S_t .

If either condition fails, either because the safety region is empty or because the confidence in parameter estimates is too low, LexTS-S takes a conservative action. Instead of playing a fixed baseline arm, it mixes the conservative baseline x_{b_t} with a randomly sampled unit vector ζ_t ,

$$x_t = (1 - \rho_1)x_{b_t} + \rho_1\zeta_t,$$

where ρ_1 balances between baseline performance and exploratory progress. This conservative play guarantees the stage-wise safety constraint even in highly uncertain situations, similar to the conservative mechanism in LexUCB-C but with a stage-wise guarantee.

After the arm x_t is played, LexTS-S observes the reward vector $[y_t^1, y_t^2, \dots, y_t^m]$ and updates the covariance matrix and estimators in the same manner as LexUCB-C. Finally, the confidence radius is computed as follows,

$$\begin{aligned} c_{t+1} &= \sigma \sqrt{d \log \left(\frac{mT(1 + t/\lambda)}{\delta} \right)} + \sqrt{\lambda}B \\ \beta_{t+1} &= c_{t+1} \cdot \sqrt{2d \log \left(\frac{8dmT}{\delta} \right)}. \end{aligned} \quad (5)$$

In summary, LexTS-S differs from LexUCB-C in two key aspects: (i) it enforces stage-wise safety, requiring that each

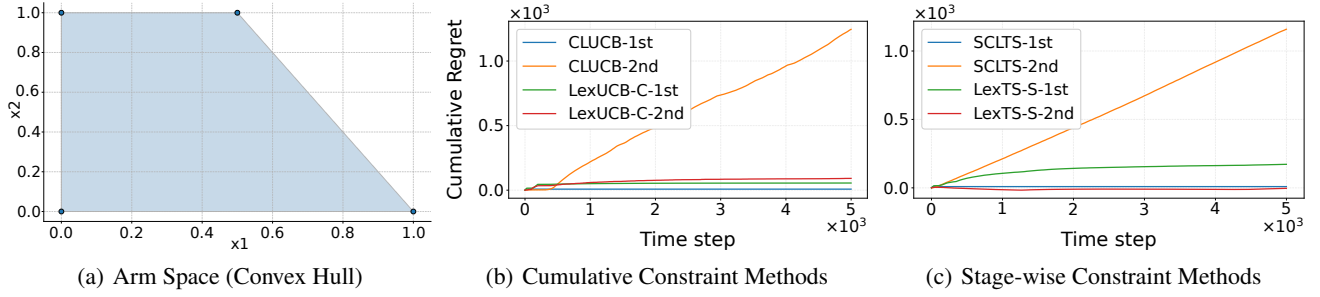


Figure 1: Visualization of arm space and regret curves under safety constraints. (a) The arm space is a convex hull of four vertices: $\{(0,0), (1,0), (0,1), (0.5,1)\}$, where the optimal arm is $(0.5,1)$. (b) Performance comparison of CLUCB and LexUCB-C under cumulative safety constraints, where LexUCB-C achieves sublinear regret on both the 1st and 2nd objectives. (c) Performance comparison of SCLTS and LexTS-S under stage-wise safety constraints, where LexTS-S achieves sublinear regret on both the 1st and 2nd objectives.

selected arm be immediately safe, rather than maintaining safety cumulatively; and (ii) it adopts a Thompson Sampling strategy based on posterior sampling of reward parameters, in contrast to the optimistic upper confidence bound approach used in LexUCB-C. These distinctions make LexTS-S more conservative in the short term, and better suited for environments where immediate safety is essential.

The regret bound of LexTS-S is provided as follows.

