
Preference-Based Dynamic Ranking Structure Recognition

Nan Lu^{1,2}
nlu99007@gmail.com

Jian Shi^{1,2,*}
jshi@iss.ac.cn

Xin-Yu Tian³
tianx@umn.edu

¹State Key Laboratory of Mathematical Sciences, Academy of Mathematics and Systems Science,
Chinese Academy of Sciences, Beijing, China

²School of Mathematical Sciences, University of Chinese Academy of Sciences, Beijing, China

³School of Statistics, University of Minnesota, Minneapolis, USA

Abstract

Preference-based data often appear complex and noisy but may conceal underlying homogeneous structures. This paper introduces a novel framework of ranking structure recognition for preference-based data. We first develop an approach to identify dynamic ranking groups by incorporating temporal penalties into a spectral estimation for the celebrated Bradley-Terry model. To detect structural changes, we introduce an innovative objective function and present a practicable algorithm based on dynamic programming. Theoretically, we establish the consistency of ranking group recognition by exploiting properties of a random ‘design matrix’ induced by a reversible Markov chain. We also tailor a group inverse technique to quantify the uncertainty in item ability estimates. Additionally, we prove the consistency of structure change recognition, ensuring the robustness of the proposed framework. Experiments on both synthetic and real-world datasets demonstrate the practical utility and interpretability of our approach.

1 Introduction

Preference-based data, where observations arise from pairwise or groupwise comparisons rather than absolute measurements, is prevalent across various domains. This form of data naturally appears in applications such as economics [1], online recommendations [2], and sports analytics [3]. In addition, the use of preference-based data in reinforcement learning from human feedback (RLHF) has led to significant improvements in the performance of large language models [4]. One major advantage of preference-based data lies in its ease of collection, as it is often more intuitive to express relative preferences rather than assign absolute scores. Many widely used datasets are inherently preference-based, making their effective modeling and analysis a pivotal research focus. To handle such data, the celebrated Bradley-Terry model [5] is widely used for inferring latent preference scores from pairwise comparisons. This model and its extensions have been extensively studied; see [6–9].

Ranking serves as a crucial tool for summarizing preference-based data, providing interpretable outcomes that facilitate decision-making in various fields. It has broad applications, including the evaluation of sports teams [10], institutions [11, 12], recommendation systems [13, 14], financial markets [15, 16], and bioinformatics [17, 18]. By leveraging ranking positions, comparison results enable the identification of top-performing entities [19, 20] while also uncovering underlying trends and patterns [21, 22]. Items to be ranked often possess latent structures due to population homogeneity, which can be reflected in phenomena such as circular comparison results. Moreover, even a slight

*Corresponding author.

Code available at <https://github.com/nanlu99/RankStruct>.

modification in comparisons can result in a different rank [23], highlighting the importance of grouped rankings. Grouping similar items can enhance robustness and reduce sensitivity to specific comparisons. For example, group rankings recognize homogeneous entities to improve interpretability and predictive accuracy [10, 24]. In the context of university rankings, Soh [25] argues that minor differences in scores should be ignored and that similar institutions should be assigned to the same group. Given time-varying comparison results, we aim to address the following questions:

- Which items exhibit similar behavior and can be categorized into the same group during a specific period?
- What is the ranking order of these groups at a particular time point?

When considering the temporal dimension, we often encounter situations where item groups evolve over time. For example, in basketball, player trades and coaching changes can significantly affect a team’s performance, potentially elevating it to a higher ranking tier. Similarly, the share prices of certain companies may surge due to emerging political, technological, or market factors. These rapidly evolving scenarios underscore the importance of detecting structural change points for long-term analysis. In this context, the term *group change* refers to shifts in group membership over time, which are crucial for accurate analysis in the dynamic grouping problem. In this work, we take a closer look at changes in group structures and aim to address another critical question:

- When does the underlying cluster structure experience significant changes?

There have been several studies on grouping methods for ranking problems in the BT model. Masarotto and Varin [10] first apply the fused lasso penalty to the maximum likelihood estimation. Vana et al. [26] utilize a similar method for journal meta-ranking, and Jeon and Choi [27] extend it to the Luce model. Tian and Shi [28] further consider the problem using the spectral method. However, it is worth noting that all these grouping methods are designed for static situations. In practice, the latent abilities of sports players and institutions may vary over time. Treating data as if it were all collected simultaneously can lead to misleading results. For example, a player in his rising period and another in a declining period may exhibit similar average performances in a game season, but they should not be classified as the same. Hence, this paper concentrates on the simultaneous ranking and grouping problem for the dynamic scenario. Li et al. [3] introduce a segmented static BT model, focusing on detecting the change points of score variation. In contrast, our approach allows scores to vary continuously, and our emphasis lies in recognizing the underlying structure and the changes in clustered groups over time.

We summarize our major contributions as follows.

- **An innovative framework for ranking group recognition.** Though item abilities can be modeled using continuous functions, the ranking positions are discrete functionals of the latent abilities, posing challenges in identifying their structure. To address this, we propose a novel workflow that nests recovering dynamic ranking groups in group change recognition.
- **A generally applicable structure change detection method.** Some works study the clustering problem of different items [10, 26, 27], while few works consider the abrupt changes of item abilities [3]. To the best of our knowledge, we are the first to consider the ranking structure changes for the BT model. We carefully design an integrated objective function, which possesses separable properties, making it permissible to develop an efficient algorithm based on dynamic programming.
- **Theoretical results on recognition consistency and estimator uncertainty.** We characterize conditions of the variability within groups that ensure the consistency recognition of groups and establish the structure recognition consistency. We quantify the uncertainty of item ability estimators using an innovative group inverse technique.
- **Ranking results with enhanced interpretability and improved accuracy.** Our method provides concise ranking results and group change information. The structured ranking results enable the identification of homogeneous items and dynamic group changes, which provide insights into the underlying structure. Simulation results also demonstrate that the grouping method integrates data effectively, yielding improved estimation accuracy.

Notations We write $a_n \lesssim b_n$ or $a_n = O(b_n)$ if there exists a constant $c > 0$ such that $a_n \leq cb_n$ for all n . We denote $a_n \asymp b_n$ if $b_n \lesssim a_n$ and $a_n \lesssim b_n$. Besides, we write $a_n = o(b_n)$ or $a_n \ll b_n$ if $\lim_{n \rightarrow \infty} a_n/b_n = 0$. We denote by $[n] = 1, \dots, n$ for any positive integer n . We let $\mathbf{I}$ represent

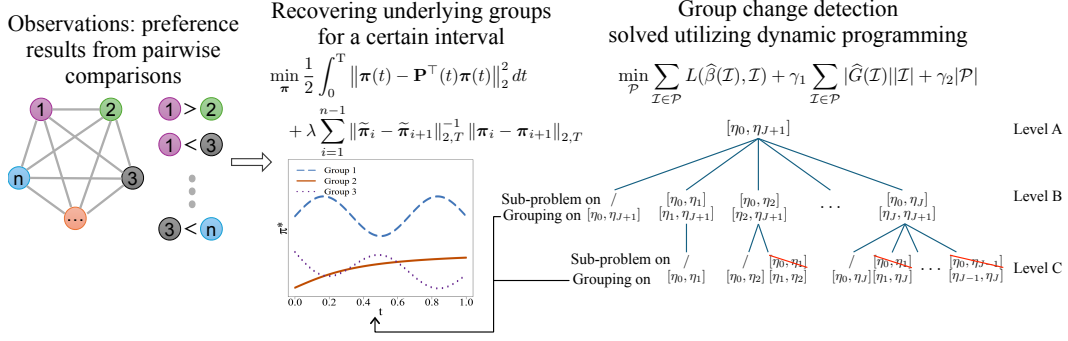


Figure 1: Workflow of the proposed method for dynamic ranking structure recognition. The left panel depicts the data format, representing preference outcomes obtained from pairwise comparisons. The middle panel illustrates the score evolution of each group over a given interval. The right panel shows the structure detection procedure, where: (1) each node can be decomposed into subproblems represented by its child nodes; (2) each child node corresponds to a subproblem and a grouping problem on intervals; (3) subproblems at level C reuse results from nodes preceding their parent node at level B, thereby avoiding redundant computation.

the identity matrix, and let e_n be an $n \times 1$ vector with each element equal to 1. $\mathbf{0}$ represents the vector or matrix composed entirely of zeros. For a vector v , $\|v\|_2$ denotes the ℓ_2 -norm. We let $\|f\|_{2,T} = (\int_0^T f(t)^2 dt)^{1/2}$, where $f(t)$ is a function on $[0, T]$.

2 Ranking Structure Recognition

2.1 Dynamic Bradley-Terry Model

In a dynamic scenario, we observe pairwise comparison results among n items, denoted as $\mathcal{Y} = \{y_{ij}(t_k), i, j \in [n], t_k \in T_{ij}\}$. The scalar $y_{ij}(t)$ represent the comparison result at time point t between items i and j , where $y_{ij}(t) = 1$ represents that item i wins over item j , and $y_{ij}(t) = 0$ indicates the opposite. We assume that the elements of $\mathcal{Y}$ are independent, and the comparison time T_{ij} of (i, j) pair is uniformly distributed over $[0, T]$. The Bradley-Terry model assigns positive scores $\pi^*(t) = (\pi_1^*(t), \pi_2^*(t), \dots, \pi_n^*(t))^\top$ to items, and presumes $y_{ij}(t) \sim \text{Bernoulli}(y_{ij}^*(t))$, where $y_{ij}^*(t) = \pi_j^*(t) / (\pi_i^*(t) + \pi_j^*(t))$ [5]. Intuitively, taking $\pi_i^*(t)$ as the ability of item i , $y_{ij}^*(t)$ represents the winning rate of item i . Notice that the BT model is invariant under the scaling of the scores, so we set $\sum_{i=1}^n \pi_i^*(t) = 1$ for all $t \in [0, T]$ to obtain a unique representation.

A straightforward approach to estimate the BT model is the maximum likelihood estimator [29, 30]. In pursuit of a computationally efficient solution, we opt for the spectral-based solver [6, 22]. Negahban et al. [6] present an insightful perspective of the spectral method, establishing a connection between the pairwise comparison results and the transition of a Markov chain. Specifically, by letting the nodes of a graph represent the items, and assigning the transformation probability from node i to j based on the frequency of item i losing to j , it is proven that the stationary distribution of this random walk corresponds to the items' abilities π^* . For a more detailed understanding, we recommend referring to Negahban et al. [6].

We adopt the kernel-based estimator. Let $K_h(t, s) = \frac{1}{h} K\left(\frac{t-s}{h}\right)$, where $K(\cdot)$ is the kernel function and h is the bandwidth. The transformation probability matrix $\mathbf{P}(t)$ is formulated as

$$\mathbf{P}_{ij}(t) = \begin{cases} \frac{\frac{1}{n} \sum_{t_k \in T_{ij}} y_{ij}(t_k) K_h(t, t_k)}{\sum_{t_k \in T_{ij}} K_h(t, t_k)} & \text{if } i \neq j, \\ 1 - \sum_{s \neq i} \mathbf{P}_{is}(t) & \text{if } i = j. \end{cases}$$

Then we have the consistent estimator $\tilde{\pi}(t)$, which is a stationary distribution of the Markov chain deduced by the stochastic matrix $\mathbf{P}$ [22].

2.2 Recognition of Dynamic Ranking Groups

We first consider recovering the ranking groups for a given interval in this section, which is a cornerstone for analyzing the changes of ranking groups for a relatively longer time period as discussed in Section 2.3.

Let B denote the number of groups. We represent these groups as $G = \{G_1, G_2, \dots, G_B\}$, forming a partition of the set $[n]$. The items within the same group possess similar scores, while those from different groups have significant score differences. Formally, we have

$$\delta_1 := \min_{\substack{k, l \in [B] \\ k \neq l}} \min_{\substack{i \in G_k \\ j \in G_l}} \|\pi_i^* - \pi_j^*\|_{2,T} \gg \max_{k \in [B]} \max_{i, j \in G_k} \|\pi_i^* - \pi_j^*\|_{2,T}.$$

Without loss of generality, we assume that $\max\{i : i \in G_k\} < \min\{i : i \in G_l\}$ for $k < l$.

