

Pairwise Comparisons without Stochastic Transitivity: Model, Theory and Applications

Sze Ming Lee

S.LEE51@LSE.AC.UK

Yunxiao Chen

Y.CHEN186@LSE.AC.UK

Department of Statistics

London School of Economics and Political Science

London WC2A 2AE, United Kingdom

Editor: Mladen Kolar

Abstract

Most statistical models for pairwise comparisons, including the Bradley-Terry (BT) and Thurstone models and many extensions, make a relatively strong assumption of stochastic transitivity. This assumption imposes the existence of an unobserved global ranking among all the players/teams/items and monotone constraints on the comparison probabilities implied by the global ranking. However, the stochastic transitivity assumption does not hold in many real-world scenarios of pairwise comparisons, especially games involving multiple skills or strategies. As a result, models relying on this assumption can have suboptimal predictive performance. In this paper, we propose a general family of statistical models for pairwise comparison data without a stochastic transitivity assumption, substantially extending the BT and Thurstone models. In this model, the pairwise probabilities are determined by a (approximately) low-dimensional skew-symmetric matrix. Likelihood-based estimation methods and computational algorithms are developed, which allow for sparse data with only a small proportion of observed pairs. Theoretical analysis shows that the proposed estimator achieves minimax-rate optimality, which adapts effectively to the sparsity level of the data. The spectral theory for skew-symmetric matrices plays a crucial role in the implementation and theoretical analysis. The proposed method's superiority against the BT model, along with its broad applicability across diverse scenarios, is further supported by simulations and real data analysis.

Keywords: Pairwise comparison, stochastic intransitivity, Bradley-Terry model, low-rank model, nuclear norm

1. Introduction

Pairwise comparison data have received intensive attention in statistics and machine learning, with diverse applications across domains. Such data often arise from tournaments, where each pairwise comparison outcome results from a match between two players or teams, or from crowdsourcing settings, where individuals are tasked with comparing two items, such as images, movies, or products. Specifically, the famous Thurstone (Thurstone, 1927) and Bradley-Terry (BT; Bradley and Terry, 1952) models have set a cornerstone in the field, followed by many extensions, including the parametric ordinal models proposed in Shah et al. (2016a), which broadens the class of parametric models. Oliveira et al. (2018) relax the assumption of a known link function and propose models that allow the

link function to belong to a broad family of functions. Nonparametric approaches have also emerged, such as the work introduced in Shah and Wainwright (2018) based on the Borda counting algorithm, and the nonparametric Bradley-Terry models studied in Chatterjee (2015) and Chatterjee and Mukherjee (2019). Additionally, pairwise comparison models have been developed for crowdsourced settings, as discussed in Chen et al. (2013) and Chen et al. (2016), among many others. The models for pairwise comparisons have received a wide range of applications, including ranking problems (Chen and Suh, 2015; Chen et al., 2019; Heckel et al., 2019; Chen et al., 2022b; Liu et al., 2023; Fan et al., 2024), predicting matches/tournaments (Cattelan et al., 2013; Tsokos et al., 2019; Macrì Demartino et al., 2024), testing the efficiency of betting markets (McHale and Morton, 2011; Lyócsa and Vÿrost, 2018; Ramirez et al., 2023), and refinement of large language models based on human evaluations (Christiano et al., 2017; Ouyang et al., 2022; Zhu et al., 2023).

While the models mentioned above have made significant contributions to the field, they rely on the assumption of stochastic transitivity, which implies a strict ranking among players/teams/items. However, this assumption may be unrealistic, particularly in settings involving multiple skills or strategies, where intransitivity naturally arises. Despite its practical importance, research on models that allow intransitivity remains limited. Some notable exceptions include the work of Chen and Joachims (2016) and Spearing et al. (2023), which extend the Bradley-Terry model by introducing additional parameters to describe intransitivity alongside parameters specifying absolute strengths based on Bradley-Terry probabilities. Spearing et al. (2023) propose a Markov chain Monte Carlo algorithm for parameter estimation under a full Bayesian framework. However, their Bayesian procedure is computationally intensive and impractical for high-dimensional settings involving many players or a relatively high latent dimension. Chen and Joachims (2016) treat the parameters as fixed quantities and estimate them by optimizing a regularized objective function. However, their objective function is non-convex, and their model is highly over-parameterized. Consequently, their optimization is still computationally intensive and does not have a convergence guarantee. Moreover, no theoretical results are established in either work for their estimator.

Motivated by these challenges, we propose a general framework for modeling intransitive pairwise comparisons, assuming an approximately low-rank structure for the pairwise comparison probability matrix. We propose an estimator for the probabilities, which can be efficiently solved by a convex optimization program. This estimator is shown to be optimal in the minimax sense, accommodating sparse data—a common issue when the number of players diverges. To our knowledge, this is the first framework to address intransitive comparisons with rigorous error analysis. The models presented in Chen and Joachims (2016) and Spearing et al. (2023), which assume a low-rank structure, can be seen as a special case of our framework. Furthermore, our method and computational algorithms scale efficiently to high-dimensional settings, making them suitable for applications with many players/teams/items. Empirical results on real-world datasets, including the e-sport *StarCraft II* and professional tennis, demonstrate the practical usefulness of our method, showing superior performance in intransitive settings and robust performance when transitivity largely holds.

Pairwise comparison data has been extensively studied in the statistics and machine learning literature, with numerous models and methods developed. We refer readers to

Cattelan (2012) for a practical overview of the field. Theoretical properties of the BT model were first established in Simons and Yao (1999). These results were later extended to likelihood-based and spectral estimators, as well as other parametric extensions, with various losses and sparsity levels (Yan et al., 2012; Shah et al., 2016a; Negahban et al., 2017; Chen et al., 2019; Han et al., 2020; Chen et al., 2022a). More recently, Han et al. (2023) propose a general framework covering most parametric models satisfying strong stochastic transitivity, establishing uniform consistency results under sparse and heterogeneous settings.

Our development is also closely related to the literature on generalized low-rank and approximate low-rank models (Cai and Zhou, 2013; Davenport et al., 2014; Cai and Zhou, 2016; Chen et al., 2020a; Chen and Li, 2022, 2024; Lee et al., 2026). While our asymptotic results and error bounds build on techniques from these works, the parameter matrix in the current work differs in that it has a skew-symmetric structure. This structure, which arises naturally from pairwise comparison data, yields dependent entries and distinguishes our setting from typical low-rank models. To address this, a tailored analysis is performed to establish rigorous theoretical results.

The rest of the paper is organized as follows. Section 2 describes the setting, introduces the general approximate low-rank model, and proposes our estimator. Section 3 establishes the theoretical properties of the proposed estimator, including results on convergence and optimality. In Section 4, we provide an algorithm for solving the optimization problem of the proposed estimator. Section 5 verifies the theoretical findings and compares the proposed model with the BT model using simulations. Section 6 applies the proposed method to two real datasets to explore the presence of intransitivity in sports and e-sports. Finally, we conclude with discussions in Section 7. Additional theoretical results are presented in Appendix A, detailed proofs of the main results are given in Appendix B, and additional simulation results are reported in Appendix C.

2. Generalized Approximate Low-rank Model for Pairwise Comparisons

2.1 Setting and Proposed Model

We consider a scenario with n subjects, such as players in a sports tournament. Let n_{ij} denote the total number of comparisons observed between subjects i and j , where $(n_{ij})_{n \times n}$ is a symmetric matrix. Let y_{ij} denote the observed counts where subject i beats subject j . Assuming no draws, we have $y_{ij} = n_{ji} - y_{ji}$ for $i, j \in \{1, \dots, n\}$.

Given the total comparisons n_{ij} , we model the observed counts y_{ij} using a Binomial distribution: $y_{ij} \sim \text{Binomial}(n_{ij}, \pi_{ij})$, where π_{ij} denotes the probability that subject i beats subject j . A fundamental property of the probabilities is that $\pi_{ij} = 1 - \pi_{ji}$ for all $i, j \in \{1, \dots, n\}$. This implies that the matrix $\Pi = (\pi_{ij})_{n \times n}$ is fully determined by its upper triangular part. Using the logistic link function $g(x) = (1 + \exp(-x))^{-1}$, we express the probabilities as $\pi_{ij} = g(m_{ij})$, where $M = (m_{ij})$ is a skew-symmetric matrix satisfying $M = -M^\top$. As a result, estimating the probabilities Π reduces to the problem of estimating M .

We say the model is stochastic transitive if there exists an unobserved global ranking among all the players, denoted by $i_1 \succ i_2 \succ \dots \succ i_n$, such that the pairwise comparison probabilities for the adjacent pairs satisfy $\pi_{i_1 i_2}, \pi_{i_2 i_3}, \dots, \pi_{i_{n-1} i_n} \geq 0.5$. In addition, $\pi_{ik} \geq \pi_{ij}$

whenever $j \succ k$, for all $i \neq j, k$. In other words, for two players, j and k , any player is more likely to win k than j if player j ranks higher than k . If stochastic transitivity does not hold, then we say a model is stochastic intransitive. For instance, stochastic intransitivity arises when there exists a triplet (i, j, k) , such that $\pi_{ik} \geq \pi_{ij}$ and $\pi_{jk} < 0.5$.

Most traditional models for pairwise comparison assume stochastic transitivity. For example, the BT model assumes $m_{ij} = u_i - u_j$, in which case, the global ranking of the players is implied by the ordering of u_i , $i = 1, \dots, n$. However, stochastic intransitivity naturally occurs in real-world competition data involving multiple strategies or skills. For example, in the professional competitions of the e-sport *StarCraft II*, players can choose from a variety of combat units with differing attributes (e.g., building cost, attack range, toughness) during the game, leading to strategic decisions that can result in intransitivity. In fact, for the best predictive model that we learned for the *StarCraft II* data, more than 70% of the (i, j, k) triplets are estimated to violate the stochastic transitivity assumption, i.e., $\pi_{ik} \geq \pi_{ij}$ and $\pi_{jk} < 0.5$; see Section 6 for the details.

From the modeling perspective, stochastic transitivity is achieved by imposing strong monotonicity constraints on the parameter matrix M . To allow for stochastic intransitivity, we need to relax such constraints. Given $Y = (y_{ij})_{n \times n}$, the log-likelihood is

$$\begin{aligned} \mathcal{L}(M) &= \sum_{i=1}^n \sum_{j=1}^n y_{ij} \log(g(m_{ij})) \\ &= \sum_{i=1}^n \sum_{j>i}^n (y_{ij} \log(g(m_{ij})) + (n_{ij} - y_{ij}) \log(1 - g(m_{ij}))). \end{aligned}$$

To prevent overfitting while accommodating stochastic intransitivity, we impose a constraint on M to reduce the size of the parameter space. Specifically, we assume that M has an approximately low-rank structure enforced through a nuclear norm constraint:

$$\|M\|_* \leq C_n n, \quad (1)$$

where $\|\cdot\|_*$ denotes the nuclear norm, and $C_n > 0$ is a constant that may vary with n . The estimator is defined as:

$$\hat{M} = \arg \max_M \mathcal{L}(M) \text{ subject to } \|M\|_* \leq C_n n, M = -M^\top. \quad (2)$$

It is easy to see that the optimization in (2) is convex; see Section 4 for its computation.

2.2 Comparison with Related Work

We compare the proposed model with existing parametric models in the literature. Han et al. (2023) introduce a general framework for analyzing pairwise comparison data under the assumption of stochastic transitivity. In the current context, their model aligns with those proposed by Shah et al. (2016b) and Heckel et al. (2019), which are expressed as

$$\pi_{ij} = \Phi(u_i - u_j), \text{ and } \pi_{ji} = 1 - \Phi(u_i - u_j).$$

Here, $\Phi(\cdot)$ is any valid symmetric cumulative distribution function specified by the user, and $\mathbf{u} = (u_1, \dots, u_n)^\top$ is a latent score vector representing the strengths of the teams.

This framework reduces to the Bradley-Terry (BT) model when $\Phi(\cdot) = g(\cdot)$, the logistic function, and to the Thurstone model when $\Phi(\cdot)$ is the cumulative distribution function of the standard normal distribution. Other models can be incorporated by specifying different forms of $\Phi(\cdot)$. The latent score $\mathbf{u}$ is treated as a fixed parameter to be estimated, enabling the framework to handle a large number of players effectively. This parametric form, however, enforces a rank-2 structure on the parameter matrix, given by

$$\Pi = \Phi(\mathbf{u}\mathbf{1}_n^\top - \mathbf{1}_n\mathbf{u}^\top),$$

where $\mathbf{1}_n$ is an n -dimensional vector of ones.

Several attempts have been made in the literature to generalize this parametric form, allowing the rank of the underlying parameter matrix to exceed two and accommodate stochastic intransitivity. We should note that, since M is a skew-symmetric matrix, its rank must be even (e.g., Horn and Johnson, 2013). For instance, Chen and Joachims (2016) proposed a blade-chest-inner model, which can be expressed as

$$\Pi = g(AB^\top - BA^\top),$$

where A and B are $n \times K$ matrices. This model allows for a general rank- $2k$ parameter matrix, with the parameters in the frequentist sense. Similar to the parametrization in Chen and Joachims (2016), Spearing et al. (2023) propose a Bayesian model for pairwise comparison under stochastic intransitivity and further develop a Markov chain Monte Carlo algorithm for its computation. Both methods lack theoretical guarantees, such as convergence results or error bounds.

Our proposed method relaxes the requirement for an exact low-rank representation by only requiring an approximate low-rank structure specified by the nuclear norm. This offers a broad parameter space that covers the models proposed in Chen and Joachims (2016) and Spearing et al. (2023), thereby providing improved robustness against model misspecification. This flexibility is particularly important for real-world applications, where comparison probabilities may be influenced by numerous weak factors alongside a few dominant ones, so that the parameter matrix need not exhibit an exact low-rank structure. In particular, if $\text{rank}(M) = 2k$ for some positive integer k , it follows that

$$\|M\|_* \leq \sqrt{2k}\|M\|_F \leq C_n n,$$

where $\|\cdot\|_F$ denotes the Frobenius norm, and C_n is a constant depending on the magnitude of the entries of M and its rank $2k$. The subscript n in C_n indicates that both the magnitude of the entries of M and its rank are allowed to grow with n . Moreover, the proposed model imposes no distributional assumptions on the parameter matrix M , and the estimator can be obtained by solving a convex optimization program, making it more scalable and computationally tractable for handling a large number of players. Furthermore, as illustrated in Section 6, practical implementation requires selecting tuning parameters for the nuclear norm constraint. Unlike the exact low-rank formulation, which requires selecting the integer-valued exact rank, the tuning parameter can be selected over a flexible grid, allowing improved accuracy through finer tuning when computational resources permit. Theoretical results, including convergence and error bounds, are presented in Section 3, with additional results provided in Appendix A. As a remark, our estimation method

and theoretical framework can be easily adapted when we replace the current assumption of the logistic form of the link function $g(\cdot)$ with other functions, such as the standard normal cumulative distribution function used in the Thurstone model.

3. Theoretical Results

We establish convergence results and lower bounds for the estimator defined in (2) under settings with different data sparsity levels. For positive sequences $\{a_n\}$ and $\{b_n\}$, we denote $a_n \lesssim b_n$ if there exists a constant $\delta > 0$ that $a_n \leq \delta b_n$ for all n . We write $a_n \asymp b_n$ if $a_n \lesssim b_n$ and $b_n \lesssim a_n$. Let $\mathcal{K}$ denote the parameter space, defined as

$$\mathcal{K} = \{M \in \mathbb{R}^{n \times n} : \|M\|_* \leq C_n n, \quad M = -M^\top\}. \quad (3)$$

We impose the following conditions:

Assumption 1 *The true parameter $M^* \in \mathcal{K}$.*

Assumption 2 *For $j = 1, \dots, n$ and $i > j$, the variables n_{ij} are independent and follow a Binomial distribution, $n_{ij} \sim \text{Binomial}(T, p_{ij,n})$, where T is a fixed integer representing the maximum possible number of comparisons between subjects, and $p_{ij,n} = p_{ji,n}$ is the comparison rate, which may vary across different pairs (i, j) . Let $0 \leq p_n \leq q_n \leq 1$ denote the minimum and maximum comparison rates, respectively, such that $p_{ij,n} \in [p_n, q_n]$ for all $i \neq j \in \{1, \dots, n\}$. We assume that $p_n \asymp q_n$ and $p_n \gtrsim \log(n)/n$.*

Assumption 1 ensures that the true parameter exhibits an approximately low-rank structure specified by our model. Assumption 2 deserves more explanations. Under this assumption, the sparsity level of the data is characterized by the rate at which the comparison rates $p_{ij,n}$ converge to 0 as n grows. The condition $p_n \gtrsim \log(n)/n$ sets a lower bound on the sparsity level, which is the best possible threshold for pairwise comparison problems. Below this bound, the comparison graph becomes disconnected with high probability (Erdős and Rényi, 1960; Han et al., 2023). The condition $p_n \asymp q_n$ imposes homogeneity on $p_{ij,n}$, a common assumption in the literature (Simons and Yao, 1999; Chen et al., 2019; Han et al., 2020). This assumption means that all pairs of items are observed with roughly equal probability. The following theorem establishes the convergence rate of the proposed estimator.

Theorem 3 *Under Assumptions 1 and 2, let $\hat{\Pi} = (\hat{\pi}_{ij})_{n \times n}$, where $\hat{\pi}_{ij} = g(\hat{m}_{ij})$. Further let $\Pi^* = g(M^*)$. Then, with probability at least $1 - \kappa_1/n$,*

$$\frac{1}{n^2 - n} \|\hat{\Pi} - \Pi^*\|_F^2 \leq \kappa_2 C_n \sqrt{\frac{1}{p_n n}},$$

where κ_1 and κ_2 are constants that do not depend on n .

As shown in Theorem 3, the convergence rate of the mean squared Frobenius error between the estimated and true pairwise comparison probabilities depends on the sample size, sampling density, and model complexity. Specifically, the error decreases as the number of subjects n increases and as the sampling becomes denser (i.e., as p_n increases). In contrast,

the upper bound increases when the underlying parameter matrix has a higher rank or a more complex structure (as reflected by a larger value of C_n).