Theorem 2. Suppose Assumptions 1 - 4 hold, and that the arm set $\mathcal{X}_t$ is convex and compact for all $t \in [T]$. Then, with probability at least $1 - \delta$, for any objective $i \in [m]$, the regret of LexTS-S is bounded as

$$R^i(T) \leq 4W^i(w) \cdot (c_T + \beta_T) \cdot (\sqrt{T} + 10 \log T \cdot \sqrt{dT}) + n_T \cdot (\Delta_u + \rho_1(\mu_u + B)),$$

where

$$\begin{aligned} W^i(w) &= 1 + w + \dots + w^{i-1}, \\ \rho_1 &= \left(\frac{\mu_\ell}{B + \mu_u} \right) \alpha, \quad h_1 = 2\rho_1(1 - \rho_1) + 2\rho_1^2, \\ n_T &\leq \left(\frac{2c_T d}{\rho_1(\Delta_\ell + \alpha\mu_\ell)} \right)^2 + \frac{2h_1^2 d^4}{\rho_1^4} \log \left(\frac{md}{\delta} \right) \\ &\quad + \frac{2h_1 c_T d^3 \sqrt{8 \log(md/\delta)}}{\rho_1^3(\Delta_\ell + \alpha\mu_\ell)}. \end{aligned}$$

Remark 2. Theorem 2 provides a high-probability regret bound for LexTS-S under stage-wise safety constraints. The leading term $\tilde{O}(W^i(w) \cdot d^{3/2} \sqrt{T})$ matches the typical regret rate of Thompson Sampling-based linear bandits [Abeille and Lazaric, 2017] up to a multiplicative factor $W^i(w)$, which arises from sequentially optimizing multiple objectives with lexicographic priorities. Similar to LexUCB-C, the bound ensures that the most important objective ($i = 1$) achieves nearly optimal regret scaling, i.e., $W^1(w) = 1$, while lower-priority objectives may incur increased regret due to the hierarchical optimization. The additional $d^3 \cdot \alpha^{-1}$ term reflects the cost of maintaining strict stage-wise safety throughout the learning process. In particular, a smaller α leads to more cautious behavior, ensuring that safety constraints are less likely to be violated, but at the expense of a larger regret penalty.

4 Experiments

In this section, we evaluate our algorithms under safety constraints on both synthetic and real-world datasets. We summarize the baselines and briefly describe the datasets here. Full experimental details are provided in Appendix B.

We compare our methods with two representative single-objective conservative linear bandit methods: (i) **CLUCB** [Kazerouni *et al.*, 2017], which optimizes only the first objective while ensuring the cumulative reward stays above $(1 - \alpha)$ times that of the baseline, and (ii) **SCLTS** [Moradipari *et al.*, 2020], which optimizes only the first objective and maintains the stage-wise reward above $(1 - \alpha)$ times the baseline.

Synthetic Dataset. The decision set (arm space) is the convex hull of four vertices: $\{(0,0), (1,0), (0,1), (0.5,1)\}$, as illustrated in Figure 1(a). The inherent vectors for the two objectives are $\theta_*^1 = [0, 1]$ and $\theta_*^2 = [1, 0]$, respectively. Under this setting, the expected reward for any arm x is given by $[\langle \theta_*^1 x \rangle, \langle \theta_*^2 x \rangle] = [x_2, x_1]$. Thus, the first objective is maximized when $x_2 = 1$, achieving the optimal value of 1. All arms lying on the top edge of the convex hull satisfy this condition. Among these arms, the second objective is maximized when $x_1 = 0.5$, which yields the highest possible second-objective reward of 0.5 under the constraint that the first objective remains optimal. Therefore, the lexicographic optimal arm is $x_t^* = [0.5, 1]$, corresponding to the expected reward vector $[1, 0.5]$. The baseline (safe) arm is set to the centroid of the convex hull, $\bar{x} = [0.375, 0.5]$.

Under the cumulative safety constraint, Figure 1(b) shows that our proposed algorithm LexUCB-C achieves performance comparable to the baseline CLUCB on the first objective. This empirically supports our theoretical guarantee that LexUCB-C preserves the performance of the most important objective. Moreover, LexUCB-C-2nd consistently yields lower cumulative regret than CLUCB-2nd across the entire time horizon, demonstrating its effectiveness in optimizing the secondary objective without compromising safety. These results highlight the strength of LexUCB-C in balancing lexicographic reward maximization with long-term cumulative safety guarantees in multi-objective decision-making.