We then recover the partition of items and present the score estimations simultaneously. Since any finite-state time-homogeneous Markov chain has at least one stationary distribution, we rewrite the estimator $\tilde{\pi}(t_0)$ as the solution of the optimization problem, $\min_{\tilde{\pi}(t_0)} \|\tilde{\pi}(t_0) - \mathbf{P}^\top(t_0)\tilde{\pi}(t_0)\|_2$ such that $\sum_{i=1}^n \tilde{\pi}_i(t_0) = 1$. Setting λ as a tuning parameter, we consider the following objective function.

$$\begin{aligned} \min_{\pi} \quad & \frac{1}{2} \int_0^T \|\pi(t) - \mathbf{P}^\top(t)\pi(t)\|_2^2 dt + \lambda \sum_{i=1}^{n-1} \|\tilde{\pi}_i - \tilde{\pi}_{i+1}\|_{2,T}^{-1} \|\pi_i - \pi_{i+1}\|_{2,T} \\ \text{s.t.} \quad & \sum_{i=1}^n \pi_i(t_k) = 1, \quad k = 1, 2, \dots, m. \end{aligned} \quad (2.1)$$

Let $t_1, t_2, \dots, t_m$ be m equidistant time points in $[0, T]$. We use the symbol with an item subscript, such as π_i , to represent the vector corresponding to the m time points $(\pi_i(t_1), \pi_i(t_2), \dots, \pi_i(t_m))^\top$. Here, the parameter m is allowed to approach infinity, allowing for the approximation of the integral. To efficiently address the constrained optimization problem, we employ a technical transformation, leading us to an unconstrained form with a well-developed optimization algorithm. Define the $n \times n$ matrix

$$\mathbf{Q}_{ij} = \begin{cases} 1 & \text{if } i = j \text{ or } i = n, \\ -1 & \text{if } i < n \text{ and } j = i + 1, \\ 0 & \text{otherwise.} \end{cases}$$

Let $\theta(t) = \mathbf{Q}(\pi(t) - \frac{1}{n}\mathbf{e}_n)$ and $\tilde{\theta} = \mathbf{Q}(\tilde{\pi}(t) - \frac{1}{n}\mathbf{e}_n)$. Let $\theta(t) = (\theta_1(t), \dots, \theta_{n-1}(t))^\top$ and $\theta = (\theta(t_1)^\top, \dots, \theta(t_m)^\top)^\top$. Let θ^* be the corresponding true value (with π substituted by π^*), and $\tilde{\theta}$ denote the counterpart induced by $\tilde{\pi}$. We define $\mathbf{X}(t) = (\mathbf{P}^\top(t) - \mathbf{I})\mathbf{Q}^{-1}$ and $\mathbf{Y}(t) = \frac{1}{n}(\mathbf{I} - \mathbf{P}^\top(t))\mathbf{e}_n$. Then we have the optimization problem (2.1) reformulated as

$$\min_{\theta} \frac{1}{2} \|\mathbf{Y} - \mathbf{X}\theta\|_2^2 + \lambda \sum_{i=1}^{n-1} \|\tilde{\theta}_i\|_2^{-1} \|\theta_i\|_2, \quad (2.2)$$

where $\mathbf{X}$ is the $mn \times m(n-1)$ matrix $\text{diag}(\mathbf{X}_{-1}(t_1), \dots, \mathbf{X}_{-1}(t_m))$, $\mathbf{X}_{-1}(t)$ is the matrix $\mathbf{X}(t)$ with its last column removed. $\mathbf{Y}$ represents the $mn \times 1$ vector, $(\mathbf{Y}(t_1), \dots, \mathbf{Y}(t_m))^\top$. This transformation directly eliminates the constraints, reducing the optimization objective to a standard adaptive group lasso problem, which possesses efficient solutions. Having obtained the solution $\hat{\theta}$, we can calculate $\hat{\pi}(t) = \mathbf{Q}^{-1}(\hat{\theta}(t)^\top, 0)^\top + \frac{1}{n}\mathbf{e}_n$. Let $\mathcal{S} = \{i : \theta_i^* \neq 0\}$, $\hat{\mathcal{S}} = \{i : \hat{\theta}_i \neq 0\}$ and $\hat{B} = |\hat{\mathcal{S}}| + 1$. We use $\tilde{\mathcal{S}}$ to denote the estimated partition points of different groups. Specifically, we let $\tilde{\mathcal{S}} = \{0\} \cup \hat{\mathcal{S}} \cup \{n\}$. Without loss of generality, we assume $\tilde{\mathcal{S}}$ is arranged in ascending order (if not, we simply reorder $\hat{\mathcal{S}}$), and $\tilde{\mathcal{S}}_i$ denotes the i -th element of $\tilde{\mathcal{S}}$. The group estimation $\hat{G} = \{\hat{G}_1, \hat{G}_2, \dots, \hat{G}_{\hat{B}}\}$ is obtained by $\hat{G}_k = \{i : \tilde{\mathcal{S}}_{k-1} < i \leq \tilde{\mathcal{S}}_k\}$.

Remark 2.1. We note that our approach and the theoretical justification presented below do not rigidly require that all items within a group possess identical scores. Instead, we establish a framework in which items within a group exhibit similar behavior, rendering practically flexibility.

Remark 2.2. Though we originally have the fused term among different items as in (2.1), the optimization objective has the same expression as the adaptive group lasso of a linear regression model in (2.2). It is worth noting that the similarity is somehow superficial since the design matrix $\mathbf{X}$ is no longer deterministic, posing difficulties for theoretical analysis.

To mitigate the issue of shrinkage in large coefficients resulting from the penalization term, a widely utilized approach is the refit procedure [31, 32]. This method entails re-estimating the coefficients after identifying the underlying structure. Upon obtaining the group estimation $\hat{G}$, we employ the refit strategy in the following manner. For $i, j \in [\hat{B}]$, define

$$\mathbf{P}_{\hat{G}ij}(t) = \begin{cases} \frac{1}{\hat{B}} \frac{\sum_{l_1 \in \hat{G}_i} \sum_{l_2 \in \hat{G}_j} \sum_{t_k \in T_{l_1 l_2}} y_{l_1 l_2}(t_k) K_h(t, t_k)}{\sum_{l_1 \in \hat{G}_i} \sum_{l_2 \in \hat{G}_j} \sum_{t_k \in T_{l_1 l_2}} K_h(t, t_k)}, & \text{if } i \neq j; \\ 1 - \sum_{s \neq i} \mathbf{P}_{\hat{G}is}(t), & \text{if } i = j. \end{cases}$$

We can obtain the stationary distribution $\hat{\pi}_{\hat{G}} = (\hat{\pi}_{\hat{G}_i})_{i \in [\hat{B}]}$ of $\mathbf{P}_{\hat{G}}(t)$. Note that we have assumed the score summation of n items to be 1 to eliminate the non-uniqueness caused by rescaling. Therefore, the refit estimator for each item is

$$\hat{\pi}_i^{rf}(t) = \frac{\hat{\pi}_{\hat{G}_l}(t)}{\sum_{k \in [\hat{G}]} |\hat{G}_k| \hat{\pi}_{\hat{G}_k}(t)}, \quad i \in \hat{G}_l. \quad (2.3)$$

Remark 2.3. We use the absolute group size for normalization because, as stated in Section 2.1, we impose the constraint that the sum of item abilities equals 1, i.e., $\sum_{i=1}^n \pi_i^* = 1$, to ensure a unique representation. Since only the ratio between scores matters in the BT model, this constraint guarantees identifiability. Furthermore, when recovering the original scores after refitting, we still expect the normalized scores to satisfy $\sum_{i=1}^n \hat{\pi}_i^{rf} = 1$. At the same time, we need to preserve the score ratios between items from different groups, meaning that for $i \in \hat{G}_l$ and $j \in \hat{G}_k$, $\hat{\pi}_i^{rf} / (\hat{\pi}_i^{rf} + \hat{\pi}_j^{rf}) = \hat{\pi}_{\hat{G}_l} / (\hat{\pi}_{\hat{G}_l} + \hat{\pi}_{\hat{G}_k})$. To satisfy both conditions simultaneously, the normalization in equation (2.3) is scaled by the absolute group size $|\hat{G}_k|$.

Remark 2.4. The refit strategy is an optional part. Treating the comparison result of items in a group as one actually compensates for more information, especially in cases where n is large and Mh is small. We also point out that refitting induces better performance, as indicated in simulations.

2.3 Recognition of Group Changes

We then focus on detecting the change points of latent clusters over an extended period. Consider a scenario with time-correlated observations occurring within the interval $[0, V]$. Still consider n entities whose structure needs to be determined. There are $J + 1$ phases, where the items' groups remain unchanged within each phase and differ between adjacent phases. In a more formal mathematical form, let $\mathbf{z}(t) = (z_1(t), \dots, z_n(t))^T$ represent the latent group of items at the time point t . There are $J + 2$ points $0 = \eta_0 < \dots < \eta_{J+1} = V$ such that $\mathbf{z}(t) \neq \mathbf{z}(s)$ for $\eta_{i-1} < t < \eta_i < s < \eta_{i+1}$, $i \in [J]$ and $\mathbf{z}(t) = \mathbf{z}(s)$ for $\eta_{i-1} < t, s < \eta_i$, $i \in [J + 1]$. The unobservable structure change points $\{\eta_i\}_{i \in [J]}$ belong to a preset candidate set $\{\xi_i\}_{i \in [U]}$. Without loss of generality, let $\{\xi_i\}_{i \in [U]}$ be in increasing order. In practice, the candidate set may be selected based on practical considerations, such as dividing points between seasons in sports games or uniformly distributed time points.

To detect changes in underlying groups, it is necessary to employ clustering methods within a subinterval $\mathcal{I} \subset [0, V]$. We utilize the clustering method proposed in Section 2.2 for dynamic ranking, by simply substituting $[0, T]$ with $\mathcal{I}$. Let $\hat{G}(\mathcal{I})$ represent the estimated group corresponding to the true structure $G(\mathcal{I})$. With a slight abuse of notations, let $\hat{\beta}(\mathcal{I})$ denote the model parameter estimations $\{\hat{\pi}_i^{rf}(t), i \in [n], t \in \mathcal{I}\}$, and let $\beta(\mathcal{I}) = \{\beta(t), t \in \mathcal{I}\}$ be the corresponding true values. Define $\bar{y}_{ij}(t) = \frac{\sum_{k=1}^M y_{ij}(t_k) K_h(t, t_k)}{\sum_{k=1}^M K_h(t, t_k)}$. We introduce the negative log-likelihood function for $\boldsymbol{\pi} = (\pi_1, \dots, \pi_n)^T$ at a time point t ,

$$l(\boldsymbol{\pi}, t) = -\frac{2}{n(n-1)} \sum_{(i,j): i \neq j} \bar{y}_{ij}(t) \log\left(\frac{\pi_j}{\pi_i + \pi_j}\right), \quad (2.4)$$

which is a natural extension of the static case. Define the function $L(\hat{\beta}(\mathcal{I}), \mathcal{I}) = \int_{t \in \mathcal{I}} l(\hat{\beta}(t), t) dt$ to measure the discrepancy between observed samples and the values expected under the grouping model.

Let $\mathcal{P}$ represents $\{[s_0, s_1], [s_1, s_2], \dots, [s_p, s_{p+1}]\}$, with $s_0 = 0$, $s_{p+1} = V$ and $\{s_i\}_{i \in [p]} \subset \{\xi_i\}_{i \in [U]}$ being a list of increasing points. We use $|\hat{G}(\mathcal{I})|$ to represent the estimated group number

and $|\mathcal{I}|$ to denote the interval length. We recover the change points of structures by considering the objective function:

$$\min_{\mathcal{P}} \sum_{\mathcal{I} \in \mathcal{P}} L(\hat{\beta}(\mathcal{I}), \mathcal{I}) + \gamma_1 \sum_{\mathcal{I} \in \mathcal{P}} |\hat{G}(\mathcal{I})| |\mathcal{I}| + \gamma_2 |\mathcal{P}|. \quad (2.5)$$

Intuitively, the first term evaluates the goodness of fit for the parameters, the second term is the penalty of groups and the last term imposes a penalty on phase changes.

We provide a brief clarification that the framework exhibits versatility. It is not confined to the ranking problem but can be applied to the general detection of group changes. As long as a clustering method designed for subintervals is provided, the framework can effectively perform the structure change detection. Specifically, it relies on $\hat{G}(\mathcal{I})$ to present the grouping results and $L(\hat{\beta}(\mathcal{I}), \mathcal{I})$ to assess the goodness of fit of the grouping method. Besides negative log-likelihood functions, $l(\cdot)$ can be residuals or other measurements, determined by the specific problem.

Note that the optimization objective (2.5) exhibits separability with respect to time and features an optimal substructure property. Specifically, the objective has an additive form across time, which allows it to be decomposed into independent subproblems separable over time. Moreover, the optimal solution to the overall problem can be constructed from the optimal solutions of its subproblems. These two properties enable the objective to be optimized recursively, forming the basis of an efficient dynamic programming solution. We can address the combinatorial problem using Algorithm 1, which provides an efficient method to estimate $\mathcal{R} = \{\hat{s}_i\}_{i \in [J]}$.

Algorithm 1 Structure Change Detection

Require: Observed data $\mathcal{Y}$, tuning parameters γ_1, γ_2 .

Ensure: Change points estimation $\mathcal{R}$.

$R = \emptyset, \mathbf{a} = -\mathbf{e}_{U+1}, \mathbf{b} = (\infty, \dots, \infty) \in \mathbb{R}^{U+1}, b_0 = 0, \xi_0 = 0, \xi_{U+1} = V.$

for r from 1 to $U + 1$ **do**

for l from 0 to $r - 1$ **do**

$b \leftarrow b_l + L(\hat{\beta}(\mathcal{I}), \mathcal{I}) + \gamma_1 |\hat{G}(\mathcal{I})| |\mathcal{I}| + \gamma_2$, where $\mathcal{I} = [\xi_l, \xi_r]$.

if $b < b_r$ **then**

$b_r \leftarrow b; \mathbf{a}_r \leftarrow l.$

$k \leftarrow U + 1$

while $k > 0$ **do**

$d \leftarrow \mathbf{a}_k; \mathcal{R} = \mathcal{R} \cup \{\xi_d\}; k \leftarrow d.$

return $\mathcal{R}$

3 Statistical Learning Theory

3.1 Consistency of Ranking Group Estimation

In this section, we present theoretical guarantees for our estimator. Specifically, we show that the probability of correctly identifying the underlying group structure approaches one as the sample size increases. We refer to this property as *group consistency*, a desirable feature that supports the reliability of the proposed estimation method. To ensure the theoretical results, we introduce the following assumptions.

Assumption 3.1. $\sup_{t \in [0, T]} \frac{\max_i \pi_i^*(t)}{\min_i \pi_i^*(t)} \leq \kappa$, where $\kappa > 0$ is a constant. $\pi_i^*(t)$ is three times continuously differentiable, $i \in [n]$.

Assumption 3.2. The kernel function is symmetric, nonnegative, and satisfies $\int_{-\infty}^{\infty} K(v) dv = 1$ and $\int_{-\infty}^{\infty} v^2 K(v) dv < \infty$.

Assumptions 3.1 and 3.2 are commonly used in the BT model and kernel methods [30, 19]. We let $|T_{ij}| = M$ for $i, j \in [n]$. We note that our method applies to the case where the number of comparisons may vary over time, and we assume a constant number of comparisons only for the simplicity of presentation. Recall that $\mathcal{S} = \{i : \theta_i^* \neq \mathbf{0}\}$. Let $\delta_2 \geq 0$ be a constant such that

$|\theta_i^*(t)| \leq \delta_2, \forall t \in T, i \in \mathcal{S}^c$. Define n_i as the number of items in G_i , and let $r_i = n_i/n$. Assume $r_i \asymp 1/B, i \in [B]$. Let $\delta = \sqrt{\frac{\log(nM)}{n^3 Mh}}$, which denotes the uniform convergence rate of the KRC estimator [19].