The following theorem addresses the optimality of Theorem 3 by providing a lower bound on the achievable estimation error, confirming that the rate in Theorem 1 cannot be improved in general.

Theorem 4 *Suppose $12 \leq C_n^2 \leq \min\{1, \kappa_3^2\}n$, where κ_3 is an absolute constant specified in Theorem 18 in Appendix B.2. Consider any algorithm which, for any $M \in \mathcal{K}$, takes as input Y and returns $\hat{M}$. Then there exists $M \in \mathcal{K}$ such that with probability at least $3/8$, $\Pi = g(M)$ and $\hat{\Pi} = g(\hat{M})$, satisfy*

$$\frac{1}{n^2 - n} \|\Pi - \hat{\Pi}\|_F^2 \geq \min \left\{ \kappa_4, \kappa_5 C_n \sqrt{\frac{1}{np_n}} \right\} \quad (4)$$

for all $n > N$. Here $\kappa_4, \kappa_5 > 0$ and N are absolute constants.

A few technical assumptions are imposed in this theorem. The condition $C_n^2 \leq \min\{1, \kappa_3^2\}n$ is mild and naturally holds for sufficiently large n , provided that the rank of M does not grow at the same rate as n . We also require $C_n^2 \geq 12$ to avoid the parameter space being too small for packing set construction, which is needed when we use Fano's inequality to establish the lower bound.

Since the rates in Theorems 3 and 4 match up to a multiplicative constant, the optimality of the proposed estimator is established.

Remark 5 *Some assumptions can be relaxed in the above theoretical analysis. Specifically, a version of Theorem 3 continues to hold with an updated convergence rate when the assumption $p_n \asymp q_n$ is dropped. Moreover, if T is not fixed but diverges, a version of Theorem 3 holds with a faster convergence rate; see Remark 15 in Appendix B.1 for details.*

Remark 6 *In this section, we establish the convergence of the matrix of pairwise comparison probabilities Π^* under the Frobenius norm. It is also possible to derive convergence results for the underlying parameter matrix M^* under additional assumptions, specifically by requiring that $\|M^*\|_\infty$ is bounded by a positive constant, as demonstrated in the proof of Lemma 19 in Appendix B.3.*

Furthermore, under an exact low-rank assumption on M^* , convergence of M^* in the element-wise maximum norm can be established by making use of the refinement procedure proposed in Chen and Li (2024). This implies the entry-wise convergence for the comparison probabilities. Such results also yield control of the two-to-infinity norm of the low-rank component, which may be particularly relevant when studying player-specific skill parameters.

A detailed discussion of the exact assumptions, along with the corresponding estimation and refinement procedures, is provided in Appendix A.1.

Remark 7 *The proposed model does not impose stochastic transitivity and therefore a global ranking is not well-defined by the true comparison probability matrix. Nevertheless, for any subset of subjects $\mathcal{S} \subseteq [n]$, one may define a relative score measuring the average strength of subject i within $\mathcal{S}$, which reduces to the score in Shah and Wainwright (2018) when $\mathcal{S} = [n]$.*

Under a suitable separation condition, we establish consistent recovery of the top- k set based on these scores. The precise formulation and theoretical results are presented in Appendix A.2.

4. Computation

To solve the nuclear-norm constrained optimization problem (2), we apply the nonmonotone spectral-projected gradient algorithm proposed by Birgin et al. (2000). Since the feasible set of M is closed and convex, this algorithm guarantees convergence to a constrained stationary point. Moreover, because $\mathcal{L}(M)$ is concave, any such stationary point is a global maximizer over the feasible set.

This section outlines the computational procedure for the proposed model. Section 4.1 reformulates the optimization problem (2) in vectorized form. Section 4.2 presents the projection algorithm used to enforce the nuclear-norm constraint. Section 4.3 describes the spectral projected gradient method for solving the constrained problem, and Section 4.4 presents the complete estimation procedure.

4.1 Reformulating the optimization problem

We begin by reformulating the matrix optimization problem into a vectorized form that is more convenient for numerical computation. Let Skew_n denote the space of $n \times n$ skew-symmetric matrices. Let $\mathcal{V}$ be the bijective linear mapping that vectorizes the upper-triangular part of any matrix in Skew_n into $\mathbb{R}^{0.5n(n-1)}$. For any $\mathbf{m} \in \mathbb{R}^{0.5n(n-1)}$, define $f(\mathbf{m}) = \mathcal{L}(\mathcal{V}^{-1}(\mathbf{m}))$. Then, solving (2) is equivalent to solving the constrained optimization problem:

$$\hat{\mathbf{m}} = \arg \max_{\mathbf{m} \in \mathbb{R}^{0.5n(n-1)}} f(\mathbf{m}) \quad \text{subject to } \|\mathcal{V}^{-1}(\mathbf{m})\|_* \leq \tau, \quad (5)$$

where $\tau = C_n n$ if C_n is known. We will later discuss an algorithm for selecting τ in practical situations where C_n is unknown.

4.2 Projection onto the nuclear-norm ball

A key step in solving (5) is to project each iterate onto the nuclear-norm ball in order to enforce the constraint. We define the orthogonal projection operator $P_\tau(\cdot)$ as

$$P_\tau(\mathbf{m}) = \arg \min_{\mathbf{x} \in \mathbb{R}^{0.5n(n-1)}} \|\mathbf{x} - \mathbf{m}\|_2 \quad \text{subject to } \|\mathcal{V}^{-1}(\mathbf{x})\|_* \leq \tau.$$

It is well known that the projection is equivalent to singular value soft-thresholding. Specifically, we determine a threshold and update each singular value by either subtracting this threshold from it or setting it to zero when the threshold exceeds the singular value, so that the resulting matrix satisfies the nuclear-norm constraint. Let $0_{n \times n}$ denote a $n \times n$ zero matrix, and $\max\{\cdot, \cdot\}$ be applied entry-wise for matrix inputs. The detailed procedure is presented in Algorithm 1, where we first apply the inverse mapping $\mathcal{V}^{-1}(\cdot)$ to recover the corresponding skew-symmetric matrix, perform singular value soft-thresholding to enforce the nuclear-norm constraint, and then apply $\mathcal{V}(\cdot)$ to return the vectorized projected output.

Algorithm 1 Projection onto the nuclear-norm ball

Input: Parameter vector $\mathbf{m}$ and nuclear norm constraint parameter τ .

Compute $M = \mathcal{V}^{-1}(\mathbf{m})$.

Perform singular value decomposition and obtain $M = U\Sigma V^\top$, where U and V are $n \times n$ orthonormal matrices and $\Sigma = \text{diag}(\sigma_1, \sigma_1, \dots, \sigma_{n/2}, \sigma_{n/2})$ if n is even and $\Sigma = \text{diag}(\sigma_1, \sigma_1, \dots, \sigma_{\lfloor n/2 \rfloor}, \sigma_{\lfloor n/2 \rfloor}, 0)$ otherwise.

Compute λ , the smallest value for which $\sum_{i=1}^{\lfloor n/2 \rfloor} 2 \max\{\sigma_i - \lambda, 0\} \leq \tau$.

Compute projected matrix $P_\tau(M) = U \max\{\Sigma - \lambda I_n, 0_{n \times n}\} V^\top$

Output: Projection outcome $P_\tau(\mathbf{m}) = \mathcal{V}(P_\tau(M))$.

In the last step, the projection outcome is defined as $P_\tau(\mathbf{m}) = \mathcal{V}(P_\tau(M))$, which is only valid provided that $P_\tau(M)$ is a skew-symmetric matrix. The following proposition ensures that this is always the case:

Proposition 8 *For any matrix $M \in \text{Skew}_n$, the projection operator satisfies $P_\tau(M) \in \text{Skew}_n$.*

4.3 Spectral projected gradient algorithm

We now introduce the spectral projected line search method. This approach combines gradient-based updates with projection onto the nuclear-norm ball, ensuring that each iterate remains within the feasible set defined by the nuclear norm constraint. The procedure is outlined in Algorithm 2.

Algorithm 2 Spectral projected line search

Input: Parameter vector from last iteration $\mathbf{m}^{(l-1)}$, Matrix of comparison outcomes Y , nuclear norm constraint parameter τ and the spectral-step length γ_{l-1}

Compute gradient: $\mathbf{g}^{(l-1)} = \nabla f(\mathbf{m}^{(l-1)})$.

Compute search direction: $\mathbf{d}^{(l-1)} = P_\tau(\mathbf{m}^{(l-1)} - \gamma_{l-1}\mathbf{g}^{(l-1)}) - \mathbf{m}^{(l-1)}$

Perform line search along the linear trajectory: $\mathbf{m}(\alpha) = \mathbf{m}^{(l-1)} + \alpha\mathbf{d}^{(l-1)}$.

if Convergence is reached **then**

Set $\mathbf{m}^{(l)}$ as the result from the line search.

else

Perform line search along the alternative trajectory:

$$\mathbf{m}^{\text{curve}}(\alpha) = P_\tau(\mathbf{m}^{(l-1)} - \alpha\gamma_{l-1}\mathbf{g}^{(l-1)}).$$

Set $\mathbf{m}^{(l)}$ as the result from the line search.

end if

Output: Updated parameter vector $\mathbf{m}^{(l)}$.

At each iteration, after computing the gradient at the current iterate, the method employs two types of line searches. The first type performs a projection once and searches along a linear trajectory $\mathbf{m}(\alpha)$. This approach is computationally efficient since the primary computational cost lies in the projection operation. If the linear search fails to converge,

the algorithm switches to a curvilinear trajectory $\mathbf{m}^{\text{curve}}(\alpha)$, which requires projecting at each step. The Spectral-step length γ_{l-1} is decided using the method from Barzilai and Borwein (1988) in each iteration.

4.4 Estimation procedure

We now summarize the full estimation procedure by integrating the reformulation, projection, and spectral projected gradient steps described above. The full procedure is detailed in Algorithm 3. The convergence criterion checks whether the optimality condition $P_\tau(\mathbf{m}^{(l)} - \nabla f(\mathbf{m}^{(l)})) = \mathbf{m}^{(l)}$ is approximately satisfied. Parts of the code are adapted from the SPGL1 package, originally implemented in MATLAB (Van Den Berg and Friedlander, 2008; Davenport et al., 2014). The proposed estimator is implemented in R, and the code for the implementation, as well as the simulation and real data analyses, is available at <https://github.com/Arthurlee51/PCWST>.

Algorithm 3 Estimation Algorithm

Input: Matrix of comparison outcomes Y , nuclear norm constraint parameter τ .

Initialization: Set $l = 0$, $\mathbf{m}^{(0)} = \mathbf{0}_{0.5n(n-1)}$, the zero vector and set the spectral step-length $\gamma_0 = 1$.

while $l = 0$ **or** convergence criterion is not satisfied **do**

 Update $l \leftarrow l + 1$.

 Update $\mathbf{m}^{(l)}$ via line search using Algorithm 2 with inputs $\mathbf{m}^{(l-1)}$, Y , τ and γ_{l-1} .

 Update γ_l as proposed by Barzilai and Borwein (1988).

end while

Output: Estimated parameter matrix $\hat{M} = \mathcal{V}^{-1}(\mathbf{m}^{(l)})$.

5. Simulation Results

This section presents simulation studies to validate the theoretical results in Section 3 and assess the numerical performance of the proposed estimator. Specifically, we evaluate the proposed model in terms of the estimation error of the pairwise comparison probability matrix and its predictive likelihood on independently generated test data. We compare its performance with that of the BT model across varying sample sizes, matrix complexities, and sparsity levels.

We consider three distinct scenarios characterized by varying levels of sparsity. Specifically, we define p_n as $n^{-1} \log(n)$, $n^{-1/2}$, and $1/4$, corresponding to sparse, less sparse, and dense data, respectively. The parameter q_n is given by $4p_n$. Each $p_{ij,n}$ is then generated from a uniform distribution with range $[p_n, q_n]$.

The parameter matrix M is constructed as $\Theta J \Theta^\top$, where Θ is an $n \times 2k$ matrix, and J is a $2k \times 2k$ block diagonal matrix of the form

$$J = \begin{pmatrix} 0 & n & 0 & \dots & 0 \\ -n & 0 & 0 & \dots & 0 \\ 0 & 0 & \ddots & \dots & 0 \\ \vdots & \vdots & \vdots & 0 & n \\ 0 & 0 & 0 & -n & 0 \end{pmatrix}.$$

The matrix Θ is orthonormal, obtained via QR decomposition of a random matrix $Z \in \mathbb{R}^{n \times 2k}$, where each entry of Z is independently sampled from a standard normal distribution $N(0, 1)$. It can be verified that $\|M\|_* = 2kn$.

We conduct 50 simulations for $n = 500, 1000, 1500$, and 2000 , with k ranging from 1 to 10. Recall that the rank of M is $2k$. Additionally, the maximum number of comparisons, T , is fixed at 5, across all settings. We set $C_n = 2k$. The loss is computed as

$$\text{Loss} = (n^2 - n)^{-1} \|\hat{\Pi} - \Pi^*\|_F^2, \quad (6)$$

and the average loss across 50 simulations is reported for each model in Figure 1, considering different values of n , k , and sparsity levels.

The results in Figure 1 show that the mean loss of the proposed estimator decreases as n increases. Moreover, the mean loss is significantly lower as the data become denser, corresponding to an increase in p_n . These observations are consistent with the results from Theorem 3.

Notably, the proposed and BT models incur higher losses as the rank parameter k increases, which is expected due to increasing complexity. However, the proposed model consistently outperforms the BT model across all settings, as reflected by the lower mean losses and the fact that its error bars lie entirely below those of the BT model. Furthermore, while the BT model’s performance remains relatively unchanged as n increases, the proposed method continues to improve, showcasing its effectiveness in handling large datasets and capturing complex structures that stochastic transitivity assumptions cannot address.

In addition to estimation error, we evaluate predictive performance using the log-likelihood computed on an independently generated test set. The results further confirm the superiority of the proposed method over the BT model, as evidenced by consistently higher average predictive likelihood values across all settings. The detailed results, together with the computational time of both methods, are presented in Appendix C.

6. Real Data Examples

In this section, we compare our model’s performance with the classical BT model using two real datasets. Section 6.1 provides a brief description of the two datasets used in the analysis, namely the StarCraft II and Tennis datasets. Section 6.2 outlines the data preparation process and describes how the nuclear norm constraint parameter $\tau = C_n n$ is decided. Section 6.3 introduces the evaluation metrics used to compare the models. Finally, Section 6.4 reports the empirical results for both the StarCraft II and Tennis datasets.

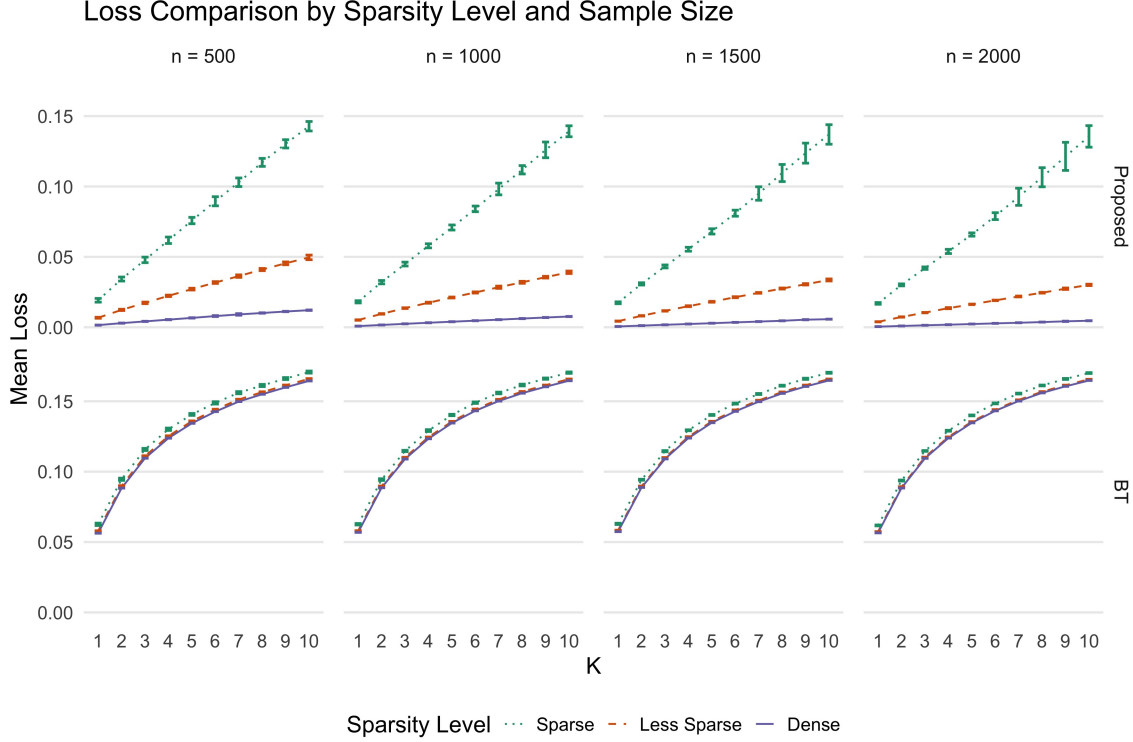


Figure 1: Comparison of loss between the proposed method and the Bradley-Terry (BT) model across different sparsity levels (sparse, less sparse, dense). Each row corresponds to a method (top: proposed; bottom: BT), and each column corresponds to a sample size ($n = 500, 1000, 1500, 2000$). The x-axis represents the rank parameter k , while the y-axis shows the mean loss, computed as the average of the losses defined in (6). Error bars indicate ± 2 standard deviations across replicates.

6.1 Description of Datasets

6.1.1 *StarCraft II* DATA

StarCraft II is a military science fiction real-time strategy game developed and published by Blizzard Entertainment. The dataset comprises match results of professional *StarCraft II* players sourced from the website aligulac.com, covering the period from 2010 to 2016. The matches follow the most common competitive format, where two players face off against each other, and each game results in either a win or loss, with no possibility of a draw.