Under the stage-wise constraint, Figure 1(c) shows that LexTS-S incurs slightly higher regret on the primary objec-

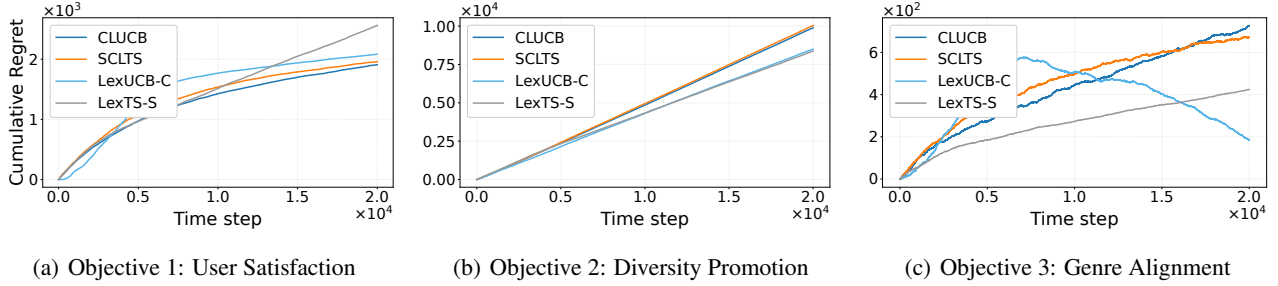


Figure 2: Cumulative regret comparison between LexUCB-C and CLUCB (cumulative safety constraint) and LexTS-S and SCLTS (stage-wise safety constraint) on MovieLens 100K dataset.

tive compared to the baseline SCLTS. This is expected, as SCLTS is designed to solely optimize the primary objective. Nonetheless, the flattened regret curves of LexTS-S on both objectives suggest that it successfully identifies the lexicographic optimal arm. In contrast, SCLTS exhibits nearly linear regret on the second objective, highlighting LexTS-S’s ability to handle multiple objectives simultaneously. These results demonstrate that LexTS-S effectively leverages the multi-objective structure to maintain strong performance even under stringent per-round safety constraints.

Real-world Dataset. We evaluate our proposed algorithms on the MovieLens 100K dataset² [Harper and Konstan, 2015], a widely used benchmark in recommender systems and bandit research [Li *et al.*, 2025]. The dataset consists of 100,000 ratings provided by 943 users on 1,682 movies, where each rating is given on a 5-star scale. Each movie is accompanied by metadata such as title, release date, and genre information covering 19 categories (Action, Adventure, Animation, Children’s, Comedy, Crime, Documentary, Drama, Fantasy, Film-Noir, Horror, Musical, Mystery, Romance, Sci-Fi, Thriller, War, and Western). We represent each movie as a 21-dimensional feature vector, consisting of a 19 one-hot encoding of its genre affiliations, a 1 normalized popularity score (ranging from 0 to 1), and a 1 bias term (set to a constant value of 1). To construct a multi-objective recommendation environment, we define three distinct objectives: User Satisfaction, Diversity Promotion, Genre Alignment.

Figures 2(a)-(c) present the cumulative regrets of all algorithms on the real-world dataset. For the first objective (Figure 2(a)), our proposed algorithms, LexUCB-C and LexTS-S, exhibit regret curves that are almost indistinguishable from those of the baselines CLUCB and SCLTS throughout the learning horizon. This observation indicates that incorporating lexicographic preferences does not compromise the optimization of the most critical objective. In the early stage, the regret of LexUCB-C slightly fluctuates due to its confidence-based exploration, while LexTS-S demonstrates smoother convergence, benefiting from the posterior sampling mechanism. As the number of rounds increases, both algorithms rapidly stabilize, achieving a comparable cumulative regret level to that of the baselines, thereby validating their ability to preserve first-priority optimality.