Theorem 3.3. Let Assumptions 3.1 and 3.2 hold. When $Mh \rightarrow \infty, n \rightarrow \infty$ and $nMh^5 \rightarrow 0$, if

1. $\max\{\delta, \frac{1}{m}, \sqrt{\frac{B}{n^3 Mh}}\} = o(\delta_1)$ and $\delta_2 = o(\sqrt{\frac{1+\cos(\frac{(n-B)\pi}{n-B+1})}{B(n-B)} \frac{1}{n^2 Mh}})$; 2. $\sqrt{\frac{m^2 B^3}{nMh}} \tilde{\delta} \ll \lambda \lesssim \frac{\delta_1 \sqrt{m}}{B\sqrt{nMh}}$, where $\tilde{\delta} = \max\{\delta, \delta_2\}$, then we have $\mathbf{P}(\hat{G} = G) \rightarrow 1$.

We have established the group consistency property. The following two remarks clarify the conditions of the theorem and highlight the distinct features of our theoretical analysis.

Remark 3.4. The first condition characterizes the requirement for δ_1 to recognize the differences among groups without being impeded by estimation errors, while the requirement of δ_2 limits the variation within each group to ensure accurate item ability estimation. The second condition requires the appropriate order of the penalized parameter λ . The term $1/m$ denotes the order of integral approximation error for $\|\tilde{\theta}_i\|_2$ and is not essential. If the midpoint approximation is replaced by the trapezoidal rule, then $1/m$ is replaced by $1/m^2$.

Remark 3.5. Unlike standard linear regression, the design matrix $\mathbf{X}$ in this context is derived from a series of transformations on the observed data y . This introduces challenges for theoretical analysis. Fortunately, $\mathbf{P}(t)$ is an approximation of a reversible Markov transition matrix $\mathbf{P}^*(t)$ (see Section C.1), where

$$\mathbf{P}_{ij}^*(t) = \begin{cases} \frac{1}{n} \frac{\pi_j^*(t)}{\pi_i^*(t) + \pi_j^*(t)} & \text{if } i \neq j; \\ 1 - \sum_{s \neq i} \mathbf{P}_{is}^*(t) & \text{if } i = j. \end{cases}$$

That plays an important role in deducing the properties of $\mathbf{X}$ and $\mathbf{P}$ and facilitates the establishment of theoretical guarantees.

3.2 Asymptotic Distribution of Item Ability Estimates

In this section, we discuss uncertainty quantification, that is, the asymptotic distribution of the item ability estimators. A well-characterized uncertainty quantification enables statistical inference tasks such as hypothesis testing and helps assess the reliability of the estimators.

Selecting one representative item from each group, $i_1 \in G_1, i_2 \in G_2, \dots, i_B \in G_B$, let $\pi_G^*(t) = (\pi_{G_1}^*(t), \dots, \pi_{G_B}^*(t))^\top = (\pi_{i_1}^*(t), \dots, \pi_{i_B}^*(t))^\top / \sum_{k=1}^B \pi_{i_k}^*(t)$ and $\hat{\pi}_G(t) = (\hat{\pi}_{i_1}^{rf}(t), \dots, \hat{\pi}_{i_B}^{rf}(t))^\top / \sum_{k=1}^B \hat{\pi}_{i_k}^{rf}(t)$. We observe that π_G^* is the stationary distribution of the $B \times B$ matrix $\mathbf{P}_G^*(t)$, where

$$\mathbf{P}_{Gij}^*(t) = \begin{cases} \frac{\pi_{Gj}^*(t)}{\pi_{Gi}^*(t) + \pi_{Gj}^*(t)}, & \text{if } i \neq j; \\ 1 - \sum_{s \neq i} \mathbf{P}_{Gis}^*(t), & \text{if } i = j. \end{cases}$$

Set $\mathbf{A}^\#(t)$ as the group inverse of $\mathbf{I} - \mathbf{P}_G^*(t)$ (see the definition of group inverse in Kirkland and Neumann [33]). We have the following result.

Theorem 3.6. Under the conditions of Theorem 3.3, if $\delta_2 = o(\frac{1}{\sqrt{n^4 Mh}})$, for a fixed B and any $t \in (0, 1)$, we have

$$\sqrt{n^2 Mh}(\hat{\pi}_G(t) - \pi_G^*(t)) \xrightarrow{\mathcal{D}} N(0, \mathbf{\Gamma}(t) \mathbf{\Lambda}(t) \mathbf{\Gamma}(t)^\top),$$

where $\mathbf{\Lambda}(t)$ is a $\frac{B(B-1)}{2}$ diagonal matrix with $\Lambda_{kl,kl}(t) = \frac{1}{r_k r_l} \frac{\pi_{G_k}^*(t) \pi_{G_l}^*(t)}{(\pi_{G_k}^*(t) + \pi_{G_l}^*(t))^2} \int K^2(v) dv$, and $\mathbf{\Gamma}(t)$ is a $B \times \frac{B(B-1)}{2}$ matrix with $\Gamma_{i,kl}(t) = (\mathbf{A}_{li}^\#(t) - \mathbf{A}_{ki}^\#(t)) \frac{(\pi_{G_k}^*(t) + \pi_{G_l}^*(t))}{B}$, $1 \leq i \leq B$, $1 \leq k < l \leq B$.

3.3 Consistency of Group Changes Detection

In this section, we focus on *structure recognition consistency*, which refers to the property that the probability of correctly identifying group structure change points approaches one as the sample size increases. This property is important as it ensures reliable detection of structural changes. We first provide a general analysis for arbitrary grouping methods in Theorem 3.9, and then specialize the discussion to the dynamic ranking setting in Corollary 3.10.

To guarantee the correctness of the estimated change points, we impose assumptions regarding the grouping accuracy within a given time interval $\mathcal{I}$.

Assumption 3.7. As sample size tends to infinity, we have $P(\widehat{G}(\mathcal{I}) = G(\mathcal{I})) \rightarrow 1$.

Assumption 3.8. $\frac{1}{|\mathcal{I}|} |L(\beta(\mathcal{I}), \mathcal{I}) - L(\widehat{\beta}(\mathcal{I}), \mathcal{I})| = O_p(\delta_3)$, where $\delta_3 \rightarrow 0$ as sample size tends to infinity.

Assumption 3.7 is intended to ensure the accurate recovery of groups. With a consistent estimator $\widehat{\beta}(\mathcal{I})$, Assumption 3.8 can be satisfied by incorporating a sufficiently smooth $l(\cdot)$. We show in Corollary 3.10 that our method in Section 2.2 is capable of satisfying these assumptions.

Theorem 3.9. Under Assumptions 3.7 and 3.8, if $\gamma_1 > \gamma_2$, and $\delta_3 = o(\gamma_2)$, then $P(\{\widehat{s}_i\}_{i=1}^{\widehat{J}} = \{\eta_i\}_{i=1}^J) \rightarrow 1$ with sample size tending to infinity.

Intuitively, the order of γ_2 should be larger than that of δ_3 to ensure efficient penalty and γ_1 is supposed to be larger than γ_2 to avoid missing change points. Based on Theorems 3.3 and 3.9, we have the following consistency guarantee for the ranking group change detection.

Corollary 3.10. Under the conditions of Theorem 3.3, if $\gamma_1 > \gamma_2$ and $\sqrt{\frac{1}{nMh}} = o(\gamma_2)$, we have $P(\{\widehat{s}_i\}_{i=1}^{\widehat{J}} = \{\eta_i\}_{i=1}^J) \rightarrow 1$ as $n \rightarrow \infty$.

4 Computational Experiments

4.1 Recognition of Dynamic Ranking Groups

We evaluate the results using the Kendall τ coefficient and the mean squared error (MSE) between the estimators $(\widehat{\pi}(t), \widehat{\pi}^{rf}(t))$ and $\pi^*(t)$ to assess the estimation accuracies of rank and value. We employ sensitivity and specificity to gauge group accuracy. Specifically, sensitivity represents the proportion of correctly identified pairs within the same group, while specificity calculates the percentage correctly distinguished between different groups. We compare our method (without and with refit strategy) with the static clustering method Group Rank Centrality (GRC) [28] and the original estimator Kernel Rank Centrality (KRC). All experiments are conducted on a machine with an 11th Gen Intel(R) Core(TM) i5-1135G7 CPU and 16GB RAM. We utilize R package sparsegl [34] for analysis. We consider two different experimental settings. Detailed configurations and parameter choices are provided in Section B.1, and the results are summarized in Table 1.

Table 1: Simulation results for simultaneously grouping and ranking.

Kendall τ				MSE				Sensitivity		Specificity	
Ours	Ours (refit)	GRC	KRC	Ours	Ours (refit)	GRC	KRC	Ours	GRC	Ours	GRC
Setting 1 (n=20, Mh=5)											
0.9998	0.9998	0.9999	0.8330	0.0586	0.0466	0.1083	0.1277	99.95%	99.97%	99.99%	100.00%
Setting 1 (n=50, Mh=10)											
1.0000	1.0000	1.0000	0.8207	0.0421	0.0375	0.1079	0.0684	100.00%	100.00%	100.00%	100.00%
Setting 2 (n=20, Mh=5)											
0.9416	0.9484	0.7977	0.7565	0.0494	0.0359	0.2168	0.1208	100.00%	100.00%	100.00%	63.64%
Setting 2 (n=50, Mh=10)											
0.9644	0.9683	0.7974	0.7742	0.0291	0.0250	0.2167	0.0606	100.00%	100.00%	100.00%	63.71%

Since the refit strategy is based on identified groups, the results of non-refit and refit estimators exhibit the same sensitivity and specificity values. The sensitivity and specificity of KRC are not listed as it does not exhibit a grouping effect. It can be observed that the refit estimator performs slightly better than the non-refit one, showing a larger Kendall τ and a smaller MSE. The Kendall τ of our method approaches one with increasing sample size, surpassing the values of the other two methods. The results of specificity highlight the necessity of a dynamic setting. It is evident that GRC cannot distinguish some groups in the second setting. Comparing the results of both settings, the Kendall τ and MSE of our method are superior to those of KRC. This suggests that our method effectively captures group information, yielding better estimation accuracy.

4.2 Recognition of Group Changes

Due to the lack of established methods for detecting dynamic ranking structural changes, we compare our method with a naive baseline that groups items within each interval between consecutive candidate change points. A dividing point is identified as a structural change point if the groupings in its adjacent intervals differ. We evaluate the experiment results using two criteria: the number of change points and the Hausdorff distance (H-dist) between the actual and estimated sets of structural change points. Detailed configurations and parameter choices are provided in Section B.2. We summarize the results for two different settings as follows.

Table 2: Simulation results for structural change detection.

	Ours				Naive			
	H-dist	$\hat{J} < J$	$\hat{J} = J$	$\hat{J} > J$	H-dist	$\hat{J} < J$	$\hat{J} = J$	$\hat{J} > J$
Setting 1								
Mh=2	0.0051	0.0%	97.4%	2.6%	0.2084	0.0%	0.0%	100.0%
Mh=4	0.0004	0.0%	99.8%	0.2%	0.1800	0.0%	0.0%	100.0%
Mh=10	0.0000	0.0%	100.0%	0.0%	0.1333	0.0%	0.8%	99.2%
Setting 2								
Mh=4	0.0492	0.0%	80.4%	19.6%	0.3258	0.0%	0.0%	100.0%
Mh=10	0.0076	0.0%	97.0%	3.0%	0.2842	0.0%	0.4%	99.6%
Mh=20	0.0004	0.0%	99.8%	0.2%	0.2434	0.0%	3.6%	96.4%

Table 2 shows that the estimated change points quickly converge to the true values as the amount of observed data increases. Compared to the naive approach, our method requires significantly fewer samples to recover the true underlying structure, demonstrating the effectiveness of the proposed objective function.

5 Empirical Analysis: Ranking Structure Recognition of NBA Teams

We analyze NBA regular season data from the 2014-2015 season to the 2018-2019 season². The candidate structure change points correspond to the season transitions and trade deadlines each season. These trade deadlines typically fall around February 20th each year and are denoted as ‘TradeDDL’. We identify two structure change points: the 2015–2016 trade deadline and the end of the 2016–2017 season, with results shown in Figure 2. For each resulting phase, we plot team win rates: alphabetically ordered on the left and ordered groups on the right. Black lines separate distinct groups. The left plot appears random, while the right displays a gradient from dark to light colors, moving from the top left to the bottom right. Items within each group exhibit similar behavior, as reflected by the color uniformity within each block, supporting the validity of the detected structure. Further details are provided in Section B.4.

Figure 3 displays the estimation of teams’ strengths using direct estimation and our ranking structure recognition method. The figure illustrates that our method provides a concise result regarding the structure of teams and the team strengths in each group.

²<https://www.nba.com/games>

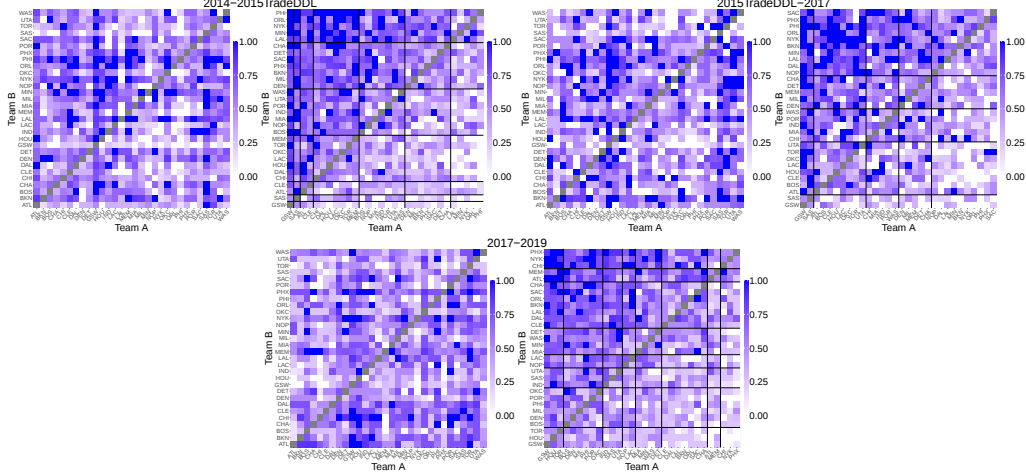


Figure 2: The winning percentage of Team A over Team B.

