

Rank Aggregation via Heterogeneous Thurstone Preference Models

Tao Jin^{*‡} and Pan Xu^{†‡} and Quanquan Gu^{§||} and Farzad Farnoud^{¶||}

Abstract

We propose the Heterogeneous Thurstone Model (HTM) for aggregating ranked data, which can take the accuracy levels of different users into account. By allowing different noise distributions, the proposed HTM model maintains the generality of Thurstone’s original framework, and as such, also extends the Bradley-Terry-Luce (BTL) model for pairwise comparisons to heterogeneous populations of users. Under this framework, we also propose a rank aggregation algorithm based on alternating gradient descent to estimate the underlying item scores and accuracy levels of different users simultaneously from noisy pairwise comparisons. We theoretically prove that the proposed algorithm converges linearly up to a statistical error which matches that of the state-of-the-art method for the single-user BTL model. We evaluate the proposed HTM model and algorithm on both synthetic and real data, demonstrating that it outperforms existing methods.

1 Introduction

Rank aggregation refers to the task of recovering the order of a set of objects given pairwise comparisons, partial rankings, or full rankings obtained from a set of users or experts. Compared to rating items, comparison is a more natural task for humans which can provide more consistent results, in part because it does not rely on arbitrary scales. Furthermore, ranked data can be obtained not only by explicitly querying users, but also through passive data collection, i.e., by observing user behavior, for example product purchases, clicks on search engine results, choice of movies in streaming services, etc. As a result, rank aggregation has a wide range of applications, from classical social choice applications (de Borda, 1781) to information retrieval (Dwork et al., 2001), recommendation systems (Baltrunas et al., 2010), and bioinformatics (Aerts et al., 2006; Kim et al., 2015).

^{*}Department of Computer Science, University of Virginia, Charlottesville, VA 22904; e-mail: taoj@virginia.edu

[†]Department of Computer Science, University of California, Los Angeles, Los Angeles, CA 90095; e-mail: panxu@cs.ucla.edu

[‡]Equal contribution

[§]Department of Computer Science, University of California, Los Angeles, Los Angeles, CA 90095; e-mail: qgu@cs.ucla.edu

[¶]Department of Electrical and Computer Engineering, University of Virginia, Charlottesville, VA 22904; e-mail: farzad@virginia.edu

^{||}Co-corresponding authors.

In aggregating rankings, the raw data is often noisy and inconsistent. One approach to arrive at a single ranking is to assume a generative model for the data whose parameters include a true score for each of the items. In particular, Thurstone’s preference model (Thurstone, 1927) assumes that comparisons or partial rankings result from comparing versions of the true scores corrupted by additive noise. Special cases of Thurstone’s model include the popular Bradley-Terry-Luce (BTL) model for pairwise comparisons and the Plackett-Luce (PL) model for partial rankings. In these settings, estimating the true scores from data will allow us to identify the true ranking of the items. Various estimation and aggregation algorithms have been developed for Thurstone’s preference model and its special cases, including (Hunter, 2004; Guiver and Snelson, 2009; Hajek et al., 2014; Chen and Suh, 2015; Vojnovic and Yun, 2016; Negahban et al., 2017).

Conventional models of ranked data and aggregation algorithms that rely on them make the assumption that the data is either produced by a single user¹ or from a set of users that are similar. In real-world datasets, however, users that provide the raw data are usually diverse with different levels of familiarity with the objects of interest, thus providing data that is not uniformly reliable and should not have equal influence on the final result. This is of particular importance in applications such as aggregating expert opinions for decision-making and aggregating annotations provided by workers in crowd sourcing settings.

In this paper, we study the problem of rank aggregation for heterogeneous populations of users. We present a generalization of Thurstone’s model, called the *heterogeneous Thurstone model* (HTM), which allows users with different noise levels, as well as a certain class of adversarial users. Unlike previous efforts on rank aggregation for heterogeneous populations such as Chen et al. (2013); Kumar and Lease (2011), the proposed model maintains the generality of Thurstone’s framework and thus also extends its special cases such as BTL and PL models. We evaluate the performance of the method using simulated data for different noise distributions. We also demonstrate that the proposed aggregation algorithm outperforms the state-of-the-art method for real datasets on evaluating the difficulty of English text and comparing the population of a set of countries.

Our Contributions: Our main contributions are summarized as follows

- We propose a general model called the heterogeneous Thurstone model (HTM) for producing ranked data based on heterogeneous sources, which reduces to the heterogeneous BTL (HBTL) model when the noise follows the Gumbel distribution and to the heterogeneous Thurstone Case V (HTCV) model when the noise follows the normal distribution respectively.
- We develop an efficient algorithm for aggregating pairwise comparisons and estimating user accuracy levels for a wide class of noise distributions based on minimizing the negative log-likelihood loss via alternating gradient descent.
- We theoretically show that the proposed algorithm converges to the unknown score vector and the accuracy vector at a locally linear rate up to a tight statistical error under mild conditions.
- For models with specific noise distributions such as the HBTL and HTCv, we prove that the proposed algorithm converges linearly to the unknown score vector and accuracy vector up to statistical errors in the order of $O(n^2 \log(mn^2)/(mk))$, where k is sample size, n is the number

¹We use the term user to refer to any entity that provides ranked data. In specific applications other terms may be more appropriate, such as voter, expert, judge, worker, and annotator.

of items and m is the number of users. When $m = 1$, the statistical error matches the error bound in the state-of-the-art work for single user BTL model (Negahban et al., 2017).

- We conduct thorough experiments on both synthetic and real world data to validate our theoretical results and demonstrate the superiority of our proposed model and algorithm.

The reminder of this paper is organized as follows. In Section 2, we review the most related work in the literature. In Section 3, we propose a family of heterogeneous Thurstone models. In Section 4, we propose an efficient algorithm for learning the ranking from pairwise comparisons. We theoretically analyze the convergence of the proposed algorithm in Section 5. Thorough experimental results are presented in Section 6 and Section 7 concludes the paper.

2 Additional Related Work

The problem of rank aggregation has a long history, dating back to the works of de Borda (1781) and de Condorcet (1785) in the 18th century, where the problems of social choice and voting were discussed. More recently, the problem of aggregating pairwise comparisons, where comparisons are incorrect with a given probability p , was studied by Braverman and Mossel (2008) and Wauthier et al. (2013). Instead of assuming the same probability for all comparisons to be incorrect, it is natural to assume that the comparison of similar items is more likely to be noisy than those items that are distinctly different. This intuition is reflected in the random utility model (RUM), also known as *Thurstone’s model* (Thurstone, 1927), where each item has a true score, and users provide rankings of subsets of items by comparing approximate version of these scores corrupted by additive noise.

When restricted to comparing pairs of items, Thurstone’s model reduces to the BTL model (Zermelo, 1929; Bradley and Terry, 1952; Luce, 1959; Hunter, 2004) if the noise follows the Gumbel distribution, and to the Thurstone Case V (TCV) model (Thurstone, 1927) if the noise is normally distributed. Recently, Negahban et al. (2012) proposed Rank Centrality, an iterative method with a random walk interpretation and showed that it performs as well as the maximum likelihood (ML) solution (Zermelo, 1929; Hunter, 2004) for BTL models and provided non asymptotic performance guarantees. Chen and Suh (2015) studied identifying the top-K candidates under the BTL model and its sample complexity.

Thurstone’s model can also be used to describe data from comparisons of multiple items. Hajek et al. (2014) provided an upper bound on the error of the ML estimator and studied its optimality when data consists of partial rankings (as opposed to pairwise comparisons) under the PL model. Yu (2000) studied order statistics under the normal noise distribution with consideration of item confusion covariance and user perception shift in a Bayesian model. Weng and Lin (2011) proposed a Bayesian approximation method for game player ranking with results from two-team matches. Guiver and Snelson (2009) studied the ranking aggregation problem with partial ranking (PL model) in a Bayesian framework. However, due to the nature of Bayesian method, above mentioned work provided few theoretical analysis. Vojnovic and Yun (2016) studied the parameter estimation problem for Thurstone models where first choices among a set of alternatives are observed. Raman and Joachims (2014, 2015) proposed the peer grading methods for solving a similar problem as ours, while the generative models to aggregate partial rankings and pairwise comparisons are completely

different. Very recently, [Zhao et al. \(2018\)](#) proposed the k -RUM model which assumes that the rank distribution has a mixture of k RUM components. They also provided the analyses of identifiability and efficiency of this model.

Almost all aforementioned works assume that all the data is provided by a single user or that all users have the same accuracy. However, this assumption is rarely satisfied in real-world datasets. The accuracy levels of different users are considered in [Kumar and Lease \(2011\)](#), which assumes that each user is correct with a certain probability and studies the problem via simulation methods such as naive Bayes and majority voting. In their pioneering work, [Chen et al. \(2013\)](#) studied rank aggregation in a crowd-sourcing environment for pairwise comparisons, modeled via the BTL or TCV model, where noisy BTL comparisons are assumed to be further corrupted. They are flipped with a probability that depends on the identity of the worker. The k -RUM model proposed by [Zhao et al. \(2018\)](#) considered a mixture of ranking distributions, without using extra information on who contributed the comparison, it may suffer from common mixture model issues.

3 Modeling Heterogeneous Ranked Data

Before introducing our Heterogeneous Thurstone Model, we start by providing some preliminaries of Thurstone’s preference model in further detail. Consider a set of n items. The score vector for the items is denoted by $\mathbf{s} = (s_1, \dots, s_n)^\top$. These items/objects are evaluated by a set of m independent users. Each user may be asked to express their preference concerning a subset of items $\{i_1, \dots, i_h\} \subseteq [n]$, where $2 \leq h \leq n$. For each item i , the user first estimates an empirical score for it as

$$z_i = s_i + \epsilon_i, \quad (3.1)$$

where ϵ_i is a random noise introduced by this evaluation process. This coarse estimate of score z_i is still implicit and cannot be queried or observed by the ranking algorithm. Instead, the user only produces a ranking of these h items by sorting the scores z_i . We thus have

$$\Pr(\pi_1 \succ \pi_2 \succ \dots \succ \pi_h) = \Pr(z_{\pi_1} > z_{\pi_2} > \dots > z_{\pi_h}), \quad (3.2)$$

where $i \succ j$ indicates that i is preferred to j by this user and $\{\pi_1, \dots, \pi_h\}$ is a permutation of $\{i_1, \dots, i_h\}$. Each time item i is compared with other items, a new score estimate z_i is produced by the user for are commonly assumed to be i.i.d. ([Braverman and Mossel, 2008](#); [Negahban et al., 2012](#); [Wauthier et al., 2013](#)).

3.1 The Heterogeneous Thurstone Model

In real-world applications, users often have different levels of expertise and some may even be adversarial. Therefore, it is natural for us to propose an extension of the Thurstone’s model presented above, referred to as the *Heterogeneous Thurstone Model* (HTM), which has the flexibility to reflect the different levels of expertise of different users. Specifically, we assume that each user has a different level of making mistakes in evaluating items, i.e., the evaluation noise of user u is

controlled by a scaling factor $\gamma_u > 0$. The proposed model is then represented as follows:

$$z_i^u = s_i + \epsilon_i / \gamma_u. \quad (3.3)$$

Based on the estimated scores of each user for each item, the probability of a certain ranking of h items provided by user u is again given by (3.2). While this extension actually applies to both pairwise comparisons and multi-item orderings, we mainly focus on pairwise comparisons in this paper.

When two items i and j are compared by user u , we denote by Y_{ij}^u the random variable representing the result,

$$Y_{ij}^u = \begin{cases} 1 & \text{if } i \succ j; \\ 0 & \text{if } i \prec j. \end{cases} \quad (3.4)$$

Let F denote the CDF of $\epsilon_j - \epsilon_i$, where ϵ_i and ϵ_j are two i.i.d. random variables. For the result Y_{ij}^u of comparison of i and j by user u , we have

$$\Pr(Y_{ij}^u = 1; s_i, s_j, \gamma_u) = \Pr(\epsilon_j - \epsilon_i < \gamma_u(s_i - s_j)) = F(\gamma_u(s_i - s_j)). \quad (3.5)$$

It is clear that the larger the value of γ_u , the more accurate the user is, since large $\gamma_u > 0$ increases the probability of preferring an item with higher score to one with lower score.

We now consider several special cases arising from specific noise distributions. First, if ϵ_i follows a Gumbel distribution with mean 0 and scale parameter 1, then we obtain the following *Heterogeneous BTL* (HBTL) model:

$$\log \Pr(Y_{ij}^u = 1; s_i, s_j, \gamma_u) = \log \frac{e^{\gamma_u s_i}}{e^{\gamma_u s_i} + e^{\gamma_u s_j}} = -\log(1 + \exp(-\gamma_u(s_i - s_j))), \quad (3.6)$$

which follows from the fact that the difference between two independent Gumbel random variables has the logistic distribution. We note that setting $\gamma_u = 1$ recovers the traditional BTL model (Bradley and Terry, 1952).

If ϵ_i follows the standard normal distribution, we obtain the following *Heterogeneous Thurstone Case V* (HTCV) model:

$$\log \Pr(Y_{ij}^u = 1; s_i, s_j, \gamma_u) = \log \Phi\left(\frac{\gamma_u(s_i - s_j)}{\sqrt{2}}\right), \quad (3.7)$$

where Φ is the CDF of the standard normal distribution. Again, when $\gamma_u = 1$, this reduces to Thurstone's Case V (TCV) model for pairwise comparisons (Thurstone, 1927).

Adversarial users: Under our heterogeneous framework, we can also model a certain class of adversarial users, whose goal is to make the estimated ranking be the opposite of the true ranking, so that, for example, an inferior item is ranked higher than the alternatives. We assume for adversarial users, the score of item i is $C - s_i$, for some constant C . Changing s_i to $C - s_i$ in (3.5) is equivalent to assuming the user has a negative accuracy γ_u . In this way, the accuracy of the user is determined by the magnitude $|\gamma_u|$ and its trustworthiness by $\text{sign}(\gamma_u)$, as illustrated in Figure 1. When adversarial users are present, this will facilitate optimizing the loss function, since instead of

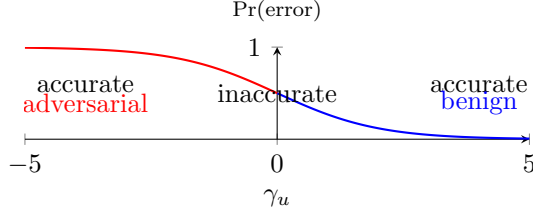


Figure 1: The effect of γ_u on the probability of error for a BTL comparison in which items have scores 0 and 1. In particular, for large negative values of γ_u , the user is accurate (with a high level of expertise) but adversarial.

solving the combinatorial optimization problem of deciding which users are adversarial, we simply optimize the value of γ_u for each user.

One relevant work to ours is the CrowdBT algorithm proposed by [Chen et al. \(2013\)](#), where they also explored the accuracy level of different users in learning a global ranking. In particular, they assume that each user has a probability η_u of making mistakes in comparing items i and j :

$$\Pr(Y_{ij}^u = 1; s_i, s_j, \eta_u) = \eta_u \Pr(i \succ j) + (1 - \eta_u) \Pr(j \succ i), \quad (3.8)$$

where $\Pr(i \succ j)$ and $\Pr(j \succ i)$ follow the BTL model. This translates to introducing a parameter in the likelihood function to quantify the reliability of each pairwise comparison. This parameterization, however, deviates from the additive noise in Thurstonian models defined as in (3.1) such as BTL and Thurstone’s Case V. Specifically, the Thurstonian model explains the noise observed in pairwise comparisons as resulting from the additive noise in estimating the latent item scores. Therefore, the natural extension of Thurstonian models to a heterogeneous population of users is to allow different noise levels for different users, as was done in (3.3). As a result, CrowdBT cannot be easily extended to settings where more than two items are compared at a time. In contrast, the model proposed here is capable to describe such generalizations of Thurstonian models, such as the PL model.