We specifically focus on matches played using the *StarCraft II: Heart of the Swarm* expansion, as different versions of the game are often treated as distinct games (Chen and Joachims, 2016). The training set includes 1,958 players, with 1.9% of all player pairs competing against each other at least once. The maximum number of matches between

any pair of players is 30. The dataset is available at <https://www.kaggle.com/datasets/alimbekovkz/starcraft-ii-matches-history>.

6.1.2 TENNIS DATA

We analyze the tennis dataset to evaluate the performance of our model in professional sports. The dataset contains the results of all men’s matches organized by the Association of Tennis Professionals (ATP) from 2000 to 2018. It includes matches from major tournaments such as the Grand Slams, the ATP World Tour Masters 1000, and other professional tennis series held during this period.

The training set consists of 723 players, with 6.4% of all player pairs having competed against each other at least once. The maximum number of matches between any pair of players is 23. The data is collected from <http://www.tennis-data.co.uk>.

6.2 Data Preparation and Parameter Tuning

The raw data consists of individual match records, with each comparison recorded as a separate entry. We reserve 30% of the match records for testing, while the remaining 70% is divided into 50% for training and 20% for validation.

The comparison data matrix is first constructed for the training set, with players absent from the training set removed from the validation set. The validation set is used to tune the nuclear constraint parameter C_n , as described in the sequel. After tuning, the training and validation sets are combined (including previously excluded entries), and the comparison data matrix is reconstructed from the combined dataset.

The test set is then evaluated against this combined dataset, excluding entries for players not present in the combined dataset. Although the proposed model can handle players who never lose or win any game, we still remove them in the training and combined dataset to ensure stabler results and a fair comparison with the BT model, as this is a common practice.

The nuclear norm of the parameter matrix M is unknown and is tuned on the training and validation sets using log-likelihood as the loss function. The nuclear constraint parameter $\tau = C_n n$ is determined by selecting C_n from 20 grid points, corresponding to powers of 10 evenly spaced between -1 and 1 . This results in $C_n = 10^{0.47} = 2.98$ for the *StarCraft II* dataset and $C_n = 10^{-0.36} = 0.43$ for the tennis dataset.

6.3 Evaluation Criteria

Let $Y^{(\text{test})} = (y_{ij}^{(\text{test})})_{n \times n}$ denote the observed comparison results from the test set. Given the estimated comparison probabilities $\hat{\Pi} = (\hat{\pi}_{ij})_{n \times n}$, we evaluate the performance of the estimates using two criteria. The first criterion is the log-likelihood, given by

$$L(Y^{(\text{test})} \mid \hat{\Pi}) = \sum_{i=1}^n \sum_{j>i}^n \left(y_{ij}^{(\text{test})} \log(\hat{\pi}_{ij}) + y_{ji}^{(\text{test})} \log(1 - \hat{\pi}_{ij}) \right),$$

where a higher log-likelihood indicates a stronger agreement between the predicted probabilities and the observed results. The second criterion is the test accuracy, given by

$$A(Y^{(\text{test})} | \hat{\Pi}) = \frac{1}{\sum_{i=1}^n \sum_{j=1}^n y_{ij}^{(\text{test})}} \sum_{i=1}^n \sum_{j>i} \left(y_{ij}^{(\text{test})} I(\hat{\pi}_{ij} \geq 0.5) + y_{ji}^{(\text{test})} I(\hat{\pi}_{ji} > 0.5) \right).$$

It measures the proportion of the comparison results correctly predicted, with higher values indicating better predictive performance. The results are presented in Table 1.

	<i>StarCraft II</i>		Tennis	
	Proposed	BT	Proposed	BT
Log-likelihood	−1,897,946	−2,137,115	−333,076	−322,483
Accuracy	0.766	0.713	0.652	0.658

Table 1: Comparison of model performance on StarCraft II and ATP datasets. The performance is evaluated using log-likelihood and accuracy for the proposed model and the BT model.

6.4 Empirical Results

6.4.1 *StarCraft II* DATA

As seen in Table 1, the proposed model achieves a higher log-likelihood of −1,897,946 compared to −2,137,115 for the BT model. This suggests that our model provides a better fit for the observed test data. The test accuracy of the proposed model is also significantly higher at 0.766, compared to 0.713 for the BT model. Among the 1,249,168,756 distinct triplets in the data, stochastic transitivity is violated in 70% of cases, as indicated by the matrix of estimated probabilities $\hat{\Pi}$ under the proposed model. Specifically, this occurs when there exists an ordering of the three players, denoted as i , j , and k , such that $\hat{\pi}_{ik} \geq \hat{\pi}_{ij}$ and $\hat{\pi}_{jk} < 0.5$.

These results are consistent with previous findings by Chen and Joachims (2016), who analyzed a similar dataset over different time frames, suggesting that a strict ranking structure may not be appropriate in e-sports. In particular, intransitivity can naturally arise from game design, such as intransitive relationships among different unit types, which provide players with significant flexibility in choosing units and strategies. Moreover, the strong performance of our method on this dataset confirms its ability to effectively handle sparsity in real-world data, aligning with both simulation and theoretical results.

6.4.2 TENNIS DATA

From Table 1, the BT model achieves a marginally better performance, with a log-likelihood of −322,483 compared to −333,076 for the proposed model, and a slightly higher test accuracy (0.658 vs 0.652). This advantage may come from the BT model’s smaller parameter space, which is more efficient when the data aligns well with the stochastic transitivity

assumption, where the level of intransitivity is minimal or absent. Nevertheless, the performance of the proposed model remains close to that of the BT model, demonstrating its robustness even in settings where transitivity holds. This flexibility is particularly useful when intransitivity is uncertain, as it maintains high accuracy without relying on strict ranking assumptions.

The lack of intransitivity in professional tennis may be due to several factors. Unlike e-sports, tennis offers limited gameplay flexibility, as adjustments to equipment like rackets and shoes have minimal impact compared to the choice of units in *StarCraft II*. Additionally, professional tennis players may be required to be well-rounded as weaknesses are quickly identified and exploited by opponents. In contrast, intransitivity may be more common at lower levels of competition, where skill imbalances are expected to be more significant. For example, a player with a strong serve but weak baseline play may be more likely to defeat one opponent while losing to another with a different style. Investigating intransitivity in lower-tier competitions remains an open question for future research.

7. Discussions

In this article, we propose a statistical framework for modeling stochastic intransitivity. The framework assumes an approximate low-rank structure in the parameter matrix, expressed through a nuclear norm constraint. Theoretical analysis demonstrates that the proposed estimator achieves optimal convergence rates under a wide range of data sparsity settings. Simulation and empirical analyses confirm that our model is superior to the Bradley-Terry model when the assumption of stochastic transitivity is violated.

Our framework stands apart from the existing literature by imposing an approximate low-rank structure. To our knowledge, all existing methods for pairwise comparison data rely on exact low-rank models, even in the limited works that allow stochastic intransitivity. By accommodating a larger parameter space, our approach offers greater flexibility and applicability to a wider range of datasets. While this may lead to slightly reduced efficiency, our analysis of the tennis dataset demonstrates that the loss of efficiency is small when stochastic transitivity largely holds. Therefore, the proposed model may predict pairwise comparison results more accurately in many real-world applications. For example, for tournament data, this could lead to more accurate predictions of the champion or the number of rounds each player can play, given historical data and the current tournament schedule.

The current research may be extended in several directions. Specifically, the current theoretical analysis focuses on the convergence of the loss $\|\hat{\Pi} - \Pi^*\|^2/(n^2 - n)$, which can be seen as a notion of convergence in an average sense (across entries of the comparison probability matrix). As shown in Remark 6, under additional structural assumptions, we further establish convergence under the matrix max-norm loss $\|\hat{\Pi} - \Pi^*\|_\infty$ by leveraging the refinement techniques proposed in Chen and Li (2024). It would be of interest to determine whether such convergence can be obtained under weaker conditions, or whether sharper convergence rates can be achieved through an improved refinement procedure. Moreover, it will be useful to further establish the asymptotic normality for each $\hat{\pi}_{ij} - \pi_{ij}^*$, which can be used to quantify the uncertainty associated with the estimated comparison probabilities.

The proposed modeling framework also needs to be extended to accommodate more complex settings of pairwise comparisons. First, covariate information can be incorporated

into the model to facilitate the prediction. For example, for many team sports tournaments (e.g., soccer and basketball), whether a team plays at their home court matters and should be included as a covariate. Second, pairwise comparison data are often collected over time, which is true for the *StarCraft II* and tennis data studied in Section 6. The current model ignores time information in data. To better predict future pairwise comparison results, it will be useful to model the comparison probabilities as a function of time. As a result, the estimation of these time-varying comparison probabilities will also differ substantially from the current procedure. Third, for pairwise comparison data produced by raters, which are commonly encountered in crowd-sourcing settings (e.g., Chen et al., 2013), characteristics of the raters, such as their reliability, affect the pairwise comparisons. In other words, the distribution of the comparison between two items depends not only on the pair of items but also on the rater who performs the comparison. In this regard, Chen et al. (2013) propose an extended version of the BT model that uses a rater-specific latent variable to account for raters’ reliability. A similar extension can be made to the current model to simultaneously account for the raters’ heterogeneity and the items’ stochastic intransitivity.

Acknowledgments

The authors would like to thank the editor and the two anonymous reviewers for their constructive and valuable comments, which substantially improved the paper.

Appendix A. Additional Theoretical Results

A.1 Assumptions and Refinement Procedure for Max-norm and Two-to-infinity Norm Convergence

In this section, we state the assumptions and refinement procedures that enable us to obtain max-norm convergence of the underlying parameter matrix. In particular, we impose the following assumptions.

Assumption 9 *M^* has rank $2K$ such that $M^* = U^* \Sigma^* J_K U^{*\top}$ with $\|U^* \Sigma^{*1/2}\|_{2 \rightarrow \infty} \leq C^*$ for some fixed constant $C^* > 0$. Moreover, we assume each singular value has multiplicity exactly 2. Let $\sigma_1^* \geq \sigma_2^* \cdots \geq \sigma_K^*$ be the largest K distinct singular values of M^* . We assume that $\sigma_K^* \asymp \sigma_1^* \asymp n$.*

Assumption 10 *Assumption 2 holds with $p_n \gtrsim n^{-1/3} \log(n)$.*

Assumption 9 requires that M^* has an exact low-rank structure, with $\|U^* \Sigma^{*1/2}\|_{2 \rightarrow \infty} \leq C^*$ serving as a standard incoherence condition to prevent spiky low-rank matrices (see, e.g., Chen et al., 2020b; Chen and Li, 2024). It is straightforward to verify that Assumption 9 implies Assumption 1, which only requires an approximate low-rank structure. Since M^* is skew-symmetric, the decomposition $M^* = U^* \Sigma^* J_K U^{*\top}$ always holds by singular value decomposition. The assumptions on distinct and comparable singular values are made for simplicity and can be relaxed.

Assumption 10 strengthens the sparsity requirement, requiring $p_n \gtrsim n^{-1/3} \log n$. This condition arises from the dependence between the estimator and the missingness pattern in the refinement step. It may be relaxed by adopting a data-splitting refinement procedure, as in Chen and Li (2024), to remove this dependence. We leave this extension to future work. Under these assumptions, we define the estimator satisfying both the nuclear norm and max norm constraint as

$$\hat{M} = \arg \max_M \mathcal{L}(M) \text{ subject to } \|M\|_* \leq C_n n \text{ and } \|M\|_\infty \leq C_n / \sqrt{2K}, M = -M^\top. \quad (7)$$

Similar to (2), the above optimisation problem can be solved by a projected line search algorithm, and the projection step can be modified following Section 4.1.3 of Davenport et al. (2014) to accommodate the additional constraint $\|M\|_\infty \leq C_n / \sqrt{2K}$.

For any $n \times 2K$ matrix $\Theta = (\theta_1, \dots, \theta_n)^\top$, define $P_{\tau, \infty}^2(\Theta) = (\check{\theta}_1, \dots, \check{\theta}_n)^\top$, where $\check{\theta}_i = \theta_i$ if $\|\theta_i\| \leq \tau$ and $\check{\theta}_i = (\tau / \|\theta_i\|) \theta_i$ otherwise. We further refine $\hat{M}$ using the following refinement procedure to obtain an estimator $\tilde{M}$ that achieves max-norm convergence.

The following theorem establishes the two-to-infinity norm convergence of $\tilde{\Theta}$ and the max-norm convergence of $\tilde{M}$ obtained from Algorithm 4.

Theorem 11 *Under Assumptions 9 and 10, with probability converging to 1, we have*

$$\|\tilde{\Theta} - \Theta^*\|_{2 \rightarrow \infty} \leq \kappa_6 n^{-1/4} p_n^{-3/4} \text{ and} \quad (8)$$

$$\|\tilde{M} - M^*\|_\infty \leq \kappa_7 n^{-1/4} p_n^{-3/4}, \quad (9)$$

where $\kappa_6, \kappa_7 > 0$ are absolute constants.

Algorithm 4 Refinement Procedure

Input: Matrix of comparison outcomes Y , nuclear norm constraint parameter τ , initial estimate $\hat{M}$ solving (7), two-to-infinity norm constraint parameter τ_2 .

Compute the best rank- $2K$ approximation of $\hat{M}$ using its top $2K$ singular values to compute $\hat{M}_K$ given by

$$\hat{M}_K = \hat{U} \hat{\Sigma} J_K \hat{U}^\top,$$

where J_K is a $2K \times 2K$ block diagonal matrix with 2×2 blocks

$$\begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix}$$

along the main diagonal.

Compute $\hat{\Theta} = P_{\tau_2}^{2,\infty}(\hat{U} \hat{\Sigma}^{1/2})$.

Calculate $\tilde{\Theta} = (\tilde{\theta}_1, \tilde{\theta}_2, \dots, \tilde{\theta}_n)^\top$, where we define $\hat{\theta}_j^{(p)} = J_K \hat{\theta}_j$ for $j = 1, \dots, n$, and $\tilde{\theta}_i$ is obtained by solving the equation $n^{-1} \sum_{j \in [n], j \neq i} \hat{\theta}_j^{(p)} \left\{ y_{ij} - n_{ij} g(\tilde{\theta}_i^\top \hat{\theta}_j^{(p)}) \right\} = 0$.

Output: $\tilde{M} = \tilde{\Theta} J_K \tilde{\Theta}^\top$.

An immediate consequence of Theorem 11 is that the same entrywise rate also holds for $\tilde{\Pi} - \Pi^*$. In particular, we have

$$\|\tilde{\Pi} - \Pi^*\|_\infty \leq \kappa_8 n^{-1/4} p_n^{-3/4}, \quad (10)$$

for some positive constant κ_8 , with probability converging to 1.

In addition to the effect of the refinement procedure discussed above, the convergence rates in the theorem also depend on the Frobenius norm error of $\hat{M}$. Therefore, sharper rates may be obtained by improving the initial convergence rate of $\hat{M}$ through a more refined exploitation of the low-rank structure. We leave this direction for future research.

A.2 Recovery of Top-k Items

For any subset of subjects $\mathcal{S} \subseteq [n]$, we may define a score within that subset by

$$r_{i,\mathcal{S}}(\Pi^*) := \frac{1}{|\mathcal{S}|} \sum_{j \in \mathcal{S}} \Pi_{ij}^*, \quad i \in \mathcal{S}.$$

This quantity can be interpreted as the probability that subject i defeats an opponent drawn uniformly at random from $\mathcal{S}$. For example, if $\mathcal{S}$ represents the set of players participating in a particular competition, then $r_{i,\mathcal{S}}(\Pi^*)$ measures the average strength of player i relative to that group. When $\mathcal{S} = [n]$, this reduces to the score considered in Shah and Wainwright (2018).

Based on these scores, we may rank the subjects in $\mathcal{S}$ and study recovery of the top- k set. Let $(1), \dots, (|\mathcal{S}|)$ denote the indices of the subjects sorted in descending order according to $r_{i,\mathcal{S}}(\Pi^*)$, and define the k -separation threshold as

$$\Delta_{k,\mathcal{S}} := r_{(k),\mathcal{S}}(\Pi^*) - r_{(k+1),\mathcal{S}}(\Pi^*).$$

The following theorem gives a sufficient condition for consistent identification of the top- k set.

Theorem 12 *Under Assumptions 9 and 10, suppose that*

$$\Delta_{k,\mathcal{S}} \geq 2\kappa_8 n^{-1/4} p_n^{-3/4},$$

where κ_8 is defined in (10). Then

$$\mathbb{P}(\tilde{\mathcal{S}}_k \neq \mathcal{S}_k^*) \rightarrow 0,$$

where $\mathcal{S}_k^*$ and $\tilde{\mathcal{S}}_k$ denote the sets of the top- k subjects in $\mathcal{S}$ ranked according to $r_{i,\mathcal{S}}(\Pi^*)$ and $r_{i,\mathcal{S}}(\tilde{\Pi})$, respectively.

The result follows directly from the entrywise convergence $\|\tilde{\Pi} - \Pi^*\|_\infty$ established in (10). In particular, the required separation level is determined by the rate of this entrywise error. Consequently, if a sharper convergence rate for $\|\tilde{\Pi} - \Pi^*\|_\infty$ is available, the separation condition on $\Delta_{k,\mathcal{S}}$ can be correspondingly weakened.