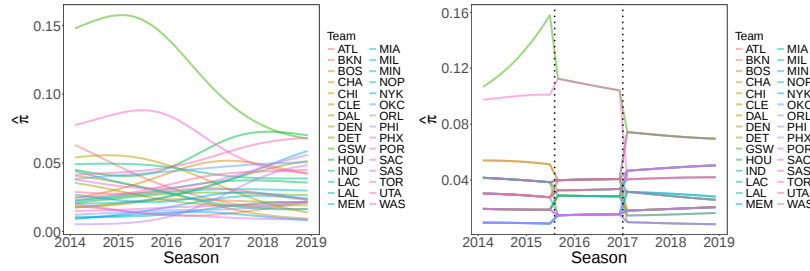


Figure 3: Estimation of team strengths using KRC (left) and our method (right).

6 Conclusion

We present a novel approach that simultaneously performs grouping and ranking based on time-varying comparisons. This offers an innovative way to analyze time-varying comparison data while generating clustered ranking results that facilitate more informed decision-making. Furthermore, we propose a combined penalty for group numbers and structure change points, allowing for the detection of long-term changes in underlying group configurations.

Several promising research directions remain open. First, extensions of the Bradley-Terry model that incorporate contextual information could be integrated into our framework to improve ranking accuracy. Second, while our current method focuses on pairwise comparisons, many practical scenarios involve comparisons among more than two candidates; extending the approach to handle such settings would broaden its applicability.

Acknowledgments and Disclosure of Funding

We thank the area chairs and reviewers for their valuable comments and suggestions, which substantially improved this paper.

References

- [1] Christopher N. Avery, Mark E. Glickman, Caroline M. Hoxby, and Andrew Metrick. A revealed preference ranking of u.s. colleges and universities *. *The Quarterly Journal of Economics*, 128(1):425–467, 12 2012. ISSN 0033-5533.
- [2] Zhou Zhao, Hanqing Lu, Deng Cai, Xiaofei He, and Yueting Zhuang. User preference learning for online social recommendation. *IEEE Transactions on Knowledge and Data Engineering*, 28(9):2522–2534, 2016.

- [3] Wanshan Li, Alessandro Rinaldo, and Daren Wang. Detecting abrupt changes in sequential pairwise comparison data. *Advances in Neural Information Processing Systems*, 35:37851–37864, 2022.
- [4] Long Ouyang, Jeffrey Wu, Xu Jiang, Diogo Almeida, Carroll Wainwright, Pamela Mishkin, Chong Zhang, Sandhini Agarwal, Katarina Slama, Alex Ray, et al. Training language models to follow instructions with human feedback. *Advances in neural information processing systems*, 35:27730–27744, 2022.
- [5] Ralph Allan Bradley and Milton E Terry. Rank analysis of incomplete block designs: I. The method of paired comparisons. *Biometrika*, 39(3/4):324–345, 1952.
- [6] Sahand Negahban, Sewoong Oh, and Devavrat Shah. Rank Centrality: Ranking from pairwise comparisons. *Operations Research*, 65(1):266–287, 2017. ISSN 0030-364X 1526-5463.
- [7] Gunther Schauburger and Gerhard Tutz. BTLLasso: a common framework and software package for the inclusion and selection of covariates in Bradley-Terry models. *Journal of Statistical Software*, 88:1–29, 2019.
- [8] Yue Liu, Ethan X. Fang, and Junwei Lu. Lagrangian inference for ranking problems. *Operations Research*, 71(1):202–223, 2023. ISSN 0030-364X 1526-5463.
- [9] Nan Lu, Ethan X Fang, and Junwei Lu. Contextual online uncertainty-aware preference learning for human feedback. *arXiv preprint arXiv:2504.19342*, 2025.
- [10] Guido Masarotto and Cristiano Varin. The ranking lasso and its application to sport tournaments. *The Annals of Applied Statistics*, 6(4):1949–1970, 2012.
- [11] Yang Zhang, Yu Xiao, Jun Wu, and Xin Lu. Comprehensive world university ranking based on ranking aggregation. *Computational Statistics*, 36(2):1139–1152, 2020. ISSN 0943-4062 1613-9658.
- [12] Jin Liu, Songyue Lin, Manling Wu, Wenjing Lyu, and Yi Su. Winning and losing relationship: A new method of university ranking in the case of countries along the belt and road. *Complexity*, 2021:1–13, 2021. ISSN 1099-0526 1076-2787.
- [13] Saúl Vargas and Pablo Castells. Rank and relevance in novelty and diversity metrics for recommender systems. In *Proceedings of the fifth ACM conference on Recommender systems*, pages 109–116, 2011.
- [14] Changhua Pei, Yi Zhang, Yongfeng Zhang, Fei Sun, Xiao Lin, Hanxiao Sun, Jian Wu, Peng Jiang, Junfeng Ge, Wenwu Ou, et al. Personalized re-ranking for recommendation. In *Proceedings of the 13th ACM conference on recommender systems*, pages 3–11, 2019.
- [15] Qiang Song, Anqi Liu, and Steve Y. Yang. Stock portfolio selection using learning-to-rank algorithms with news sentiment. *Neurocomputing*, 264:20–28, 2017. ISSN 09252312.
- [16] Fuli Feng, Moxin Li, Cheng Luo, Ritchie Ng, and Tat-Seng Chua. Hybrid learning to rank for financial event ranking. In *Proceedings of the 44th International ACM SIGIR Conference on Research and Development in Information Retrieval*, pages 233–243, 2021.
- [17] Shili Lin. Rank aggregation methods. *WIREs Computational Statistics*, 2(5):555–570, 2010. ISSN 1939-5108 1939-0068.
- [18] M. Kim, F. Farnoud, and O. Milenkovic. HyDRA: gene prioritization via hybrid distance-score rank aggregation. *Bioinformatics*, 31(7):1034–43, 2015. ISSN 1367-4811 (Electronic) 1367-4803 (Linking).
- [19] Nan Lu, Jian Shi, Xin-Yu Tian, and Kai Song. Tests on dynamic ranking. *Statistica Sinica*, 2024. doi: 10.5705/ss.202024.0153.
- [20] Rose D. Baker and Ian G. McHale. An empirical Bayes model for time-varying paired comparisons ratings: Who is the greatest women’s tennis player? *European Journal of Operational Research*, 258(1):328–333, 2017. ISSN 03772217.
- [21] Gerardo Iñiguez, Carlos Pineda, Carlos Gershenson, and Albert-László Barabási. Dynamics of ranking. *Nature Communications*, 13(1):1646, 2022.
- [22] Xinyu Tian, Jian Shi, Xiaotong Shen, and Kai Song. A spectral approach for the dynamic Bradley–Terry model. *Stat*, 13(3):e722, 2024.
- [23] Luca Faramondi, Gabriele Oliva, Roberto Setola, and Sándor Bozóki. Robustness to rank reversal in pairwise comparison matrices based on uncertainty bounds. *European Journal of Operational Research*, 304(2):676–688, 2023.

- [24] Gerhard Tutz and Gunther Schauberger. Extended ordered paired comparison models with application to football data from German Bundesliga. *AStA Advances in Statistical Analysis*, 99:209–227, 2015.
- [25] Kaycheng Soh. The seven deadly sins of world university ranking: A summary from several papers. *Journal of Higher Education Policy and Management*, 39(1):104–115, 2017.
- [26] Laura Vana, Ronald Hochreiter, and Kurt Hornik. Computing a journal meta-ranking using paired comparisons and adaptive lasso estimators. *Scientometrics*, 106(1):229–251, 2015. ISSN 0138-9130 1588-2861.
- [27] Jong-June Jeon and Hosik Choi. The sparse Luce model. *Applied Intelligence*, 48(8):1953–1964, 2016.
- [28] Xin-Yu Tian and Jian Shi. Grouped rank centrality: Ranking and grouping from pairwise comparisons simultaneously. *Stat*, 12(1):e626, 2023.
- [29] Heejong Bong, Wanshan Li, Shamindra Shrotriya, and Alessandro Rinaldo. Nonparametric estimation in the dynamic Bradley-Terry model. In *International Conference on Artificial Intelligence and Statistics*, pages 3317–3326. PMLR, 2020.
- [30] Chao Gao, Yandi Shen, and Anderson Y Zhang. Uncertainty quantification in the Bradley-Terry-Luce model. *Information and Inference: A Journal of the IMA*, 12(2):1073–1140, 2023.
- [31] Alexandre Belloni and Victor Chernozhukov. Least squares after model selection in high-dimensional sparse models. *Bernoulli*, 19(2):521–547, 2013.
- [32] Charles-Alban Deledalle, Nicolas Papadakis, Joseph Salmon, and Samuel Vaïter. CLEAR: Covariant least-square refitting with applications to image restoration. *SIAM Journal on Imaging Sciences*, 10(1): 243–284, 2017.
- [33] Stephen J Kirkland and Michael Neumann. *Group inverses of M-matrices and their applications*. CRC Press, 2012.
- [34] Xiaoxuan Liang, Aaron Cohen, Anibal Sólón Heinsfeld, Franco Pestilli, and Daniel J McDonald. sparsegl: An r package for estimating sparse group lasso. *Journal of Statistical Software*, 110:1–23, 2024.
- [35] Tao Wang and LiXing Zhu. A distribution-based lasso for a general single-index model. *Science China Mathematics*, 58:109–130, 2015.
- [36] Ming Yuan and Yi Lin. Model selection and estimation in regression with grouped variables. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 68(1):49–67, 2006.
- [37] Hansheng Wang and Chenlei Leng. A note on adaptive group lasso. *Computational Statistics & Data Analysis*, 52(12):5277–5286, 2008.
- [38] Caiya Zhang and Yanbiao Xiang. On the oracle property of adaptive group lasso in high-dimensional linear models. *Statistical Papers*, 57:249–265, 2016.
- [39] Yuxin Chen, Yuejie Chi, Jianqing Fan, Cong Ma, et al. Spectral methods for data science: A statistical perspective. *Foundations and Trends in Machine Learning*, 14(5):566–806, 2021.
- [40] Lucas Maystre, Victor Kristof, and Matthias Grossglauser. Pairwise comparisons with flexible time-dynamics. In *Proceedings of the 25th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*, pages 1236–1246, 2019.
- [41] Wei Xiong, Hanze Dong, Chenlu Ye, Ziqi Wang, Han Zhong, Heng Ji, Nan Jiang, and Tong Zhang. Iterative preference learning from human feedback: Bridging theory and practice for RLHF under KL-constraint. In *Forty-first International Conference on Machine Learning*, 2024.
- [42] Huiying Zhong, Zhun Deng, Weijie J Su, Zhiwei Steven Wu, and Linjun Zhang. Provable multi-party reinforcement learning with diverse human feedback. *arXiv preprint arXiv:2403.05006*, 2024.
- [43] Banghua Zhu, Michael Jordan, and Jiantao Jiao. Principled reinforcement learning with human feedback from pairwise or k-wise comparisons. In *International Conference on Machine Learning*, pages 43037–43067. PMLR, 2023.
- [44] Eglantine Karlé and Hemant Tyagi. Dynamic ranking with the BTL model: a nearest neighbor based rank centrality method. *Journal of Machine Learning Research*, 24(269):1–57, 2023.

NeurIPS Paper Checklist

1. Claims

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

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

Justification: The main claims made in the abstract and introduction accurately reflect the paper's contributions and scope.

Guidelines:

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

2. Limitations

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

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

Justification: Yes, please refer to our Section 6.

Guidelines:

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

3. Theory assumptions and proofs

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

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

Justification: Our Appendix provides the detailed proof of all theoretical results.

Guidelines:

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

4. Experimental result reproducibility

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

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

Justification: Please refer to experiment settings in Sections 4 and B. The codes for the proposed algorithm and experiments will also be made publicly available upon publication.

Guidelines:

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

5. Open access to data and code

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

Answer: [Yes]

Justification: The codes for the proposed algorithm and experiments will be made publicly available upon publication.

Guidelines:

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

6. Experimental setting/details

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

Answer: [Yes]

Justification: Yes, please refer to our Sections 4 and B.

Guidelines:

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

7. Experiment statistical significance

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

Answer: [Yes]

Justification: Yes, please refer to our Sections 4 and B.

Guidelines:

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

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

8. Experiments compute resources

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

Answer: [Yes]

Justification: Yes, please refer to our Section 4.

Guidelines:

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

9. Code of ethics

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

Answer: [Yes]

Justification: The research conducted in the paper adheres to the NeurIPS Code of Ethics.

Guidelines:

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

10. Broader impacts

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

Answer: [NA]

Justification: This study purely contributes to the technical advancement of dynamic ranking problem and does not have any societal impacts.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.

- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

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

Answer: [NA]

Justification: This paper poses no such risks.

Guidelines:

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

12. Licenses for existing assets

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

Answer: [Yes]

Justification: Yes, please refer to our Section 5.

Guidelines:

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

- If this information is not available online, the authors are encouraged to reach out to the asset’s creators.

13. **New assets**

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

Answer: [NA]

Justification: This paper does not release new assets.

Guidelines:

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

14. **Crowdsourcing and research with human subjects**

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

Answer: [NA]

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

Guidelines:

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

15. **Institutional review board (IRB) approvals or equivalent for research with human subjects**

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

Answer: [NA]

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

Guidelines:

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

16. **Declaration of LLM usage**

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigorousness, or originality of the research, declaration is not required.