4 Optimization and Rank Aggregation

In this section, we define the pairwise comparison loss function for the population of users and propose an efficient and effective optimization algorithm to minimize it. We denote by $\mathbf{Y}^u$ the matrix containing all pairwise comparisons Y_{ij}^u of user u on items i and j . The entries of $\mathbf{Y}^u$ are 0/1/?, where ? indicates that the pair was not compared by the user. Furthermore, let $\mathcal{D}_u$ denote the set of all pairs (i, j) compared by user u . We define the loss function for each user u as

$$\begin{aligned} \mathcal{L}_u(\mathbf{s}, \gamma_u; \mathbf{Y}^u) &= -\frac{1}{k_u} \sum_{(i,j) \in \mathcal{D}_u} \log \Pr(Y_{ij}^u = 1 | s_i, s_j, \gamma_u) \\ &= -\frac{1}{k_u} \sum_{(i,j) \in \mathcal{D}_u} \log F(\gamma_u(s_i - s_j)), \end{aligned}$$

where $k_u = |\mathcal{D}_u|$ is the number of comparisons by user u . Then, the total loss function for m users is

$$\mathcal{L}(\mathbf{s}, \boldsymbol{\gamma}; \mathbf{Y}) = \frac{1}{m} \sum_{u=1}^m \mathcal{L}_u(\mathbf{s}, \boldsymbol{\gamma}_u; \mathbf{Y}^u), \quad (4.1)$$

where $\boldsymbol{\gamma} = (\gamma_1, \dots, \gamma_m)^\top$ and $\mathbf{Y} = (\mathbf{Y}^1, \dots, \mathbf{Y}^m)$. We denote the unknown true score vector as $\mathbf{s}^*$ and the true accuracy vector as $\boldsymbol{\gamma}^*$. Given observation $\mathbf{Y}$, our goal is to recover $\mathbf{s}^*$ and $\boldsymbol{\gamma}^*$ via minimizing the loss function in (4.1). To ensure the identifiability of $\mathbf{s}^*$, we follow [Negahban et al. \(2017\)](#) to assume that $\mathbf{1}^\top \mathbf{s}^* = \sum_{i=1}^n s_i^* = 0$, where $\mathbf{1} \in \mathbb{R}^n$ is the all one vector. The following proposition shows that the loss function $\mathcal{L}$ is convex in $\mathbf{s}$ and in $\boldsymbol{\gamma}$ separately if the PDF of ϵ_i is log-concave.

Proposition 4.1. If the distribution of the noise ϵ_i in (3.3) is log-concave, then the loss function $\mathcal{L}(\mathbf{s}, \boldsymbol{\gamma}; \mathbf{Y})$ given in (4.1) is convex in $\mathbf{s}$, and in $\boldsymbol{\gamma}$ respectively.

The log-concave family includes many well-known distributions such as normal, exponential, Gumbel, gamma and beta distributions. In particular, the noise distributions used in BTL and Thurstone's Case V (TCV) models fall into this category. Although the loss function $\mathcal{L}$ is non convex with respect to the joint variable $(\mathbf{s}, \boldsymbol{\gamma})$, Proposition 4.1 inspires us to perform alternating gradient descent ([Jain et al., 2013](#)) on $\mathbf{s}$ and $\boldsymbol{\gamma}$ to minimize the loss function. As is shown in Algorithm 1, we alternating perform gradient descent update on $\mathbf{s}$ (or $\boldsymbol{\gamma}$) while fixing $\boldsymbol{\gamma}$ (or $\mathbf{s}$) at each iteration. In addition to the alternating gradient descent steps, we shift $\mathbf{s}^{(t)}$ in Line 4 of Algorithm 1 such that $\mathbf{1}^\top \mathbf{s}^{(t)} = 0$ to avoid the aforementioned identifiability issue of $\mathbf{s}^*$. After T iterations, given the output $\mathbf{s}^{(T)}$, the estimated ranking of the items is obtained by sorting $\{s_1^{(T)}, \dots, s_n^{(T)}\}$ in descending order (item with the highest score in $\mathbf{s}^{(T)}$ is the most preferred).

Algorithm 1 HTMs with Alternating Gradient Descent

- 1: **input:** learning rates $\eta_1, \eta_2 > 0$, initial points $\mathbf{s}^{(0)}$ and $\boldsymbol{\gamma}^{(0)}$ satisfying $\|\mathbf{s}^{(0)} - \mathbf{s}^*\|_2^2 + \|\boldsymbol{\gamma}^{(0)} - \boldsymbol{\gamma}^*\|_2^2 \leq r$, number of iteration T , comparison results by users $\mathbf{Y}$.
 - 2: **for** $t = 0, \dots, T - 1$ **do**
 - 3: $\tilde{\mathbf{s}}^{(t+1)} = \mathbf{s}^{(t)} - \eta_1 \nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}^{(t)}, \boldsymbol{\gamma}^{(t)}; \mathbf{Y})$
 - 4: $\mathbf{s}^{(t+1)} = (\mathbf{I} - \mathbf{1}\mathbf{1}^\top/n) \tilde{\mathbf{s}}^{(t+1)}$
 - 5: $\boldsymbol{\gamma}^{(t+1)} = \boldsymbol{\gamma}^{(t)} - \eta_2 \nabla_{\boldsymbol{\gamma}} \mathcal{L}(\mathbf{s}^{(t)}, \boldsymbol{\gamma}^{(t)}; \mathbf{Y})$
 - 6: **end for**
 - 7: **output:** $\mathbf{s}^{(T)}, \boldsymbol{\gamma}^{(T)}$.
-

As we will show in the next section, the convergence of Algorithm 1 to the optimal points $\mathbf{s}^*$ and $\boldsymbol{\gamma}^*$ is guaranteed if an initialization such that $\mathbf{s}^{(0)}$ and $\boldsymbol{\gamma}^{(0)}$ are close to the unknown parameters is available. In practice, to initialize $\mathbf{s}$, we can use the solution provided by the rank centrality algorithm ([Negahban et al., 2012](#)) or start from uniform or random scores. In this paper, we initialize $\mathbf{s}$ and $\boldsymbol{\gamma}$, as $\mathbf{s}^{(0)} = \mathbf{1}$ and $\boldsymbol{\gamma}^{(0)} = \mathbf{1}$. We note that multiplying $\mathbf{s}$ or $\boldsymbol{\gamma}$ by a negative constant does not alter the loss but reverses the estimated ranking. Implicit in our initialization is the assumption that the majority of the users are trustworthy and thus have positive γ . When data is sparse, there

may be subsets of items that are not compared directly or indirectly. In such cases, regularization may be necessary, which is discussed in further detail in Section 6.

5 Theoretical Analysis of the Proposed Algorithm

In this section, we provide the convergence analysis of Algorithm 1 for the general loss function defined in (4.1). Without loss of generality, we assume the number of observations $k_u = k$ for all users $u \in [m]$ throughout our analysis. Since there's no specific requirement on the noise distributions in the general HTM model, to derive the linear convergence rate, we need the following conditions on the loss function $\mathcal{L}$, which are standard in the literature of alternating minimization (Jain et al., 2013; Zhu et al., 2017; Xu et al., 2017b,a; Zhang et al., 2018; Chen et al., 2018). Note that all these conditions can actually be verified once we specify the noise distribution in specific models. We provide the justifications of these conditions in the appendix.

Condition 5.1 (Strong Convexity). $\mathcal{L}$ is μ_1 -strongly convex with respect to $\mathbf{s} \in \mathbb{R}^n$ and μ_2 -strongly convex with respect to $\boldsymbol{\gamma} \in \mathbb{R}^m$. In particular, there is a constant $\mu_1 > 0$ such that for all $\mathbf{s}, \mathbf{s}' \in \mathbb{R}^n$,

$$\mathcal{L}(\mathbf{s}, \boldsymbol{\gamma}) \geq \mathcal{L}(\mathbf{s}', \boldsymbol{\gamma}) + \langle \nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}', \boldsymbol{\gamma}), \mathbf{s} - \mathbf{s}' \rangle + \mu_1/2 \|\mathbf{s} - \mathbf{s}'\|_2^2.$$

And there is a constant $\mu_2 > 0$ such that for all $\boldsymbol{\gamma}, \boldsymbol{\gamma}' \in \mathbb{R}^m$, it holds

$$\mathcal{L}(\mathbf{s}, \boldsymbol{\gamma}) \geq \mathcal{L}(\mathbf{s}, \boldsymbol{\gamma}') + \langle \nabla_{\boldsymbol{\gamma}} \mathcal{L}(\mathbf{s}, \boldsymbol{\gamma}'), \boldsymbol{\gamma} - \boldsymbol{\gamma}' \rangle + \mu_2/2 \|\boldsymbol{\gamma} - \boldsymbol{\gamma}'\|_2^2.$$

Condition 5.2 (Smoothness). $\mathcal{L}$ is L_1 -smooth with respect to $\mathbf{s} \in \mathbb{R}^n$ and L_2 -smooth with respect to $\boldsymbol{\gamma} \in \mathbb{R}^m$. In particular, there is a constant $L_1 > 0$ such that for all $\mathbf{s}, \mathbf{s}' \in \mathbb{R}^n$, it holds

$$\mathcal{L}(\mathbf{s}, \boldsymbol{\gamma}) \leq \mathcal{L}(\mathbf{s}', \boldsymbol{\gamma}) + \langle \nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}', \boldsymbol{\gamma}), \mathbf{s} - \mathbf{s}' \rangle + L_1/2 \|\mathbf{s} - \mathbf{s}'\|_2^2.$$

And there is a constant $L_2 > 0$ such that for all $\boldsymbol{\gamma}, \boldsymbol{\gamma}' \in \mathbb{R}^m$, it holds

$$\mathcal{L}(\mathbf{s}, \boldsymbol{\gamma}) \leq \mathcal{L}(\mathbf{s}, \boldsymbol{\gamma}') + \langle \nabla_{\boldsymbol{\gamma}} \mathcal{L}(\mathbf{s}, \boldsymbol{\gamma}'), \boldsymbol{\gamma} - \boldsymbol{\gamma}' \rangle + L_2/2 \|\boldsymbol{\gamma} - \boldsymbol{\gamma}'\|_2^2.$$

The next condition is a variant of the usual Lipschitz gradient condition. It is worth noting that the gradient is derived with respect to $\mathbf{s}$ (or $\boldsymbol{\gamma}$), while the upper bound is the difference of $\boldsymbol{\gamma}$ (or $\mathbf{s}$). This condition is commonly imposed and verified in the analysis of expectation-maximization algorithms (Wang et al., 2015) and alternating minimization (Jain et al., 2013).

Condition 5.3 (First-order Stability). There are constants $M_1, M_2 > 0$ such that $\mathcal{L}$ satisfies

$$\begin{aligned} \|\nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}, \boldsymbol{\gamma}) - \nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}, \boldsymbol{\gamma}')\|_2 &\leq M_1 \|\boldsymbol{\gamma} - \boldsymbol{\gamma}'\|_2, \\ \|\nabla_{\boldsymbol{\gamma}} \mathcal{L}(\mathbf{s}, \boldsymbol{\gamma}) - \nabla_{\boldsymbol{\gamma}} \mathcal{L}(\mathbf{s}', \boldsymbol{\gamma})\|_2 &\leq M_2 \|\mathbf{s} - \mathbf{s}'\|_2, \end{aligned}$$

for all $\mathbf{s}, \mathbf{s}' \in \mathbb{R}^n$ and $\boldsymbol{\gamma}, \boldsymbol{\gamma}' \in \mathbb{R}^m$.

Note that the loss function in (4.1) is defined based on finitely many samples of observations. The next condition shows how close the gradient of the sample loss function is to the expected loss function.

Condition 5.4. Denote $\bar{\mathcal{L}}$ as the expected loss, where the expectation of $\mathcal{L}$ is taken over the random choice of the comparison pairs and the observation $\mathbf{Y}$. With probability at least $1 - 1/n$, we have

$$\begin{aligned}\|\nabla_{\mathbf{s}}\mathcal{L}(\mathbf{s}, \gamma) - \nabla_{\mathbf{s}}\bar{\mathcal{L}}(\mathbf{s}, \gamma)\|_2 &\leq \epsilon_1(k, n), \\ \|\nabla_{\gamma}\mathcal{L}(\mathbf{s}, \gamma) - \nabla_{\gamma}\bar{\mathcal{L}}(\mathbf{s}, \gamma)\|_2 &\leq \epsilon_2(k, n),\end{aligned}$$

where n is the number of items and k is the number of observations for each user. In addition, $\epsilon_1(k, n)$ and $\epsilon_2(k, n)$ will go to zero when sample size k goes to infinity.

$\epsilon_1(k, n)$ and $\epsilon_2(k, n)$ in Condition 5.4 are also called the statistical errors (Wang et al., 2015; Xu et al., 2017a) between the sample version gradient and the expected (population) gradient.

Now we deliver our main theory on the linear convergence of Algorithm 1 for general HTM models. Full proofs can be found in the appendix.

Theorem 5.5. For a general HTM model, assume Conditions 5.1, 5.2, 5.3 and 5.4 hold and that $M_1, M_2 \leq \sqrt{\mu_1\mu_2}/4$. Denote that $\|\mathbf{s}^*\|_\infty = s_{\max}$ and $\|\gamma^*\|_\infty = \gamma_{\max}$. Suppose the initialization guarantees that $\|\mathbf{s}^{(0)} - \mathbf{s}^*\|_2^2 + \|\gamma^{(0)} - \gamma^*\|_2^2 \leq r^2$, where $r = \min\{\mu_1/(2M_1), \mu_2/(2M_2)\}$. If we set the step size $\eta_1 = \eta_2 = \mu/(12(L^2 + M^2))$, where $L = \max\{L_1, L_2\}$, $\mu = \min\{\mu_1, \mu_2\}$ and $M = \max\{M_1, M_2\}$, then the output of Algorithm 1 satisfies

$$\|\mathbf{s}^{(T)} - \mathbf{s}^*\|_2^2 + \|\gamma^{(T)} - \gamma^*\|_2^2 \leq r^2 \rho^T + \frac{\epsilon_1(k, n)^2 + \epsilon_2(k, n)^2}{\mu^2}$$

with probability at least $1 - 1/n$, where the contraction parameter is $\rho = 1 - \mu^2/(48(L^2 + M^2))$.

Remark 5.6. Theorem 5.5 establishes the linear convergence of Algorithm 1 when the initial points are close to the unknown parameters. The first term on the right-hand side is called the optimization error, which goes to zero as iteration number t goes to infinity. The second term is called the statistical error of the HTM model, which goes to zero when sample size mk goes to infinity. Hence, the estimation error of our proposed algorithm converges to the order of $O((\epsilon_1(k, n)^2 + \epsilon_2(k, n)^2)/\mu^2)$ after $t = O(\log((\epsilon_1(k, n)^2 + \epsilon_2(k, n)^2)/\mu^2 r^2)/\log \rho)$ iterations.

Note that the results in Theorem 5.5 hold for any general HTM models with Algorithm 1 as a solver. In particular, if we run the alternating gradient descent algorithm on the HBTL and HTCVC models proposed in Section 3, we will also obtain linear convergence rate to the true parameters up to a statistical error in the order of $O(n^2 \log(mn^2)/(mk))$, which matches the state-of-the-art statistical error for such models (Negahban et al., 2017). We provide the implications of Theorem 5.5 on specific models in the supplementary material.

6 Experiments

In this section, we present experimental results to show the performance of the proposed algorithm on heterogeneous populations of users. The experiments are conducted on both synthetic and real data with both benign users and adversarial users. We use the Kendall's tau correlation (Kendall, 1948) between the estimated and true rankings to measure the similarity between rankings, which

is defined as $\tau = \frac{2(c-d)}{n(n-1)}$, where c and d are the number of pairs on which the two rankings agree and disagree, respectively. Pairs that are tied in at least one of the rankings are not counted in c or d .

Baseline methods: In Gumbel noise setting, we compare Algorithm 1 based on our proposed HBTL model with (i) the BTL model that can be optimized through iterative maximum-likelihood methods (Negahban et al., 2012) or spectral methods such as Rank Centrality (Negahban et al., 2017); and (ii) the CrowdBT algorithm (Chen et al., 2013), which is a variation of BTL that allows users with different levels of accuracy. In the normal noise setting, we compare Algorithm 1 based on our proposed HTCv model with TCV model. We also implemented a TCV equivalent of CrowdBT and report its performance as CrowdTCV.

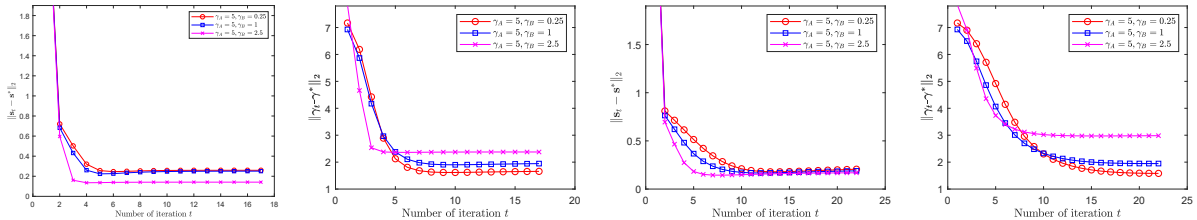
6.1 Experimental Results on Synthetic Data

We set number of items $n = 20$, number of users $m = 9$ and set the ground truth score vector $\mathbf{s}$ to be uniformly distributed in $[0, 1]$. The m users are divided into groups A and B , consisting of 3 and 6 users respectively. These two groups of users generate heterogeneous data in the sense that users in group A are more accurate than those in group B . We vary γ_A in the range of $\{2.5, 5, 10\}$ and γ_B in the range of $\{0.25, 1, 2.5\}$, which leads to in total 9 configurations of data generation. For each configuration, we conduct the experiment under the following two settings:

- (1) **Benign:** $\gamma_1, \dots, \gamma_3 = \gamma_A$ (Group A); $\gamma_4, \dots, \gamma_9 = \gamma_B$ (Group B).
- (2) **Adversarial:** $\gamma_1 = -\gamma_A$, $\gamma_2, \gamma_3 = \gamma_A$ (Group A); $\gamma_4, \gamma_5 = -\gamma_B$, $\gamma_6, \dots, \gamma_9 = \gamma_B$ (Group B).