In contrast, for the second and third objectives (Figures 2(b)-(c)), the advantage of our lexicographic design becomes evident. Both LexUCB-C and LexTS-S achieve smaller cumulative regrets compared to CLUCB and SCLTS, with the performance gap widening over time. This demonstrates that, after optimizing the dominant objective, our algorithms effectively exploit the remaining exploration capacity to improve outcomes for subsequent objectives. The hierarchical update rules in LexUCB-C and LexTS-S enable the agent to adaptively focus on secondary and tertiary objectives without violating the lexicographic priority constraints. Overall, these results confirm that the proposed algorithms not only maintain competitiveness on the leading objective but also substantially enhance performance on subordinate ones, thereby realizing true lexicographic optimization in practice.

5 Conclusion and Future Work

In this paper, we formalize and investigate the problem of safe lexicographic MOSLB, which integrates lexicographic preferences with two widely adopted safety constraints: cumulative and stage-wise constraints. We design two algorithms tailored to these settings. For the cumulative-safe setting, we develop LexUCB-C, a UCB-based algorithm that respects lexicographic order while ensuring cumulative safety. For the stage-wise-safe setting, we offer LexTS-S, a Thompson Sampling-based method that guarantees per-round safety. We provide rigorous regret analysis for both algorithms, showing that their performance matches that of existing single-objective safe bandit algorithms [Kazerouni *et al.*, 2017; Moradipari *et al.*, 2020] in terms of dependence on the dimension d , time horizon T , and safety threshold α . Notably, we demonstrate that optimizing lower-priority objectives incurs minimal additional regret to the highest-priority objective.

This work opens several promising directions for future research. First, extending our framework to nonlinear reward models or more general action spaces would broaden its applicability. Second, while our analysis provides worst-case guarantees, exploring problem-dependent algorithms that leverage structural properties to improve empirical performance is an important direction. Finally, eliminating the reliance on prior knowledge assumed in Assumptions 1-4 would enhance the practical utility of our algorithms.

²<https://movielens.org/>

Acknowledgments

The work described in this paper was supported by the Research Grants Council of the Hong Kong Special Administrative Region [GRF Project No. CityU 9043546] and the National Natural Science Foundation of China (NSFC) under Grant No. 62276223.