Appendix B. Proof of Theoretical Results

This section presents the proofs of the theoretical results, where Section B.1 proves Theorem 3, Section B.2 proves Theorem 4, Section B.3 proves Theorem 11, Section B.4 proves Theorem 12 and Section B.5 proves Proposition 8. Throughout this section, $\delta_0, \delta_1, \dots$ denote positive constants that do not depend on n . For two probability distributions $\mathcal{P}$ and Q on a finite set A , $D(\mathcal{P}||Q)$ will denote the Kullback-Leibler (KL) divergence,

$$D(\mathcal{P}||Q) = \sum_{x \in A} \mathcal{P}(x) \log \left(\frac{\mathcal{P}(x)}{Q(x)} \right).$$

B.1 Proof of Theorem 3

For two scalars $x, z \in [0, 1]$, define the Hellinger distance as

$$d_H^2(x, z) = (\sqrt{x} - \sqrt{z})^2 + (\sqrt{1-x} - \sqrt{1-z})^2.$$

For $n \times n$ matrices $X = (x_{ij})_{n \times n}$ and $Z = (z_{ij})_{n \times n}$ where $X, Z \in [0, 1]^{n \times n}$, define

$$d_H^2(X, Z) = \frac{1}{n^2} \sum_{i=1}^n \sum_{j=1}^n d_H^2(x_{ij}, z_{ij}).$$

It is straightforward to show that $d_H^2(X, Z) \gtrsim \|X - Z\|_F^2 / (n^2 - n)$. Moreover, let $\|X\|_\infty = \max_{i,j} |x_{ij}|$ denotes the entry-wise infinity norm of X . We will first prove the theorem under an additional constraint that $\|M^*\|_\infty \leq \gamma$ and $\|\hat{M}\|_\infty \leq \gamma$ for some $\gamma > 0$, then send $\gamma \rightarrow \infty$ to recover Theorem 3. Formally, we prove the following theorem:

Theorem 13 *Under the conditions in Theorem 3, suppose in addition that $\|M^*\|_\infty \leq \gamma$. Let $\hat{M}$ be a solution to (2) under the additional constraint that $\|\hat{M}\|_\infty \leq \gamma$. Then with probability at least $1 - \delta_1/n$,*

$$d_H^2(\hat{\Pi}, \Pi^*) \leq \delta_2 C_n \sqrt{\frac{1}{p_n n}},$$

where δ_1 and δ_2 are absolute constants.

Proof Define $\bar{\mathcal{L}}(M) = \mathcal{L}(M) - \mathcal{L}(0_{n \times n})$. The following lemma is essential to proving Theorem 13:

Lemma 14 *Under the conditions in Theorem 13, we have*

$$P \left(\frac{1}{n^2} \sup_{M \in \mathcal{G}} |\bar{\mathcal{L}}(M) - E(\bar{\mathcal{L}}(M))| \geq \delta_0 C_n \sqrt{\frac{T q_n}{n}} \right) \leq \frac{\delta_1}{n},$$

where δ_0 is an absolute constant, and $\mathcal{G} \subset \mathbb{R}^{n \times n}$ is defined as

$$\mathcal{G} = \{M \in \mathbb{R}^{n \times n} : \|M\|_* \leq C_n n, \|M\|_\infty \leq \gamma, M = -M^\top\}.$$

Before proving the lemma, we first show how Lemma 14 implies Theorem 13. For two scalars $x, z \in [0, 1]$, we abuse the notation of $D(\cdot \|\cdot)$ and define the divergence measure as

$$D(x \| z) = x \log \left(\frac{x}{z} \right) + (1 - x) \log \left(\frac{1 - x}{1 - z} \right).$$

Similarly, for two matrices $X, Z \in [0, 1]^{n \times n}$, define

$$D(X \| Z) = \sum_{i=1}^n \sum_{j=1}^n D(x_{ij} \| z_{ij}).$$

For any choice of $M \in \mathcal{G}$, we have

$$\begin{aligned} & E(\bar{\mathcal{L}}(M) - \bar{\mathcal{L}}(M^*)) \\ &= E(\mathcal{L}(M) - \mathcal{L}(M^*)) \\ &= \sum_{i=1}^n \sum_{j>i} E \left(y_{ij} \log \left(\frac{g(m_{ij})}{g(m_{ij}^*)} \right) + (n_{ij} - y_{ij}) \log \left(\frac{1 - g(m_{ij})}{1 - g(m_{ij}^*)} \right) \right) \\ &= \sum_{i=1}^n \sum_{j>i} E \left(n_{ij} g(m_{ij}^*) \log \left(\frac{g(m_{ij})}{g(m_{ij}^*)} \right) + n_{ij} (1 - g(m_{ij}^*)) \log \left(\frac{1 - g(m_{ij})}{1 - g(m_{ij}^*)} \right) \right) \\ &= -T \sum_{i=1}^n \sum_{j>i} p_{ij,n} D(g(m_{ij}^*) \| g(m_{ij})) \\ &\leq -0.5 T p_n D(\Pi^* \| \Pi). \end{aligned}$$

Note that $M^* \in \mathcal{G}$ by assumption. Therefore, for any $M \in \mathcal{G}$, we have

$$\begin{aligned} \bar{\mathcal{L}}(M) - \bar{\mathcal{L}}(M^*) &= E(\bar{\mathcal{L}}(M) - \bar{\mathcal{L}}(M^*)) + (\bar{\mathcal{L}}(M) - E(\bar{\mathcal{L}}(M))) - (\bar{\mathcal{L}}(M^*) - E(\bar{\mathcal{L}}(M^*))) \\ &\leq E(\bar{\mathcal{L}}(M) - \bar{\mathcal{L}}(M^*)) + 2 \sup_{X \in \mathcal{G}} |\bar{\mathcal{L}}(X) - E(\bar{\mathcal{L}}(X))| \\ &\leq -0.5Tp_n D(\Pi^* \|\Pi) + 2 \sup_{X \in \mathcal{G}} |\bar{\mathcal{L}}(X) - E(\bar{\mathcal{L}}(X))|. \end{aligned}$$

Moreover, from the definition of $\hat{M}$, we have $\hat{M} \in \mathcal{G}$ and $\mathcal{L}(\hat{M}) \geq \mathcal{L}(M^*)$. Therefore, we obtain

$$0 \leq -0.5Tp_n D(\Pi^* \|\hat{\Pi}) + 2 \sup_{M \in \mathcal{G}} |\bar{\mathcal{L}}(M) - E(\bar{\mathcal{L}}(M))|.$$

Applying Lemma 14, then with probability at least $1 - \delta_1/n$, we have

$$0 \leq \frac{-0.5Tp_n D(\Pi^* \|\hat{\Pi})}{n^2} + 2\delta_0 C_n \sqrt{\frac{Tq_n}{n}}.$$

This implies that

$$\frac{D(\Pi^* \|\hat{\Pi})}{n^2} \leq \frac{4\delta_0 C_n}{Tp_n} \sqrt{\frac{Tq_n}{n}} \lesssim \frac{4\delta_0 C_n}{\sqrt{Tp_n}} \sqrt{\frac{1}{n}}$$

by Assumption 2. Note that $d_H^2(\hat{\Pi}, \Pi^*) \leq n^{-2} D(\Pi^* \|\hat{\Pi})$ by Jensen's inequality combined with the fact that $(1-x) \leq \log(x)$. Hence Theorem 13 is proved. Theorem 3 then follows by the fact that $d_H^2(\hat{\Pi}, \Pi^*) \gtrsim \|\hat{\Pi} - \Pi^*\|_F^2 / (n^2 - n)$ and taking the limit as $\gamma \rightarrow \infty$. $\blacksquare$

We now begin to prove Lemma 14.

Proof For any $h > 0$, using Markov's inequality, we have

$$\begin{aligned} &P\left(\frac{1}{n^2} \sup_{M \in \mathcal{G}} |\bar{\mathcal{L}}(M) - E(\bar{\mathcal{L}}(M))| \geq \delta_0 C_n \sqrt{Tq_n/n}\right) \\ &= P\left(\sup_{M \in \mathcal{G}} |\bar{\mathcal{L}}(M) - E(\bar{\mathcal{L}}(M))|^h \geq \left(\delta_0 C_n n^{1.5} \sqrt{Tq_n}\right)^h\right) \\ &\leq \frac{E\left(\sup_{M \in \mathcal{G}} |\bar{\mathcal{L}}(M) - E(\bar{\mathcal{L}}(M))|^h\right)}{(\delta_0 C_n n^{1.5} \sqrt{Tq_n})^h}. \end{aligned} \tag{11}$$

The bound in Lemma 14 will be established by combining (11), deriving an upper bound on $E\left(\sup_{M \in \mathcal{G}} |\bar{\mathcal{L}}(M) - E(\bar{\mathcal{L}}(M))|^h\right)$ and setting $h = \log(n)$. Note that we can write $\bar{\mathcal{L}}(M)$ as

$$\bar{\mathcal{L}}(M) = \sum_{i=1}^n \sum_{j>i} y_{ij} \log\left(\frac{g(m_{ij})}{g(0)}\right) + (n_{ij} - y_{ij}) \log\left(\frac{1 - g(m_{ij})}{1 - g(0)}\right).$$

By a symmetrization argument (Lemma 6.3 in Ledoux and Talagrand (1991)), we have

$$\begin{aligned} &E\left(\sup_{M \in \mathcal{G}} |\bar{\mathcal{L}}(M) - E(\bar{\mathcal{L}}(M))|^h\right) \\ &\leq 2^h E\left(\sup_{M \in \mathcal{G}} \left|\sum_{i=1}^n \sum_{j>i} \epsilon_{ij} \left\{y_{ij} \log\left(\frac{g(m_{ij})}{g(0)}\right) + (n_{ij} - y_{ij}) \log\left(\frac{1 - g(m_{ij})}{1 - g(0)}\right)\right\}\right|^h\right), \end{aligned}$$

where $\epsilon_{i,j}$ are i.i.d. Rademacher random variables for $i, j = 1, \dots, n$. To bound the latter term, we apply a contraction principle (Theorem 4.12 in Ledoux and Talagrand (1991)). From the assumption that $\|M\|_\infty \leq \gamma$, conditional on n_{ij} , for $n_{ij} \geq 1$,

$$n_{ij}^{-1} \left(y_{ij} \log \left(\frac{g(m_{ij})}{g(0)} \right) + (n_{ij} - y_{ij}) \log \left(\frac{1 - g(m_{ij})}{1 - g(0)} \right) \right)$$

is a contraction that vanish at 0. Thus, we have

$$\begin{aligned} E \left(\sup_{M \in \mathcal{G}} |\bar{\mathcal{L}}(M) - E(\bar{\mathcal{L}}(M))|^h \right) &\leq (2^h)(2^h) E \left(\sup_{M \in \mathcal{G}} \left| \sum_{i=1}^n \sum_{j>i} n_{ij} \epsilon_{ij} m_{ij} \right|^h \right) \\ &= 4^h E \left(\sup_{M \in \mathcal{G}} \left| \sum_{i=1}^n \sum_{j>i} n_{ij} \epsilon_{ij} m_{ij} \right|^h \right). \end{aligned} \quad (12)$$

To bound $E \left(\sup_{M \in \mathcal{G}} \left| \sum_{i=1}^n \sum_{j>i} n_{ij} \epsilon_{ij} m_{ij} \right|^h \right)$, we apply the skew-symmetric property of M and the fact that $n_{ij} = n_{ji}$ for $i, j \in \{1, \dots, n\}$. For any $M \in \mathcal{G}$, we have

$$\sum_{i=1}^n \sum_{j=1}^n n_{ij} \epsilon_{ij} m_{ij} = \sum_{i=1}^n \sum_{j>i} n_{ij} (\epsilon_{ij} - \epsilon_{ji}) m_{ij}.$$

On the other hand, for $h > 1$, by the convexity of $|\cdot|^h$, we have

$$\begin{aligned} \left| \sum_{i=1}^n \sum_{j>i} n_{ij} \epsilon_{ij} m_{ij} \right|^h &= \left| 0.5 \left\{ \sum_{i=1}^n \sum_{j>i} n_{ij} (\epsilon_{ij} - \epsilon_{ji}) m_{ij} + \sum_{i=1}^n \sum_{j>i} n_{ij} (\epsilon_{ij} + \epsilon_{ji}) m_{ij} \right\} \right|^h \\ &\leq 0.5 \left(\left| \sum_{i=1}^n \sum_{j>i} n_{ij} (\epsilon_{ij} - \epsilon_{ji}) m_{ij} \right|^h + \left| \sum_{i=1}^n \sum_{j>i} n_{ij} (\epsilon_{ij} + \epsilon_{ji}) m_{ij} \right|^h \right). \end{aligned}$$

Since ϵ_{ji} and $-\epsilon_{ji}$ have identical distribution, after taking expectation, we have

$$\begin{aligned} E \left(\sup_{M \in \mathcal{G}} \left| \sum_{i=1}^n \sum_{j>i} n_{ij} \epsilon_{ij} m_{ij} \right|^h \right) &\leq E \left(\sup_{M \in \mathcal{G}} \left| \sum_{i=1}^n \sum_{j>i} n_{ij} (\epsilon_{ij} - \epsilon_{ji}) m_{ij} \right|^h \right) \\ &= E \left(\sup_{M \in \mathcal{G}} \left| \sum_{i=1}^n \sum_{j=1}^n n_{ij} \epsilon_{ij} m_{ij} \right|^h \right) \\ &= E \left(\sup_{M \in \mathcal{G}} |\langle \mathcal{E} \circ \mathcal{N}, M \rangle|^h \right). \end{aligned} \quad (13)$$

Here, $\mathcal{E} = (\epsilon_{ij})_{n \times n}$, $\mathcal{N} = (n_{ij})_{n \times n}$ and $\mathcal{E} \circ \mathcal{N}$ represents the hadamard product between $\mathcal{E}$ and $\mathcal{N}$, and $\langle X, Z \rangle = \sum_{i=1}^n \sum_{j=1}^n x_{ij} z_{ij}$ for any $n \times n$ matrices X and Z . Note that $|\langle X, Z \rangle| \leq \|X\|_{op} \|Z\|_*$, where $\|\cdot\|_{op}$ is the Euclidean operator norm. Hence we have

$$\begin{aligned} E \left(\sup_{M \in \mathcal{G}} |\langle \mathcal{E} \circ \mathcal{N}, M \rangle|^h \right) &\leq E \left(\sup_{M \in \mathcal{G}} \|\mathcal{E} \circ \mathcal{N}\|_{op}^h \|M\|_*^h \right) \\ &\leq (C_n n)^h E \left(\|\mathcal{E} \circ \mathcal{N}\|_{op}^h \right). \end{aligned} \quad (14)$$

We can write $\mathcal{E} \circ \mathcal{N} = \sum_{i=1}^n \sum_{j=1}^n \epsilon_{ij} n_{ij} E_{ij}$, where E_{ij} is a $n \times n$ matrix with 1 at the (i, j) th entry and 0 otherwise. Following arguments similar to Section 4.3 of Tropp et al. (2015), and applying Theorem 4.1.1 in Tropp et al. (2015), for $t > 0$, we set $s = -t^{2/h} / (2 \max_j \{\sum_{i=1}^n n_{ij}^2\})$ such that

$$\begin{aligned} E(\|\mathcal{E} \circ \mathcal{N}\|_{op}^h | \mathcal{N}) &= \left(\int_0^\infty P(\|\mathcal{E} \circ \mathcal{N}\|_{op}^h \geq t) dt \right) \\ &\leq \left(\int_0^\infty 2n \exp \left(\frac{-t^{2/h}}{2 \max_j \{\sum_{i=1}^n n_{ij}^2\}} \right) dt \right) \\ &= \left(\int_0^\infty 2n \left(\frac{h}{2} (2 \max_j \{\sum_{i=1}^n n_{ij}^2\})^{h/2} s^{h/2-1} \right) \exp(-s) ds \right) \\ &= \left((nh) (2 \max_j \{\sum_{i=1}^n n_{ij}^2\})^{h/2} \int_0^\infty s^{h/2-1} \exp(-s) ds \right) \\ &= nh \Gamma(h/2) (2 \max_j \{\sum_{i=1}^n n_{ij}^2\})^{h/2}, \end{aligned}$$

where $\Gamma(\cdot)$ is the gamma function. Taking expectation, we have

$$E(\|\mathcal{E} \circ \mathcal{N}\|_{op}^h) \leq nh \Gamma(h/2) 2^{h/2} E(\max_j \{\sum_{i=1}^n n_{ij}^2\}^{h/2}). \quad (15)$$

We aim to find a bound for $E(\max_j \{\sum_{i=1}^n n_{ij}^2\}^{h/2})$. Using Bernstein's inequality, for each j and all $t > 0$, we have

$$\begin{aligned} P \left(\left| \sum_{i=1}^n (n_{ij}^2 - E(n_{ij}^2)) \right| > t \right) &\leq 2 \exp \left(\frac{-t^2/2}{\sum_{i=1}^n \{E(n_{ij}^4) - (E(n_{ij}^2))^2\} + T^2 t/3} \right) \\ &\leq 2 \exp \left(\frac{-t^2/2}{nT^4 q_n + T^2 t/3} \right). \end{aligned}$$

In particular, for $t \geq 6nT^2 q_n$, we have

$$P \left(\left| \sum_{i=1}^n (n_{ij}^2 - E(n_{ij}^2)) \right| > t \right) \leq 2 \exp(-t/T^2) = 2P(U_j > t/T^2),$$

where $U_1, \dots, U_n$ are independent and identically distributed exponential random variables. Hence, we have

$$\begin{aligned}
 & \left(E \left(\max_j \left\{ \sum_{i=1}^n n_{ij}^2 \right\}^{h/2} \right) \right)^{1/h} \\
 &= \left(E \left(\max_j \left| \sum_{i=1}^n n_{ij}^2 - E(n_{ij}^2) + E(n_{ij}^2) \right|^{h/2} \right) \right)^{1/h} \\
 &\leq 2 \left(E \left(\max_j \left| \sum_{i=1}^n n_{ij}^2 - E(n_{ij}^2) \right|^{h/2} \right) \right)^{1/h} + 2 \left(E \left(\max_j \left| \sum_{i=1}^n E(n_{ij}^2) \right|^{h/2} \right) \right)^{1/h} \\
 &\leq 2\sqrt{nT^2q_n} + 2 \left(E \left(\max_j \left| \sum_{i=1}^n n_{ij}^2 - E(n_{ij}^2) \right|^h \right) \right)^{1/2h} \\
 &= 2\sqrt{nT^2q_n} + 2 \left(\int_0^\infty P \left(\max_j \left| \sum_{i=1}^n n_{ij}^2 - E(n_{ij}^2) \right| \geq t \right) dt \right)^{1/2h} \\
 &\leq 2\sqrt{nT^2q_n} + 2 \left\{ (6nT^2q_n)^h + \int_{(6nT^2q_n)^h}^\infty P \left(\max_j \left| \sum_{i=1}^n n_{ij}^2 - E(n_{ij}^2) \right| \geq t \right) dt \right\}^{1/2h} \\
 &\leq 2\sqrt{nT^2q_n} + 2 \left\{ (6nT^2q_n)^h + 2 \int_{(6nT^2q_n)^h}^\infty P \left(\max_j \{U_j\}^h \geq t/T^{2h} \right) dt \right\}^{1/2h} \\
 &\leq 2\sqrt{nT^2q_n} + 2 \left\{ (6nT^2q_n)^h + 2E \left(\max_j \{T^2U_j\}^h \right) \right\}^{1/2h} \\
 &= 2\sqrt{nT^2q_n} + 2 \left\{ (6nT^2q_n)^h + 2T^{2h}E \left((\max_j \{U_j\})^h \right) \right\}^{1/2h}.
 \end{aligned}$$