We also test on various densities of compared pairs, which effectively controls the sample size. In particular, we choose 4 sets of α , which denote the portion of all possible pairs that are compared. The larger the value, the more pairs are compared by each user. The simulation process is as follows: we first generate $n(n-1)$ ordered pairs of items, where n is the number of items. This is equivalent to comparing each unique pair of items twice. Then for each pair of items, response from every annotator had a probability of α to be recorded and used for training the model. And α is chosen from $\{0.2, 0.4, 0.6, 0.8\}$ to make up for four runs. Each experiment is repeated 100 times with different random seeds.

Under setting (1), we plot the estimation error of Algorithm 1 v.s. number of iterations for HBTL and HTCv model in Figures 2(a)-2(b) and 2(c)-2(d) respectively. In all settings, our algorithm enjoys a linear convergence rate to the true parameters up to statistical errors, which is well aligned with the theoretical results in Theorem 5.5.



(a) Estimation error for $\mathbf{s}^*$ (b) Estimation error for γ^* (c) Estimation error for $\mathbf{s}^*$ (d) Estimation error for γ^*

Figure 2: Evolution of estimation errors vs. number of iterations t for HBTL model. (c)-(d): Evolution of estimation errors vs. number of iterations t for HTCv model.

When there is no adversarial users in the system, the ranking results for Gumbel noises under different configurations of γ_A and γ_B are shown in Table 1 and the ranking results for normal noises under different configurations of γ_A and γ_B are shown in Table 2. In both tables, each cell presents the Kendall’s tau correlation between the aggregated ranking and the ground truth, averaged over 100 trials. For each experimental setting, we use the bold text to denote the method which achieved highest performance. We also underline the highest score whenever there is a tie. It can be observed that in almost all cases, HBTL provides much more accurate rankings than BTL and HTCv significantly outperforms TCV as well. In particular, the larger the difference between γ_A and γ_B is, the more significant the improvement is. The only exception is when $\gamma_A = \gamma_B = 2.5$, in which case the data is not heterogeneous and our HTM model has no advantage. Nevertheless, our method still achieve comparable performance as BTL for non-heterogeneous data. It can also be observed that HBTL generally outperforms CrowdBT. But the advantage is not large, as CrowdBT also includes the different accuracy levels of different users. Importantly, however, as discussed in Section 3.1, CrowdBT is not compatible with the additive noise in Thurstonian models and cannot be extended in a natural way to ranked data other than pairwise comparison. In addition, unlike CrowdBT, our method enjoys strong theoretical guarantees while maintaining a good performance. Tables 1 and 2 also illustrate an important fact: If there are users with high accuracy, the presence of low quality data does not significantly impact the performance of Algorithm 1.

When there are a portion of adversarial users as stated in setting (2), we consider adversarial users whose accuracy level γ_u may take negative values as discussed above. The results for Gumbel and normal noises under setting (2) are shown in Table 3 and Table 4 respectively. It can be seen that in this case, the difference between the methods is even more pronounced.

6.2 Experimental Results on Real-World Data

We evaluate our method on two real-world datasets. The first one named “Reading Level” (Chen et al., 2013) contains English text excerpts whose reading difficulty level is compared by workers. 624 workers annotated 490 excerpts which resulting in a total of 12,728 pairwise comparisons. We also used Mechanical Turk to collect another dataset named “Country Population”. In this crowdsourcing task, we asked workers to compare the population between two countries and pick the one which has more population. Since the population ranking of countries has a universal consensus, which can be obtained by looking up demographic data, it is a better choice than those movie rankings which subjects to personal preferences. There were 15 countries as shown in Table 5 which made up to 105 pairwise comparisons. The values were collected according to the latest demography statistics on Wikipedia for each country as of March 2019. Each user was asked 16 pairs randomly selected from all those 105 pairs. A total of 199 workers provided response to this task through Mechanical Turk. These two datasets were both collected in online crowdsourcing environments so that we can expect varying worker accuracy where effectiveness of our approach can be demonstrated.

In real-world datasets, it may happen that two items from two subsets are never compared with each other, directly or indirectly. In such cases, the ranking will not be unique. Furthermore, if

Table 1: Kendall’s tau correlation for different method under Gumbel noise. Group A users all have the accuracy level γ_A and Group B users all have the accuracy level γ_B . α represents the portion of all possible pairwise comparisons each annotator labeled in the simulation. The **bold** number highlights the highest performance and the underlined number indicates a tie.

Observ. Ratio	γ_B	Methods	γ_A		
			2.5	5	10
$\alpha = 0.8$	0.25	BTL	0.767 $\pm$ 0.055	0.836 $\pm$ 0.043	0.879 $\pm$ 0.032
		CrowdBT	0.847 $\pm$ 0.042	0.928 $\pm$ 0.023	0.962 $\pm$ 0.016
		HBTL	0.850 $\pm$ 0.041	0.930 $\pm$ 0.024	0.964 $\pm$ 0.015
	1.0	BTL	0.863 $\pm$ 0.036	0.896 $\pm$ 0.028	0.923 $\pm$ 0.026
		CrowdBT	<u>0.875</u> $\pm$ 0.033	<u>0.930</u> $\pm$ 0.024	0.967 $\pm$ 0.018
		HBTL	<u>0.875</u> $\pm$ 0.033	<u>0.930</u> $\pm$ 0.024	0.969 $\pm$ 0.017
	2.5	BTL	0.933 $\pm$ 0.022	0.946 $\pm$ 0.019	0.959 $\pm$ 0.018
		CrowdBT	0.931 $\pm$ 0.024	0.947 $\pm$ 0.019	0.967 $\pm$ 0.017
		HBTL	0.931 $\pm$ 0.025	0.948 $\pm$ 0.021	0.972 $\pm$ 0.015
$\alpha = 0.6$	0.25	BTL	0.743 $\pm$ 0.064	0.814 $\pm$ 0.048	0.853 $\pm$ 0.037
		CrowdBT	0.823 $\pm$ 0.050	0.909 $\pm$ 0.034	0.954 $\pm$ 0.018
		HBTL	0.824 $\pm$ 0.051	0.908 $\pm$ 0.033	0.955 $\pm$ 0.018
	1.0	BTL	0.837 $\pm$ 0.036	0.872 $\pm$ 0.033	0.903 $\pm$ 0.033
		CrowdBT	0.853 $\pm$ 0.035	0.911 $\pm$ 0.031	0.955 $\pm$ 0.018
		HBTL	0.851 $\pm$ 0.033	0.913 $\pm$ 0.028	0.958 $\pm$ 0.017
	2.5	BTL	0.913 $\pm$ 0.032	0.931 $\pm$ 0.024	0.948 $\pm$ 0.021
		CrowdBT	0.910 $\pm$ 0.028	0.935 $\pm$ 0.020	0.961 $\pm$ 0.016
		HBTL	0.912 $\pm$ 0.029	0.936 $\pm$ 0.022	0.967 $\pm$ 0.017
$\alpha = 0.4$	0.25	BTL	0.671 $\pm$ 0.062	0.761 $\pm$ 0.053	0.812 $\pm$ 0.048
		CrowdBT	0.764 $\pm$ 0.065	0.872 $\pm$ 0.037	0.933 $\pm$ 0.024
		HBTL	0.769 $\pm$ 0.061	0.873 $\pm$ 0.034	0.934 $\pm$ 0.022
	1.0	BTL	0.791 $\pm$ 0.051	0.844 $\pm$ 0.043	0.866 $\pm$ 0.035
		CrowdBT	0.798 $\pm$ 0.050	0.889 $\pm$ 0.029	0.934 $\pm$ 0.027
		HBTL	0.806 $\pm$ 0.051	0.891 $\pm$ 0.031	0.936 $\pm$ 0.026
	2.5	BTL	0.882 $\pm$ 0.034	0.910 $\pm$ 0.030	0.919 $\pm$ 0.027
		CrowdBT	0.879 $\pm$ 0.034	0.912 $\pm$ 0.026	0.943 $\pm$ 0.022
		HBTL	0.880 $\pm$ 0.032	0.916 $\pm$ 0.028	0.945 $\pm$ 0.020
$\alpha = 0.2$	0.25	BTL	0.575 $\pm$ 0.095	0.663 $\pm$ 0.078	0.712 $\pm$ 0.069
		CrowdBT	0.644 $\pm$ 0.094	0.798 $\pm$ 0.055	0.884 $\pm$ 0.035
		HBTL	0.665 $\pm$ 0.090	0.805 $\pm$ 0.051	0.882 $\pm$ 0.034
	1.0	BTL	0.708 $\pm$ 0.073	0.768 $\pm$ 0.057	0.804 $\pm$ 0.039
		CrowdBT	0.696 $\pm$ 0.081	0.813 $\pm$ 0.052	0.876 $\pm$ 0.034
		HBTL	0.702 $\pm$ 0.079	0.819 $\pm$ 0.052	0.882 $\pm$ 0.034
	2.5	BTL	0.820 $\pm$ 0.044	<u>0.861</u> $\pm$ 0.043	0.883 $\pm$ 0.033
		CrowdBT	0.803 $\pm$ 0.048	0.857 $\pm$ 0.037	0.898 $\pm$ 0.030
		HBTL	0.807 $\pm$ 0.049	<u>0.861</u> $\pm$ 0.038	0.904 $\pm$ 0.029

Table 2: Kendall’s tau correlation for different methods under noise from the normal distribution. Group A users all have the accuracy level γ_A and Group B users all have the accuracy level γ_B . α represents the portion of all possible pairwise comparisons each annotator labeled in the simulation. The **bold** number highlights the highest performance and the underlined number indicates a tie.

Observ. Ratio	γ_B	Methods	γ_A		
			2.5	5	10
$\alpha = 0.8$	0.25	TCV	0.811±0.048	0.860±0.040	0.885±0.036
		CrowdTCV	0.881±0.032	<u>0.943</u> ±0.021	<u>0.971</u> ±0.014
		HTCV	0.882 ±0.030	<u>0.943</u> ±0.021	<u>0.971</u> ±0.015
	1.0	TCV	0.885±0.036	0.910±0.027	0.925±0.029
		CrowdTCV	<u>0.897</u> ±0.030	<u>0.944</u> ±0.020	0.973±0.015
		HTCV	<u>0.897</u> ±0.033	<u>0.944</u> ±0.020	0.975 ±0.013
	2.5	TCV	<u>0.945</u> ±0.021	0.956±0.018	0.965±0.018
		CrowdTCV	<u>0.945</u> ±0.021	0.954±0.019	0.976±0.014
		HTCV	0.944±0.021	0.959 ±0.017	0.981 ±0.014
$\alpha = 0.6$	0.25	TCV	0.763±0.059	0.830±0.043	0.850±0.041
		CrowdTCV	0.845±0.038	0.926 ±0.023	<u>0.961</u> ±0.020
		HTCV	0.846 ±0.040	0.925±0.025	<u>0.961</u> ±0.020
	1.0	TCV	0.862±0.038	0.892±0.034	0.912±0.025
		CrowdTCV	0.870±0.035	0.930±0.028	0.962±0.019
		HTCV	0.875 ±0.033	0.932 ±0.027	0.963 ±0.018
	2.5	TCV	0.927 ±0.027	0.943±0.021	0.955±0.019
		CrowdTCV	0.925±0.027	0.946±0.026	0.968±0.015
		HTCV	0.925±0.027	0.952 ±0.022	0.974 ±0.013
$\alpha = 0.4$	0.25	TCV	0.691±0.073	0.790±0.047	0.809±0.048
		CrowdTCV	0.804±0.050	0.901±0.028	0.946 ±0.022
		HTCV	0.808 ±0.049	0.904 ±0.028	0.945±0.022
	1.0	TCV	0.821±0.047	0.859±0.036	0.875±0.036
		CrowdTCV	0.832±0.044	0.900±0.035	0.946±0.020
		HTCV	0.836 ±0.043	0.904 ±0.032	0.947 ±0.020
	2.5	TCV	0.901 ±0.027	0.921±0.029	0.935±0.026
		CrowdTCV	0.895±0.031	0.923±0.028	0.950±0.019
		HTCV	0.895±0.030	0.926 ±0.025	0.957 ±0.018
$\alpha = 0.2$	0.25	TCV	0.599±0.088	0.688±0.077	0.738±0.060
		CrowdTCV	0.689±0.080	0.826±0.046	0.899 ±0.031
		HTCV	0.693 ±0.082	0.828 ±0.049	0.898±0.034
	1.0	TCV	0.733±0.070	0.791±0.055	0.815±0.041
		CrowdTCV	0.729±0.074	0.836±0.043	0.904 ±0.033
		HTCV	0.740 ±0.072	0.841 ±0.038	0.901±0.031
	2.5	TCV	0.856 ±0.041	0.878±0.036	0.888±0.032
		CrowdTCV	0.844±0.048	0.873±0.035	0.905±0.027
		HTCV	0.848±0.041	0.881 ±0.036	0.913 ±0.026

Table 3: Kendall’s tau correlation for different methods under noise from the Gumbel distribution when a third of the users are *adversarial*. The **bold** number highlights the highest performance and the underlined number indicates a tie.

Observ. Ratio	γ_B	Methods	γ_A		
			2.5	5	10
$\alpha = 0.8$	0.25	BTL	0.443±0.107	0.569±0.096	0.614±0.085
		CrowdBT	<u>0.852±0.044</u>	0.925±0.023	0.967±0.017
		HBTL	<u>0.852±0.045</u>	0.926±0.023	0.966±0.017
	1.0	BTL	0.575±0.089	0.663±0.071	0.710±0.074
		CrowdBT	0.873±0.037	0.931±0.023	0.967±0.014
		HBTL	0.875±0.037	0.932±0.024	0.966±0.017
	2.5	BTL	0.725±0.057	0.780±0.046	0.798±0.047
		CrowdBT	<u>0.931±0.025</u>	0.948±0.019	0.966±0.016
		HBTL	<u>0.931±0.025</u>	0.951±0.019	0.973±0.015
$\alpha = 0.6$	0.25	BTL	0.384±0.122	0.491±0.107	0.557±0.095
		CrowdBT	0.822±0.046	0.908±0.030	0.953±0.019
		HBTL	0.824±0.044	0.910±0.028	0.954±0.018
	1.0	BTL	0.546±0.097	0.627±0.078	0.670±0.080
		CrowdBT	0.852±0.037	0.911±0.029	0.954±0.018
		HBTL	0.854±0.037	0.914±0.028	0.956±0.019
	2.5	BTL	0.684±0.078	0.736±0.064	0.755±0.062
		CrowdBT	0.910±0.028	0.934±0.025	0.960±0.016
		HBTL	0.912±0.029	0.936±0.024	0.965±0.017
$\alpha = 0.4$	0.25	BTL	0.323±0.130	0.405±0.132	0.485±0.109
		CrowdBT	0.742±0.169	<u>0.877±0.033</u>	0.934±0.025
		HBTL	0.766±0.059	<u>0.877±0.035</u>	0.933±0.024
	1.0	BTL	0.448±0.118	0.544±0.096	0.583±0.094
		CrowdBT	0.810±0.044	0.886±0.031	<u>0.934±0.026</u>
		HBTL	0.819±0.045	0.891±0.031	<u>0.934±0.029</u>
	2.5	BTL	0.627±0.087	0.660±0.075	0.698±0.063
		CrowdBT	0.879±0.034	0.913±0.027	0.939±0.023
		HBTL	0.880±0.032	0.914±0.029	0.948±0.022
$\alpha = 0.2$	0.25	BTL	0.246±0.145	0.305±0.151	0.361±0.143
		CrowdBT	0.613±0.235	0.712±0.356	<u>0.848±0.256</u>
		HBTL	0.614±0.263	0.709±0.380	<u>0.848±0.249</u>
	1.0	BTL	0.336±0.154	0.407±0.127	0.452±0.132
		CrowdBT	0.644±0.282	0.795±0.176	0.878±0.038
		HBTL	0.650±0.281	0.807±0.172	0.888±0.040
	2.5	BTL	0.498±0.106	0.548±0.103	0.571±0.098
		CrowdBT	0.803±0.049	0.858±0.039	0.897±0.032
		HBTL	0.807±0.049	0.865±0.039	0.900±0.029

Table 4: Kendall tau correlation for different methods under noise from the normal distribution when a third of the users are *adversarial*. The **bold** number highlights the highest performance and the underlined number indicates a tie.