References

- [Abbasi-yadkori *et al.*, 2011] Yasin Abbasi-yadkori, Dávid Pál, and Csaba Szepesvári. Improved algorithms for linear stochastic bandits. In *Advances in Neural Information Processing Systems 24*, pages 2312–2320, 2011.
- [Abeille and Lazaric, 2017] Marc Abeille and Alessandro Lazaric. Linear thompson sampling revisited. In *Proceedings of the 20th International Conference on Artificial Intelligence and Statistics*, pages 176–184, 2017.
- [Abernethy *et al.*, 2024] Jacob A. Abernethy, Robert Schapire, and Umar Syed. Lexicographic optimization: Algorithms and stability. In *Proceedings of The 27th International Conference on Artificial Intelligence and Statistics*, pages 2503–2511, 2024.
- [Auer *et al.*, 2002] Peter Auer, Nicolò Cesa-Bianchi, and Paul Fischer. Finite-time analysis of the multiarmed bandit problem. *Machine Learning*, 47(2–3):235–256, 2002.
- [Auer *et al.*, 2016] Peter Auer, Chao-Kai Chiang, Ronald Ortner, and Madalina Drugan. Pareto front identification from stochastic bandit feedback. In *Proceedings of the 19th International Conference on Artificial Intelligence and Statistics*, pages 939–947, 2016.
- [Auer, 2002] Peter Auer. Using confidence bounds for exploitation-exploration trade-offs. *Journal of Machine Learning Research*, 3(11):397–422, 2002.
- [Bubeck and Cesa-Bianchi, 2012] Sébastien Bubeck and Nicolò Cesa-Bianchi. Regret analysis of stochastic and nonstochastic multi-armed bandit problems. *Foundations and Trends in Machine Learning*, 5(1):1–122, 2012.
- [Chapelle and Li, 2011] Olivier Chapelle and Lihong Li. An empirical evaluation of thompson sampling. In *Advances in Neural Information Processing Systems 24*, pages 2249–2257, 2011.
- [Cheng *et al.*, 2026] Ji Cheng, Song Lai, Shunyu Yao, and Bo Xue. Offline multi-objective bandits: From logged data to pareto-optimal policies. *Proceedings of the 40th Annual AAAI Conference on Artificial Intelligence*, pages 36636–36644, 2026.
- [Dani *et al.*, 2008] Varsha Dani, Thomas P. Hayes, and Sham M. Kakade. Stochastic linear optimization under bandit feedback. In *Proceedings of the 21st Annual Conference on Learning*, pages 355–366, 2008.
- [Drugan and Nowe, 2013] Madalina M. Drugan and Ann Nowe. Designing multi-objective multi-armed bandits algorithms: A study. In *The 2013 International Joint Conference on Neural Networks*, pages 1–8, 2013.
- [Ehrgott, 2005] Matthias Ehrgott. *Multicriteria Optimization*. Springer-Verlag, Berlin, Heidelberg, 2005.
- [Harper and Konstan, 2015] F. Maxwell Harper and Joseph A. Konstan. The movielens datasets: History and context. *ACM Transactions on Interactive Intelligent Systems*, 5(4):1–19, 2015.
- [Hüyük and Tekin, 2021] Alihan Hüyük and Cem Tekin. Multi-objective multi-armed bandit with lexicographically ordered and satisficing objectives. *Machine Learning*, 110(6):1233–1266, 2021.
- [Jun and Kim, 2024] Kwang-Sung Jun and Jungtaek Kim. Noise-adaptive confidence sets for linear bandits and application to Bayesian optimization. In *Proceedings of the 41st International Conference on Machine Learning*, pages 22643–22671, 2024.
- [Katariya *et al.*, 2019] Sumeet Katariya, Branislav Kveton, Zheng Wen, and Vamsi K. Potluru. Conservative exploration using interleaving. In *Proceedings of the 22nd International Conference on Artificial Intelligence and Statistics*, pages 954–963, 2019.
- [Kazerouni *et al.*, 2017] Abbas Kazerouni, Mohammad Ghavamzadeh, Yasin Abbasi-Yadkori, and Benjamin Van Roy. Conservative contextual linear bandits. In *Advances in Neural Information Processing Systems 30*, page 3913–3922, 2017.
- [Lai and Robbins, 1985] Tze Leung Lai and Herbert Robbins. Asymptotically efficient adaptive allocation rules. *Advances in Applied Mathematics*, 6(1):4–22, 1985.
- [Lattimore and Szepesvári, 2020] Tor Lattimore and Csaba Szepesvári. *Bandit Algorithms*. Cambridge University Press, 2020.
- [Li *et al.*, 2010] Lihong Li, Wei Chu, John Langford, and Robert E. Schapire. A contextual-bandit approach to personalized news article recommendation. In *Proceedings of the 19th International Conference on World Wide Web*, pages 661–670, 2010.
- [Li *et al.*, 2025] Zhuohua Li, Maoli Liu, Xiangxiang Dai, and John C.S. Lui. Towards efficient conversational recommendations: Expected value of information meets bandit learning. In *Proceedings of the ACM on Web Conference 2025*, pages 4226–4238, 2025.
- [Lin *et al.*, 2022] Jiabin Lin, Xian Yeow Lee, Talukder Jubery, Shana Moothedath, Soumik Sarkar, and Baskar Ganapathysubramanian. Stochastic conservative contextual linear bandits. In *IEEE 61st Conference on Decision and Control*, pages 7321–7326, 2022.
- [Lu *et al.*, 2019] Shiyin Lu, Guanghui Wang, Yao Hu, and Lijun Zhang. Multi-objective generalized linear bandits. In *Proceedings of the 28th International Joint Conference on Artificial Intelligence*, pages 3080–3086, 2019.
- [Magureanu *et al.*, 2014] Stefan Magureanu, Richard Combes, and Alexandre Proutiere. Lipschitz bandits: Regret lower bound and optimal algorithms. In *Proceedings of the 27th Conference on Learning Theory*, pages 975–999, 2014.