By standard computations for exponential random variables, we can obtain the inequality $E((\max_j \{U_j\})^h) \leq 2h! + \log^h(n)$. Thus, we have

$$\begin{aligned}
 \left(E \left(\max_j \left\{ \sum_{i=1}^n n_{ij}^2 \right\}^{h/2} \right) \right)^{1/h} &\leq 2\sqrt{nT^2q_n} + 2 \left\{ (6nT^2q_n)^h + 2T^{2h}(2h! + \log^h(n)) \right\}^{1/2h} \\
 &\leq 2T(1 + \sqrt{6})\sqrt{nq_n} + 2T(2)^{1/2h}(\sqrt{\log(n)} + 2^{1/2h}\sqrt{h}) \\
 &\leq 2T(1 + \sqrt{6})\sqrt{nq_n} + 2T(2 + \sqrt{2})\sqrt{\log(n)}
 \end{aligned}$$

using the choice $h = \log(n)$ in the final line. Combining this result with (15), we have

$$\begin{aligned}
 E(\|\mathcal{E} \circ \mathcal{N}\|_{op}^h)^{1/h} &\leq (nh\Gamma(h/2))^{1/h}\sqrt{2}\{2T(1 + \sqrt{6})\sqrt{nq_n} + 2T(2 + \sqrt{2})\sqrt{\log(n)}\} \\
 &\leq \delta_3 T \sqrt{nq_n}
 \end{aligned}$$

for some constant $\delta_3 > 0$ by Assumption 2. Combining this with (12), (13) and (14), we obtain

$$E \left(\sup_{M \in \mathcal{G}} |\bar{\mathcal{L}}(M) - E(\bar{\mathcal{L}}(M))|^h \right)^{1/h} \leq (4T) (C_n n) (\delta_3) \sqrt{n q_n}.$$

Plugging this into (11), the probability in (11) is upper bounded by

$$\left\{ \frac{(4T) (C_n n) (\delta_3) \sqrt{n q_n}}{\delta_0 C_n n^{1.5} \sqrt{T q_n}} \right\}^h \leq \left(\frac{4\sqrt{T} \delta_3}{\delta_0} \right)^{\log(n)} \leq \frac{\delta_1}{n},$$

provided that $\delta_0 \geq 4\sqrt{T} \delta_3 / e$, which establishes the lemma. $\blacksquare$

Remark 15 *The assumption $p_n \asymp q_n$ can be relaxed based on the above proof. In particular, without imposing this condition, one can show that with probability at least $1 - \kappa_1/n$,*

$$\frac{1}{n^2 - n} \|\hat{\Pi} - \Pi^*\|_F^2 \leq \kappa_2 C_n \sqrt{\frac{q_n}{p_n^2 n}},$$

under the remaining assumptions of Theorem 3. If T diverges rather than remains fixed, a similar argument, following the proof of Theorem 13 using the rate established in Lemma 14, yields

$$\frac{1}{n^2 - n} \|\hat{\Pi} - \Pi^*\|_F^2 \leq \kappa_2 C_n \sqrt{\frac{1}{T p_n n}}$$

with probability converging to one. Under these relaxed settings, it is unclear whether the rate is optimal, and we leave this question for future work.

B.2 Proof of Theorem 4

We first quote the following lemma from Davenport et al. (2014):

Lemma 16 *Suppose $x, z \in (0, 1)$. Then*

$$D(x||z) \leq \frac{(x - z)^2}{z(1 - z)}.$$

The following lemma constructs a packing set $\mathcal{X} \subset \mathcal{K}$ such that, for any distinct $X^{(a)}, X^{(b)} \in \mathcal{X}$, $\|X^{(a)} - X^{(b)}\|_F^2$ is large:

Lemma 17 *Let $\mathcal{K}$ be defined as in (3), and k a positive integer. Let $\gamma \leq 1$ be such that k/γ^2 is an integer, and suppose $k/\gamma^2 \leq n$. Then, there exists a set $\mathcal{X} \subset \mathcal{K}$ satisfying*

$$|\mathcal{X}| \geq \exp \left(\frac{kn}{25600\gamma^2} \right)$$

with the following properties:

1. For all $X = (x_{ij})_{n \times n} \in \mathcal{X}$, each entry of X satisfies $|x_{ij}| \leq C_n \gamma / \sqrt{2k}$.
2. For all $X^{(a)} \neq X^{(b)} \in \mathcal{X}$,

$$\|X^{(a)} - X^{(b)}\|_F^2 > \frac{C_n^2 \gamma^2 n^2}{16k}.$$

Proof We use a probabilistic argument. The set will be constructed by drawing

$$|\mathcal{X}| = \left\lceil \exp \left(\frac{kn}{25600\gamma^2} \right) \right\rceil$$

matrices independently from the following distribution. Set $B = k/\gamma^2$. Each matrix in $\mathcal{X}$ is constructed of the form $S - S^\top$, where $S = (s_{ij})_{n \times n}$ consists of blocks of dimension $B \times n$, stacked vertically. The entries of the first block are independent and identically distributed symmetric random variables taking values $\pm C_n \gamma / (2\sqrt{2k})$. Then S is filled out by copying this block as many times as it fits. That is,

$$s_{ij} = s_{i'j}, \text{ where } i' = i \pmod{B} + 1.$$

Now we argue that with nonzero probability, this set will have all the desired properties. For $X \in \mathcal{X}$, it is easy to verify that $X = -X^\top$. Moreover, we have

$$\|X\|_\infty \leq 2\{C_n \gamma / (2\sqrt{2k})\} \leq C_n / \sqrt{2k}.$$

Further, since $\text{rank}(X) \leq 2\text{rank}(S) \leq 2B$,

$$\|X\|_* \leq \sqrt{2B} \|X\|_F \leq \sqrt{2k/\gamma^2} n (C_n \gamma / \sqrt{2k}) = C_n n.$$

Thus $\mathcal{X} \subset \mathcal{K}$, and it remains to show that $\mathcal{X}$ satisfies property 2 in Lemma 17. Let $p = \lfloor n/B \rfloor$. Consider the submatrix of S containing the first B rows, denoted by $S_{[1:B, :]}$. This can be written as

$$S_{[1:B, :]} = (S_1, S_2, \dots, S_p, S_{p+1}),$$

where $S_1, \dots, S_p$ are matrices of dimension $B \times B$, and S_{p+1} accounts for the remaining part of $S_{[1:B, :]}$. If n is divisible by B , then S_{p+1} is an empty matrix. For $X^{(a)} = S^{(a)} - (S^{(a)})^\top$ and $X^{(b)} = S^{(b)} - (S^{(b)})^\top$, drawn from the above distribution, define

$$\Theta_i = \frac{\sqrt{2k}}{C_n \gamma} \left(S_i^{(a)} - S_i^{(b)} \right), \quad \text{for } i = 1, \dots, p.$$

Each Θ_i is a $B \times B$ matrix, and we write $\Theta_i = (\theta_{i,sl})_{B \times B}$, where each $\theta_{i,sl}$ is independent and identically distributed random variables such that for each $s, l \in \{1, \dots, B\}$, we have

$$P(\theta_{i,sl} = 1) = P(\theta_{i,sl} = -1) = 0.25 \text{ and } P(\theta_{i,sl} = 0) = 0.5.$$

Hence we can write

$$\begin{aligned}
 \|X^{(a)} - X^{(b)}\|_F^2 &\geq \sum_{i=1}^p \sum_{j=1}^p \|S_i^{(a)} - (S_j^{(a)})^\top - S_i^{(b)} + (S_j^{(b)})^\top\|_F^2 \\
 &= \frac{C_n^2 \gamma^2}{2k} \sum_{i=1}^p \sum_{j=1}^p \|\Theta_i - \Theta_j^\top\|_F^2 \\
 &= \frac{C_n^2 \gamma^2}{2k} \sum_{i=1}^p \sum_{j=1}^p (\|\Theta_i\|_F^2 + \|\Theta_j^\top\|_F^2 - 2\text{tr}(\Theta_i \Theta_j)) \\
 &= \frac{C_n^2 \gamma^2}{2k} \left\{ 2p \sum_{i=1}^p (\|\Theta_i\|_F^2) - 2\text{tr} \left(\left(\sum_{i=1}^p \Theta_i \right) \left(\sum_{i=1}^p \Theta_i \right) \right) \right\}.
 \end{aligned}$$

The trace can be expanded as:

$$\begin{aligned}
 \text{tr} \left(\left(\sum_{i=1}^p \Theta_i \right) \left(\sum_{i=1}^p \Theta_i \right) \right) &= \sum_{s=1}^B \sum_{k=1}^B \left(\sum_{i=1}^p \theta_{i,sl} \right) \left(\sum_{i=1}^p \theta_{i,ls} \right) \\
 &= 2 \sum_{s=1}^B \sum_{k>s}^B \left(\sum_{i=1}^p \theta_{i,sl} \right) \left(\sum_{i=1}^p \theta_{i,ls} \right) + \sum_{s=1}^B \left(\sum_{i=1}^p \theta_{i,ss} \right)^2.
 \end{aligned}$$

Hence we can write

$$\begin{aligned}
 &2p \sum_{i=1}^p (\|\Theta_i\|_F^2) - 2\text{tr} \left(\left(\sum_{i=1}^p \Theta_i \right) \left(\sum_{i=1}^p \Theta_i \right) \right) \\
 &= 2p \sum_{i=1}^p \sum_{s=1}^B \sum_{k=1}^B \theta_{i,sl}^2 - 4 \sum_{s=1}^B \sum_{k>s}^B \left(\sum_{i=1}^p \theta_{i,sl} \right) \left(\sum_{i=1}^p \theta_{i,ls} \right) - 2 \sum_{s=1}^B \left(\sum_{i=1}^p \theta_{i,ss} \right)^2 \\
 &= 2 \sum_{s=1}^B \{ p \sum_{i=1}^p (\theta_{i,ss}^2) - (\sum_{i=1}^p \theta_{i,ss})^2 \} + 2 \sum_{s=1}^B \sum_{k>s}^B \{ p \sum_{i=1}^p (\theta_{i,sl}^2 + \theta_{i,ls}^2) - 2 \left(\sum_{i=1}^p \theta_{i,sl} \right) \left(\sum_{i=1}^p \theta_{i,ls} \right) \}.
 \end{aligned}$$

Taking expectation, we have

$$\begin{aligned}
 E \left(2p \sum_{i=1}^p (\|\Theta_i\|_F^2) - 2\text{tr} \left(\left(\sum_{i=1}^p \Theta_i \right) \left(\sum_{i=1}^p \Theta_i \right) \right) \right) &= 2B \{ p(0.5p) - 0.5p \} + \frac{2B(B-1)p^2}{2} \\
 &= Bp^2 - Bp + B^2p^2 - Bp^2 \\
 &= B^2p^2 - Bp.
 \end{aligned}$$

Using the fact that $p = \lfloor n/B \rfloor \geq n/2B$ and $p \leq n/B$, we have

$$\begin{aligned}
 & P \left(2p \sum_{i=1}^p (\|\Theta_i\|_F^2) - 2tr \left(\left(\sum_{i=1}^p \Theta_i \right) \left(\sum_{i=1}^p \Theta_i \right) \right) \leq n^2/8 \right) \\
 &= P \left(-2p \sum_{i=1}^p (\|\Theta_i\|_F^2) + 2tr \left(\left(\sum_{i=1}^p \Theta_i \right) \left(\sum_{i=1}^p \Theta_i \right) \right) + B^2 p^2 - Bp \geq -n^2/8 + B^2 p^2 - Bp \right) \\
 &\leq P \left(-2p \sum_{i=1}^p (\|\Theta_i\|_F^2) + 2tr \left(\left(\sum_{i=1}^p \Theta_i \right) \left(\sum_{i=1}^p \Theta_i \right) \right) + B^2 p^2 - Bp \geq -n^2/8 + n^2/4 - n \right) \\
 &\leq P \left(-2p \sum_{i=1}^p (\|\Theta_i\|_F^2) + 2tr \left(\left(\sum_{i=1}^p \Theta_i \right) \left(\sum_{i=1}^p \Theta_i \right) \right) + B^2 p^2 - Bp \geq n^2/16 \right),
 \end{aligned}$$

where the last inequality holds as long as $n \geq 16$. Using McDiarmid's inequality, we can obtain the bound

$$\begin{aligned}
 P \left(2p \sum_{i=1}^p (\|\Theta_i\|_F^2) - 2tr \left(\left(\sum_{i=1}^p \Theta_i \right) \left(\sum_{i=1}^p \Theta_i \right) \right) \leq n^2/8 \right) &\leq \exp \left(-\frac{2(n^2/16)^2}{\sum_{s=1}^B \sum_{k=1}^B \sum_{i=1}^p (10p)^2} \right) \\
 &= \exp \left(-\frac{n^4}{12800B^2p^3} \right) \\
 &\leq \exp \left(-\frac{nB}{12800} \right).
 \end{aligned}$$

Using Union bound, we have that

$$P \left(\min_{X^{(a)} \neq X^{(b)} \in \mathcal{X}} 2p \sum_{i=1}^p (\|\Theta_i\|_F^2) - 2tr \left(\left(\sum_{i=1}^p \Theta_i \right) \left(\sum_{i=1}^p \Theta_i \right) \right) \leq n^2/8 \right) \leq \binom{|\mathcal{X}|}{2} \exp \left(-\frac{nB}{12800} \right),$$

which is less than 1 given the size of $\mathcal{X}$. Thus the event that

$$2p \sum_{i=1}^p (\|\Theta_i\|_F^2) - 2tr \left(\left(\sum_{i=1}^p \Theta_i \right) \left(\sum_{i=1}^p \Theta_i \right) \right) > n^2/8$$

for all $X^{(a)} \neq X^{(b)} \in \mathcal{X}$ has non-zero probability. In this event,

$$\|X^{(a)} - X^{(b)}\|_F^2 > \frac{C_n^2 \gamma^2}{2k} (n^2/8) = \frac{C_n^2 \gamma^2 n^2}{16k}.$$

The proof of the lemma is thus complete. ■

We now proceed to prove the following theorem, which concerns the lower bound treating n_{ij} as given.

Theorem 18 Suppose $12 \leq C_n^2 \leq \min\{1, (\kappa_3')^2/T\}n$. For any given n_{ij} , $i, j \in \{1, \dots, n\}$, $j > i$, consider any algorithm which, for any $M \in \mathcal{K}$, takes as input Y and returns $\hat{M}$. Then there exists $M \in \mathcal{K}$ such that with probability at least $3/4$, $\Pi = g(M)$ and $\hat{\Pi} = g(\hat{M})$ satisfy

$$\frac{1}{n^2 - n} \|\Pi - \hat{\Pi}\|_F^2 \geq \min \left\{ \kappa_4, \kappa_3' C_n \sqrt{\frac{n}{\sum_{i=1}^n \sum_{j=1}^n y_{ij}}} \right\}. \quad (16)$$

for all $n > N$. Here $\kappa_3', \kappa_4 > 0$ and N are absolute constants. We set $\kappa_3 = \kappa_3'/\sqrt{T}$.