Answer: [NA]

Justification: The core method development in this research does not involve LLMs.

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.

A Notations

We include Table 3 to summarize the key notations used throughout the paper.

Table 3: Notations

Symbol	Description
$[n]$	Set of integers $\{1, \dots, n\}$
$\mathbf{I}$	Identity matrix
$\mathbf{e}_n$	$n \times 1$ vector with all elements equal to 1
$\mathbf{0}$	Zero vector or matrix
$\ \mathbf{v}\ _2$	ℓ_2 -norm of vector $\mathbf{v}$
$\ f\ _{2,T}$	L_2 norm of function $f(\cdot)$ over the interval $[0, T]$
$\mathcal{Y}$	Set of pairwise comparison results
$y_{ij}(t)$	Comparison result at time t between items i and j
$\boldsymbol{\pi}^*(t)$	Score vector of items at time t in the Bradley-Terry model
$y_{ij}^*(t)$	Winning probability between items i and j at time t
$K(\cdot)$	Kernel function
$\mathbf{P}(t)$	Transformation probability matrix for the Markov chain
B	Number of groups
G	Group partition of items
δ_1	Minimum pairwise score difference between groups
δ_2	Maximum pairwise score difference within group
λ	Tuning parameter for group recognition
$\mathbf{Q}$	Constant matrix used for transformations
$\boldsymbol{\theta}(t)$	Transformation of the score vector $\boldsymbol{\pi}(t)$
$\mathbf{X}(t)$	Matrix after transformation used in optimization objective for group recognition
$\mathbf{Y}(t)$	Vector after transformation used in optimization objective for group recognition
$\mathbf{z}(t)$	Latent group of items at time point t
η_i	Time points where the structure changes
ξ_i	Candidate structure change points
$l(\boldsymbol{\pi}, t)$	Negative log-likelihood function at time t for score vector $\boldsymbol{\pi}$
$\mathcal{P}$	Partition of time interval
γ_1, γ_2	Regularization parameters for the objective function
$\mathbf{A}^\#$	Group inverse matrix

B Supplementary to numerical results

B.1 Experiment settings for ranking group recognition

For the experiments in Section 4.1, we consider two settings to evaluate our methods, as illustrated in Figure 4. We set $T = 1$, $B = 3$, and assign the number of items in each group as 3:3:4.

Define the first set as follows:

$$\boldsymbol{\pi}_i^*(t) - \text{pert}_i(t) = \begin{cases} \frac{1}{n} (2 + 0.3 \sin(6\pi t)) & i \in G_1, \\ \frac{1}{n} (1 - 0.2 \sin(6\pi t)) & i \in G_2, \\ \frac{1}{n} (0.25 - 0.075 \sin(6\pi t)) & i \in G_3. \end{cases}$$

The $\text{pert}_i(t)$ is a perturbation term whose absolute value is less than $0.01/n$. For the second set, define it as:

$$\boldsymbol{\pi}_i^*(t) - \text{pert}_i(t) = \begin{cases} \frac{1}{n} (1.9 + 0.5 \sin(3\pi t)) & i \in G_1, \\ \frac{1}{n} (0.1 + 0.6 \arctan(\pi t)) & i \in G_2, \\ \frac{1}{n} (1 - 0.375 \sin(3\pi t) - 0.45 \arctan(\pi t)) & i \in G_3. \end{cases}$$

The point ϵ used for order estimation is 0.001. The first setting represents a simple case, while the second setting is more complex with intersections of scores among different groups.

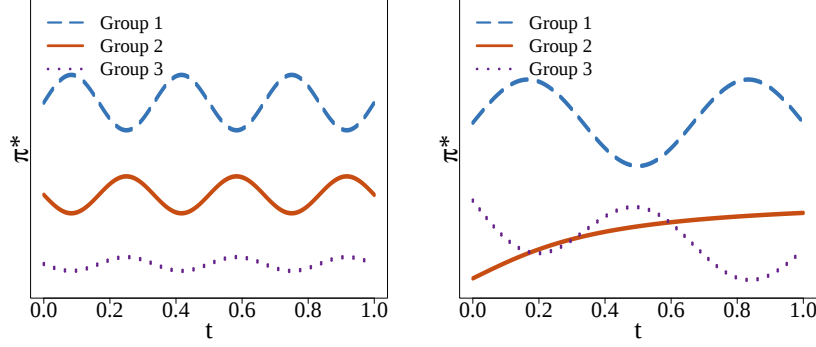


Figure 4: $\pi^*(t)$ of each group in simulations.

Set $h = 0.05$, $m = 30$ and vary n and M . We repeat each setting 500 times and use the extended BIC (EBIC) criterion [35, 22] to choose the tuning parameter λ . We note that cross-validation can be used here, but with the EBIC criteria, the computational cost is much lower. Specifically, we have

$$\text{EBIC}(\lambda) = nm \log\left(\frac{\text{RSS}(\lambda)}{nm}\right) + c_0 \text{Var}(\mathbf{Y}) + \log(nm) \lceil df(\lambda) \rceil.$$

Here, $c_0 = 0.1$, $\lceil \cdot \rceil$ is the round down function. We let $\text{RSS}(\lambda) = \|\mathbf{Y} - \mathbf{X}\hat{\boldsymbol{\theta}}(\lambda)\|_2^2$ and

$$df(\lambda) = \sum_{i=1}^{n-1} \mathbb{1}\{\|\hat{\boldsymbol{\theta}}_i(\lambda)\|_2 > 0\} + \sum_{i=1}^{n-1} \frac{\|\hat{\boldsymbol{\theta}}_i(\lambda)\|_2}{\|\boldsymbol{\theta}_i\|_2} (m-1),$$

which is commonly used, for example, in [36, 37].

B.2 Experiment settings for group changes recognition

For the problem of structural change detection in Section 4.2, we consider two settings: one with three stages and another with two stages. The true abilities are depicted in Figures 5 and 6, respectively. We set $h = 0.02$ and denote the observation times within each phase as M . The experiment is repeated 500 times. We employ the widely-used 10-fold cross-validation for the choice of tuning parameters γ_1 and γ_2 .

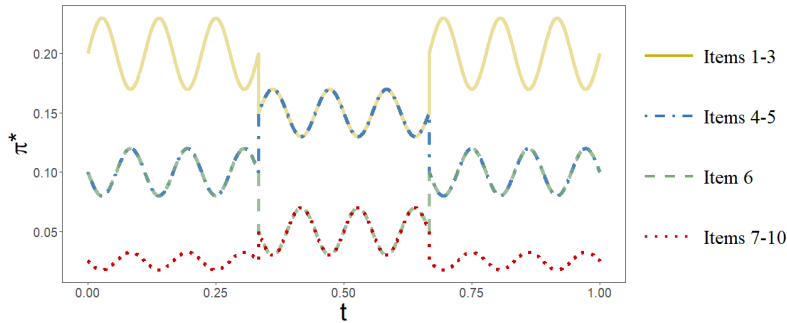


Figure 5: $\pi^*(t)$ in the first setting.

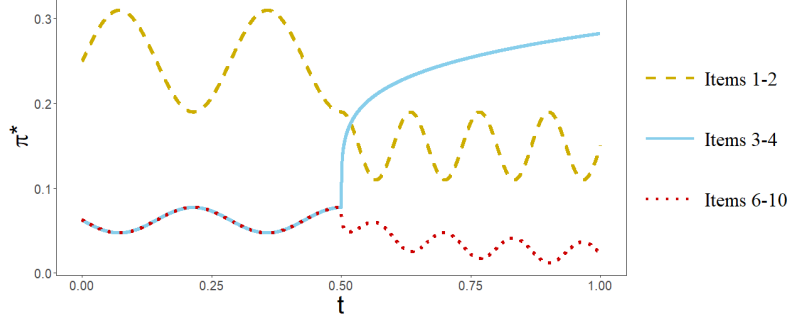


Figure 6: $\pi^*(t)$ in the second setting.

Specifically, in setting 1, three phases are considered. The score functions in phases I and III remain the same, with a proportion of items being 3:3:4.

$$\pi_i^*(t) = \begin{cases} 0.2 + 0.03 \sin(18\pi t) & i = 1, 2, 3, \\ 0.1 - 0.02 \sin(18\pi t) & i = 4, 5, 6, \\ 0.025 - 0.0075 \sin(18\pi t) & i = 7, 8, 9, 10. \end{cases}$$

For phase II, the item proportion is 1:1.

$$\pi_i^*(t) = \begin{cases} 0.15 + 0.02 \sin(18\pi t) & i = 1, \dots, 5, \\ 0.05 - 0.02 \sin(18\pi t) & i = 6, \dots, 10. \end{cases}$$

In the second setting, during phase I, the item proportion is 1:4.

$$\pi_i^*(t) = \begin{cases} 0.25 + 0.06 \sin(7\pi t) & i = 1, 2, \\ 0.0625 - 0.015 \sin(7\pi t) & i = 3, \dots, 10. \end{cases}$$

For phase II, the item proportion is 1:1:3.

$$\pi_i^*(t) = \begin{cases} 0.15 - 0.04 \sin(15\pi t), & i = 1, 2, \\ 0.065 + 0.25(t - \frac{1}{2})^{1/10}, & i = 3, 4, \\ 0.1425 + 0.02 \sin(15\pi t) - 0.0625(t - \frac{1}{2})^{1/10}, & i = 5, \dots, 10. \end{cases}$$

B.3 Sensitivity of the hyperparameter choice

Hyperparameters play a crucial role in the performance of the proposed method, and their selection can be challenging. Ideally, computationally effective guidelines, such as those based on information criteria, would facilitate hyperparameter tuning. However, due to the novel optimization objective in our case, establishing such rules is non-trivial and warrants further investigation.

In this study, we employ cross-validation, which yields good empirical performance. To evaluate the sensitivity of the method to hyperparameter choices, we conduct additional experiments using parameter grids for the two experimental settings described in Section 4.2. Specifically, we repeat each combination of hyperparameters 50 times and calculate the average number of detected change points for each combination. The results are shown in Tables 4 and 5.

From the experimental results, we observe that the number of estimated change points varies with different values of γ_1 and γ_2 . In particular, we find that γ_1 should be greater than γ_2 to ensure effective change point detection, consistent with our theoretical analysis in Theorem 3.9. Moreover, for a fixed γ_1 , the estimated number of change points decreases as γ_2 increases. This aligns with our intuition, as larger values of γ_2 impose a higher penalty for each additional change point. Overall, while the tuning parameters affect the change point estimation, the estimated number of change points remains relatively stable even under substantial variations in the scales of γ_1 and γ_2 .

B.4 Supplementary to empirical analysis

For the empirical study, we set the bandwidth to match the season length, and other parameters remain consistent with those used in the simulations. Tuning parameters chosen through cross-validation

Table 4: Average change point number for different values of γ_1 and γ_2 for setting 1.

γ_1	0.02	0.04	0.06	0.08	0.1	0.2	0.4	0.6	0.8	1	2	4	6	8	10
$\gamma_2=0.002$	2.12	2.1	2.22	2.22	2.22	2.04	1.58	1.54	1.54	1.54	1.54	1.54	1.54	1.54	1.54
$\gamma_2=0.004$	2.1	2.08	2.2	2.2	2.2	2.02	1.58	1.54	1.54	1.54	1.54	1.54	1.54	1.54	1.54
$\gamma_2=0.006$	2.1	2.08	2.2	2.2	2.2	2.02	1.58	1.54	1.54	1.54	1.54	1.54	1.54	1.54	1.54
$\gamma_2=0.008$	2.1	2.08	2.2	2.2	2.2	2.02	1.58	1.54	1.54	1.54	1.54	1.54	1.54	1.54	1.54
$\gamma_2=0.01$	2.1	2.08	2.2	2.2	2.2	2.02	1.58	1.54	1.54	1.54	1.54	1.54	1.54	1.54	1.54
$\gamma_2=0.02$	0	2.06	2.18	2.18	2.18	1.96	1.54	1.54	1.54	1.54	1.54	1.54	1.54	1.54	1.54
$\gamma_2=0.04$	0	0	2.1	2.14	2.14	1.9	1.52	1.52	1.52	1.52	1.52	1.52	1.52	1.52	1.52
$\gamma_2=0.06$	0	0	0	2.14	2.14	1.86	1.52	1.52	1.52	1.52	1.52	1.52	1.52	1.52	1.52
$\gamma_2=0.08$	0	0	0	0	2.12	1.82	1.52	1.52	1.52	1.52	1.52	1.52	1.52	1.52	1.52
$\gamma_2=0.1$	0	0	0	0	0	1.82	1.52	1.52	1.52	1.52	1.52	1.52	1.52	1.52	1.52
$\gamma_2=0.2$	0	0	0	0	0	0	1.52	1.52	1.52	1.52	1.52	1.52	1.52	1.52	1.52
$\gamma_2=0.4$	0	0	0	0	0	0	0	1.52	1.52	1.52	1.52	1.52	1.52	1.52	1.52
$\gamma_2=0.6$	0	0	0	0	0	0	0	0	1.52	1.52	1.52	1.52	1.52	1.52	1.52
$\gamma_2=0.8$	0	0	0	0	0	0	0	0	0	1.52	1.52	1.52	1.52	1.52	1.52
$\gamma_2=1$	0	0	0	0	0	0	0	0	0	0	1.52	1.52	1.52	1.52	1.52

Table 5: Average change point number for different values of γ_1 and γ_2 for setting 2.