- [Moradipari *et al.*, 2020] Ahmadreza Moradipari, Christos Thrampoulidis, and Mahnoosh Alizadeh. Stage-wise conservative linear bandits. In *Advances in Neural Information Processing Systems 33*, pages 11191–11201, 2020.
- [Rodriguez *et al.*, 2012] Mario Rodriguez, Christian Posse, and Ethan Zhang. Multiple objective optimization in recommender systems. In *Proceedings of the 6th ACM Conference on Recommender Systems*, pages 11–18, 2012.
- [Skalse *et al.*, 2022] Joar Skalse, Lewis Hammond, Charlie Griffin, and Alessandro Abate. Lexicographic multi-objective reinforcement learning. In *Proceedings of the 31st International Joint Conference on Artificial Intelligence*, pages 3430–3436, 2022.
- [Tekin, 2019] Cem Tekin. The biobjective multiarmed bandit: learning approximate lexicographic optimal allocations. *Turkish Journal of Electrical Engineering and Computer Sciences*, 27(2):1065–1080, 2019.
- [Thompson, 1933] William R. Thompson. On the likelihood that one unknown probability exceeds another in view of the evidence of two samples. *Biometrika*, 25:285–294, 1933.
- [Wanigasekara *et al.*, 2019] Nirandika Wanigasekara, Yuxuan Liang, Siong Thye Goh, Ye Liu, Joseph Jay Williams, and David S. Rosenblum. Learning multi-objective rewards and user utility function in contextual bandits for personalized ranking. In *Proceedings of the 28th International Joint Conference on Artificial Intelligence*, pages 3835–3841, 2019.
- [Wu *et al.*, 2016] Yifan Wu, Roshan Shariff, Tor Lattimore, and Csaba Szepesvári. Conservative bandits. In *Proceedings of the 33rd International Conference on International Conference on Machine Learning*, page 1254–1262, 2016.
- [Xu *et al.*, 2023] Ruitu Xu, Yifei Min, and Tianhao Wang. Noise-adaptive thompson sampling for linear contextual bandits. In *Advances in Neural Information Processing Systems 36*, pages 23630–23657, 2023.
- [Xue *et al.*, 2020] Bo Xue, Guanghui Wang, Yimu Wang, and Lijun Zhang. Nearly optimal regret for stochastic linear bandits with heavy-tailed payoffs. In *Proceedings of the 29th International Joint Conference on Artificial Intelligence*, pages 2936–2942, 2020.
- [Xue *et al.*, 2023] Bo Xue, Yimu Wang, Yuanyu Wan, Jinfeng Yi, and Lijun Zhang. Efficient algorithms for generalized linear bandits with heavy-tailed rewards. In *Advances in Neural Information Processing Systems 36*, pages 70880–70891, 2023.
- [Xue *et al.*, 2024] Bo Xue, Ji Cheng, Fei Liu, Yimu Wang, and Qingfu Zhang. Multiobjective lipschitz bandits under lexicographic ordering. *Proceedings of the 38th AAAI Conference on Artificial Intelligence*, pages 16238–16246, 2024.
- [Xue *et al.*, 2025a] Bo Xue, Xi Lin, Yuanyu Wan, and Qingfu Zhang. Problem-dependent regret for lexicographic multi-armed bandits with adversarial corruptions. In *Proceedings of the 35th International Joint Conference on Artificial Intelligence*, pages 6776–6784, 2025.
- [Xue *et al.*, 2025b] Bo Xue, Xi Lin, Xiaoyuan Zhang, and Qingfu Zhang. Multiple trade-offs: An improved approach for lexicographic linear bandits. *Proceedings of the 39th AAAI Conference on Artificial Intelligence*, pages 21850–21858, 2025.
- [Xue *et al.*, 2026] Bo Xue, Yuanyu Wan, Zhichao Lu, and Qingfu Zhang. Beyond the lower bound: Bridging regret minimization and best arm identification in lexicographic bandits. *Proceedings of the 40th Annual AAAI Conference on Artificial Intelligence*, pages 27414–27422, 2026.
- [Zhang *et al.*, 2023] Hengrui Zhang, Youfang Lin, Sheng Han, and Kai Lv. Lexicographic actor-critic deep reinforcement learning for urban autonomous driving. *IEEE Transactions on Vehicular Technology*, 72(4):4308–4319, 2023.