Observ. Ratio	γ_B	Methods	γ_A		
			2.5	5	10
$\alpha = 0.8$	0.25	TCV	0.471±0.105	0.590±0.095	0.640±0.075
		CrowdTCV	<u>0.882</u> ±0.034	0.938 ±0.023	0.972±0.017
		HTCV	<u>0.882</u> ±0.033	0.937±0.023	0.973 ±0.016
	1.0	TCV	0.642±0.083	0.694±0.068	0.722±0.064
		CrowdTCV	0.893±0.030	0.945±0.020	0.973±0.016
		HTCV	0.895 ±0.031	0.947 ±0.019	0.975 ±0.017
	2.5	TCV	0.772±0.055	0.804±0.045	0.821±0.050
		CrowdTCV	0.945 ±0.021	0.956±0.019	0.978±0.014
		HTCV	0.944±0.021	0.960 ±0.019	0.982 ±0.013
$\alpha = 0.6$	0.25	TCV	0.416±0.129	0.527±0.107	0.552±0.099
		CrowdTCV	<u>0.847</u> ±0.039	0.924±0.025	<u>0.960</u> ±0.020
		HTCV	<u>0.847</u> ±0.039	0.925 ±0.023	<u>0.960</u> ±0.020
	1.0	TCV	0.569±0.086	0.648±0.066	0.686±0.080
		CrowdTCV	0.866±0.036	0.930±0.024	<u>0.966</u> ±0.018
		HTCV	0.870 ±0.036	0.932 ±0.025	<u>0.966</u> ±0.018
	2.5	TCV	0.718±0.060	0.762±0.045	0.786±0.055
		CrowdTCV	0.926 ±0.027	0.949±0.023	0.969±0.014
		HTCV	0.925±0.027	0.952 ±0.020	0.972 ±0.014
$\alpha = 0.4$	0.25	TCV	0.359±0.119	0.472±0.116	0.514±0.103
		CrowdTCV	0.797±0.053	0.893±0.034	0.942 ±0.022
		HTCV	0.799 ±0.048	0.896 ±0.031	0.938±0.022
	1.0	TCV	0.487±0.116	0.577±0.088	0.587±0.088
		CrowdTCV	0.842±0.049	0.898±0.029	0.945 ±0.021
		HTCV	0.843 ±0.046	0.902 ±0.027	0.944±0.022
	2.5	TCV	0.648±0.073	0.704±0.071	0.718±0.066
		CrowdTCV	<u>0.895</u> ±0.031	0.925±0.031	0.951±0.021
		HTCV	<u>0.895</u> ±0.030	0.929 ±0.028	0.957 ±0.018
$\alpha = 0.2$	0.25	TCV	0.259±0.147	0.349±0.135	0.382±0.133
		CrowdTCV	0.600±0.340	0.826±0.044	0.895 ±0.038
		HTCV	0.636 ±0.282	0.828 ±0.044	0.893±0.036
	1.0	TCV	0.397±0.119	0.436±0.115	0.469±0.100
		CrowdTCV	0.721±0.065	0.834 ±0.043	0.901±0.033
		HTCV	0.736 ±0.066	0.832±0.046	0.905 ±0.032
	2.5	TCV	0.518±0.102	0.577±0.098	0.600±0.077
		CrowdTCV	0.843±0.049	0.873±0.037	0.908±0.030
		HTCV	0.848 ±0.041	0.880 ±0.036	0.917 ±0.028

data is sparse, the estimates may suffer from overfitting. To address these issues, regularization is often used. While this can be done in a variety of ways, for the sake of comparison with CrowdBT, we use virtual node regularization (Chen et al., 2013). Specifically, it is assumed that there is a virtual item of utility $s_0 = 0$ which is compared to all other items by a virtual user. This leads to the loss function $\mathcal{L} + \lambda_0 \mathcal{L}_0$, where $\mathcal{L}_0 = -\sum_{i \in [n]} \log F(s_0 - s_i) - \sum_{i \in [n]} \log F(s_i - s_0)$ and $\lambda_0 \geq 0$ is a tuning parameter.

We evaluate the performance of the methods for $\lambda_0 = 0, 1, 5, 10$. For different values of λ_0 , HBTL performs best more often than any other method and, in particular, it performs best for $\lambda_0 = 0$. Table 6 reports the best performance of each method across different regularization values for the two real-world data experiment. It can be observed that HBTL and HTCv outperform their counterparts, CrowdBT and CrowdTCV, as well as the uniform models, BTL and TCV.

Table 5: Ground truth for “Country Population” dataset.

Country	Population (million)
China	1410
India	1340
United States	324
Indonesia	264
Brazil	209
Pakistan	197
Nigeria	191
Bangladesh	165
Russia	144
Mexico	129
Japan	127
Ethiopia	105
Philippines	104.9
Egypt	97.6
Vietnam	95.5

Table 6: Performance of ranking algorithms on real-world dataset. The **bold** number highlights the highest performance.

Dataset	BTL	TCV	CrowdBT	CrowdTCV	HBTL	HTCV
Reading Level	0.3472	0.3452	0.3737	0.3672	0.3763	0.3729
Country Population	0.7524	0.7524	0.7714	0.7714	0.7905	0.7714

Table 7: Performance of ranking algorithms for the “Reading Level” dataset with different regularization parameters. The **bold** number highlights the highest performance.

	$\lambda_0 = 0$	$\lambda_0 = 1$	$\lambda_0 = 5$	$\lambda_0 = 10$
BTL	0.3299	0.3433	0.3472	0.3402
TCV	0.3294	0.3423	0.3452	0.3375
CrowdBT	0.3490	0.3737	0.3648	0.3535
CrowdTCV	0.3512	0.3672	0.3511	0.3388
HBTL	0.3608	0.3660	0.3719	0.3763
HTCV	0.3578	0.3696	0.3729	0.3680

Table 8: Performance of ranking algorithms for the “Country Population” dataset with different regularization parameters. The **bold** number highlights the highest performance.

	$\lambda_0 = 0$	$\lambda_0 = 1$	$\lambda_0 = 5$	$\lambda_0 = 10$
BTL	0.7524	0.7524	0.7524	0.7524
TCV	0.7524	0.7524	0.7524	0.7524
CrowdBT	0.7714	0.7714	0.7714	0.7524
CrowdTCV	0.7714	0.7714	0.7714	0.7524
HBTL	0.7905	0.7905	0.7524	0.7524
HTCV	0.7714	0.7714	0.7524	0.7524

6.3 Analysis on regularization effects

Detailed result with various regularization settings can be found in Table 7 and Table 8. The reported values are Kendall’s tau correlation. It shows that without regularization our method outperforms other methods. And with virtual node trick, it shows relative amount of improvement in the final ranking result, yet not essential. However, this method needs to tune another parameter λ_0 . If no gold/ground-truth comparison is given, there will be no validation standard to tune this parameter. Furthermore, the performance of the proposed methods is less dependent on the regularization parameter, which facilitates their application to real data. It is also interesting to see that our method is less prone to be affected by the regularization parameter.

7 Conclusions and Future Work

In this paper, we propose the heterogeneous Thurstone model for pairwise comparisons and partial rankings when data is produced by a population of users with diverse levels of expertise, as is often the case in real-world applications. The proposed model maintains the generality of Thurstone’s framework and thus also extends common models such as Bradley-Terry-Luce, Thurstone’s Case V,

and Plackett-Luce. We also developed an alternating gradient descent algorithm to estimate the score vector and expertise level vector simultaneously. We prove the local linear convergence of our algorithm for general HTM models satisfying mild conditions. We also prove the convergence of our algorithm for the two most common noise distributions, which leads to the HBTL and HTCVC models. Experiments on both synthetic and real data show that our proposed model and algorithm generally outperforms the competing methods, sometimes by a significant margin.

There are several interesting future directions that could be explored. First, it would be of great importance to devise a provable initialization algorithm since our current analysis relies on certain initialization methods that are guaranteed to be close to the true values. Another direction is extending the algorithm and analysis to the case of partial ranking such as the Plackett-Luce model. Finally, lower bounds on the estimation error would enable better evaluating algorithms for rank aggregation in heterogeneous Thurstone models.

A Implications of Specific Models

Our Theorem 5.5 is for general HTM models that satisfy Conditions 5.1, 5.2, 5.3 and 5.4. In this subsection, we will show that the linear convergence rate of Algorithm 1 can also be attained for specific models without assuming these conditions when the random noise ϵ_i in (3.3) follows the Gumbel distribution and the Gaussian distribution respectively.

A.1 Heterogeneous BTL model

We first consider the model with Gumbel noise. Specifically, $\{\epsilon_i\}_{i=1,\dots,n}$ follow the Gumbel distribution with mean 0 and scale parameter 1. Then we obtain the HBTL model defined in (3.6). The following corollary states the convergence result of Algorithm 1 for HBTL models.

Corollary A.1. Consider the HBTL model in (3.6) and assume the sample size $k \geq n^2 \log(mn)/m^2$. Let $\|\mathbf{s}^*\|_\infty = s_{\max}$, $\max_u |\gamma^{*u}| = \gamma_{\max}$ and $\min_u |\gamma^{*u}| = \gamma_{\min}$. Assume $\gamma_{\max} s_{\max} = C_0$ for a constant $C_0 \geq 1/2$ and

$$s_{\max} \leq \frac{\sqrt{m}\|\mathbf{s}^*\|_2}{n} \cdot \frac{\gamma_{\min} e^{5C_0}}{32\sqrt{2}\gamma_{\max}(1 + e^{5C_0})^2}.$$

Suppose the initialization points $\mathbf{s}^{(0)}$ and $\boldsymbol{\gamma}^{(0)}$ satisfy that $\|\mathbf{s}^{(0)} - \mathbf{s}^*\|_2^2 + \|\boldsymbol{\gamma}^{(0)} - \boldsymbol{\gamma}^*\|_2^2 \leq r^2$, where $r = \min\{\|\mathbf{s}^*\|_2/2, \gamma_{\min}/2, s_{\max}, \sqrt{\gamma_{\max} s_{\max}}\}$. If we set the step size small enough such that

$$\eta_1 = \eta_2 < \frac{mne^{5C_0}\Gamma_1^2}{6(1 + e^{5C_0})^2(m\Gamma_2^4 + 32n^2C_0^2)},$$

where $\Gamma_1 = \min\{\gamma_{\min}/2, \|\mathbf{s}^*\|_2\}$ and $\Gamma_2 = \max\{2\gamma_{\max}, 2\|\mathbf{s}^*\|_2\}$, then the output of Algorithm 1 satisfies

$$\|\mathbf{s}^{(T)} - \mathbf{s}^*\|_2^2 + \|\boldsymbol{\gamma}^{(T)} - \boldsymbol{\gamma}^*\|_2^2 \leq r^2 \rho^T + \frac{\Lambda n^2 \log(4mn^2)}{mk}$$

with probability at least $1 - 1/n$, where $\rho = 1 - \eta(\mu - 6\eta(\Gamma_2^4/n^2 + 32C_0^2/m))/2$ and Λ is a constant which only depends on $C_0, \gamma_{\max}$ and Γ_1 .

Remark A.2. According to Corollary A.1, when the initial points $\mathbf{s}^{(0)}$ and $\boldsymbol{\gamma}^{(0)}$ lie in a small neighborhood of the unknown parameter $\mathbf{s}^*, \boldsymbol{\gamma}^*$, the proposed algorithm converges linearly fast to a term in the order of $O(n^2 \log(mn^2)/(mk))$, which is called the statistical error of the HBTL model. Note that when $m = 1$, the statistical error reduces to $O(n^2 \log(n)/k)$, which matches the state-of-the-art estimation error bound for single user BTL model (Negahban et al., 2017). In addition, we assumed that $\|\mathbf{s}^*\|_\infty \lesssim O(\sqrt{m}/n \|\mathbf{s}^*\|_2)$ in order to derive the linear convergence of Algorithm 1. When m is in the same order of n , the requirement reduces to $\|\mathbf{s}^*\|_\infty \lesssim O(\|\mathbf{s}^*\|_2/\sqrt{n})$. This assumption is similar to the spikiness assumption in Agarwal et al. (2012); Negahban and Wainwright (2012), which ensures that there are not too many items that have zero or nearly zero scores.

A.2 Heterogeneous Thurstone Case V model

Now we consider the HTM model with Gaussian noise. Assume that $\{\epsilon_i\}_{i=1,\dots,n}$ are i.i.d. from $N(0, 1)$. Then the general HTM model becomes HTCv model defined in (3.7), which generalizes the single user TCV model (Thurstone, 1927). Before we present the convergence results of Algorithm 1 for this model, we first remark some notations of the normal distribution to simplify the presentation. In particular, let $\Phi(x)$ be the CDF of standard normal distribution. We define $H(x) = (\Phi'(x)^2 - \Phi(x)\Phi''(x))/\Phi(x)^2$, which can be verified to be a monotonically decreasing function.

Corollary A.3. Consider the HTCv model in (3.7) and assume the sample size $k \geq n^2 \log(mn)/m^2$. $s_{\max}, \gamma_{\max}, \gamma_{\min}$ and C_0 are defined the same as in Corollary A.1. Assume $s_{\max}$ satisfies

$$s_{\max} \leq \frac{\sqrt{m}\|\mathbf{s}^*\|_2}{n} \cdot \frac{\gamma_{\min}H(5C_0)}{30\gamma_{\max}(\Phi(-5C_0)^{-1} + H(-5C_0))}.$$

Suppose the initialization points $\mathbf{s}^{(0)}$ and $\boldsymbol{\gamma}^{(0)}$ satisfy that $\|\mathbf{s}^{(0)} - \mathbf{s}^*\|_2^2 + \|\boldsymbol{\gamma}^{(0)} - \boldsymbol{\gamma}^*\|_2^2 \leq r^2$, where $r = \min\{\|\mathbf{s}^*\|_2/2, \gamma_{\min}/2, s_{\max}, \sqrt{\gamma_{\max}s_{\max}}\}$. If we set the step size

$$\eta_1 = \eta_2 < \frac{mn\Gamma_1^2 H(5C_0)}{6(m\Gamma_2^4 + 50n^2C_0^2)H(-5C_0)^2},$$

where $\Gamma_1 = \min\{\gamma_{\min}/2, \|\mathbf{s}^*\|_2\}$ and $\Gamma_2 = \max\{2\gamma_{\max}, 2\|\mathbf{s}^*\|_2\}$, then the output of Algorithm 1 satisfies

$$\|\mathbf{s}^{(T)} - \mathbf{s}^*\|_2^2 + \|\boldsymbol{\gamma}^{(T)} - \boldsymbol{\gamma}^*\|_2^2 \leq r^2 \rho^T + \frac{\Lambda' n^2 \log(4mn^2)}{mk}$$

with probability at least $1 - 1/n$, where $\rho = 1 - \eta(\mu - 6\eta(\Gamma_2^4/n^2 + 32C_0^2/m))/2$ and Λ' is a constant which only depends on $C_0, \gamma_{\max}$ and Γ_1 .

Remark A.4. Corollary A.3 suggests that under suitable initialization, Algorithm 1 enjoys a linear convergence rate when the random noise follows the standard normal distribution. The statistical error for the HTCv model is in the order of $O(n^2 \log(mn^2)/(mk))$. We again need the ‘spikiness’

assumption on the unknown score vector $\mathbf{s}^*$ in order to ensure the algorithm to find the true parameter. The results are almost the same as those of the HBTL model presented in Corollary A.1 except that the constants in the HTCv model depends on the normal CDF Φ and its first and second derivatives.

B Proof of the Generic Model

In this section, we provide the proof of Theorem 5.5 for general heterogeneous Thurstone models.