γ_1	0.02	0.04	0.06	0.08	0.1	0.2	0.4	0.6	0.8	1	2	4	6	8	10
$\gamma_2=0.002$	1.92	1.86	1.94	1.92	1.92	1.9	1.9	1.9	1.86	1.78	1.76	1.76	1.76	1.76	1.76
$\gamma_2=0.004$	1.82	1.84	1.92	1.9	1.9	1.88	1.88	1.88	1.84	1.76	1.74	1.74	1.74	1.74	1.74
$\gamma_2=0.006$	1.8	1.84	1.92	1.9	1.9	1.88	1.88	1.86	1.84	1.76	1.74	1.74	1.74	1.74	1.74
$\gamma_2=0.008$	1.68	1.78	1.88	1.86	1.86	1.84	1.84	1.86	1.84	1.76	1.74	1.74	1.74	1.74	1.74
$\gamma_2=0.01$	1.68	1.78	1.88	1.86	1.86	1.84	1.84	1.86	1.84	1.76	1.74	1.74	1.74	1.74	1.74
$\gamma_2=0.02$	0	1.76	1.88	1.86	1.86	1.84	1.82	1.84	1.84	1.76	1.74	1.74	1.74	1.74	1.74
$\gamma_2=0.04$	0	0	1.84	1.82	1.82	1.8	1.78	1.8	1.8	1.72	1.7	1.7	1.7	1.7	1.7
$\gamma_2=0.06$	0	0	0	1.82	1.8	1.78	1.74	1.8	1.8	1.72	1.7	1.7	1.7	1.7	1.7
$\gamma_2=0.08$	0	0	0	0	1.8	1.78	1.74	1.8	1.8	1.72	1.7	1.7	1.7	1.7	1.7
$\gamma_2=0.1$	0	0	0	0	0	1.78	1.74	1.78	1.78	1.7	1.68	1.68	1.68	1.68	1.68
$\gamma_2=0.2$	0	0	0	0	0	0	1.68	1.78	1.78	1.7	1.68	1.68	1.68	1.68	1.68
$\gamma_2=0.4$	0	0	0	0	0	0	0	1.76	1.78	1.7	1.68	1.68	1.68	1.68	1.68
$\gamma_2=0.6$	0	0	0	0	0	0	0	0	1.74	1.7	1.68	1.68	1.68	1.68	1.68
$\gamma_2=0.8$	0	0	0	0	0	0	0	0	0	1.66	1.68	1.68	1.68	1.68	1.68
$\gamma_2=1$	0	0	0	0	0	0	0	0	0	0	1.68	1.68	1.68	1.68	1.68

are $\gamma_1 = 0.04$ and $\gamma_2 = 0.006$. For the learning results demonstrated in Figure 2, in the first period, the GSW and SAS teams occupy the first two groups due to their extremely high winning rates. In the second phase, the first group includes GSW and SAS, while the second group consists of items from the top groups in the previous stage, with some changes. For instance, MEM and DAL shift to weaker groups. From the second to the third stage, a major change in the leading teams is notable: ATL, CLE, SAS disappear from the top groups, and teams like DEN, MIL, PHI emerge in the top 2 groups.

It is important to note that employing a static method yields significantly different results. For instance, in the initial phase, while also identifying seven groups, the static method GRC categorizes CHI into a single group, amalgamates POR into the MEM group, and groups all items from BOS to CHA (in the order presented in Figure 2, excluding POR) as a unified entity. Moreover, in the third phase, the static method recognizes three groups. The first group remains unchanged, with the subsequent ten items (excluding NOP but including LAC) forming the second group, while all other teams constitute the third group.

C Technical proofs

C.1 Proof of Theorem 3.3

Proof. Let $\mu(\mathbf{A})$ represent the eigenvalue of matrix $\mathbf{A}$. Let $\mathbf{A}_S$ be the submatrix of $\mathbf{A}$ consisting of columns that correspond to items in S for matrix $\mathbf{A}$, and x_S be the subvector of x comprising components corresponding to S for vector x . Before presenting the main theorem, we present the following lemma on the properties of $\mathbf{X}$ and $\mathbf{P}$.

Lemma C.1. Let Assumptions 3.1 and 3.2 hold. If $Mh \rightarrow \infty$, $n \rightarrow \infty$ and $nMh^5 \rightarrow 0$, we have $\|\mathbf{I} - \mathbf{P}(t)\|_2 = O_p(1)$, $\|(\mathbf{Q}^{-1})_{S^c}\|_2 \lesssim \sqrt{\frac{B}{1 + \cos \frac{(n-B)\pi}{n-B+1}}}$, $\mu_{\min}((\mathbf{X}_S)^\top(t)\mathbf{X}_S(t)) \gtrsim \frac{n}{B}$ and $\|(\mathbf{X}_S^\top(t)\mathbf{X}_S(t))^{-1}\mathbf{X}_S^\top(t)\|_2 \lesssim \sqrt{\frac{B}{n}}$ with probability tending to 1.

We first show the consistency of $\tilde{\boldsymbol{\theta}}$. From Theorem S1 of Lu et al. [19], we have $\|\tilde{\boldsymbol{\pi}}(t) - \boldsymbol{\pi}^*(t)\|_\infty = O_p(\delta)$. Combining the definition of $\boldsymbol{\theta}$, $\|\tilde{\boldsymbol{\theta}}(t) - \boldsymbol{\theta}^*(t)\|_\infty = O_p(\delta)$.

Notice that

$$\hat{\boldsymbol{\theta}} = \arg \min_{\boldsymbol{\theta}} Q(\boldsymbol{\theta}) = \arg \min_{\boldsymbol{\theta}} \frac{1}{2} \|\mathbf{Q} - \mathbf{X}\boldsymbol{\theta}\|_2^2 + \lambda \sum_{i=1}^{n-1} \|\tilde{\boldsymbol{\theta}}_i\|_2^{-1} \|\boldsymbol{\theta}_i\|_2. \quad (\text{C.1})$$

Following the proof of Theorem 2.1 in [38], $Q(\boldsymbol{\theta})$ is a strictly convex function. Lemma 4.1 in [38] points out that, (C.1) is equivalent to

$$-\mathbf{X}_j^\top (\mathbf{Y} - \mathbf{X}\hat{\boldsymbol{\theta}}) + \lambda \|\tilde{\boldsymbol{\theta}}_j\|_2^{-1} \frac{\hat{\boldsymbol{\theta}}_j}{\|\hat{\boldsymbol{\theta}}_j\|_2} = \mathbf{0}, \quad \forall \hat{\boldsymbol{\theta}}_j \neq \mathbf{0},$$

and

$$\|\mathbf{X}_j^\top (\mathbf{Y} - \mathbf{X}\hat{\boldsymbol{\theta}})\|_2 \leq \lambda \|\tilde{\boldsymbol{\theta}}_j\|_2^{-1}, \quad \forall \hat{\boldsymbol{\theta}}_j = \mathbf{0},$$

where $\mathbf{X}_j$ represents the columns of $\mathbf{X}$ corresponding to $\boldsymbol{\theta}_j$. Therefore, it is sufficient to prove $\exists \boldsymbol{\theta}_0, \forall j \in S, \boldsymbol{\theta}_{0j} \neq \mathbf{0}$, and $\forall j \notin S, \boldsymbol{\theta}_{0j} = \mathbf{0}$, such that

$$-\mathbf{X}_j^\top (\mathbf{Y} - \mathbf{X}_S \boldsymbol{\theta}_{0S}) + \lambda \|\tilde{\boldsymbol{\theta}}_j\|_2^{-1} \frac{\boldsymbol{\theta}_{0j}}{\|\boldsymbol{\theta}_{0j}\|_2} = \mathbf{0}, \quad \forall j \in S, \quad (\text{C.2})$$

and

$$\|\mathbf{X}_j^\top (\mathbf{Y} - \mathbf{X}_S \boldsymbol{\theta}_{0S})\|_2 < \lambda \|\tilde{\boldsymbol{\theta}}_j\|_2^{-1}, \quad \forall j \notin S. \quad (\text{C.3})$$

Using (C.2), we have

$$-\mathbf{X}_S^\top (\mathbf{Y} - \mathbf{X}_S \boldsymbol{\theta}_{0S}) + \lambda \beta_{0S} = \mathbf{0},$$

where $\beta_{0S} = (\frac{\boldsymbol{\theta}_{0j}^\top}{\|\tilde{\boldsymbol{\theta}}_j\|_2 \|\boldsymbol{\theta}_{0j}\|_2})_{j \in S}^\top$. Notice that $\mathbf{X}_S^\top \mathbf{X}_S$ is invertible. Hence,

$$\begin{aligned} \boldsymbol{\theta}_{0S} &= (\mathbf{X}_S^\top \mathbf{X}_S)^{-1} \mathbf{X}_S^\top \mathbf{Y} - \lambda (\mathbf{X}_S^\top \mathbf{X}_S)^{-1} \beta_{0S} \\ &= \boldsymbol{\theta}_S^* + ((\mathbf{X}_S^\top \mathbf{X}_S)^{-1} \mathbf{X}_S^\top \mathbf{Y} - \boldsymbol{\theta}_S^*) - \lambda (\mathbf{X}_S^\top \mathbf{X}_S)^{-1} \beta_{0S}. \end{aligned} \quad (\text{C.4})$$

As for the second term,

$$\begin{aligned} \|(\mathbf{X}_S^\top \mathbf{X}_S)^{-1} \mathbf{X}_S^\top \mathbf{Y} - \boldsymbol{\theta}_S^*\|_\infty &= \|(\mathbf{X}_S^\top \mathbf{X}_S)^{-1} \mathbf{X}_S^\top (\mathbf{Y} - \mathbf{X}_S \boldsymbol{\theta}_S^*)\|_\infty \\ &\leq \sup_{k \in [m]} \|(\mathbf{X}_S^\top(t_k) \mathbf{X}_S(t_k))^{-1} \mathbf{X}_S^\top(t_k) (\mathbf{Y}(t_k) - \mathbf{X}_S(t_k) \boldsymbol{\theta}_S^*(t_k))\|_2 \\ &\leq \sup_{k \in [m]} \|(\mathbf{X}_S^\top(t_k) \mathbf{X}_S(t_k))^{-1} \mathbf{X}_S^\top(t_k)\|_2 \|\mathbf{Y}(t_k) - \mathbf{X}_S(t_k) \boldsymbol{\theta}_S^*(t_k)\|_2, \end{aligned} \quad (\text{C.5})$$

where

$$\begin{aligned} \|\mathbf{Y}(t) - \mathbf{X}_{\mathcal{S}}(t)\boldsymbol{\theta}_{\mathcal{S}}^*(t)\|_2 &\leq \|\mathbf{Y}(t) - \mathbf{X}_{-1}(t)\boldsymbol{\theta}^*(t)\|_2 + \|\mathbf{X}_{\mathcal{S}^c}(t)\boldsymbol{\theta}_{\mathcal{S}^c}^*(t)\|_2 \\ &= \|(\mathbf{P}^{*\top}(t) - \mathbf{P}^\top(t))\boldsymbol{\pi}^*(t)\|_2 + \|\mathbf{X}_{\mathcal{S}^c}(t)\boldsymbol{\theta}_{\mathcal{S}^c}^*(t)\|_2. \end{aligned} \quad (\text{C.6})$$

The proof of Theorem 1 in [22] implies

$$\|(\mathbf{P}^{*\top}(t) - \mathbf{P}^\top(t))\boldsymbol{\pi}^*(t)\|_2 = O_p(\sqrt{\frac{1}{n^2 M h}}).$$

For the second term in (C.6),

$$\|\mathbf{X}_{\mathcal{S}^c}(t)\boldsymbol{\theta}_{\mathcal{S}^c}^*(t)\|_2 \leq \|\mathbf{X}_{\mathcal{S}^c}(t)\|_2 \|\boldsymbol{\theta}_{\mathcal{S}^c}^*(t)\|_2 \leq \|\mathbf{P}^\top(t) - \mathbf{I}\|_2 \|(Q^{-1})_{\mathcal{S}^c}\|_2 \|\boldsymbol{\theta}_{\mathcal{S}^c}^*(t)\|_2. \quad (\text{C.7})$$

From Lemma C.1, we have (C.7) $\lesssim \sqrt{\frac{B}{1+\cos \frac{(n-B)\pi}{n-B+1}}}(n-B)\delta_2$. If $\delta_2 = o_p(\sqrt{\frac{1+\cos \frac{(n-B)\pi}{n-B+1}}{B(n-B)} \frac{1}{n^2 M h}})$, then (C.7) is $o_p(\sqrt{\frac{1}{n^2 M h}})$. The first term in (C.5) is $O_p(\sqrt{\frac{B}{n}})$ using Lemma C.1. Hence, (C.5) is $O_p(\sqrt{\frac{B}{n^3 M h}})$.

For the third term in (C.4), we have

$$\begin{aligned} \|\lambda(\mathbf{X}_{\mathcal{S}}^\top \mathbf{X}_{\mathcal{S}})^{-1} \beta_{0\mathcal{S}}\|_\infty &\leq \sup_{k \in [m]} \|\lambda(\mathbf{X}_{\mathcal{S}}^\top(t_k) \mathbf{X}_{\mathcal{S}}(t_k))^{-1} \beta_{0\mathcal{S}}(t_k)\|_\infty \\ &\leq \sup_{k \in [m]} \|\lambda(\mathbf{X}_{\mathcal{S}}^\top(t_k) \mathbf{X}_{\mathcal{S}}(t_k))^{-1} \beta_{0\mathcal{S}}(t_k)\|_2 \leq \sup_{k \in [m]} \lambda\|(\mathbf{X}_{\mathcal{S}}^\top(t_k) \mathbf{X}_{\mathcal{S}}(t_k))^{-1}\|_2 \|\beta_{0\mathcal{S}}(t_k)\|_2 \\ &\leq \lambda \frac{B}{n} \frac{\sqrt{B}}{\min_{i \in \mathcal{S}} \|\tilde{\boldsymbol{\theta}}_i\|_2}. \end{aligned} \quad (\text{C.8})$$

Note that

$$\begin{aligned} \sqrt{\frac{1}{m}} \min_{i \in \mathcal{S}} \|\tilde{\boldsymbol{\theta}}_i\|_2 &\geq \sqrt{\frac{1}{m}} \min_{i \in \mathcal{S}} \|\boldsymbol{\theta}_i^*\|_2 - \sqrt{\frac{1}{m}} \max_{i \in \mathcal{S}} \|\tilde{\boldsymbol{\theta}}_i - \boldsymbol{\theta}_i^*\|_2 \\ &\gtrsim \min_{i \in \mathcal{S}} \|\boldsymbol{\theta}_i^*(t)\|_{2,T} + O(\frac{1}{m}) + O_p(\delta) \gtrsim \delta_1. \end{aligned}$$

Hence, (C.8) $\lesssim \frac{\lambda \sqrt{B^3}}{n \sqrt{m} \delta_1}$.