Proof Let $c = g'(-1) = g(-1)(1 - g(-1))$, and let $c' = g(-1)$. Note that for all $x \in [-1, 1]$, we have $g'(x) \geq c$ and $c' \leq g(x) \leq 1 - c'$. We begin by choosing ϵ so that

$$\epsilon^2 = \min \left\{ \frac{c}{64}, \kappa_3' C_n \sqrt{\frac{n}{\sum_{i=1}^n \sum_{j=1}^n y_{ij}}} \right\}, \quad (17)$$

where κ_3' is an absolute constant to be determined. Let $k = 6$ and choose γ so that k/γ^2 is an integer and

$$4\sqrt{2} \frac{\epsilon\sqrt{2k}}{C_n c} \leq \gamma \leq \frac{8\epsilon\sqrt{2k}}{C_n c}.$$

This is possible since by assumption $C_n \geq \sqrt{12}$, $\epsilon \leq c/8$ and $c = 0.197$. One can check that γ satisfies the assumptions of Lemma 17. Note that for $X^{(i)} \neq X^{(j)} \in \mathcal{X}$,

$$\|g(X^{(i)}) - g(X^{(j)})\|_F^2 \geq c^2 \|X^{(i)} - X^{(j)}\|_F^2 > c^2 C_n^2 \gamma^2 n^2 / 16k \geq 4\epsilon^2 n^2. \quad (18)$$

Now suppose for the sake of a contradiction that there exists an algorithm such that for any $X \in \mathcal{K}$, it returns an $\hat{X}$ such that

$$\frac{1}{n^2} \|g(X) - g(\hat{X})\|_F^2 \leq \epsilon^2. \quad (19)$$

with probability at least $1/4$. Define

$$X^* = \arg \min_{X^{(a)} \in \mathcal{X}} \frac{1}{n^2} \|g(X^{(a)}) - g(\hat{X})\|_F^2.$$

If (19) holds, then (18) implies that $X^* = X$. Thus, if (19) holds with probability at least $1/4$ then

$$P(X \neq X^*) \leq 3/4.$$

However, by a variant of Fano's inequality, we have

$$P(X \neq X^*) \geq 1 - \frac{n^2 \max_{X^{(a)} \neq X^{(b)}} D(Y | X^{(a)} \| Y | X^{(b)}) + 1}{\log |\mathcal{X}|}. \quad (20)$$

Since $y_{ij} + y_{ji} = n_{ij}$ (with n_{ij} given), the value of y_{ji} is determined by y_{ij} . Moreover, y_{ij} are independent for $i = 1, \dots, n, j > i$. Therefore,

$$D(Y | X^{(a)} \| Y | X^{(b)}) = \sum_{i=1}^n \sum_{j>i} D(y_{ij} | x_{ij}^{(a)} \| y_{ij} | x_{ij}^{(b)}).$$

Using Lemma 16, we have

$$\begin{aligned} D(y_{ij} \mid x_{ij}^{(a)} \parallel y_{ij} \mid x_{ij}^{(b)}) &\leq \frac{(g(C_n\gamma/\sqrt{2k}) - g(-C_n\gamma/\sqrt{2k}))^2}{g(C_n\gamma/\sqrt{2k})(1 - g(C_n\gamma/\sqrt{2k}))} \\ &\leq \frac{4(g'(\xi))^2 C_n^2 \gamma^2 / (2k)}{g(C_n\gamma/\sqrt{2k})(1 - g(C_n\gamma/\sqrt{2k}))} \\ &= \frac{4\{g(\xi)(1 - g(\xi))\}^2 C_n^2 \gamma^2 / (2k)}{g(C_n\gamma/\sqrt{2k})(1 - g(C_n\gamma/\sqrt{2k}))} \end{aligned}$$

for some $|\xi| \leq C_n\gamma/\sqrt{2k}$. Since $c' < g(x) < 1 - c'$ for $|x| < 1$, $g(\xi) \leq g(C_n\gamma/\sqrt{2k})$, and that

$$C_n\gamma/\sqrt{2k} \leq C_n \frac{8\epsilon\sqrt{2k}}{C_n c\sqrt{2k}} = \frac{8\epsilon}{c} \leq 1,$$

we have

$$D(y_{ij} \mid x_{ij}^{(a)} \parallel y_{ij} \mid x_{ij}^{(b)}) \leq \frac{4(1 - c')}{c'} \frac{64\epsilon^2}{c^2} = \delta_4 \epsilon^2,$$

where $\delta_4 = 256(1 - c')/(c' c^2)$. Thus, from (20), we have

$$\begin{aligned} \frac{1}{4} &\leq \frac{\delta_4(\sum_{i=1}^n \sum_{j=1}^n y_{ij})\epsilon^2 + 1}{\log(|\mathcal{X}|)} \leq \frac{25600\gamma^2}{kn} \left\{ \delta_4 \left(\sum_{i=1}^n \sum_{j=1}^n y_{ij} \right) \epsilon^2 + 1 \right\} \\ &\leq \frac{3276800}{c^2} \epsilon^2 \left(\frac{\delta_4(\sum_{i=1}^n \sum_{j=1}^n y_{ij})\epsilon^2 + 1}{nC_n^2} \right). \end{aligned}$$

We now argue that this leads to a contradiction. Specifically, if $\delta_4(\sum_{i=1}^n \sum_{j=1}^n y_{ij})\epsilon^2 \leq 1$, then together with (17) implies that $nC_n^2 \leq 409600/c$. Since $C_n^2 \geq 2k$ by assumption, if we set $N > 204800/(kc)$, this would lead to a contradiction. Thus, suppose now that $\delta_4(\sum_{i=1}^n \sum_{j=1}^n y_{ij})\epsilon^2 > 1$, in which case we have

$$\epsilon^2 \geq \frac{cC_n\sqrt{n}}{5120\sqrt{\delta_4(\sum_{i=1}^n \sum_{j=1}^n y_{ij})}}.$$

Thus setting $\kappa_3' \leq c/(5120\sqrt{\delta_4})$ in (17) leads to a contradiction, and hence (19) must fail to hold with probability at least $3/4$, which completes the proof. $\blacksquare$

We now apply Theorem 18 to prove Theorem 4. For any $\epsilon > 0$, Hoeffding's inequality allows us to derive that

$$\begin{aligned}
 & P \left(\sqrt{\frac{n}{\sum_{i=1}^n \sum_{j=1}^n y_{ij}}} \geq \epsilon \sqrt{\frac{1}{np_n}} \right) \\
 &= P \left(\sum_{i=1}^n \sum_{j=1}^n y_{ij} \leq \frac{n^2 p_n}{\epsilon^2} \right) \\
 &= 1 - P \left(\sum_{i=1}^n \sum_{j>i} (n_{ij} - T p_{ij,n}) \geq \frac{n^2 p_n}{\epsilon^2} - T \sum_{i=1}^n \sum_{j>i} p_{ij,n} \right) \\
 &\geq 1 - P \left(\sum_{i=1}^n \sum_{j>i} (n_{ij} - T p_{ij,n}) \geq \frac{n^2 p_n}{\epsilon^2} - \frac{T n(n-1) q_n}{2} \right) \\
 &\geq 1 - \exp \left(\frac{-2[(n^2 p_n / \epsilon^2) - \{T n(n-1) q_n\} / 2]^2}{T^2 n(n-1) / 2} \right) \\
 &= 1 - \exp \left(\frac{-\{(2n^2 p_n / \epsilon^2) - T n(n-1) q_n\}^2}{T^2 n(n-1)} \right).
 \end{aligned}$$

To apply Theorem 18, it suffices to find ϵ such that

$$1 - \exp \left(\frac{-\{(2n^2 p_n / \epsilon^2) - T n(n-1) q_n\}^2}{T^2 n(n-1)} \right) \geq 0.5$$

for sufficiently large n . From Assumption 2, we have $p_n \asymp q_n$ and $q_n \gtrsim \log(n)/n$. Consequently, there exists $\delta_5, \delta_6 > 0$ such that $p_n \geq \delta_5 q_n$ and $q_n \geq \delta_6 \log(n)/n$. Taking $\epsilon = \sqrt{\delta_5 / T}$, we have

$$\begin{aligned}
 P \left(\sqrt{\frac{n}{\sum_{i=1}^n \sum_{j=1}^n y_{ij}}} \geq \sqrt{\frac{\delta_5}{T}} \sqrt{\frac{1}{np_n}} \right) &\geq P \left(\sqrt{\frac{n}{\sum_{i=1}^n \sum_{j=1}^n y_{ij}}} \geq \sqrt{\frac{p_n}{T q_n}} \sqrt{\frac{1}{np_n}} \right) \\
 &\geq 1 - \exp \left(-\frac{(2T n^2 q_n - T n(n-1) q_n)^2}{T^2 n(n-1)} \right) \\
 &= 1 - \exp \left(-\frac{T^2 n^2 q_n^2 (n+1)^2}{T^2 n(n-1)} \right) \\
 &\geq 1 - \exp(-q_n^2 (n+1)^2) \\
 &\geq 1 - \exp(-\delta_6^2 (\log(n))^2) \\
 &\geq 1/2
 \end{aligned}$$

for sufficiently large n . Therefore, the proof of Theorem 4 is complete by setting $\kappa_5 = \kappa_3 \sqrt{\delta_5 / T}$.

B.3 Proof of Theorem 11

We first derive a bound for $\|\tilde{\Theta} - \Theta^*\|_F$.

Lemma 19 *Under Assumptions 9 and 10, for $\hat{M}$ solving (7), with probability converging to 1, we have*

$$\frac{1}{n^2} \|\hat{M} - M^*\|_F^2 \leq \delta_7 C_n \sqrt{\frac{1}{p_n n}}, \quad (21)$$

$$\frac{1}{n} \|\hat{\Theta} - \Theta^*\|_F^2 \leq \delta_8 C_n \sqrt{\frac{1}{p_n n}}. \quad (22)$$

Here $\Theta^* = U^* \hat{O} \Sigma^{*1/2}$, where $\hat{O} = \text{diag}(\hat{O}_1, \dots, \hat{O}_K)$ is a block-diagonal matrix such that $\hat{O}_k$ is orthogonal with $\det(\hat{O}_k) = 1$ for $k \in [K]$.

Proof (21) follows from Theorem 13 and the argument in the proof of Theorem 1 of Davenport et al. (2014), noting that $\|M^*\|_\infty$ is bounded by a positive constant from Assumption 9. Hence, we now focus on showing (22).

Note that as M^* is skew-symmetric, we can write the singular value decomposition of M^* as $M^* = U^* \Sigma^* J_K U^{*\top}$, where U^* is orthogonal and Σ^* is a diagonal matrix containing the singular values. Write $U^* = (\mathbf{u}_1^*, \dots, \mathbf{u}_{2K}^*)$. We can write

$$M^* = \sum_{k=1}^K \sigma_k^* (\mathbf{u}_{2k-1}^* \mathbf{u}_{2k}^{*\top} - \mathbf{u}_{2k}^* \mathbf{u}_{2k-1}^{*\top}).$$

Note that we have $(n^2)^{-1} \|\hat{M} - M^*\|_F^2 \lesssim \delta_7 C_n (p_n n)^{-1/2}$ with probability converging to 1 by (21). Let $\hat{M}_K = \hat{U} \hat{\Sigma} J_K \hat{U}^\top$ denote the best rank- $2K$ approximation of $\hat{M}$ obtained by singular value decomposition. We have

$$\|\hat{M}_K - \hat{M}\|_F \leq \|M^* - \hat{M}\|_F.$$

Therefore, we have

$$\|\hat{M}_K - M^*\|_F \leq \|\hat{M}_K - \hat{M}\|_F + \|\hat{M} - M^*\|_F \leq 2\|M^* - \hat{M}\|_F \quad (23)$$

Using Davis-Kahan theorem, we have

$$\sin(\hat{U}, U^*) \leq 2 \frac{\|\hat{M}_K - M^*\|_F}{\sigma_K^*}.$$

By (21), (23) and the assumption that $\sigma_K \asymp n$, we have

$$\sin(\hat{U}, U^*) \lesssim \sqrt{C_n (p_n n)^{-1/4}}$$

with probability converging to 1. Under this event, we have

$$\|\hat{U} - U^* \hat{O}\|_F \lesssim \sqrt{C_n (p_n n)^{-1/4}},$$

where $\hat{O}$ is an $2K \times 2K$ orthogonal matrix. By Theorem 2 of Yu et al. (2015), we can further show that for $k = 1, \dots, K$,

$$\sin((\hat{\mathbf{u}}_{2k-1}, \hat{\mathbf{u}}_{2k}), (\mathbf{u}_{2k-1}^*, \mathbf{u}_{2k}^*)) \leq 2 \frac{\|\hat{M}_K^\top \hat{M}_K - M^{*\top} M^*\|_F}{\min\{\sigma_{k-1}^{*2} - \sigma_k^{*2}, \sigma_k^{*2} - \sigma_{k+1}^{*2}\}},$$

where $\sigma_{K+1} = 0$ by definition. Define $E = \hat{M}_K - M^*$, we have

$$\|\hat{M}_K^\top \hat{M}_K - M^{*\top} M^*\|_F = \|M^{*\top} E + E^\top M^* + E^\top E\|_F \leq 2\|M^*\|_{op}\|E\|_F + \|E\|_F^2$$

and thus $\|(\hat{\mathbf{u}}_{2k-1}, \hat{\mathbf{u}}_{2k}) - (\mathbf{u}_{2k-1}^*, \mathbf{u}_{2k}^*)\hat{O}_k\|_F \lesssim \sqrt{C_n}(p_n n)^{-1/4}$ with probability converging to 1, where $\hat{O}_{2k}$ is a 2×2 orthogonal matrix. In this event, we have

$$\|\hat{U} - U^* \text{diag}(\hat{O}_1, \dots, \hat{O}_K)\|_F \lesssim \sqrt{C_n}(p_n n)^{-1/4}.$$

We now show that $\det(O_k) = 1$ for $k = 1, \dots, K$. We first show it when $K = 1$. Note that $\det(O_1) = 1$ or -1 as O_1 is orthogonal. If $\det(O_1) = -1$, we have

$$\begin{aligned} U^* \hat{O}_1 \Sigma^* J_1 (U^* \hat{O}_1)^\top &= U^* \hat{O}_1 \Sigma^* J_1 \hat{O}_1^\top U^{*\top} \\ &= \sigma_1^* \det(\hat{O}_1) U^* J_1 U^{*\top} \\ &= \det(\hat{O}_1) U^* \Sigma^* J_1 U^{*\top} \\ &= \det(\hat{O}_1) M^* \\ &= -M^*. \end{aligned}$$

However this implies that we can write $n^{-1}\|M^* + \hat{M}\|_F = n^{-1}\| -M^* - \hat{M}\|_F$ as

$$\begin{aligned} &\frac{1}{n} \|U^* \hat{O}_1 \Sigma^* J_1 (U^* \hat{O}_1)^\top - \hat{U} \hat{\Sigma} J_1 (\hat{U})^\top\|_F \\ &\leq \frac{1}{n} \|U^* \hat{O}_1 \Sigma^* J_1 (U^* \hat{O}_1 - \hat{U})^\top\|_F + \frac{1}{n} \|(U^* \hat{O}_1 \Sigma^* - \hat{U} \hat{\Sigma}) J_1 (\hat{U})^\top\|_F \\ &\leq \frac{1}{n} \|\Sigma^*\|_2 \|U^* \hat{O}_1 - \hat{U}\|_F + \frac{1}{n} \|(U^* \hat{O}_1 - \hat{U}) \Sigma^*\|_F + \frac{1}{n} \|\hat{U} (\Sigma^* - \hat{\Sigma})\|_F \\ &\leq \frac{2\sigma_1^*}{n} \|U^* \hat{O}_1 - \hat{U}\|_F + \frac{1}{n} \|\Sigma^* - \hat{\Sigma}\|_F \\ &\lesssim \sqrt{C_n}(p_n n)^{-1/4}. \end{aligned}$$

This implies that $n^{-1}\|2M^*\|_F \lesssim n^{-1}\|M^* - \hat{M}\|_F + \|M^* + \hat{M}\|_F \lesssim \sqrt{C_n}(p_n n)^{-1/4}$. However, it leads to contradiction as $n^{-1}\|2M^*\|_F \geq n^{-1}\|M^*\|_2 = \sigma_1/n$, which does not vanishes. Hence we are forced to have $\det(O_k) = 1$. The argument for general K is similar. Specifically, suppose there are at least one $k \in [K]$ such that $\det(O_k) = -1$. Let U_{sub}^* be the matrix containing the columns of U^* corresponding to σ_k^* such that $\det(O_k) = -1$, and similarly let Σ_{sub}^* be the diagonal matrix containing those pairs of σ_k^* . Define $M_{sub}^* = U_{sub}^* \Sigma_{sub}^* J_{K_{sub}} (U^*)^\top$, where K_{sub} is the number of k with $\det(O_k) = -1$. We can then show that $n^{-1}\|M_{sub}^* + \hat{M}_{sub}\|_F \lesssim \sqrt{C_n}(p_n n)^{-1/4}$ and $n^{-1}\|M_{sub}^* - \hat{M}_{sub}\|_F \lesssim \sqrt{C_n}(p_n n)^{-1/4}$, which leads to contradiction.

Recall that we defined $\hat{\Theta} = P_{\tau_2}^{2,\infty}(\hat{U}\hat{\Sigma}^{1/2})$ and $\Theta^* = U^*\hat{O}\Sigma^{*1/2}$. Take $\tau_2 = C^*$. Since $\|\Theta^*\|_{2 \rightarrow \infty} \leq C^*$ by assumption, we have

$$\begin{aligned}
& n^{-1/2} \|\hat{\Theta} - \Theta^*\|_F \\
&= n^{-1/2} \|P_{C^*}^{2,\infty}(\hat{U}\hat{\Sigma}^{1/2}) - U^*\hat{O}\Sigma^{*1/2}\|_F \\
&\leq n^{-1/2} \|P_{C^*}^{2,\infty}(\hat{U}\hat{\Sigma}^{1/2}) - \hat{U}\hat{\Sigma}^{1/2}\|_F + \|\hat{U}\hat{\Sigma}^{1/2} - U^*\hat{O}\Sigma^{*1/2}\|_F \\
&\leq 2n^{-1/2} \|\hat{U}\hat{\Sigma}^{1/2} - U^*\hat{O}\Sigma^{*1/2}\|_F \\
&\leq n^{-1/2} \|\hat{U}(\hat{\Sigma}^{1/2} - \Sigma^{*1/2})\|_F + n^{-1} \|(\hat{U} - U^*\hat{O})\Sigma^{*1/2}\|_F \\
&\leq n^{-1/2} \|\hat{\Sigma}^{1/2} - \Sigma^{*1/2}\|_F + n^{-1/2} \sqrt{\sigma_1} \|\hat{U} - U^*\hat{O}\|_F \\
&\lesssim n^{-1/2} \|\hat{M} - M^*\|_F + \|\hat{U} - U^*\hat{O}\|_F \\
&\lesssim \sqrt{C_n} (p_n n)^{-1/4},
\end{aligned}$$

which completes the proof of Lemma 19. $\blacksquare$

Let $\Theta = (\theta_1, \dots, \theta_n)^\top$, and let $\theta_j^{(p)} = J_K \theta_j$ be the permuted θ_j for $j = 1, \dots, n$. We further let Θ_{-i} represent Θ with the i th row removed. We define the log-likelihood associated with the i -th subject and its derivative as

$$\begin{aligned}
l_i(\theta, \Theta_{-i}) &= \sum_{j \in [n], j \neq i} \{y_{ij} \log(g(\theta_i^\top \theta_j^{(p)})) + (n_{ij} - y_{ij}) \log(1 - g(\theta_i^\top \theta_j^{(p)}))\} \text{ and} \\
S_{1,i}(\theta; \Theta_{-i}) &= \frac{\partial}{\partial \theta_i} l_i(\theta, \Theta_{-i}) = \sum_{j \in [n], j \neq i} \theta_j^{(p)} \left\{ y_{ij} - n_{ij} g(\theta_i^\top \theta_j^{(p)}) \right\}.
\end{aligned}$$

Let $\omega_{ij} = I(n_{ij} \geq 1)$ be the missing indicator. Define the functions $\alpha_1(u) = \sup_{|x| \leq u} |g'(u)|$ and $\alpha_2(u) = \sup_{|x| \leq u} |g''(u)|$. The following Lemma provides a non-probabilistic bound for the solution to $S_{1,i}(\theta; \Theta_{-i}) = 0$.