Proof of Theorem 5.5. According to the update in Algorithm 1 and the fact that $\mathbf{1}^\top \mathbf{s}^* = 0$, we have

$$\begin{aligned}\|\mathbf{s}^{(t+1)} - \mathbf{s}^*\|_2^2 &= \|(\mathbf{I} - \mathbf{1}\mathbf{1}^\top/n)(\tilde{\mathbf{s}}^{(t+1)} - \mathbf{s}^*)\|_2^2 \\ &\leq \|\tilde{\mathbf{s}}^{(t+1)} - \mathbf{s}^*\|_2^2 \\ &= \|\mathbf{s}^{(t)} - \mathbf{s}^*\|_2^2 + \eta_1^2 \|\nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}^{(t)}, \gamma^{(t)})\|_2^2 - 2\eta_1 \langle \nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}^{(t)}, \gamma^{(t)}), \mathbf{s}^{(t)} - \mathbf{s}^* \rangle,\end{aligned}$$

where the inequality comes from the fact that $\|\mathbf{I} - \mathbf{1}\mathbf{1}^\top/n\|_2 \leq 1$. We first bound the second term on the right hand side above

$$\begin{aligned}\|\nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}^{(t)}, \gamma^{(t)})\|_2^2 &\leq 3\|\nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}^{(t)}, \gamma^{(t)}) - \nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}^{(t)}, \gamma^*)\|_2^2 + 3\|\nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}^{(t)}, \gamma^*) - \nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}^*, \gamma^*)\|_2^2 \\ &\quad + 3\|\nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}^*, \gamma^*) - \nabla_{\mathbf{s}} \bar{\mathcal{L}}(\mathbf{s}^*, \gamma^*)\|_2^2 \\ &\leq 3M_1^2 \|\gamma^{(t)} - \gamma^*\|_2^2 + 3L_1^2 \|\mathbf{s}^{(t)} - \mathbf{s}^*\|_2^2 + 3\epsilon_1(k, n)^2,\end{aligned}$$

where the first inequality is due to $\nabla_{\mathbf{s}} \bar{\mathcal{L}}(\mathbf{s}^*, \gamma^*) = \mathbf{0}$ and the second inequality is due to Conditions 5.2, 5.3, and 5.4. Now we bound the inner product term. Note that

$$\begin{aligned}\langle \nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}^{(t)}, \gamma^{(t)}), \mathbf{s}^{(t)} - \mathbf{s}^* \rangle &= \langle \nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}^{(t)}, \gamma^{(t)}) - \nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}^*, \gamma^{(t)}), \mathbf{s}^{(t)} - \mathbf{s}^* \rangle + \langle \nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}^*, \gamma^{(t)}) - \nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}^*, \gamma^*), \mathbf{s}^{(t)} - \mathbf{s}^* \rangle \\ &\quad + \langle \nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}^*, \gamma^*) - \nabla_{\mathbf{s}} \bar{\mathcal{L}}(\mathbf{s}^*, \gamma^*), \mathbf{s}^{(t)} - \mathbf{s}^* \rangle.\end{aligned}$$

By strong convexity (Condition 5.1) of $\mathcal{L}$ we have

$$\langle \nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}^{(t)}, \gamma^{(t)}) - \nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}^*, \gamma^{(t)}), \mathbf{s}^{(t)} - \mathbf{s}^* \rangle \geq \mu_1 \|\mathbf{s}^{(t)} - \mathbf{s}^*\|_2^2. \quad (\text{B.1})$$

Applying Young's inequality and Condition 5.3, we obtain

$$\begin{aligned}|\langle \nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}^*, \gamma^{(t)}) - \nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}^*, \gamma^*), \mathbf{s}^{(t)} - \mathbf{s}^* \rangle| &\leq \|\nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}^*, \gamma^{(t)}) - \nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}^*, \gamma^*)\|_2 \cdot \|\mathbf{s}^{(t)} - \mathbf{s}^*\|_2 \\ &\leq \frac{\alpha M_1^2}{2} \|\gamma^{(t)} - \gamma^*\|_2^2 + \frac{1}{2\alpha} \|\mathbf{s}^{(t)} - \mathbf{s}^*\|_2^2,\end{aligned} \quad (\text{B.2})$$

where $\alpha > 0$ is an arbitrarily chosen constant. In addition, by Condition 5.4 and Young's inequality we have

$$|\langle \nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}^*, \gamma^*) - \nabla_{\mathbf{s}} \bar{\mathcal{L}}(\mathbf{s}^*, \gamma^*), \mathbf{s}^{(t)} - \mathbf{s}^* \rangle| \leq \|\nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}^*, \gamma^*) - \nabla_{\mathbf{s}} \bar{\mathcal{L}}(\mathbf{s}^*, \gamma^*)\|_2 \cdot \|\mathbf{s}^{(t)} - \mathbf{s}^*\|_2$$

$$\leq \frac{1}{2\mu_1}\epsilon_1(k, n)^2 + \frac{\mu_1}{2}\|\mathbf{s}^{(t)} - \mathbf{s}^*\|_2^2. \quad (\text{B.3})$$

Combining (B.1), (B.2) and (B.3), we have

$$\langle \nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}^{(t)}, \boldsymbol{\gamma}^{(t)}), \mathbf{s}^{(t)} - \mathbf{s}^* \rangle \geq \frac{\mu_1 \alpha - 1}{2\alpha} \|\mathbf{s}^{(t)} - \mathbf{s}^*\|_2^2 - \frac{\alpha M_1^2}{2} \|\boldsymbol{\gamma}^{(t)} - \boldsymbol{\gamma}^*\|_2^2 - \frac{1}{2\mu_1} \epsilon_1(k, n)^2.$$

Therefore, we have

$$\begin{aligned} \|\mathbf{s}^{(t+1)} - \mathbf{s}^*\|_2^2 &\leq \left(1 + 3L_1^2\eta_1^2 - \eta_1\left(\mu_1 - \frac{1}{\alpha}\right)\right) \|\mathbf{s}^{(t)} - \mathbf{s}^*\|_2^2 + M_1^2(3\eta_1^2 + \alpha\eta_1) \|\boldsymbol{\gamma}^{(t)} - \boldsymbol{\gamma}^*\|_2^2 \\ &\quad + (3\eta_1^2 + \eta_1/\mu_1)\epsilon_1(k, n)^2. \end{aligned} \quad (\text{B.4})$$

Similarly, we can bound $\|\boldsymbol{\gamma}^{(t+1)} - \boldsymbol{\gamma}^*\|_2^2$ as follows

$$\begin{aligned} \|\boldsymbol{\gamma}^{(t+1)} - \boldsymbol{\gamma}^*\|_2^2 &\leq \left(1 + 3L_2^2\eta_2^2 - \eta_2\left(\mu_2 - \frac{1}{\beta}\right)\right) \|\boldsymbol{\gamma}^{(t)} - \boldsymbol{\gamma}^*\|_2^2 + M_2^2(3\eta_2^2 + \beta\eta_2) \|\mathbf{s}^{(t)} - \mathbf{s}^*\|_2^2 \\ &\quad + (3\eta_2^2 + \eta_2/\mu_2)\epsilon_2(k, n)^2, \end{aligned} \quad (\text{B.5})$$

where $\beta > 0$ are arbitrarily chosen constants. In particular, set $\alpha = \mu_2/(4M_1^2)$, $\beta = \mu_1/(4M_2^2)$ and $\eta_1 = \eta_2 = \eta$. When $M_1, M_2 \leq \sqrt{\mu_1\mu_2}/4$, we have

$$\begin{aligned} \|\mathbf{s}^{(t+1)} - \mathbf{s}^*\|_2^2 + \|\boldsymbol{\gamma}^{(t+1)} - \boldsymbol{\gamma}^*\|_2^2 &\leq (1 + 3(L_1^2 + M_2^2)\eta^2 - \mu_1\eta/2) \|\mathbf{s}^{(t)} - \mathbf{s}^*\|_2^2 \\ &\quad + (1 + 3(L_2^2 + M_1^2)\eta^2 - \mu_2\eta/2) \|\boldsymbol{\gamma}^{(t)} - \boldsymbol{\gamma}^*\|_2^2 \\ &\quad + (3\eta^2 + \eta/\mu_1)\epsilon_1(k, n)^2 + (3\eta^2 + \eta/\mu_2)\epsilon_2(k, n)^2 \\ &\leq (1 + 3(L^2 + M^2)\eta^2 - \mu\eta/2) (\|\mathbf{s}^{(t)} - \mathbf{s}^*\|_2^2 + \|\boldsymbol{\gamma}^{(t)} - \boldsymbol{\gamma}^*\|_2^2) \\ &\quad + (3\eta^2 + \eta/\mu)(\epsilon_1(k, n)^2 + \epsilon_2(k, n)^2), \end{aligned} \quad (\text{B.6})$$

where $L = \max\{L_1, L_2\}$, $M = \max\{M_1, M_2\}$ and $\mu = \min\{\mu_1, \mu_2\}$. Note that we have $\|\mathbf{s}_0 - \mathbf{s}^*\|_2^2 + \|\boldsymbol{\gamma}_0 - \boldsymbol{\gamma}^*\|_2^2 \leq r^2$ by some initialization process. We can prove that $\|\mathbf{s}^{(t)} - \mathbf{s}^*\|_2^2 + \|\boldsymbol{\gamma}^{(t)} - \boldsymbol{\gamma}^*\|_2^2 \leq r^2$ for all $t \geq 0$ by induction. Specifically, assume it holds for t , then it suffices to ensure

$$(3\eta + 1/\mu)(\epsilon_1(k, n)^2 + \epsilon_2(k, n)^2) \leq r^2(\mu/2 - 3(L^2 + M^2)\eta), \quad (\text{B.7})$$

which holds when k is sufficiently large. Choosing η to be sufficiently small, we can ensure that $1 + 3(L^2 + M^2)\eta^2 - \mu\eta/2 \leq 1$. In particular, we can set $\eta = \mu/(12(L^2 + M^2))$, which implies

$$\begin{aligned} \|\mathbf{s}^{(t+1)} - \mathbf{s}^*\|_2^2 + \|\boldsymbol{\gamma}^{(t+1)} - \boldsymbol{\gamma}^*\|_2^2 &\leq \rho (\|\mathbf{s}^{(t)} - \mathbf{s}^*\|_2^2 + \|\boldsymbol{\gamma}^{(t)} - \boldsymbol{\gamma}^*\|_2^2) \\ &\quad + (3\eta^2 + \eta/\mu)(\epsilon_1(k, n)^2 + \epsilon_2(k, n)^2), \end{aligned}$$

with $\rho = 1 - \mu^2/(48(L^2 + M^2))$. Therefore, we have

$$\|\mathbf{s}^{(t)} - \mathbf{s}^*\|_2^2 + \|\boldsymbol{\gamma}^{(t)} - \boldsymbol{\gamma}^*\|_2^2 \leq \rho^t (\|\mathbf{s}_0 - \mathbf{s}^*\|_2^2 + \|\boldsymbol{\gamma}_0 - \boldsymbol{\gamma}^*\|_2^2) + \frac{3\eta^2 + \eta/\mu}{1 - \rho} (\epsilon_1(k, n)^2 + \epsilon_2(k, n)^2)$$

$$\leq r^2 \rho^t + \frac{\epsilon_1(k, n)^2 + \epsilon_2(k, n)^2}{\mu^2},$$

which completes the proof. $\square$

C Proofs of Specific Examples

In this section, we will provide the convergence analysis of Algorithm 1 for two specific examples with different noise distributions. In particular, we will show that Conditions 5.1 and 5.2 can be verified under these specific distributions. Recall the log-likelihood function

$$\mathcal{L}(\mathbf{s}, \gamma; \mathbf{Y}) = -\frac{1}{mk} \sum_{u=1}^m \sum_{(i,j) \in \mathcal{D}_u} \log F(\gamma_u(s_i - s_j); Y_{ij}^u). \quad (\text{C.1})$$

For the ease of presentation, we will omit $\mathbf{Y}$ in the rest of the proof and assume that the observation set $\mathcal{D}_u$ is parametrized by $k = |\mathcal{D}_u|$ and vectors $\mathbf{a}_{l,u} \in \mathbb{R}^n$ for $l = 1, \dots, k$, where each $\mathbf{a}_{l,u} = \mathbf{e}_{i_l} - \mathbf{e}_{j_l}$ for some pair of items (i_l, j_l) that is compared by user u and $\mathbf{e}_i$ is the natural basis. Then, we can rewrite the loss function in terms of vector $\mathbf{s}$ as follows

$$\mathcal{L}(\mathbf{s}, \gamma) = -\frac{1}{mk} \sum_{u=1}^m \sum_{l=1}^k \log F(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s}; Y_{ij_l}^u). \quad (\text{C.2})$$

Denote $g(x) = -\log F(x)$ for $x \in \mathbb{R}$. Then we can calculate the gradient of loss function $\mathcal{L}$ with respect to $\mathbf{s}$ and γ .

$$\begin{aligned} \nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}, \gamma) &= \frac{1}{mk} \sum_{u=1}^m \sum_{l=1}^k g'(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s}) \gamma_u \mathbf{a}_{l,u}, \\ \nabla_{\gamma} \mathcal{L}(\mathbf{s}, \gamma) &= \frac{1}{mk} \begin{bmatrix} \sum_{l=1}^k g'(\gamma_1 \mathbf{a}_{l,1}^\top \mathbf{s}) \mathbf{a}_{l,1}^\top \mathbf{s} \\ \vdots \\ \sum_{l=1}^k g'(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s}) \mathbf{a}_{l,u}^\top \mathbf{s} \\ \vdots \end{bmatrix}. \end{aligned} \quad (\text{C.3})$$

And the Hessian matrix can be calculated as

$$\begin{aligned} \nabla_{\mathbf{s}}^2 \mathcal{L}(\mathbf{s}, \gamma) &= \frac{1}{mk} \sum_{u=1}^m \sum_{l=1}^k g''(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s}) (\gamma_u)^2 \mathbf{a}_{l,u} \mathbf{a}_{l,u}^\top, \\ \nabla_{\gamma}^2 \mathcal{L}(\mathbf{s}, \gamma) &= \frac{1}{mk} \text{diag} \begin{bmatrix} \sum_{l=1}^k g''(\gamma_1 \mathbf{a}_{l,1}^\top \mathbf{s}) \mathbf{a}_{l,1}^\top \mathbf{s} \mathbf{a}_{l,1}^\top \mathbf{s} \\ \vdots \\ \sum_{l=1}^k g''(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s}) \mathbf{a}_{l,u}^\top \mathbf{s} \mathbf{a}_{l,u}^\top \mathbf{s} \\ \vdots \end{bmatrix}, \end{aligned} \quad (\text{C.4})$$

where $\text{diag}(\mathbf{x})$ is the diagonal matrix with diagonal entries given by $\mathbf{x}$.

C.1 Proof of Heterogeneous BTL model

Recall the definition in (4.1). The loss function can be written as

$$\mathcal{L}(\mathbf{s}, \gamma) = \frac{1}{mk} \sum_{u=1}^m \sum_{l=1}^k g\left(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s}; Y_{i_l j_l}^u\right), \quad (\text{C.5})$$

where $g(\cdot)$ is defined as

$$g(x; Y_{i_l j_l}^u) = -\log \frac{\exp(Y_{i_l j_l}^u x)}{1 + \exp(x)}. \quad (\text{C.6})$$

Therefore, the loss function of the HBTL model can be rewritten as follows:

$$\mathcal{L}(\mathbf{s}, \gamma) = \frac{1}{mk} \sum_{u=1}^m \sum_{l=1}^k \log\left(1 + \exp(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s})\right) - Y_{i_l j_l}^u \gamma_u \mathbf{a}_{l,u}^\top \mathbf{s}. \quad (\text{C.7})$$

Recall the gradients and Hessian matrices calculated in (C.3) and (C.4). We need to calculate $g'(\cdot)$ and $g''(\cdot)$. In particular, we have

$$g'(x; Y) = \frac{-Y + (1 - Y) \exp(x)}{1 + \exp(x)}, \quad g''(x; Y) = \frac{\exp(x)}{(1 + \exp(x))^2}. \quad (\text{C.8})$$

It is easy to verify that $g'(x)$ is monotonically increasing on $\mathbb{R}$. For any $|x| \leq \theta$, we have

$$\frac{-1}{1 + e^{-\theta}} \leq g'(x; Y = 1) \leq \frac{-1}{1 + e^{\theta}}, \quad \frac{e^{-\theta}}{1 + e^{-\theta}} \leq g'(x; Y = 0) \leq \frac{e^{\theta}}{1 + e^{\theta}}. \quad (\text{C.9})$$

Furthermore, $g''(x) = g''(-x)$, $g''(x)$ is increasing on $(-\infty, 0]$ and decreasing on $[0, \infty)$. Hence, for all $|x| \leq \theta$, we have

$$e^{\theta}/(1 + e^{\theta})^2 \leq g''(x) \leq g''(0) = 1/4. \quad (\text{C.10})$$

We can further show that the following lemmas hold, which validates Conditions 5.1, 5.2, 5.3 and 5.4 used in the convergence analysis.

The first two lemmas verify the strong convexity and smoothness of $\mathcal{L}$ with respect to $\mathbf{s}$ and γ respectively.