From (C.4), $\forall j \in \mathcal{S}$,

$$\|\boldsymbol{\theta}_{0j}\|_2 \gtrsim \sqrt{m} \delta_1 - \sqrt{m} \sqrt{\frac{B}{n^3 M h}} - \frac{\lambda \sqrt{B^3}}{n \delta_1}.$$

If $\sqrt{\frac{B}{n^3 M h \delta_1^2}} = o(1)$ and $\frac{\lambda B^{3/2}}{n m^{1/2} \delta_1^2} = o(1)$, then with probability tending to 1, we have $\forall j \in \mathcal{S}, \boldsymbol{\theta}_{0j} \neq \mathbf{0}$. Actually, we have proved that if $\frac{\lambda B \sqrt{n M h}}{\sqrt{m} \delta_1} = O(1)$, then $\|\boldsymbol{\theta}_{0\mathcal{S}} - \boldsymbol{\theta}_{\mathcal{S}}^*\|_\infty = O_p(\sqrt{\frac{B}{n^3 M h}})$.

Then we prove (C.3). Assume $j \notin \mathcal{S}$.

$$\begin{aligned} \|\mathbf{X}_j^\top (\mathbf{Y} - \mathbf{X}_{\mathcal{S}} \boldsymbol{\theta}_{0\mathcal{S}})\|_2 &\leq \|\mathbf{X}_j^\top (\mathbf{Y} - \mathbf{X}_{\mathcal{S}} \boldsymbol{\theta}_{\mathcal{S}}^*)\|_2 + \|\mathbf{X}_j^\top (\mathbf{X}_{\mathcal{S}} \boldsymbol{\theta}_{\mathcal{S}}^* - \mathbf{X}_{\mathcal{S}} \boldsymbol{\theta}_{0\mathcal{S}})\|_2 \\ &\lesssim \sup_t \sqrt{m} \|\mathbf{X}_{\mathcal{S}^c}^\top(t) (\mathbf{Y}(t) - \mathbf{X}_{\mathcal{S}}(t) \boldsymbol{\theta}_{\mathcal{S}}^*(t))\|_\infty \\ &\quad + \sqrt{m} \|(\mathbf{X}_{\mathcal{S}^c}^\top(t) (\mathbf{X}_{\mathcal{S}}(t) \boldsymbol{\theta}_{\mathcal{S}}^*(t) - \mathbf{X}_{\mathcal{S}}(t) \boldsymbol{\theta}_{0\mathcal{S}}(t)))\|_\infty. \end{aligned} \quad (\text{C.9})$$

Similar to the proof of Theorem 1 in [28], we can obtain

$$\begin{aligned} (\text{C.9}) &\leq \sup_t \sqrt{m} \max_i \|(\mathbf{X}_{\mathcal{S}^c}^\top(t))_{i\cdot}\|_2 \|(\mathbf{Y}(t) - \mathbf{X}_{\mathcal{S}}(t) \boldsymbol{\theta}_{\mathcal{S}}^*(t))\|_2 \\ &\quad + \sqrt{m} \max_i \|(\mathbf{X}_{\mathcal{S}^c}^\top(t) \mathbf{X}_{\mathcal{S}}(t))_{i\cdot}\|_2 \|\boldsymbol{\theta}_{\mathcal{S}}^*(t) - \boldsymbol{\theta}_{0\mathcal{S}}(t)\|_2 \\ &\leq \sup_t \sqrt{m n} \max_{i,k} |(\mathbf{X}_{\mathcal{S}^c}^\top(t))_{ik}| \|(\mathbf{P}^{*\top}(t) - \mathbf{P}^\top(t))\boldsymbol{\pi}^*(t)\|_2 \\ &\quad + \sqrt{m B} \max_{i,k} |(\mathbf{X}_{\mathcal{S}^c}^\top(t) \mathbf{X}_{\mathcal{S}}(t))_{ik}| \|\boldsymbol{\theta}_{\mathcal{S}}^*(t) - \boldsymbol{\theta}_{0\mathcal{S}}(t)\|_2. \end{aligned} \quad (\text{C.10})$$

The second inequality is gotten using (C.6). From the proof of Theorem 1 in [28], we have $\max_{i,k} |(\mathbf{X}_{SC}^\top(t))_{ik}| = O(1)$ and $\max_{i,k} |(\mathbf{X}_{SC}^\top(t)\mathbf{X}_S(t))_{ik}| = O(n)$. Therefore, (C.10) = $O_p(\sqrt{\frac{mB^3}{nMh}})$.

On the other hand, since $\sqrt{\frac{1}{m}}\|\tilde{\boldsymbol{\theta}}_j\|_2 \leq \sqrt{\frac{1}{m}}\|\tilde{\boldsymbol{\theta}}_j - \boldsymbol{\theta}_j^*\|_2 + \sqrt{\frac{1}{m}}\|\boldsymbol{\theta}_j^*\|_2 \lesssim \delta_2 + \delta$, we have $\min_{j \notin \mathcal{S}} \lambda \|\tilde{\boldsymbol{\theta}}_j\|_2^{-1} \gtrsim \frac{\lambda}{m^{1/2}\delta}$. Since $\lambda \gtrsim \sqrt{\frac{m^2 B^3}{nMh}}\tilde{\delta}$, the theorem is proved. $\square$

C.2 Proof of Theorem 3.6

Proof. Let $\mathbf{P}_G$ take the following form:

$$\mathbf{P}_{Gij}(t) = \begin{cases} \frac{1}{B} \frac{\sum_{l_1 \in G_i} \sum_{l_2 \in G_j} \sum_{t_k \in T_{l_1 l_2}} y_{l_1 l_2}(t_k) K_h(t, t_k)}{\sum_{l_1 \in G_i} \sum_{l_2 \in G_j} \sum_{t_k \in T_{l_1 l_2}} K_h(t, t_k)}, & \text{if } i \neq j; \\ 1 - \sum_{s \neq i} \mathbf{P}_{Gis}(t), & \text{if } i = j. \end{cases}$$

Set $\tilde{\pi}_G$ as the stationary distribution of $\mathbf{P}_G$. Define $\mathbf{Y}_G(t) = B\mathbf{P}_G$. Let $\mathbf{Y}_G^* = B\mathbf{P}_G^*$. Similar to the proof of Theorem 1 in [28], using the derivative of stationary distribution, we can obtain

$$\frac{\partial \pi_G^{*\top}(t)}{\partial \mathbf{Y}_{Gij}^*(t)} = \pi_G^{*\top}(t) \frac{\partial \mathbf{P}_G^{*\top}(t)}{\partial \mathbf{Y}_{Gij}^*(t)} A^\#(t), \quad (\text{C.11})$$

where

$$\left(\frac{\partial \mathbf{P}_G^{*\top}(t)}{\partial \mathbf{Y}_{Gij}^*(t)}\right)_{ij} = \left(\frac{\partial \mathbf{P}_G^{*\top}(t)}{\partial \mathbf{Y}_{Gij}^*(t)}\right)_{jj} = -\left(\frac{\partial \mathbf{P}_G^{*\top}(t)}{\partial \mathbf{Y}_{Gij}^*(t)}\right)_{ji} = -\left(\frac{\partial \mathbf{P}_G^{*\top}(t)}{\partial \mathbf{Y}_{Gij}^*(t)}\right)_{ii} = \frac{1}{B},$$

and other elements of the derivative matrix are zero.

Then we consider the difference between $\mathbf{Y}_G(t)$ and $\mathbf{Y}_G^*(t)$. For $i \neq j$,

$$\begin{aligned} & \mathbf{Y}_{Gij}(t) - \mathbf{Y}_{Gij}^*(t) \\ &= \left(\frac{\sum_{l_1 \in G_i} \sum_{l_2 \in G_j} \sum_{t_k \in T_{l_1 l_2}} y_{l_1 l_2}(t_k) K_h(t, t_k)}{\sum_{l_1 \in G_i} \sum_{l_2 \in G_j} \sum_{t_k \in T_{l_1 l_2}} K_h(t, t_k)} - \frac{\sum_{l_1 \in G_i} \sum_{l_2 \in G_j} \sum_{t_k \in T_{l_1 l_2}} y_{l_1 l_2}^*(t_k) K_h(t, t_k)}{\sum_{l_1 \in G_i} \sum_{l_2 \in G_j} \sum_{t_k \in T_{l_1 l_2}} K_h(t, t_k)} \right) \\ &+ \left(\frac{\sum_{l_1 \in G_i} \sum_{l_2 \in G_j} \sum_{t_k \in T_{l_1 l_2}} y_{l_1 l_2}^*(t_k) K_h(t, t_k)}{\sum_{l_1 \in G_i} \sum_{l_2 \in G_j} \sum_{t_k \in T_{l_1 l_2}} K_h(t, t_k)} - \frac{1}{n_i n_j} \sum_{l \in G_i} \sum_{k \in G_j} \frac{\pi_k^*(t)}{\pi_l^*(t) + \pi_k^*(t)} \right) \\ &+ \left(\frac{1}{n_i n_j} \sum_{l \in G_i} \sum_{k \in G_j} \frac{\pi_k^*(t)}{\pi_l^*(t) + \pi_k^*(t)} - \mathbf{Y}_{Gij}^*(t) \right) \\ &=: (\mathbf{Y}_{Gij}(t) - \mathbf{Y}_{wij}^*(t)) + (\mathbf{Y}_{wij}^*(t) - \tilde{\mathbf{Y}}_{Gij}^*(t)) + (\tilde{\mathbf{Y}}_{Gij}^*(t) - \mathbf{Y}_{Gij}^*(t)). \end{aligned} \quad (\text{C.12})$$

Using central limit theorem, we have

$$\sqrt{n^2 M h}(\mathbf{Y}_{Gij} - \mathbf{Y}_{wij}^*) \xrightarrow{\mathcal{D}} N(0, \frac{1}{r_i r_j} \frac{\pi_{G_i}^*(t) \pi_{G_j}^*(t)}{(\pi_{G_i}^*(t) + \pi_{G_j}^*(t))^2} \int K^2(v) dv). \quad (\text{C.13})$$

Notice that when $nMh^5 \rightarrow 0$,

$$\begin{aligned} & \sqrt{n^2 M h}(\mathbf{Y}_{wij}^* - \tilde{\mathbf{Y}}_{Gij}^*(t)) \\ &= \sqrt{n^2 M h} \left(\frac{\sum_{l_1 \in G_i} \sum_{l_2 \in G_j} \sum_{t_k \in T_{l_1 l_2}} y_{l_1 l_2}^*(t_k) K_h(t, t_k)}{\sum_{l_1 \in G_i} \sum_{l_2 \in G_j} \sum_{t_k \in T_{l_1 l_2}} K_h(t, t_k)} - \frac{1}{n_i n_j} \sum_{l \in G_i} \sum_{k \in G_j} \frac{\pi_k^*(t)}{\pi_l^*(t) + \pi_k^*(t)} \right) \\ &\rightarrow \sqrt{n^2 M h} \left(\frac{h^2 \sum_{l \in G_i} \sum_{k \in G_j} \ddot{y}_{lk}^*(t) \int v^2 K(v) dv}{2n_i n_j} \right) \rightarrow 0. \end{aligned} \quad (\text{C.14})$$

Besides,

$$\begin{aligned} & \sqrt{n^2 M h} |\tilde{\mathbf{Y}}_G^*(t) - \mathbf{Y}_G^*(t)| \\ &= \sqrt{n^2 M h} \left| \frac{1}{n_i n_j} \sum_{l \in G_i} \sum_{k \in G_j} \frac{\pi_k^*(t)}{\pi_l^*(t) + \pi_k^*(t)} - \frac{\pi_{G_j}^*(t)}{\pi_{G_i}^*(t) + \pi_{G_j}^*(t)} \right| \\ &\lesssim \sqrt{n^4 M h} \delta_2 \rightarrow 0. \end{aligned} \quad (\text{C.15})$$

Combining (C.11), (C.12), (C.13), (C.14) and (C.15), we have,

$$\sqrt{n^2 M h}(\tilde{\pi}_G(t) - \pi_G^*(t)) \xrightarrow{\mathcal{D}} \Gamma(t)N(0, \Lambda(t)). \quad (\text{C.16})$$

Set $T_n(G) = \tilde{\pi}_G(t) - \pi_G^*(t)$ and $T_n(\hat{G}) = \hat{\pi}_G(t) - \pi_G^*(t)$. From Theorem 3.3, $P(\hat{G} = G) \rightarrow 1$. Therefore, for every $A \subset \mathbb{R}^B$, from

$$P(T_n(\hat{G}) \in A) = P(T_n(\hat{G}) \in A | \hat{G} = G)P(\hat{G} = G) + P(T_n(\hat{G}) \in A | \hat{G} \neq G)P(\hat{G} \neq G),$$

we can obtain

$$\lim_{n \rightarrow \infty} P(T_n(\hat{G}) \in A) = \lim_{n \rightarrow \infty} P(T_n(\hat{G}) \in A | \hat{G} = G) = \lim_{n \rightarrow \infty} P(T_n(G) \in A). \quad (\text{C.17})$$

From (C.16) and (C.17), we have

$$\sqrt{n^2 M h}(\hat{\pi}_G(t) - \pi_G^*(t)) \xrightarrow{\mathcal{D}} \Gamma(t)N(0, \Gamma(t)\Lambda(t)\Gamma(t)^\top).$$