Lemma 20 *Let $\text{diag}(\Omega_i) = \text{diag}(\{\omega_{ij}\}_{j \in [n], j \neq i})$ be a diagonal matrix of ω_{ij} for $j = 1, \dots, n, j \neq i$, and $\mathbf{z}_i = (z_{ij})_{j \in [n], j \neq i}$, where $z_{ij} = y_{ij} - n_{ij} g(m_{ij}^*)$. Under Assumptions 9 and 10, if there exists $\xi > 0$ such that*

$$\begin{aligned}
& 2\sigma_{2K}^{-1}(\mathcal{I}_{1,i}(\Theta_{-i})) \{ \|\mathbf{z}_i \text{diag}(\Omega_i) \Theta_{-i}^{(p)}\| + \|\mathcal{B}_{1,i}(\Theta_{-i})\| + \beta_{1,i}(\Theta_{-i}) \alpha_2((C^* + \xi)C^*) \} \\
& \leq \xi \leq 0.5 \{ \gamma_{1,i}(\Theta_{-i}) \alpha_2((C^* + \xi)C^*) \}^{-1} \sigma_{2K}(\mathcal{I}_{1,i}(\Theta_{-i})),
\end{aligned}$$

where we define

$$\begin{aligned}
\mathcal{B}_{1,i}(\Theta_{-i}) &= \sum_{j \in [n], j \neq i} n_{ij} g'(m_{ij}^*) \theta_j^{(p)} (\theta_j^{(p)} - \theta_j^{*(p)})^\top \theta_i^*, \\
\mathcal{I}_{1,i}(\Theta_{-i}) &= \sum_{j \in [n], j \neq i} n_{ij} g'(m_{ij}^*) \theta_j^{(p)} (\theta_j^{(p)})^\top, \\
\beta_{1,i}(\Theta_{-i}) &= \sup_{\|u\|=1} \sum_{j \in [n], j \neq i} n_{ij} ((\theta_j^{(p)} - \theta_j^{*(p)})^\top \theta_i^*)^2 |\theta_j^{(p)\top} \mathbf{u}|, \text{ and} \\
\gamma_{1,i}(\Theta_{-i}) &= \sup_{\|u\|=1} \sum_{j \in [n], j \neq i} n_{ij} |\theta_j^{(p)\top} \mathbf{u}|^3.
\end{aligned}$$

Then, there is $\check{\boldsymbol{\theta}}_i$ such that $\|\check{\boldsymbol{\theta}}_i - \boldsymbol{\theta}_i^*\| \leq \xi$ and $S_{1,i}(\check{\boldsymbol{\theta}}_i; \Theta_{-i}) = 0$.

Proof Let $\boldsymbol{\theta}$ be a vector such that $\|\boldsymbol{\theta} - \boldsymbol{\theta}_i^*\| \leq \xi$ and let $m_{ij} = \boldsymbol{\theta}^\top \boldsymbol{\theta}_j^{(p)}$. Consider the Taylor expansion of $S_{1,i}(\boldsymbol{\theta}; \Theta_{-i})$:

$$\begin{aligned} S_{1,i}(\boldsymbol{\theta}; \Theta_{-i}) &= \sum_{j \in [n], j \neq i} \omega_{ij} \{y_{ij} - n_{ij}g(m_{ij}^*)\} \boldsymbol{\theta}_j^{(p)} - \sum_{j \in [n], j \neq i} n_{ij} \{g(m_{ij}^*) - g(m_{ij})\} \boldsymbol{\theta}_j^{(p)} \\ &= \Theta_{-i}^{(p)\top} \text{diag}(\Omega_i) \mathbf{z}_i^\top - \sum_{j \in [n], j \neq i} n_{ij} g'(m_{ij}^*) (m_{ij} - m_{ij}^*) \boldsymbol{\theta}_j^{(p)} \\ &\quad - 0.5 \sum_{j \in [n], j \neq i} n_{ij} g''(\tilde{m}_{ij}) (m_{ij} - m_{ij}^*)^2 \boldsymbol{\theta}_j^{(p)}. \end{aligned}$$

for some $\tilde{m}_{ij}$ between m_{ij}^* and m_{ij} . Plugging $m_{ij} - m_{ij}^* = (\boldsymbol{\theta}_j^{(p)})^\top (\boldsymbol{\theta} - \boldsymbol{\theta}_i^*) + (\boldsymbol{\theta}_j^{(p)} - \boldsymbol{\theta}_j^{*(p)})^\top \boldsymbol{\theta}_i^*$ into the above display, we obtain

$$\begin{aligned} S_{1,i}(\boldsymbol{\theta}; \Theta_{-i}) &= \Theta_{-i}^{(p)\top} \text{diag}(\Omega_i) \mathbf{z}_i^\top - \sum_{j \in [n], j \neq i} n_{ij} g'(m_{ij}^*) \boldsymbol{\theta}_j^{(p)} (\boldsymbol{\theta}_j^{(p)})^\top (\boldsymbol{\theta} - \boldsymbol{\theta}_i^*) \\ &\quad - \sum_{j \in [n], j \neq i} n_{ij} g'(m_{ij}^*) \boldsymbol{\theta}_j^{(p)} (\boldsymbol{\theta}_j^{(p)} - \boldsymbol{\theta}_j^{*(p)})^\top \boldsymbol{\theta}_i^* \\ &\quad - 0.5 \sum_{j \in [n], j \neq i} n_{ij} g''(\tilde{m}_{ij}) (m_{ij} - m_{ij}^*)^2 \boldsymbol{\theta}_j^{(p)}. \end{aligned}$$

Multiplying $(\boldsymbol{\theta} - \boldsymbol{\theta}_i^*)^\top$ on both sides, we obtain

$$\begin{aligned} &(\boldsymbol{\theta} - \boldsymbol{\theta}_i^*)^\top S_{1,i}(\boldsymbol{\theta}; \Theta_{-i}) \\ &= (\boldsymbol{\theta} - \boldsymbol{\theta}_i^*)^\top \Theta_{-i}^{(p)\top} \text{diag}(\Omega_i) \mathbf{z}_i^\top - (\boldsymbol{\theta} - \boldsymbol{\theta}_i^*)^\top \sum_{j \in [n], j \neq i} n_{ij} g'(m_{ij}^*) \boldsymbol{\theta}_j^{(p)} (\boldsymbol{\theta}_j^{(p)})^\top (\boldsymbol{\theta} - \boldsymbol{\theta}_i^*) \\ &\quad - (\boldsymbol{\theta} - \boldsymbol{\theta}_i^*)^\top \sum_{j \in [n], j \neq i} n_{ij} g'(m_{ij}^*) \boldsymbol{\theta}_j^{(p)} (\boldsymbol{\theta}_j^{(p)} - \boldsymbol{\theta}_j^{*(p)})^\top \boldsymbol{\theta}_i^* \\ &\quad - 0.5 (\boldsymbol{\theta} - \boldsymbol{\theta}_i^*)^\top \sum_{j \in [n], j \neq i} n_{ij} g''(\tilde{m}_{ij}) (m_{ij} - m_{ij}^*)^2 \boldsymbol{\theta}_j^{(p)}. \end{aligned} \tag{24}$$

Recall that $\|\boldsymbol{\theta} - \boldsymbol{\theta}_i^*\| = \xi$. Using inequalities about matrix products and singular values, we have the following upper bounds for the first three terms on the right-hand side of the above display/

$$|(\boldsymbol{\theta} - \boldsymbol{\theta}_i^*)^\top \Theta_{-i}^{(p)\top} \text{diag}(\Omega_i) \mathbf{z}_i^\top| \leq |\xi| \|\Theta_{-i}^{(p)\top} \text{diag}(\Omega_i) \mathbf{z}_i^\top\| = \xi \|\mathbf{z}_i \text{diag}(\Omega_i) \Theta_{-i}^{(p)}\|, \tag{25}$$

$$-(\boldsymbol{\theta} - \boldsymbol{\theta}_i^*)^\top \sum_{j \in [n], j \neq i} n_{ij} g'(m_{ij}^*) \boldsymbol{\theta}_j^{(p)} (\boldsymbol{\theta}_j^{(p)} - \boldsymbol{\theta}_j^{*(p)})^\top \boldsymbol{\theta}_i^* \leq -\xi^2 \sigma_{2K}(\mathcal{I}_{1,i}(\Theta_{-i})), \tag{26}$$

where $\sigma_{2K}(\mathcal{I}_{1,i}(\Theta_{-i}))$ denotes the $2K$ th largest singular value of $\mathcal{I}_{1,i}(\Theta_{-i})$, and

$$|(\boldsymbol{\theta} - \boldsymbol{\theta}_i^*)^\top \sum_{j \in [n], j \neq i} n_{ij} g'(m_{ij}^*) \boldsymbol{\theta}_j^{(p)} (\boldsymbol{\theta}_j^{(p)} - \boldsymbol{\theta}_j^{*(p)})^\top \boldsymbol{\theta}_i^*| = \|(\boldsymbol{\theta} - \boldsymbol{\theta}_i^*)^\top \mathcal{B}_{1,i}(\Theta_{-i})\| \leq \xi \|\mathcal{B}_{1,i}(\Theta_{-i})\|. \tag{27}$$

Now we analyze the last term in (24). Note that $|\tilde{m}_{ij}| \leq |m_{ij}^*| \vee |m_{ij}| \leq (C^* + \xi)C^*$ and $m_{ij} - m_{ij}^* = (\boldsymbol{\theta}_j^{(p)})^\top (\boldsymbol{\theta} - \boldsymbol{\theta}_i^*) + (\boldsymbol{\theta}_j^{(p)} - \boldsymbol{\theta}_j^{*(p)})^\top \boldsymbol{\theta}_i^*$. We have

$$\begin{aligned}
 & 0.5(\boldsymbol{\theta} - \boldsymbol{\theta}_i^*)^\top \sum_{j \in [n], j \neq i} n_{ij} g''(\tilde{m}_{ij})(m_{ij} - m_{ij}^*)^2 \boldsymbol{\theta}_j^{(p)} \\
 & \leq 0.5\alpha_2 ((C^* + \xi)C^*) \xi \sup_{\|u\|=1} \sum_{j \in [n], j \neq i} n_{ij} ((\boldsymbol{\theta}_j^{(p)} - \boldsymbol{\theta}_j^{*(p)})^\top \boldsymbol{\theta}_i^* + \xi(\boldsymbol{\theta}_j^{(p)})^\top u)^2 |\boldsymbol{\theta}_j^{(p)}^\top \mathbf{u}| \\
 & \leq \alpha_2 ((C^* + \xi)C^*) \left\{ \xi \sup_{\|u\|=1} \sum_{j \in [n], j \neq i} n_{ij} ((\boldsymbol{\theta}_j^{(p)} - \boldsymbol{\theta}_j^{*(p)})^\top \boldsymbol{\theta}_i^*)^2 |\boldsymbol{\theta}_j^{(p)}^\top \mathbf{u}| \right. \\
 & \quad \left. + \xi^3 \sup_{\|u\|=1} \sum_{j \in [n], j \neq i} n_{ij} |\boldsymbol{\theta}_j^{(p)}^\top \mathbf{u}|^3 \right\} \\
 & = \alpha_2 ((C^* + \xi)C^*) (\xi\beta_{1,i}(\Theta_{-i}) + \xi^3\gamma_{1,i}(\Theta_{-i})).
 \end{aligned}$$

Combining the analysis with (25) and (27), we obtain

$$\begin{aligned}
 (\boldsymbol{\theta} - \boldsymbol{\theta}_i^*)^\top S_{1,i}(\boldsymbol{\theta}; \Theta_{-i}) & \leq -\sigma_{2K}(\mathcal{I}_{1,i}(\Theta_{-i}))\xi^2 + \gamma_{1,i}(\Theta_{-i})\alpha_2 ((C^* + \xi)C^*)\xi^3 \\
 & \quad + \{\|\mathbf{z}_i \text{diag}(\Omega_i)\Theta_{-i}^{(p)}\| + \|\mathcal{B}_{1,i}(\Theta_{-i})\| + \beta_{1,i}(\Theta_{-i})\alpha_2 ((C^* + \xi)C^*)\}\xi.
 \end{aligned}$$

The rest of the proof follows from the argument in the proof of Lemma 15 in Chen and Li (2024). $\blacksquare$

We now simplify the result of Lemma 15 to a more user-friendly version.

Lemma 21 *Under Assumptions 9 and 10, if*

$$\begin{aligned}
 & \|\mathbf{z}_i \text{diag}(\Omega_i)\Theta_{-i}^{(p)}\| + \|\mathcal{B}_{1,i}(\Theta_{-i})\| + \beta_{1,i}(\Theta_{-i})\alpha_2 (3C^{*3}) \\
 & \leq \min\{0.25\gamma_{1,i}(\Theta_{-i})^{-1}(\alpha_2 (3C^{*3}))^{-1}\sigma_{2K}^2(\mathcal{I}_{1,i}(\Theta_{-i})), 0.5\sigma_{2K}(\mathcal{I}_{1,i}(\Theta_{-i}))C^*\}.
 \end{aligned}$$

Then, there is $\check{\boldsymbol{\theta}}_i$ such that $S_{1,i}(\check{\boldsymbol{\theta}}_i; \Theta_{-i}) = 0$, and

$$\|\check{\boldsymbol{\theta}}_i - \boldsymbol{\theta}_i^*\| \leq 2\sigma_{2K}^{-1}(\mathcal{I}_{1,i}(\Theta_{-i}))\{\|\mathbf{z}_i \text{diag}(\Omega_i)\Theta_{-i}^{(p)}\| + \|\mathcal{B}_{1,i}(\Theta_{-i})\| + \beta_{1,i}(\Theta_{-i})\alpha_2 (3C^{*3})\}.$$

The proof follows from the argument in Lemma 16 of Chen and Li (2024). We omit the details. We now analyze each term in Lemma 21 with $\Theta_{-i} = \hat{\Theta}_{-i}$. Let $p_{\max} = \max_{i \in [n]} \sum_{j \in [n], j \neq i} n_{ij}$ denotes the maximum number of comparisons in each row.

Lemma 22 *(Upper bound for $\|\mathbf{z}_i \text{diag}(\Omega_i)\hat{\Theta}_{-i}^{(p)}\|$). Under Assumptions 9 and 10, with probability at least $1 - (nK)^{-1}$,*

$$\max_{i \in [n]} \|\mathbf{z}_i \text{diag}(\Omega_i)\hat{\Theta}_{-i}^{(p)}\| \leq 16C^*\sqrt{2K} \log(nK)p_{\max}^{1/2} + \sqrt{T}p_{\max}^{1/2}\|\hat{\Theta}^{(p)} - \Theta^{*(p)}\|_F.$$

Proof Note that

$$\|\mathbf{z}_i \text{diag}(\Omega_i) \hat{\Theta}_{-i}^{(p)}\| \leq \|\mathbf{z}_i \text{diag}(\Omega_i) \Theta_{-i}^{*(p)}\| + \|\mathbf{z}_i \text{diag}(\Omega_i) (\hat{\Theta}_{-i}^{(p)} - \Theta_{-i}^{*(p)})\|. \quad (28)$$

It is easy to see that $|y_{ij} - n_{ij}g(m_{ij}^*)| \leq n_{ij}$ for all i, j . Hence we have

$$\begin{aligned} \|\mathbf{z}_i \text{diag}(\Omega_i) (\hat{\Theta}_{-i}^{(p)} - \Theta_{-i}^{*(p)})\| &= \left\| \sum_{j \in [n], j \neq i} (y_{ij} - n_{ij}g(m_{ij}^*)) (\hat{\theta}_j - \theta_j^*) \right\| \\ &\leq \sum_{j \in [n], j \neq i} n_{ij} \|\hat{\theta}_j - \theta_j^*\| \\ &\leq \sqrt{\sum_{j \in [n], j \neq i} n_{ij}^2} \|\hat{\Theta}_{-i}^{(p)} - \Theta_{-i}^{*(p)}\|_F \\ &\leq \sqrt{T} p_{\max}^{1/2} \|\hat{\Theta}_{-i}^{(p)} - \Theta_{-i}^{*(p)}\|_F \\ &\leq \sqrt{T} p_{\max}^{1/2} \|\hat{\Theta}^{(p)} - \Theta^{*(p)}\|_F. \end{aligned} \quad (29)$$

Combining (28) and (29) and taking maximum over $i \in [n]$, we have

$$\max_{i \in [n]} \|\mathbf{z}_i \text{diag}(\Omega_i) \hat{\Theta}_{-i}^{(p)}\| \leq \max_{i \in [n]} \|\mathbf{z}_i \text{diag}(\Omega_i) \Theta_{-i}^{*(p)}\| + \sqrt{T} p_{\max}^{1/2} \|\hat{\Theta}^{(p)} - \Theta^{*(p)}\|_F.$$

For the first term, we can show that for each $k \in \{1, \dots, 2K\}$ and positive t , since $y_{ij} - n_{ij}g(m_{ij}^*)$ is mean zero given n_{ij} , we can apply Bernstein inequality and get

$$\begin{aligned} &P\left(\left| \sum_{j \in [n], j \neq i} (y_{ij} - n_{ij}g(m_{ij}^*)) \theta_{jk}^{*(p)} \right| \geq t \mid \{n_{ij}\}_{j \in [n], j \neq i}\right) \\ &\leq \exp\left(-\frac{0.5t^2}{C^*(\sum_{j \in [n], j \neq i} n_{ij}g(m_{ij}^*)(1 - g(m_{ij}^*)) + \frac{1}{3}Tt)}\right) \\ &\leq \exp\left(-\frac{0.5t^2}{C^*(\sum_{j \in [n], j \neq i} n_{ij} + Tt)}\right). \end{aligned}$$