Lemma C.1. Suppose the noise ϵ follows the Gumbel distribution and the sample size $mk \geq 64(\gamma_{\max} + r)^2/(\gamma_{\min} - r)^2 n \log n$. Let $r \leq \min\{s_{\max}, \sqrt{\gamma_{\max} s_{\max}}\}$, for all $\mathbf{s}, \mathbf{s}' \in \mathbb{R}^n, \gamma \in \mathbb{R}^m$ such that $\|\mathbf{s} - \mathbf{s}^*\|_2 \leq r, \|\mathbf{s}' - \mathbf{s}^*\|_2 \leq r$ and $\|\gamma - \gamma^*\|_2 \leq r$, we have

$$\begin{aligned} \mathcal{L}(\mathbf{s}, \gamma) &\geq \mathcal{L}(\mathbf{s}', \gamma) + \langle \nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}', \gamma), \mathbf{s} - \mathbf{s}' \rangle + \frac{\mu_1}{2} \|\mathbf{s} - \mathbf{s}'\|_2^2, \\ \mathcal{L}(\mathbf{s}, \gamma) &\leq \mathcal{L}(\mathbf{s}', \gamma) + \langle \nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}', \gamma), \mathbf{s} - \mathbf{s}' \rangle + \frac{L_1}{2} \|\mathbf{s} - \mathbf{s}'\|_2^2, \end{aligned}$$

where the coefficients are defined as

$$\mu_1 = \frac{(\gamma_{\min} - r)^2 e^{5\gamma_{\max} s_{\max}}}{n(1 + e^{5\gamma_{\max} s_{\max}})^2}, \quad L_1 = \frac{(\gamma_{\max} + r)^2}{n}.$$

Lemma C.2. Suppose the noise ϵ follows the Gumbel distribution and the sample size satisfies $k \geq 18(s_{\max} + r)^4 n^2 / (m^2 (\|\mathbf{s}^*\|_2 + r)^4 \log(mn))$. Let $r \leq \min\{s_{\max}, \sqrt{\gamma_{\max} s_{\max}}\}$, for all $\mathbf{s} \in \mathbb{R}^n, \gamma, \gamma' \in \mathbb{R}^m$ such that $\|\mathbf{s} - \mathbf{s}^*\|_2 \leq r, \mathbf{s}^\top \mathbf{1} = 0$, and $\|\gamma - \gamma^*\|_2 \leq r, \|\gamma' - \gamma^*\|_2 \leq r$, we have with probability at least $1 - 1/n$ that

$$\begin{aligned} \mathcal{L}(\mathbf{s}, \gamma) &\geq \mathcal{L}(\mathbf{s}, \gamma') + \langle \nabla_\gamma \mathcal{L}(\mathbf{s}, \gamma'), \gamma - \gamma' \rangle + \frac{\mu_2}{2} \|\gamma - \gamma'\|_2^2, \\ \mathcal{L}(\mathbf{s}, \gamma) &\leq \mathcal{L}(\mathbf{s}, \gamma') + \langle \nabla_\gamma \mathcal{L}(\mathbf{s}, \gamma'), \gamma - \gamma' \rangle + \frac{L_2}{2} \|\gamma - \gamma'\|_2^2, \end{aligned}$$

where the coefficients are defined as

$$\mu_2 = \frac{(\|\mathbf{s}^*\|_2 + r)^2 e^{5\gamma_{\max} s_{\max}}}{n(1 + e^{5\gamma_{\max} s_{\max}})^2}, \quad L_2 = \frac{(\|\mathbf{s}^*\|_2 + r)^2}{n}.$$

Lemma C.3. Let $r \leq \min\{s_{\max}, \sqrt{\gamma_{\max} s_{\max}}\}$, for all $\mathbf{s} \in \mathbb{R}^n, \gamma \in \mathbb{R}^m$ such that $\|\mathbf{s} - \mathbf{s}^*\|_2 \leq r, \|\mathbf{s}' - \mathbf{s}^*\|_2 \leq r$ and $\|\gamma - \gamma^*\|_2 \leq r, \|\gamma' - \gamma^*\|_2 \leq r$, we have

$$\begin{aligned} \|\nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}, \gamma) - \nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}, \gamma')\|_2 &\leq \frac{\sqrt{2}(1 + 2\gamma_{\max} s_{\max})}{\sqrt{m}} \|\gamma - \gamma'\|_2, \\ \|\nabla_\gamma \mathcal{L}(\mathbf{s}, \gamma) - \nabla_\gamma \mathcal{L}(\mathbf{s}', \gamma)\|_2 &\leq \frac{\sqrt{2}(1 + 2\gamma_{\max} s_{\max})}{\sqrt{m}} \|\mathbf{s} - \mathbf{s}'\|_2. \end{aligned}$$

Lemma C.4. Let $r \leq \min\{s_{\max}, \sqrt{\gamma_{\max} s_{\max}}\}$, for all $\mathbf{s} \in \mathbb{R}^n, \gamma \in \mathbb{R}^m$ such that $\|\mathbf{s} - \mathbf{s}^*\|_2 \leq r$ and $\|\gamma - \gamma^*\|_2 \leq r$. Denote $\bar{\mathcal{L}}$ as the expected loss which takes expectation of $\mathcal{L}$ over the random choice of comparison pair. We have

$$\begin{aligned} \|\nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}, \gamma) - \nabla_{\mathbf{s}} \bar{\mathcal{L}}(\mathbf{s}, \gamma)\|_2 &\leq \epsilon_1(k, n) := \frac{2(\gamma_{\max} + r)}{1 + e^{-5\gamma_{\max} s_{\max}}} \sqrt{\frac{2 \log(2n)}{mk}}, \\ \|\nabla_\gamma \mathcal{L}(\mathbf{s}, \gamma) - \nabla_\gamma \bar{\mathcal{L}}(\mathbf{s}, \gamma)\|_2 &\leq \epsilon_2(k, n) := \frac{10\gamma_{\max} s_{\max}}{1 + e^{5\gamma_{\max} s_{\max}}} \sqrt{\frac{2 \log(2mn)}{mk}}, \end{aligned}$$

holds with probability at least $1 - 1/n$.

Proof of Corollary A.1. Now we prove the convergence of Algorithm 1 for Gumbel noise. Our proof will be similar to that of Theorem 5.5. In particular, we only need to verify that Conditions 5.1, 5.2, 5.3 and 5.4 hold when the noise follows a Gumbel distribution. According to Lemmas C.1 and C.2, we know that $\mathcal{L}(\mathbf{s}, \gamma)$ is μ_1 -strongly convex and L_1 -smooth with respect to $\mathbf{s}$, and is μ_2 -strongly convex and L_2 -smooth with respect to γ . More specifically, when $mk \geq 64n \log(n)$, we have

$$\mu_1 \geq (\gamma_{\min} - r)^2 e^{5C_0} / (n(1 + e^{5C_0})^2), \quad L_1 \leq (\gamma_{\max} + r)^2 / n, \quad (\text{C.11})$$

where we use the fact that $\gamma_{\max} s_{\max} = C_0$. In addition, note that $s_{\max} \leq \sqrt{m}/n \|\mathbf{s}^*\|_s$ and

$\|\mathbf{s}^{(t)} - \mathbf{s}^*\| \leq r$. Hence if $mk \geq 18 \log(mn)$, we have

$$\mu_2 \geq (\|\mathbf{s}^*\|_2 + r)^2 e^{5C_0} / (n(1 + e^{5C_0})^2), \quad L_2 \leq (\|\mathbf{s}^*\|_2 + r)^2 / n \quad (\text{C.12})$$

By Lemma C.3 and the assumption that $C_0 \geq 1/2$, we know that $\mathcal{L}$ satisfies the first-order stability (Condition 5.3) with $M_1 = M_2 = 4\sqrt{2}\gamma_{\max}s_{\max}/\sqrt{m}$. Note that by assumption, we have

$$s_{\max} \leq \frac{\gamma_{\min} e^{2C_0}}{16\sqrt{2}\gamma_{\max}(1 + e^{2C_0})^2} \frac{\sqrt{m}\|\mathbf{s}^*\|_2}{n}.$$

This immediately implies that $M = M_1 = M_2 \leq \sqrt{\mu_1\mu_2}/4$. Therefore, by similar arguments as in the proof of Theorem 5.5, we need to set step sizes $\eta_1 = \eta_2 = \eta < \mu/(6(L^2 + M^2))$, where $\mu = \min\{\mu_1, \mu_2\}$, $L = \max\{L_1, L_2\}$. In fact, it suffices to set

$$\eta < \frac{mne^{5C_0}\Gamma_1^2}{6(1 + e^{5C_0})^2(m\Gamma_2^4 + 32n^2C_0^2)},$$

with $\Gamma_1 = \min\{\gamma_{\min}/2, \|\mathbf{s}^*\|_2\}$ and $\Gamma_2 = \max\{2\gamma_{\max}, 2\|\mathbf{s}^*\|_2\}$. We thus obtain

$$\|\mathbf{s}^{(t)} - \mathbf{s}^*\|_2^2 + \|\boldsymbol{\gamma}^{(t)} - \boldsymbol{\gamma}^*\|_2^2 \leq r^2 \rho^t + \frac{\epsilon_1(k, n)^2 + \epsilon_2(k, n)^2}{\mu^2} \leq r^2 \rho^t + \frac{\Lambda n^2 \log(4mn^2)}{mk},$$

where $\rho = 1 - \eta(\mu - 6\eta(\Gamma_2^4/n^2 + 32C_0^2/m))/2$ and the last inequality comes from Lemma C.4 with the constant Λ defined as follows:

$$\Lambda = \max \left\{ \frac{200C_0^2(1 + e^{5C_0})^2}{\Gamma_1^4 e^{10C_0}}, \frac{8(\gamma_{\max} + r)^2(1 + e^{5C_0})^4}{\Gamma_1^4(1 + e^{-5C_0})^2} \right\}.$$

This completes the proof. $\square$

C.2 Proof of Heterogeneous Thurstone Case V model

In this subsection, we provide the analysis of our algorithm when the noise ϵ_i follows a Gaussian distribution, which results in the Thurstone model. The log-likelihood function can be written as

$$\mathcal{L}(\mathbf{s}, \boldsymbol{\gamma}; \mathbf{Y}) = \frac{1}{mk} \sum_{u=1}^m \sum_{l=1}^k g(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s}; Y_{i_{lj}}^u). \quad (\text{C.13})$$

with $g(\cdot)$ defined as $g(x) = -\log \Phi(x)$ with $\Phi(\cdot)$ be the CDF of the standard normal distribution. Note that $\Pr(Y_{i_{lj}}^u = 1) = \Phi(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s})$ and $\Pr(Y_{i_{lj}}^u = 0) = 1 - \Phi(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s}) = \Phi(-\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s})$. Thus we can write $g(\cdot)$ as $g(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s}; Y_{i_{lj}}^u) = -\log \Phi((2Y_{i_{lj}}^u - 1)\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s})$. Note that $(2Y - 1)^2 = 1$, we have

$$g'(x; Y) = -\frac{(2Y - 1)\Phi'(x)}{\Phi(x)}, \quad g''(x; Y) = \frac{\Phi'(x)^2 - \Phi(x)\Phi''(x)}{\Phi(x)^2}.$$

In order to bound $g'(x)$ and $g''(x)$, we first calculate the derivatives of $\Phi(x)$ as follows:

$$\Phi(x) = \int_{-\infty}^x \frac{1}{\sqrt{2\pi}} e^{-\frac{z^2}{2}} dz, \quad \Phi'(x) = \frac{1}{\sqrt{2\pi}} e^{-\frac{x^2}{2}}, \quad \Phi''(x) = \frac{-x}{\sqrt{2\pi}} e^{-\frac{x^2}{2}}. \quad (\text{C.14})$$

For any $\theta > 0$ such that $|x| \leq \theta$, we have

$$\frac{e^{-\theta^2/2}}{\sqrt{2\pi}\Phi(\theta)} \leq |g'(x)| \leq \frac{1}{\sqrt{2\pi}\Phi(-\theta)}.$$

We can verify that $g''(x)$ is monotonically decreasing on $\mathbb{R}^d$ and $g''(x) > 0$ also always hold. Thus for all $|x| \leq \theta$, we have $g''(\theta) \leq g''(x) \leq g''(-\theta)$.

Proof of Corollary A.3. Recall the derivation of the gradient in (C.3) and the Hessian in (C.4) of the loss function $\mathcal{L}$. In order to verify Conditions 5.1, 5.2, 5.3 and 5.4, we only need the upper and lower bounds of $g'(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s}; Y_{i_l j_l}^u)$ for all $u = 1, \dots, m$ and $l = 1, \dots, k$. Therefore, using exactly the same proof techniques as in Section C.1, we can also establish strong convexity, smoothness, first-order stability and the statistical error bound for sample loss function $\mathcal{L}$ when the noise ϵ follows the standard normal distribution. We omit the proof since it is the same as that of the Gumbel case. We can verify that $\mathcal{L}$ is μ_1 -strongly convex and L_1 -smooth with respect to $\mathbf{s}$, and is μ_2 -strongly convex and L_2 -smooth with respect to γ . The coefficient parameters are defined as $\mu_1 = (\gamma_{\min} - r)^2 H(5C_0)/n$, $L_1 = (\gamma_{\max} + r)^2 H(-5C_0)/n$, $\mu_2 = (\|\mathbf{s}^*\|_2 + r)^2 H(5C_0)/n$ and $L_2 = (\|\mathbf{s}^*\|_2 + r)^2 H(-5C_0)/n$. Note that $H(x)$ is a function defined based on the normal CDF $\Phi(\cdot)$:

$$H(x) = [\Phi'(x)^2 - \Phi(x)\Phi''(x)]/\Phi(x)^2,$$

where Φ, Φ', Φ'' are defined in (C.14). The loss function $\mathcal{L}$ also satisfies Condition 5.3 with $M = M_1 = M_2 = (1/\Phi(-5C_0) + 5\sqrt{2\pi}H(-5C_0)\gamma_{\max}s_{\max})/\sqrt{m\pi}$. In order to make sure that $M \leq \sqrt{\mu_1\mu_2}/4$, we only need $s_{\max} \leq \sqrt{\pi}\gamma_{\min}H(5C_0)/[4\gamma_{\max}(2/\Phi(-5C_0)) + 5\sqrt{2\pi}H(-5C_0)] \cdot \sqrt{m}\|\mathbf{s}^*\|_2/n$. Therefore, by Theorem 5.5, if we choose step sizes $\eta_1 = \eta_2 = \eta$ such that

$$\eta < \frac{mn\Gamma_1^2 H(5C_0)}{6(m\Gamma_2^4 + 50n^2C_0^2)H(-5C_0)^2},$$

with $\Gamma_2 = \min\{\gamma_{\min}/2, \|\mathbf{s}^*\|_2\}$, $\Gamma_2 = \max\{2\gamma_{\max}, 2\|\mathbf{s}^*\|_2\}$,

then we are able to obtain the following convergence result:

$$\|\mathbf{s}^{(t)} - \mathbf{s}^*\|_2^2 + \|\gamma^{(t)} - \gamma^*\|_2^2 \leq r^2 \rho^t + \frac{\epsilon_1(k, n)^2 + \epsilon_2(k, n)^2}{\mu^2}, \quad (\text{C.15})$$

where $\mu = \Gamma_1^2 H(5C_0)/n$, $\rho = 1 - \eta(\mu - 6\eta(\Gamma_2^4/n^2 + 32C_0^2/m))/2$ and $\epsilon_1(k, n), \epsilon_2(k, n)$ are the statistical error bounds. Similar to the proof of Lemma C.4, we know that $\epsilon_1(k, n) = (\gamma_{\max} + r)/(\sqrt{\pi}\Phi(-5C_0))\sqrt{2\log(2n)/(mk)}$ and $\epsilon_2(k, n) = 10\gamma_{\max}s_{\max}/(\sqrt{\pi}\Phi(-5C_0))\sqrt{\log(2mn)/(mk)}$. Plugging these two bounds into (C.15) yields

$$\|\mathbf{s}^{(t)} - \mathbf{s}^*\|_2^2 + \|\gamma^{(t)} - \gamma^*\|_2^2 \leq r^2 \rho^t + \frac{\Lambda' n^2 \log(4mn^2)}{mk},$$

which holds with probability at least $1 - 1/n$, where Λ' is a constant defines as follows.

$$\Lambda' = \frac{2 \max\{(\gamma_{\max} + r)^2, 50C_0^2\}}{\pi \Gamma_1^4 H(5C_0)^2 \Phi(-5C_0)^2}.$$

This completes the proof. $\square$

D Proofs of Technical Lemmas

In this section, we provide the proofs of technical lemmas used in the previous section.

D.1 Proof of Lemma C.1

We first lay down the following useful lemma.

Lemma D.1. (Tropp, 2012) Consider a sequence of i.i.d. random matrices $\{\mathbf{X}_k\}$ in $\mathbb{R}^{d \times d}$ with $\mathbb{E}[\mathbf{X}_k] = \mathbf{0}$ and $\|\mathbf{X}_k\|_2 \leq R$. Then for all $t \geq 0$

$$\Pr\left(\left\|\sum_k \mathbf{X}_k\right\| \geq t\right) \leq d \exp\left(-\frac{t^2}{2\sigma^2 + 2Rt/3}\right),$$

where $\sigma^2 = \|\sum_k \mathbb{E}[\mathbf{X}_k^2]\|_2$.