□

C.3 Proof of Theorem 3.9

Proof. We first prove that $P(\{\hat{s}_i\}_{i=1}^{\hat{J}} \supset \{\eta_i\}_{i=1}^J) \rightarrow 1$. Suppose there are change points in the interval $(\hat{s}_i, \hat{s}_{i+1})$, and assume the first one be η_j . Let $\mathcal{J}_1 = (\hat{s}_i, \eta_j)$ and $\mathcal{J}_2 = (\eta_j, \hat{s}_{i+1})$.

$$\begin{aligned} & \left(\sum_{i=1}^2 (L(\hat{\beta}(\mathcal{J}_i), \mathcal{J}_i) + \gamma_1 |\hat{G}(\mathcal{J}_i)| |\mathcal{J}_i|) + \gamma_2 \right) - \left(L(\hat{\beta}(\mathcal{J}_1 \cup \mathcal{J}_2), \mathcal{J}_1 \cup \mathcal{J}_2) + \gamma_1 |\hat{G}(\mathcal{J}_1 \cup \mathcal{J}_2)| |\mathcal{J}_1 \cup \mathcal{J}_2| \right) \\ &= \gamma_1 \left(|\hat{G}(\mathcal{J}_1)| |\mathcal{J}_1| + |\hat{G}(\mathcal{J}_2)| |\mathcal{J}_2| - |\hat{G}(\mathcal{J}_1 \cup \mathcal{J}_2)| |\mathcal{J}_1 \cup \mathcal{J}_2| \right) + \gamma_2 + O_p(\delta_3). \end{aligned} \quad (\text{C.18})$$

Notice that

$$\begin{aligned} & |\hat{G}(\mathcal{J}_1)| |\mathcal{J}_1| + |\hat{G}(\mathcal{J}_2)| |\mathcal{J}_2| - |\hat{G}(\mathcal{J}_1 \cup \mathcal{J}_2)| |\mathcal{J}_1 \cup \mathcal{J}_2| \\ & \rightarrow |G^*(\mathcal{J}_1)| |\mathcal{J}_1| + |G^*(\mathcal{J}_2)| |\mathcal{J}_2| - |G^*(\mathcal{J}_1 \cup \mathcal{J}_2)| |\mathcal{J}_1 \cup \mathcal{J}_2|, \end{aligned}$$

and with the fact that $\mathcal{J}_1 \subset (\eta_{j-1}, \eta_j)$ and the right endpoint of $\mathcal{J}_1$ is exactly η_j , we have

$$|G^*(\mathcal{J}_1 \cup \mathcal{J}_2)| - |G^*(\mathcal{J}_1)| \geq 1.$$

Hence, with probability tending to 1,

$$(C.18) \leq \gamma_2 - \gamma_1 |\mathcal{J}_1| \leq \gamma_2 - \gamma_1 < 0,$$

which is contradictory to the definition of $\hat{\mathcal{P}}$. Therefore, there are no change points between the estimated ones. In another word, $P(\{\hat{s}_i\}_{i=1}^{\hat{J}} \supset \{\eta_i\}_{i=1}^J) \rightarrow 1$.

The above proof classifies that all change points are in the estimation set, and we then prove that all points in the set are change points. Suppose that $\hat{s}_i \notin \{\eta_i\}_{i=1}^J$, then there exists j such that $\hat{s}_i \in (\eta_j, \eta_{j+1})$. Let $\mathcal{J}_1 = (\eta_j, \hat{s}_i)$ and $\mathcal{J}_2 = (\hat{s}_i, \eta_{j+1})$.

Then we have

$$\begin{aligned} & \left(\sum_{i=1}^2 (L(\hat{\beta}(\mathcal{J}_i), \mathcal{J}_i) + \gamma_1 |\hat{G}(\mathcal{J}_i)| |\mathcal{J}_i|) + \gamma_2 \right) - \left(L(\hat{\beta}(\mathcal{J}_1 \cup \mathcal{J}_2), \mathcal{J}_1 \cup \mathcal{J}_2) + \gamma_1 |\hat{G}(\mathcal{J}_1 \cup \mathcal{J}_2)| |\mathcal{J}_1 \cup \mathcal{J}_2| \right) \\ &= \gamma_1 \left(|\hat{G}(\mathcal{J}_1)| |\mathcal{J}_1| + |\hat{G}(\mathcal{J}_2)| |\mathcal{J}_2| - |\hat{G}(\mathcal{J}_1 \cup \mathcal{J}_2)| |\mathcal{J}_1 \cup \mathcal{J}_2| \right) + \gamma_2 + O_p(\delta_3) \rightarrow \gamma_2 > 0 \end{aligned} \quad (\text{C.19})$$

holds with probability tending to 1. Note that (C.19) means that removing $\hat{s}_i$ in $\hat{\mathcal{P}}$ leads to a strictly smaller value of (2.5), which is contradictory to the definition of $\hat{\mathcal{P}}$. □

C.4 Proof of Corollary 3.10

Proof. First notice that $\frac{\|\pi^*(t) - \tilde{\pi}^{rf}(t)\|_2}{\|\pi^*(t)\|_2} = O_p(\sqrt{\frac{1}{nMh}})$, $\forall t \in [0, V]$ by noticing the consistency of group estimation by Theorem 3.3 and convergence rate results in [22]. Then we only need to prove that if $\frac{\|\pi(t) - \tilde{\pi}(t)\|_2}{\|\pi(t)\|_2} = O_p(\zeta)$, $\forall t \in [0, V]$, then $\frac{1}{|\mathcal{I}|} |L(\pi, \mathcal{I}) - L(\tilde{\pi}, \mathcal{I})| = O_p(\zeta)$. In fact, we can obtain

$$\begin{aligned} \frac{1}{|\mathcal{I}|} |L(\pi, \mathcal{I}) - L(\tilde{\pi}, \mathcal{I})| &= \frac{1}{|\mathcal{I}|} \left| \int_{t \in \mathcal{I}} l(\pi(t)) - l(\tilde{\pi}(t)) dt \right| \\ &\leq \frac{1}{|\mathcal{I}|} \int_{t \in \mathcal{I}} \frac{2}{n(n-1)} \sum_{(i,j): i \neq j} \bar{y}_{ij}(t) \left| \log\left(\frac{\pi_j(t)}{\pi_i(t) + \pi_j(t)}\right) - \log\left(\frac{\tilde{\pi}_j(t)}{\tilde{\pi}_i(t) + \tilde{\pi}_j(t)}\right) \right| dt, \end{aligned} \quad (\text{C.20})$$

where

$$\left| \log\left(\frac{\pi_j(t)}{\pi_i(t) + \pi_j(t)}\right) - \log\left(\frac{\tilde{\pi}_j(t)}{\tilde{\pi}_i(t) + \tilde{\pi}_j(t)}\right) \right| = \log\left(1 + \frac{\tilde{\pi}_i(t)/\tilde{\pi}_j(t) - \pi_i(t)/\pi_j(t)}{1 + \pi_i(t)/\pi_j(t)}\right)$$

is $O_p(\zeta)$ using Taylor expansion. Then (C.20) is $O_p(\zeta)$. $\square$

C.5 Proof of Lemma C.1

Proof. We first prove $\|\mathbf{I} - \mathbf{P}^*(t)\|_2 = O(1)$. We omit t in this part for simplicity. Let π_0 be the stationary distribution of $\mathbf{P}^*$, i.e., $\pi_0^\top \mathbf{P}^* = \pi_0$. Let $\Pi = \text{diag}(\pi_0)$. Then under Assumption 3.1,

$$\sqrt{\frac{\min_i \pi_{0i}}{\max_i \pi_{0i}}} \frac{\|\Pi^{1/2}(\mathbf{I} - \mathbf{P}^*)x\|_2}{\|\Pi^{1/2}x\|_2} \leq \frac{\|(\mathbf{I} - \mathbf{P}^*)x\|_2}{\|x\|_2} \leq \sqrt{\frac{\max_i \pi_{0i}}{\min_i \pi_{0i}}} \frac{\|\Pi^{1/2}(\mathbf{I} - \mathbf{P}^*)x\|_2}{\|\Pi^{1/2}x\|_2}.$$

Hence, we have

$$\|\mathbf{I} - \mathbf{P}^*\|_2 \asymp \max_{x \neq 0} \frac{\|\Pi^{1/2}(\mathbf{I} - \mathbf{P}^*)x\|_2}{\|\Pi^{1/2}x\|_2} \asymp \max_{x \neq 0} \frac{\|\Pi^{1/2}(\mathbf{I} - \mathbf{P}^*)\Pi^{-1/2}\Pi^{1/2}x\|_2}{\|\Pi^{1/2}x\|_2},$$

which is the maximum singular value of $\Pi^{1/2}(\mathbf{I} - \mathbf{P}^*)\Pi^{-1/2}$. From 2.4.3 of [39], it is a symmetric matrix and its spectral norm equals $1 - \mu_{\min}(\mathbf{P}^*)$, which is $O(1)$. Further,

$$\|\mathbf{I} - \mathbf{P}\|_2 \leq \|\mathbf{I} - \mathbf{P}^*\|_2 + \|\mathbf{P} - \mathbf{P}^*\|_2 = O_p(1)$$

using Lemma 5 of [22].

Let W be the inversion of $(\mathbf{I}^{-1})_{\mathcal{S}^c}(\mathbf{Q}^{-1})_{\mathcal{S}^c}$. Similar to Lemma 3 in [28], W is a $(n-B) \times (n-B)$ symmetric tridiagonal matrix, and for a vector x ,

$$x^\top W x \gtrsim \frac{1}{B} x_1^2 + \sum_{i=2}^{n-B} \frac{1}{B} (x_i - x_{i-1})^2 + \frac{1}{B} x_{n-B}^2,$$

which is corresponding to a diagonal-constant matrix with the minimum eigenvalue being $\frac{2}{B}(1 + \cos \frac{(n-B)\pi}{n-B+1})$. Hence, $\|(\mathbf{Q}^{-1})_{\mathcal{S}^c}\|_2 \lesssim \sqrt{\frac{B}{1 + \cos \frac{(n-B)\pi}{n-B+1}}}$.

Since $\mathbf{X}^\top \mathbf{X}$ is a matrix with diagonal elements $\mathbf{X}_{-1}(t)^\top \mathbf{X}_{-1}(t)$, it is sufficient to prove the same for $\mathbf{X}_{-1}(t)$. Then we conclude the results using Lemma 4 in [28]. $\square$

D Additional discussions

D.1 Independence assumption of pairwise comparisons

We assume that pairwise comparisons, the elements of $\mathcal{Y}$, are independent in Section 2.1. The independence assumption among pairwise comparisons is a standard and widely adopted simplification in the BT model. In many real-world applications, such as sports tournaments and recommender systems, comparisons are typically collected independently across individuals or time, which renders the assumption approximately valid in practice.

For example, Masarotto and Varin [10] analyze outcomes from the National Football League and American College Hockey under the independence assumption. Similarly, Maystre et al. [40] apply a time-varying BT model to datasets from the NBA and the Association of Tennis Professionals (ATP), also assuming independence between matches. In the context of recommender systems, Liu et al. [8] adopt the BT framework for product ranking based on independent pairwise comparisons. All of these analyses demonstrate strong performance in practical settings.

Furthermore, recent studies in reinforcement learning from human feedback (RLHF) continue to employ the BT model with the independence assumption, which has proven effective. For instance, Xiong et al. [41], Zhong et al. [42], and Zhu et al. [43] develop RLHF algorithms using pairwise comparisons assumed to be independent, and report successful outcomes across various tasks.

D.2 Smoothness condition of Assumption 3.1

We clarify that the smoothness condition in Assumption 3.1 is standard in dynamic settings and can be relaxed. In Section 2.2, it suffices for the score functions to be Lipschitz continuous to ensure consistent group identification (Theorem 3.3). The stronger smoothness assumption is only needed to derive the asymptotic distribution of the estimators (Theorem 3.6). Importantly, the Lipschitz condition is widely adopted in theoretical analyses of dynamic ranking problems, such as Assumption 5.2 of Bong et al. [29] and Assumption 1 of Karlé and Tyagi [44].

Furthermore, in Section 2.3, the general theoretical results in Theorem 3.9 rely only on Assumptions 3.7 and 3.8, which do not require the ability trajectories to be smooth. The smoothness condition in Corollary 3.10 is imposed solely to facilitate the application of Theorem 3.3, but as noted above, this can be weakened to a Lipschitz condition. Alternatively, one may use a segmented estimation strategy over a gridded time interval to estimate $\hat{G}(\mathcal{I})$, and then apply our proposed change detection framework without requiring any smoothness assumption.

Finally, we would like to emphasize that our method performs well even in the presence of abrupt changes. This is supported by simulation results in Section 4.2, where the underlying score trajectories (shown in Figures 5 and 6 in Section B.2) feature nonsmooth and abrupt changes, yet our method maintains strong performance.

D.3 Optimality of Theorem 3.6

Theorem 3.6 shows that the estimation error satisfies $\|\hat{\pi}_G(t) - \pi_G^*(t)\|_\infty = O_p((n^2 Mh)^{-1/2})$, which matches the optimal convergence rate. Specifically, for the refitted estimator $\hat{\pi}_i^{rf}$ defined in equation (2.3), we obtain the relative error rate $\|\hat{\pi}^{rf}(t) - \pi^*(t)\|_\infty / \|\pi^*(t)\|_\infty = O_p((n^2 Mh)^{-1/2})$. We analyze the result based on the effective sample size. On average, there are n/B items per group, each compared against roughly $(B-1)n$ others. Assuming the Epanechnikov kernel with bandwidth h , each pair contributes about $2Mh$ effective observations. Hence, the total number of comparisons used to estimate each $\hat{\pi}_i^{rf}(t)$ is of order $n^2 Mh$. Since we pool comparisons across all items within the same group to estimate each ability score, the estimation procedure leverages this aggregated information. This matches the optimal order $L^{-1/2}$ established in Chen et al. [39] and Karlé and Tyagi [44], where L denotes the average number of comparisons per item.