Proof of Lemma C.1. Using Taylor expansion, we have

$$\mathcal{L}(\mathbf{s}, \boldsymbol{\gamma}) = \mathcal{L}(\mathbf{s}', \boldsymbol{\gamma}) + \langle \nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}', \boldsymbol{\gamma}), \mathbf{s} - \mathbf{s}' \rangle + \frac{1}{2}(\mathbf{s} - \mathbf{s}')^\top \nabla_{\mathbf{s}}^2 \mathcal{L}(\tilde{\mathbf{s}}, \boldsymbol{\gamma})(\mathbf{s} - \mathbf{s}'),$$

where $\tilde{\mathbf{s}} = \mathbf{s} + \theta(\mathbf{s}' - \mathbf{s})$ for some $\theta \in (0, 1)$. In order to show the strong convexity and smoothness of $\mathcal{L}$, we need to bound the minimal and maximum eigenvalues of $\nabla_{\mathbf{s}}^2 \mathcal{L}(\mathbf{s}, \boldsymbol{\gamma})$. Note that $\mathbf{s}, \boldsymbol{\gamma}$ lie in a neighborhood with radius r of the true parameters $\mathbf{s}^*, \boldsymbol{\gamma}^*$ respectively. When $r \leq \min\{s_{\max}, \sqrt{\gamma_{\max} s_{\max}}\}$, we have

$$|\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s}| \leq |(\gamma_u - \gamma_u^*) \mathbf{a}_{l,u}^\top (\mathbf{s} - \mathbf{s}^*)| + |\gamma_u^* \mathbf{a}_{l,u}^\top (\mathbf{s} - \mathbf{s}^*)| + |\gamma_u^* \mathbf{a}_{l,u}^\top \mathbf{s}^*| \leq 5\gamma_{\max} s_{\max}. \quad (\text{D.1})$$

For any $\boldsymbol{\Delta} \in \mathbb{R}^n$, we have

$$\begin{aligned} \frac{1}{mk} \sum_{u=1}^m \sum_{l=1}^k \frac{(\gamma_u)^2 \exp(5\gamma_{\max} s_{\max})}{(1 + \exp(5\gamma_{\max} s_{\max}))^2} \boldsymbol{\Delta}^\top \mathbf{a}_{l,u} \mathbf{a}_{l,u}^\top \boldsymbol{\Delta} &\leq \boldsymbol{\Delta}^\top \nabla_{\mathbf{s}}^2 \mathcal{L}(\mathbf{s}, \boldsymbol{\gamma}) \boldsymbol{\Delta} \\ &= \frac{1}{mk} \sum_{u=1}^m \sum_{l=1}^k g''\left(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s}\right) (\gamma_u)^2 \boldsymbol{\Delta}^\top \mathbf{a}_{l,u} \mathbf{a}_{l,u}^\top \boldsymbol{\Delta} \\ &\leq \frac{1}{4mk} \sum_{u=1}^m \sum_{l=1}^k (\gamma_u)^2 \boldsymbol{\Delta}^\top \mathbf{a}_{l,u} \mathbf{a}_{l,u}^\top \boldsymbol{\Delta}, \end{aligned}$$

where we used the monotonicity of g'' . Since $\mathbf{a}_{l,u} = \mathbf{e}_{i_l} - \mathbf{e}_{j_l}$ and i_l, j_l are uniformly distributed, we have $\mathbb{E}[\mathbf{a}_{l,u}\mathbf{a}_{l,u}^\top] = \mathbb{E}[\mathbf{e}_{i_l}\mathbf{e}_{i_l}^\top + \mathbf{e}_{j_l}\mathbf{e}_{j_l}^\top - \mathbf{e}_{i_l}\mathbf{e}_{j_l}^\top - \mathbf{e}_{j_l}\mathbf{e}_{i_l}^\top] = 2/n\mathbf{I} - 2/n(\mathbf{1}\mathbf{1}^\top/n)$. We define

$$\mathbf{X}_{l,u} = (\gamma^u)^2 \left[\mathbf{a}_{l,u}\mathbf{a}_{l,u}^\top - \frac{2(\mathbf{I} - \mathbf{1}\mathbf{1}^\top/n)}{n} \right], \quad \mathbf{L} = \frac{2(\mathbf{I} - \mathbf{1}\mathbf{1}^\top/n)}{n}. \quad (\text{D.2})$$

Thus we have $\mathbb{E}[\mathbf{X}_{l,u}] = \mathbf{0}$. Furthermore, we have $\|\mathbf{X}_{l,u}\|_2 \leq 2(\gamma_{\max} + r)^2$ and $\mathbb{E}[\mathbf{X}_{l,u}^2] \leq 4(\gamma_{\max} + r)^4(n-1)/n^2(\mathbf{I} - \mathbf{1}\mathbf{1}^\top/n)$. Applying Lemma D.1 yields

$$\begin{aligned} \Pr \left(\left\| \frac{1}{mk} \sum_{u=1}^m \sum_{l=1}^k \mathbf{X}_{l,u} \right\|_2 \geq t \right) &\leq 2n \exp \left(\frac{-t^2}{8(\gamma_{\max} + r)^4(n-1)/(n^2mk) + 4t(\gamma_{\max} + r)^2/(3mk)} \right) \\ &\leq 2n \exp \left(\frac{-t^2}{8(\gamma_{\max} + r)^4/(nmk) + 4t(\gamma_{\max} + r)^2/(3mk)} \right), \end{aligned}$$

which implies that

$$\begin{aligned} \left\| \frac{1}{mk} \sum_{u=1}^m \sum_{l=1}^k \mathbf{X}_{l,u} \right\|_2 &\leq \frac{8(\gamma_{\max} + r)^2 \log n}{3mk} + 4(\gamma_{\max} + r)^2 \sqrt{\frac{\log n}{nmk}} \\ &\leq 8(\gamma_{\max} + r)^2 \sqrt{\frac{\log n}{nmk}} \end{aligned}$$

holds with probability at least $1 - 1/n$, where the last inequality holds when $mk \geq 4/9n \log n$. Therefore, we have

$$\|\nabla_{\mathbf{s}}^2 \mathcal{L}(\mathbf{s}, \gamma)\|_2 \leq (\gamma_{\max} + r)^2 \left(\frac{1}{2n} + 2\sqrt{\frac{\log n}{nmk}} \right) \leq \frac{(\gamma_{\max} + r)^2}{n}.$$

On the other hand, for any $\Delta \in \mathbb{R}^n$ such that $\Delta^\top \mathbf{1} = 0$, we have

$$\frac{1}{mk} \sum_{u=1}^m \sum_{l=1}^k \Delta^\top \mathbf{X}_{l,u} \Delta \geq -8\gamma_{\max}^2 \sqrt{\frac{\log n}{nmk}} \|\Delta\|_2^2,$$

which implies

$$\Delta^\top \nabla_{\mathbf{s}}^2 \mathcal{L}(\mathbf{s}, \gamma) \Delta \geq \left(\frac{2(\gamma_{\min} - r)^2}{n} - 8(\gamma_{\max} + r)^2 \sqrt{\frac{\log n}{nmk}} \right) \|\Delta\|_2^2.$$

Therefore, when k is sufficiently large such that $mk \geq 64(\gamma_{\max} + r)^2/(\gamma_{\min} - r)^2 n \log n$, we have

$$\lambda_{\min}(\nabla_{\mathbf{s}}^2 \mathcal{L}(\mathbf{s}, \gamma)) \geq \frac{(\gamma_{\max} - r)^2 e^{5\gamma_{\max} s_{\max}}}{n(1 + e^{5\gamma_{\max} s_{\max}})^2}.$$

This completes the proof. $\square$

D.2 Proof of Lemma C.2

Proof. Using Taylor expansion, we get

$$\mathcal{L}(\mathbf{s}, \gamma) = \mathcal{L}(\mathbf{s}, \gamma') + \langle \nabla_{\gamma} \mathcal{L}(\mathbf{s}, \gamma'), \gamma - \gamma' \rangle + \frac{1}{2}(\gamma - \gamma')^{\top} \nabla_{\gamma}^2 \mathcal{L}(\mathbf{s}, \tilde{\gamma})(\gamma - \gamma'), \quad (\text{D.3})$$

where $\tilde{\gamma} = \gamma + \theta(\gamma' - \gamma)$ for some $\theta \in (0, 1)$. Recall the Hessian matrix with respect to γ :

$$\nabla_{\gamma}^2 \mathcal{L}(\mathbf{s}, \gamma) = \frac{1}{mk} \text{diag} \begin{bmatrix} \sum_{l=1}^k g'' \left(\gamma_1 \mathbf{a}_{l,1}^{\top} \mathbf{s} \right) \mathbf{a}_{l,1}^{\top} \mathbf{s} \mathbf{a}_{l,1}^{\top} \mathbf{s} \\ \vdots \\ \sum_{l=1}^k g'' \left(\gamma_u \mathbf{a}_{l,u}^{\top} \mathbf{s} \right) \mathbf{a}_{l,u}^{\top} \mathbf{s} \mathbf{a}_{l,u}^{\top} \mathbf{s} \\ \vdots \end{bmatrix}.$$

For any fixed u , we denote $X_{l,u} = \mathbf{a}_{l,u}^{\top} \mathbf{s} \mathbf{a}_{l,u}^{\top} \mathbf{s} - \mathbf{s}^{\top} \mathbf{L} \mathbf{s}$, where $\mathbf{L}$ is defined as in (D.2). Recall the calculation of g'' in (C.8), (C.10) and that $|\gamma_u \mathbf{a}_{l,u}^{\top} \mathbf{s}| \leq 5\gamma_{\max} s_{\max}$ by (D.1), we have

$$\frac{e^{5\gamma_{\max} s_{\max}}}{(1 + e^{5\gamma_{\max} s_{\max}})^2} \leq g'' \left(\gamma_u \mathbf{a}_{l,u}^{\top} \mathbf{s} \right) = \frac{\exp(\gamma_u \mathbf{a}_{l,u}^{\top} \mathbf{s})}{(1 + \exp(\gamma_u \mathbf{a}_{l,u}^{\top} \mathbf{s}))^2} \leq \frac{1}{4}.$$

Since $\nabla_{\gamma}^2 \mathcal{L}(\mathbf{s}, \gamma)$ is a diagonal matrix, the eigenvalues of $\nabla_{\gamma}^2 \mathcal{L}(\mathbf{s}, \gamma)$ can be bounded by

$$\begin{aligned} \frac{e^{5\gamma_{\max} s_{\max}}}{(1 + e^{5\gamma_{\max} s_{\max}})^2} \min_u \frac{1}{mk} \sum_{l=1}^k (\mathbf{a}_{l,u}^{\top} \mathbf{s})^2 &\leq \lambda_{\min}(\nabla_{\gamma}^2 \mathcal{L}(\gamma)) \\ &\leq \lambda_{\max}(\nabla_{\gamma}^2 \mathcal{L}(\gamma)) \\ &\leq \frac{1}{4} \max_u \frac{1}{mk} \sum_{l=1}^k (\mathbf{a}_{l,u}^{\top} \mathbf{s})^2. \end{aligned} \quad (\text{D.4})$$

Since $\mathbf{s}^{\top} \mathbf{1} = 0$, it is easy to verify $\mathbb{E}[X_{l,u}] = \mathbb{E}[\mathbf{s}^{\top} (\mathbf{a}_{l,u} \mathbf{a}_{l,u}^{\top} - \mathbf{L}) \mathbf{s}] = 0$ and $|X_{l,u}| \leq 6(s_{\max} + r)^2$. For any fixed u , applying Hoeffding's inequality yields

$$\Pr \left(-\frac{1}{mk} \sum_{l=1}^k X_{l,u} \geq t \right) = \Pr \left(\frac{1}{mk} \sum_{l=1}^k X_{l,u} \geq t \right) \leq \exp \left(-\frac{m^2 t^2 k}{18(s_{\max} + r)^4} \right).$$

Further applying union bound, we have

$$\Pr \left(\max_u \frac{1}{mk} \sum_{l=1}^k X_{l,u} \geq t \right) \leq \sum_u \Pr \left(\frac{1}{k} \sum_{l=1}^k X_{l,u} \geq mt \right) \leq m \exp \left(-\frac{m^2 t^2 k}{18(s_{\max} + r)^4} \right),$$

which immediately implies that

$$\begin{aligned}
\lambda_{\max}(\nabla_{\gamma}^2 \mathcal{L}(\mathbf{s}, \gamma)) &\leq \frac{1}{4} \max_u \frac{1}{mk} \sum_{l=1}^k \mathbf{a}_{l,u}^{\top} \mathbf{s} \mathbf{a}_{l,u}^{\top} \mathbf{s} \\
&\leq \frac{(\|\mathbf{s}^*\|_2 + r)^2}{2n} + \frac{3(s_{\max} + r)^2}{4m} \sqrt{\frac{2 \log(mn)}{k}} \\
&\leq \frac{(\|\mathbf{s}^*\|_2 + r)^2}{n}
\end{aligned} \tag{D.5}$$

holds with probability at least $1 - 1/n$, where the last inequality is true when the sample size satisfies $k \geq 5(s_{\max} + r)^4 n^2 / (m^2 (\|\mathbf{s}^*\|_2 + r)^4 \log(mn))$. On the other hand, we also have

$$\Pr \left(\max_u -\frac{1}{mk} \sum_{l=1}^k X_{l,u} \geq t \right) \leq \sum_u \Pr \left(-\frac{1}{k} \sum_{l=1}^k X_{l,u} \geq mt \right) \leq m \exp \left(-\frac{m^2 t^2 k}{18(s_{\max} + r)^4} \right),$$

which leads to the conclusion that

$$\begin{aligned}
\lambda_{\min}(\nabla_{\gamma}^2 \mathcal{L}(\mathbf{s}, \gamma)) &\geq \frac{e^{5\gamma_{\max} s_{\max}}}{(1 + e^{5\gamma_{\max} s_{\max}})^2} \max_u \frac{1}{mk} \sum_{l=1}^k \mathbf{a}_{l,u}^{\top} \mathbf{s} \mathbf{a}_{l,u}^{\top} \mathbf{s} \\
&\geq \frac{e^{5\gamma_{\max} s_{\max}}}{(1 + e^{5\gamma_{\max} s_{\max}})^2} \left(\frac{2(\|\mathbf{s}^*\|_2 + r)^2}{n} - \frac{3(s_{\max} + r)^2}{m} \sqrt{\frac{2 \log(mn)}{k}} \right) \\
&\geq \frac{(\|\mathbf{s}^*\|_2 + r)^2 e^{5\gamma_{\max} s_{\max}}}{n(1 + e^{5\gamma_{\max} s_{\max}})^2}
\end{aligned} \tag{D.6}$$

holds with probability at least $1 - 1/n$, where the last inequality is due to $k \geq 18(s_{\max} + r)^4 n^2 / (m^2 (\|\mathbf{s}^*\|_2 + r)^4 \log(mn))$. $\square$

D.3 Proof of Lemma C.3

Proof. Recall the gradient of $\mathcal{L}$ with respect to $\mathbf{s}$ in (C.3). It holds that

$$\begin{aligned}
\|\nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}, \gamma) - \nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}, \gamma')\|_2 &= \left\| \frac{1}{mk} \sum_{u=1}^m \sum_{l=1}^k (g'(\gamma_u \mathbf{a}_{l,u}^{\top} \mathbf{s}) \gamma_u - g'(\gamma'_u \mathbf{a}_{l,u}^{\top} \mathbf{s}) \gamma'_u) \mathbf{a}_{l,u} \right\|_2 \\
&\leq \frac{1}{mk} \sum_{u=1}^m \sum_{l=1}^k [|g'(\gamma_u \mathbf{a}_{l,u}^{\top} \mathbf{s}) (\gamma_u - \gamma'_u)| \\
&\quad + |(g'(\gamma_u \mathbf{a}_{l,u}^{\top} \mathbf{s}) - g'(\gamma'_u \mathbf{a}_{l,u}^{\top} \mathbf{s})) \gamma'_u|] \|\mathbf{a}_{l,u}\|_2.
\end{aligned}$$

Note that we have $|g'(\gamma_u \mathbf{a}_{l,u}^{\top} \mathbf{s})| \leq 1$ and $\|\mathbf{a}_{l,u}\|_2 = \sqrt{2}$. In addition, by the mean value theorem we have

$$g'(\gamma_u \mathbf{a}_{l,u}^{\top} \mathbf{s}) - g'(\gamma'_u \mathbf{a}_{l,u}^{\top} \mathbf{s}) = g''(x)(\gamma_u - \gamma'_u) \mathbf{a}_{l,u}^{\top} \mathbf{s},$$

where $x = t\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s} + (1-t)\gamma'_u \mathbf{a}_{l,u}^\top \mathbf{s}$ for some $t \in (0, 1)$. By plugging the range of γ_u and $\mathbf{s}$, we have $|x| \leq 5\gamma_{\max} s_{\max}$ by (D.1) and hence $|g''(x)| = |e^x/(1+e^x)^2| \leq 1/4$. Now we can bound $\|\nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}, \gamma) - \nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}, \gamma')\|_2$ as follows:

$$\begin{aligned} \|\nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}, \gamma) - \nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}, \gamma')\|_2 &\leq \frac{1}{mk} \sum_{u=1}^m \sum_{l=1}^k \sqrt{2}(1 + 2\gamma_{\max} s_{\max}) |\gamma_u - \gamma'_u| \\ &\leq \frac{\sqrt{2}(1 + 2\gamma_{\max} s_{\max})}{\sqrt{m}} \|\gamma - \gamma'\|_2. \end{aligned}$$

Now we prove the upper bound of $\|\nabla_{\gamma} \mathcal{L}(\mathbf{s}, \gamma) - \nabla_{\gamma} \mathcal{L}(\mathbf{s}', \gamma)\|_2$. First, we have by (C.3) that

$$\nabla_{\gamma} \mathcal{L}(\mathbf{s}, \gamma) - \nabla_{\gamma} \mathcal{L}(\mathbf{s}', \gamma) = \frac{1}{mk} \begin{bmatrix} \sum_{l=1}^k \mathbf{a}_{l,1}^\top (g'(\gamma_1 \mathbf{a}_{l,1}^\top \mathbf{s}) \mathbf{s} - g'(\gamma_1 \mathbf{a}_{l,1}^\top \mathbf{s}') \mathbf{s}') \\ \vdots \\ \sum_{l=1}^k \mathbf{a}_{l,u}^\top (g'(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s}) \mathbf{s} - g'(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s}') \mathbf{s}') \\ \vdots \end{bmatrix}.$$

Note that for each u , we have

$$\begin{aligned} &\mathbf{a}_{l,u}^\top (g'(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s}) \mathbf{s} - g'(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s}') \mathbf{s}') \\ &= \mathbf{a}_{l,u}^\top [g'(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s}) (\mathbf{s} - \mathbf{s}') + (g'(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s}) - g'(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s}')) \mathbf{s}']. \end{aligned} \quad (\text{D.7})$$

For the first term in (D.7), we have

$$|\mathbf{a}_{l,u}^\top g'(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s}) (\mathbf{s} - \mathbf{s}')| \leq \sqrt{2} \|\mathbf{s} - \mathbf{s}'\|_2.$$

For the second term in (D.7), applying the mean value theorem yields

$$|\mathbf{a}_{l,u}^\top (g'(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s}) - g'(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s}')) \mathbf{s}'| = |g''(x) \gamma_u \mathbf{a}_{l,u}^\top (\mathbf{s} - \mathbf{s}') \mathbf{a}_{l,u}^\top \mathbf{s}'| \leq \frac{5\sqrt{2}\gamma_{\max} s_{\max}}{4} \|\mathbf{s} - \mathbf{s}'\|_2,$$

where $x = t\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s} + (1-t)\gamma^u \mathbf{a}_{l,u}^\top \mathbf{s}'$ for some $t \in (0, 1)$. Therefore, we have

$$\|\nabla_{\gamma} \mathcal{L}(\mathbf{s}, \gamma) - \nabla_{\gamma} \mathcal{L}(\mathbf{s}', \gamma)\|_2 \leq \frac{\sqrt{2}(1 + 2\gamma_{\max} s_{\max})}{\sqrt{m}} \|\mathbf{s} - \mathbf{s}'\|_2,$$

which completes our proof. $\square$

D.4 Proof of Lemma C.4

Proof. According to (C.3), the gradient of $\mathcal{L}$ with respect to $\mathbf{s}$ is

$$\nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}, \gamma) = \frac{1}{mk} \sum_u \sum_l \frac{(-Y + (1-Y) \exp(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s})) \gamma_u \mathbf{a}_{l,u}}{1 + \exp(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s})}.$$

By assumption we have $|\gamma_u| \leq (\gamma_{\max} + r)$ and $|\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s}| \leq 5\gamma_{\max} s_{\max}$ by (D.1). In addition, we have $\|\gamma_u \mathbf{a}_{l,u} / (1 + \exp(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s}))\|_2 \leq \sqrt{2}(\gamma_{\max} + r) / (1 + e^{-5\gamma_{\max} s_{\max}})$. Applying Hoeffding's inequality, we have

$$\Pr(\|\nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}, \gamma) - \nabla_{\mathbf{s}} \bar{\mathcal{L}}(\mathbf{s}, \gamma)\|_2 \geq t) \leq 2 \exp\left(\frac{-(1 + e^{-5\gamma_{\max} s_{\max}})^2 m k t^2}{8(\gamma_{\max} + r)^2}\right),$$

which implies that

$$\|\nabla_{\mathbf{s}} \mathcal{L}(\mathbf{s}, \gamma) - \nabla_{\mathbf{s}} \bar{\mathcal{L}}(\mathbf{s}, \gamma)\|_2 \leq \frac{2(\gamma_{\max} + r)}{1 + e^{-5\gamma_{\max} s_{\max}}} \sqrt{\frac{2 \log(2n)}{m k}}$$

holds with probability at least $1 - 1/n$. Recall the calculation in (C.3), the gradient of $\mathcal{L}$ with respect to γ is

$$\nabla_{\gamma} \mathcal{L}(\mathbf{s}, \gamma) = \frac{1}{m k} \begin{bmatrix} \sum_{l=1}^k g'(\gamma_1 \mathbf{a}_{l,1}^\top \mathbf{s}) \mathbf{a}_{l,1}^\top \mathbf{s} \\ \vdots \\ \sum_{l=1}^k g'(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s}) \mathbf{a}_{l,u}^\top \mathbf{s} \\ \vdots \end{bmatrix}.$$

The squared statistical error is

$$\|\nabla_{\gamma} \mathcal{L}(\mathbf{s}, \gamma) - \nabla_{\gamma} \bar{\mathcal{L}}(\mathbf{s}, \gamma)\|_2 = \frac{1}{m k} \sqrt{\sum_u \left[\sum_l (g'(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s}) \mathbf{a}_{l,u} - \mathbb{E}[g'(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s}) \mathbf{a}_{l,u}])^\top \mathbf{s} \right]^2},$$

which implies for all $t \geq 0$

$$\begin{aligned} & \Pr(\|\nabla_{\gamma} \mathcal{L}(\mathbf{s}, \gamma) - \nabla_{\gamma} \bar{\mathcal{L}}(\mathbf{s}, \gamma)\|_2 \geq t) \\ & \leq \Pr\left(\max_u \frac{1}{k} \sum_l (g'(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s}) \mathbf{a}_{l,u} - \mathbb{E}[g'(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s}) \mathbf{a}_{l,u}])^\top \mathbf{s} \geq \sqrt{m t}\right) \\ & \leq \sum_u \Pr\left(\frac{1}{k} \sum_l (g'(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s}) \mathbf{a}_{l,u} - \mathbb{E}[g'(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s}) \mathbf{a}_{l,u}])^\top \mathbf{s} \geq \sqrt{m t}\right), \end{aligned}$$

where the last inequality is due to union bound. For each user u , we have

$$|(g'(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s}) \mathbf{a}_{l,u} - \mathbb{E}[g'(\gamma_u \mathbf{a}_{l,u}^\top \mathbf{s}) \mathbf{a}_{l,u}])^\top \mathbf{s}| \leq \frac{10\gamma_{\max} s_{\max}}{1 + e^{-5\gamma_{\max} s_{\max}}}.$$

Applying Hoeffding's inequality yields

$$\Pr(\|\nabla_{\gamma} \mathcal{L}(\mathbf{s}, \gamma) - \nabla_{\gamma} \bar{\mathcal{L}}(\mathbf{s}, \gamma)\|_2 \geq t) \leq 2m \exp\left(\frac{-(1 + e^{-5\gamma_{\max} s_{\max}})^2 t^2 m k}{100\gamma_{\max}^2 s_{\max}^2}\right),$$

which immediately leads to the conclusion that

$$\|\nabla_{\gamma}\mathcal{L}(\mathbf{s}, \gamma) - \nabla_{\gamma}\bar{\mathcal{L}}(\mathbf{s}, \gamma)\|_2 \leq \frac{10\gamma_{\max}s_{\max}}{1 + e^{-5\gamma_{\max}s_{\max}}} \sqrt{\frac{2\log(2mn)}{mk}}$$

holds with probability at least $1 - 1/n$. This completes the proof. $\square$

D.5 Proof of Proposition 4.1

Proof. Since the PDF g of the noise terms ϵ_i is log-concave, and because the convolution of log-concave functions is log-concave [Merkle \(1998\)](#), the CDF F of $\epsilon_j - \epsilon_i$ for any pair i, j is also log-concave. Hence $h(x) = -\log F(x)$ is convex. The loss function is the sum of terms of the form $h_{iju} = h(\gamma_u(s_i - s_j))$. Fix i, j , and u . We have

$$\nabla_{\mathbf{s}}^2 h_{iju} = h''(\gamma_u(s_i - s_j))(\gamma_u)^2(\mathbf{e}_i - \mathbf{e}_j)(\mathbf{e}_i - \mathbf{e}_j)^{\top},$$

where $\mathbf{e}_i$ is the standard unit vector for coordinate i in $\mathbb{R}^n$. By the convexity of h and the fact that $(\mathbf{e}_i - \mathbf{e}_j)(\mathbf{e}_i - \mathbf{e}_j)^{\top}$ is positive-definite, the loss function is convex in $\mathbf{s}$. Similarly, it is easy to show that it is convex in γ . $\square$

Acknowledgment

We would like to thank the anonymous reviewers for their helpful comments. We would like to thank Ashish Kumar for the collection of “Country Population” dataset. PX and QG are supported in part by the NSF grants CIF-1911168, III-1904183 and CAREER Award 1906169. TJ and FF are supported in part by the NSF grant CIF-1908544. The views and conclusions contained in this paper are those of the authors and should not be interpreted as representing any funding agencies.

References

- AERTS, S., LAMBRECHTS, D., MAITY, S., VAN LOO, P., COESSENS, B., DE SMET, F., TRANCHEVENT, L.-C., DE MOOR, B., MARYNEN, P., HASSAN, B., CARMELIET, P. and MOREAU, Y. (2006). Gene prioritization through genomic data fusion. *Nature Biotechnology* **24** 537–544.
- AGARWAL, A., NEGAHBAN, S., WAINWRIGHT, M. J. ET AL. (2012). Noisy matrix decomposition via convex relaxation: Optimal rates in high dimensions. *The Annals of Statistics* **40** 1171–1197.
- BALTRUNAS, L., MAKCINSKAS, T. and RICCI, F. (2010). Group Recommendations with Rank Aggregation and Collaborative Filtering. In *Proceedings of the Fourth ACM Conference on Recommender Systems*. RecSys ’10, ACM, New York, NY, USA.
- BRADLEY, R. A. and TERRY, M. E. (1952). Rank Analysis of Incomplete Block Designs: I. The Method of Paired Comparisons. *Biometrika* **39** 324–345.

- BRAVERMAN, M. and MOSSEL, E. (2008). Noisy Sorting Without Resampling. In *ACM-SIAM Symp. Discrete Algorithms (SODA)*. Society for Industrial and Applied Mathematics, San Francisco, California.
- CHEN, J., XU, P., WANG, L., MA, J. and GU, Q. (2018). Covariate adjusted precision matrix estimation via nonconvex optimization. In *International Conference on Machine Learning*.
- CHEN, X., BENNETT, P. N., COLLINS-THOMPSON, K. and HORVITZ, E. (2013). Pairwise ranking aggregation in a crowdsourced setting. In *Proceedings of the sixth ACM international conference on Web search and data mining*. ACM.
- CHEN, Y. and SUH, C. (2015). Spectral MLE: Top-k rank aggregation from pairwise comparisons. In *International Conference on Machine Learning*.
- DE BORDA, J.-C. (1781). Mémoire sur les élections au scrutin. *Histoire de l'Académie royale des sciences*.
- DE CONDORCET, M. (1785). *Essai Sur l'application de l'analyse à La Probabilité Des Décisions Rendues à La Pluralité Des Voix*. L'imprimerie royale.
- DWORK, C., KUMAR, R., NAOR, M. and SIVAKUMAR, D. (2001). Rank aggregation methods for the web. In *Proc. 10th Int. Conf. World Wide Web*. ACM.
- GUIVER, J. and SNELSON, E. (2009). Bayesian Inference for Plackett-Luce Ranking Models. In *Proceedings of the 26th Annual International Conference on Machine Learning*. ICML '09, ACM, New York, NY, USA.
- HAJEK, B., OH, S. and XU, J. (2014). Minimax-optimal Inference from Partial Rankings. In *Advances in Neural Information Processing Systems 27*.
- HUNTER, D. R. (2004). MM algorithms for generalized Bradley-Terry models. *The Annals of Statistics* **32** 384–406.
- JAIN, P., NETRAPALLI, P. and SANGHAVI, S. (2013). Low-rank matrix completion using alternating minimization. In *Proceedings of the forty-fifth annual ACM symposium on Theory of computing*. ACM.
- KENDALL, M. (1948). *Rank Correlation Methods*. London: Griffin.
- KIM, M., FARNOUD, F. and MILENKOVIC, O. (2015). HyDRA: Gene prioritization via hybrid distance-score rank aggregation. *Bioinformatics* **31** 1034–1043.
- KUMAR, A. and LEASE, M. (2011). Learning to Rank from a Noisy Crowd. In *Proceedings of the 34th International ACM SIGIR Conference on Research and Development in Information Retrieval*. SIGIR '11, ACM, New York, NY, USA.
- LUCE, R. D. (1959). *Individual Choice Behavior: A Theoretical Analysis*. John Wiley & Sons, Inc., New York.

- MERKLE, M. (1998). Convolutions of logarithmically concave functions. *Publikacije Elektrotehničkog fakulteta. Serija Matematika* 113–117.
- NEGAHBAN, S., OH, S. and SHAH, D. (2012). Iterative ranking from pair-wise comparisons. In *Advances in Neural Information Processing Systems* 25.
- NEGAHBAN, S., OH, S. and SHAH, D. (2017). Rank Centrality: Ranking from Pairwise Comparisons. *Operations Research* **65** 266–287.
- NEGAHBAN, S. and WAINWRIGHT, M. J. (2012). Restricted strong convexity and weighted matrix completion: Optimal bounds with noise. *Journal of Machine Learning Research* **13** 1665–1697.
- RAMAN, K. and JOACHIMS, T. (2014). Methods for ordinal peer grading. In *Proceedings of the 20th ACM SIGKDD international conference on Knowledge discovery and data mining*. ACM.
- RAMAN, K. and JOACHIMS, T. (2015). Bayesian ordinal peer grading. In *Proceedings of the Second (2015) ACM Conference on Learning@ Scale*. ACM.
- THURSTONE, L. L. (1927). A law of comparative judgment. *Psychological Review* **34** 273–286.
- TROPP, J. A. (2012). User-friendly tail bounds for sums of random matrices. *Foundations of computational mathematics* **12** 389–434.
- VOJNOVIC, M. and YUN, S. (2016). Parameter Estimation for Generalized Thurstone Choice Models. In *PMLR*.
- WANG, Z., GU, Q., NING, Y. and LIU, H. (2015). High dimensional em algorithm: Statistical optimization and asymptotic normality. In *Advances in neural information processing systems*.
- WAUTHIER, F., JORDAN, M. and JOJIC, N. (2013). Efficient Ranking from Pairwise Comparisons. In *PMLR*.
- WENG, R. C. and LIN, C.-J. (2011). A Bayesian approximation method for online ranking. *Journal of Machine Learning Research* **12** 267–300.
- XU, P., MA, J. and GU, Q. (2017a). Speeding up latent variable Gaussian graphical model estimation via nonconvex optimization. In *Advances in Neural Information Processing Systems*.
- XU, P., ZHANG, T. and GU, Q. (2017b). Efficient algorithm for sparse tensor-variate Gaussian graphical models via gradient descent. In *Artificial Intelligence and Statistics*.
- YU, P. L. H. (2000). Bayesian analysis of order-statistics models for ranking data. *Psychometrika* **65** 281–299.
- ZERMELO, E. (1929). Die Berechnung der Turnier-Ergebnisse als ein Maximumproblem der Wahrscheinlichkeitsrechnung. *Mathematische Zeitschrift* **29** 436–460.
- ZHANG, X., WANG, L. and GU, Q. (2018). A unified framework for nonconvex low-rank plus sparse matrix recovery. In *International Conference on Artificial Intelligence and Statistics*.

- ZHAO, Z., VILLAMIL, T. and XIA, L. (2018). Learning mixtures of random utility models. In *Thirty-Second AAAI Conference on Artificial Intelligence*.
- ZHU, R., WANG, L., ZHAI, C. and GU, Q. (2017). High-dimensional variance-reduced stochastic gradient expectation-maximization algorithm. In *Proceedings of the 34th International Conference on Machine Learning-Volume 70*. JMLR. org.