This implies that

$$\begin{aligned} &P(\|\mathbf{z}_i \text{diag}(\Omega_i) \Theta_{-i}^{*(p)}\| \geq t \mid \{n_{ij}\}_{j \in [n], j \neq i}) \\ &\leq \sum_{k=1}^{2K} P(|\sum_{j \in [n], j \neq i} (y_{ij} - n_{ij}g(m_{ij}^*)) \theta_{jk}^{*(p)}| \geq t/\sqrt{2K} \mid \{n_{ij}\}_{j \in [n], j \neq i}) \\ &\leq 2K \exp\left(-\frac{0.25t^2}{KC^*(\sum_{j \in [n], j \neq i} n_{ij} + Tt/\sqrt{2K})}\right). \end{aligned}$$

Combining results for different i with a union bound, we have

$$P(\max_{i \in [n]} \|\mathbf{z}_i \text{diag}(\Omega_i) \Theta_{-i}^{*(p)}\| \geq t \mid \{n_{ij}\}_{j \in [n], j \neq i}) \leq 2Kn \exp\left(-\frac{0.25t^2}{KC^*(p_{\max} + Tt/\sqrt{2K})}\right).$$

In particular, for $t = 16C^*\sqrt{2K}\log(nK)p_{\max}^{1/2}$, we can show that $\max_{i \in [n]} \|\mathbf{z}_i \text{diag}(\Omega_i) \Theta_{-i}^{*(p)}\| \leq 16C^*\sqrt{2K}\log(nK)p_{\max}^{1/2}$ with probability at least $1 - (nK)^{-1}$, and the proof is complete. ■

Lemma 23 (*Upper bound for $\|\mathcal{B}_{1,i}(\hat{\Theta}_{-i})\|$*) Under Assumptions 9 and 10,

$$\max_{i \in [n]} \|\mathcal{B}_{1,i}(\hat{\Theta}_{-i})\| \leq C^{*2} \sqrt{T} p_{\max}^{1/2} \|\hat{\Theta} - \Theta^*\|_F.$$

Proof Recall that $|g'(x)| = |g(x)(1 - g(x))| \leq 1$. Hence we have

$$\begin{aligned} \|\mathcal{B}_{1,i}(\hat{\Theta}_{-i})\| &= \left\| \sum_{j \in [n], j \neq i} n_{ij} g'(m_{ij}^*) \hat{\boldsymbol{\theta}}_j^{(p)} (\hat{\boldsymbol{\theta}}_j^{(p)} - \boldsymbol{\theta}_j^{*(p)})^\top \boldsymbol{\theta}_i^* \right\| \\ &\leq C^{*2} \sum_{j \in [n], j \neq i} n_{ij} \|\hat{\boldsymbol{\theta}}_j^{(p)} - \boldsymbol{\theta}_j^{*(p)}\| \\ &\leq C^{*2} \sqrt{T} p_{\max}^{1/2} \|\hat{\Theta}_{-i} - \Theta_{-i}^*\| \\ &\leq C^{*2} \sqrt{T} p_{\max}^{1/2} \|\hat{\Theta} - \Theta^*\|. \end{aligned}$$

The proof is completed by taking maximum for $i \in [n]$. ■

Lemma 24 (*Bound for $\beta_{1,i}(\hat{\Theta}_{-i})$*) Under Assumptions 9 and 10,

$$\max_{i \in [n]} \beta_{1,i}(\hat{\Theta}_{-i}) \leq C^{*3} T \|\hat{\Theta} - \Theta^*\|_F^2.$$

Proof For each $i \in [n]$

$$\begin{aligned} \beta_{1,i}(\hat{\Theta}_{-i}) &= \sup_{\|u\|=1} \sum_{j \in [n], j \neq i} n_{ij} ((\hat{\boldsymbol{\theta}}_j^{(p)} - \boldsymbol{\theta}_j^{*(p)})^\top \boldsymbol{\theta}_i^*)^2 |(\hat{\boldsymbol{\theta}}_j^{(p)})^\top \mathbf{u}| \\ &\leq C^{*3} \sup_{\|u\|=1} \sum_{j \in [n], j \neq i} n_{ij} \|\hat{\boldsymbol{\theta}}_j^{(p)} - \boldsymbol{\theta}_j^{*(p)}\|^2 \leq C^{*3} T \|\hat{\Theta} - \Theta^*\|_F^2 \end{aligned}$$

■

Lemma 25 (*Upper bound for $p_{\max}$*). Under Assumptions 9 and 10, if $nq_n \geq 6\log(n)$, then

$$P(p_{\max} \geq 2Tnq_n) \leq 1/n.$$

Proof It follows from Lemma 23 in Chen and Li (2024) that $P(\max_{i \in [n]} \sum_{j \in [n], j \neq i} \omega_{ij} \geq 2pq_n) \leq 1/n$, which implies $P(\max_{i \in [n]} \sum_{j \in [n], j \neq i} T\omega_{ij} \geq 2Tpq_n) \leq 1/n$. The proof then complete as $P(p_{\max} \geq 2Tpq_n) \leq P(\max_{i \in [n]} \sum_{j \in [n], j \neq i} T\omega_{ij} \geq 2Tpq_n)$. ■

Lemma 26 (Bound for $\gamma_{1,i}(\hat{\Theta}_{-i})$). Under Assumptions 9 and 10, if $nq_n \geq 6 \log(n)$, then with probability at least $1 - 1/n$,

$$\gamma_{1,i}(\hat{\Theta}_{-i}) \leq 2Tnq_n C^{*3}.$$

Proof The lemma follow by Lemma 25 and the following inequality

$$\gamma_{1,i}(\Theta_{-i}) = \sup_{\|u\|=1} \sum_{j \in [n], j \neq i} n_{ij} |\theta_j^{(p)\top} \mathbf{u}|^3 \leq C^{*3} p_{\max}$$

■

The next two lemmas give a lower bound for $\sigma_{2K}(\mathcal{I}_{1,i}(\hat{\Theta}_{-i}))$.

Lemma 27 Define $\mu(y) = \inf_{|x| \leq y} g'(x)$. Under Assumptions 9 and 10, if $\| \text{diag}(\Omega_i)(\hat{\Theta}_{-i}^{(p)} - \Theta_{-i}^{*(p)}) \|_2 \leq 2^{-1} \sigma_{2K}(\text{diag}(\Omega_i) \Theta_{-i}^{*(p)})$, then

$$\sigma_{2K}(\mathcal{I}_{1,i}(\hat{\Theta}_{-i})) \geq 2^{-2} \mu(C^{*2}) \sigma_{2K}^2(\text{diag}(\Omega_i) \Theta_{-i}^{*(p)}).$$

Proof Note that for any $\mathbf{u} \in \mathbb{R}^{2K}$ such that $\|\mathbf{u}\| = 1$,

$$\begin{aligned} \mathcal{I}_{1,i}(\hat{\Theta}_{-i}) &= \sum_{j \in [n], j \neq i} n_{ij} g'(m_{ij}^*) (\mathbf{u}^\top \hat{\theta}_j^{(p)})^2 \\ &\geq \mu(C^{*2}) \sum_{j \in [n]} \omega_{ij} (\mathbf{u}^\top \hat{\theta}_j^{(p)})^2 \\ &\geq \mu(C^{*2}) \sigma_{2K}^2(\text{diag}(\Omega_i) \hat{\Theta}_{-i}^{(p)}). \end{aligned}$$

The rest of the proof follows from the argument in Lemma 25 of Chen and Li (2024). ■

Lemma 28 Let $P_{1,i} = \text{diag}(\{p_{ij}\}_{j \in [n], j \neq i}) = E(\text{diag}(\Omega_i))$ and $\lambda_{i,\min}^* = \lambda_{2K}((\Theta_{-i}^*)^\top P_{1,i} \Theta_{-i}^*)$, where $\lambda_{2K}(\cdot)$ denotes the $2K$ th largest eigenvalue of a symmetric matrix. If $\lambda_{\min}^* := \min_{i \in [n]} \lambda_{i,\min}^* \geq 16 \|\Theta^*\|_{2 \rightarrow \infty}^2 \log(2Kn)$, then

$$P(\min_{i \in [n]} \sigma_{2K}^2(\text{diag}(\Omega_i) \Theta_{-i}^{*(p)}) \leq 0.5 \lambda_{\min}^*) \leq 1/(2nK).$$

Moreover, if $p_n \sigma_{2K}^2(\Theta^*) \geq 32 \|\Theta^*\|_{2 \rightarrow \infty}^2 \log(n)$, then

$$P(\min_{i \in [n]} \sigma_{2K}^2(\text{diag}(\Omega_i) \Theta_{-i}^{*(p)}) \leq 0.5 p_n \sigma_{2K}^2(\Theta^*)) \leq 1/(2nK).$$

The proof follows from the argument in the proof of Lemma 26 in Chen and Li (2024), noting that $\text{diag}(\Omega_i) \Theta_{-i}^{*(p)} = \text{diag}(\bar{\Omega}_i) \Theta_{-i}^{*(p)}$, where $\bar{\Omega}_i = \text{diag}(\omega_{i1}, \dots, \omega_{in})$. We omit the details. We now turn to asymptotic analysis of $\tilde{\Theta}$.

Lemma 29 *Under Assumptions 9 and 10, with probability converging to 1, there is $\tilde{\Theta} = (\tilde{\theta}_{ik})_{n \times 2K}$ such that $S_{1,i}(\tilde{\theta}_i, \hat{\Theta}_{-i}) = 0$ for all $i \in [n]$, $\|\tilde{\Theta} - \Theta^*\|_{2 \rightarrow \infty} \leq C^*$, and*

$$\|\tilde{\Theta} - \Theta^*\|_{2 \rightarrow \infty} \lesssim n^{1/2}(nq_n)^{1/4}(np_n)^{-1} = n^{-1/4}q_n^{-3/4}.$$

Moreover, $\tilde{\theta}_i$ is the unique solution to the optimization problem $\max_{\theta_i \in \mathbb{R}^{2K}} l_i(\theta_i, \hat{\Theta}_{-i})$ for all $i \in [n]$.

Proof Throughout the proof, we restrict the analysis on the event $\{p_{\max} \leq 2Tnq_n\}$, which has probability converging to 1 by Lemma 25. On this event, we have that with probability at least $1 - 1/(nK)^{-1}$,

$$\max_{i \in [n]} \|\mathbf{z}_i \text{diag}(\Omega_i) \hat{\Theta}_{-i}^{(p)}\| \leq 32C^*K \log(nK)(Tnq_n)^{1/2} + \sqrt{2T}(Tnq_n)^{1/2} \|\hat{\Theta} - \Theta^*\|_F \quad (30)$$

according to Lemma 22 and the fact that $\|\hat{\Theta}^{(p)} - \Theta^{*(p)}\|_F = \|\hat{\Theta} - \Theta^*\|_F$. Next, according to Lemma 23, we have

$$\max_{i \in [n]} \|\mathcal{B}_{1,i}(\hat{\Theta}_{-i})\| \leq C^{*2} \sqrt{T} p_{\max}^{1/2} \|\hat{\Theta} - \Theta^*\|_F$$

Thus with probability converging to one, we have

$$\max_{i \in [n]} \|\mathcal{B}_{1,i}(\hat{\Theta}_{-i})\| \leq C^{*2} \sqrt{T} \sqrt{2Tnq_n} \|\hat{\Theta} - \Theta^*\|_F. \quad (31)$$

Next, according to Lemma 24, we have

$$\max_{i \in [n]} \beta_{1,i}(\hat{\Theta}_{-i}) \leq C^{*3} T \|\hat{\Theta} - \Theta^*\|_F^2. \quad (32)$$

Combining (22), (30), (31) and (32), we have

$$\begin{aligned} & \max_{i \in [n]} \{ \|\mathbf{z}_i \text{diag}(\Omega_i) \Theta_{-i}^{(p)}\| + \|\mathcal{B}_{1,i}(\Theta_{-i})\| + \beta_{1,i}(\Theta_{-i}) \alpha_2(3C^{*3}) \} \\ & \leq 32C^*K \log(nK)(Tnq_n)^{1/2} + \sqrt{2T}(Tnq_n)^{1/2} \|\hat{\Theta} - \Theta^*\|_F + C^{*2} \sqrt{T} \sqrt{2Tnq_n} \|\hat{\Theta} - \Theta^*\|_F \\ & \quad + C^{*3} T \|\hat{\Theta} - \Theta^*\|_F^2 \\ & \lesssim K \log(nK)(Tnq_n)^{1/2} + \sqrt{T}(Tnq_n)^{1/2} \sqrt{nC_n}(p_n n)^{-1/4} + TC_n n(p_n n)^{-1/2} \\ & \asymp (nq_n)^{1/2} + n^{1/2}(nq_n)^{1/4} + n^{1/2}q_n^{-1/2} \\ & \leq n^{1/2}(nq_n)^{1/4} \\ & \leq nq_n \end{aligned} \quad (33)$$

for n large enough as $q_n \gtrsim n^{-1/3} \log(n)$ by Assumption 10. Next, we find a lower bound for $\sigma_{2K}^2(\mathcal{I}_{1,i}(\hat{\Theta}_{-i}))$. Note that $\sigma_{2K}^2(\Theta^*) \asymp n$ and $\|\Theta^*\|_{2 \rightarrow \infty} \leq C^*$ by Assumption 9. Hence we have $p_n \sigma_{2K}^2(\Theta^*) \geq 32\|\Theta^*\|_{2 \rightarrow \infty} \log(2Kn)$ for n large enough. According to Lemma 28, with probability at least $1 - 1/(2nK)$,

$$\min_{i \in [n]} \sigma_{2K}^2(\text{diag}(\Omega_i) \Theta_{-i}^{*(p)}) \geq 0.5 p_n \sigma_{2K}^2(\Theta^*).$$

Also, we can verify that

$$\begin{aligned}
 \max_{i \in [n]} \|\text{diag}(\Omega_i)(\hat{\Theta}_{-i}^{(p)} - \Theta_{-i}^{*(p)})\|_2^2 &\leq \|\hat{\Theta}^{(p)} - \Theta^{*(p)}\|_F^2 \\
 &\lesssim nC_n(p_n n)^{-1/2} \\
 &\leq 0.5 \min_{i \in [n]} \sigma_{2K}^2(\text{diag}(\Omega_i)\Theta_{-i}^{*(p)})
 \end{aligned}$$

under the assumption that $p_n \gtrsim n^{-1/3} \log(n)$. Therefore, by Lemma 27, with probability converging to 1, we have

$$\min_{i \in [n]} \sigma_{2K}(\mathcal{I}_{1,i}(\hat{\Theta}_{-i})) \geq \delta_9 p_n n, \quad (34)$$

for some constant $\delta_9 > 0$. Next, we verify conditions of Lemma 21. According to Lemma 26, we have

$$\max_{i \in [n]} \gamma_{1,i}(\hat{\Theta}_{-i}) \leq 2Tnq_n C^{*3} \lesssim nq_n$$

for n large enough. Therefore, with probability tending to 1, we have

$$\min_{i \in [n]} \gamma_{1,i}(\Theta_{-i})^{-1} (\alpha_2 (3C^{*3}))^{-1} \sigma_{2K}^2(\mathcal{I}_{1,i}(\Theta_{-i})) \gtrsim (nq_n)^{-1} (p_n n)^2 \asymp nq_n. \quad (35)$$

Therefore, by (33) and (35), we have

$$\begin{aligned}
 &\max_{i \in [n]} \{\|\mathbf{z}_i \text{diag}(\Omega_i) \hat{\Theta}_{-i}^{(p)}\| + \|\mathcal{B}_{1,i}(\hat{\Theta}_{-i})\| + \beta_{1,i}(\hat{\Theta}_{-i}) \alpha_2 (3C^{*3})\} \\
 &\leq \min_{i \in [n]} 0.25 \gamma_{1,i}(\hat{\Theta}_{-i})^{-1} (\alpha_2 (3C^{*3}))^{-1} \sigma_{2K}^2(\mathcal{I}_{1,i}(\hat{\Theta}_{-i}))
 \end{aligned} \quad (36)$$

with probability converging to 1. Similarly, according to 34, we have

$$\begin{aligned}
 &\max_{i \in [n]} \{\|\mathbf{z}_i \text{diag}(\Omega_i) \hat{\Theta}_{-i}^{(p)}\| + \|\mathcal{B}_{1,i}(\hat{\Theta}_{-i})\| + \beta_{1,i}(\hat{\Theta}_{-i}) \alpha_2 (3C^{*3})\} \\
 &\leq 0.5 \sigma_{2K}(\mathcal{I}_{1,i}(\hat{\Theta}_{-i})) C^*
 \end{aligned} \quad (37)$$

with probability converging to 1. Thus, conditions of Lemma 21 are satisfied. By (33) and (34) we have $\|\tilde{\Theta} - \Theta^*\|_{2 \rightarrow \infty} \leq C^*$ and

$$\|\tilde{\Theta} - \Theta^*\|_{2 \rightarrow \infty} \lesssim n^{1/2} (nq_n)^{1/4} (np_n)^{-1} = n^{-1/4} p_n^{-3/4}.$$

Moreover, from (34), the optimization problem $\max_{\boldsymbol{\theta}_i \in \mathbb{R}^{2K}} \sum_{j \in [n], j \neq i} \{y_{ij} \log(g(\boldsymbol{\theta}_i^\top \hat{\boldsymbol{\theta}}_j^{(p)})) + (n_{ij} - y_{ij}) \log(1 - g(\boldsymbol{\theta}_i^\top \hat{\boldsymbol{\theta}}_j^{(p)}))\}$ is strictly convex. Thus, $\tilde{\boldsymbol{\theta}}_i$ is the unique solution to this optimization problem, and the proof is complete. $\blacksquare$

With the above lemma, we have shown that (8) in Theorem 11 holds. It only remains to prove that (9) holds. Note that $\tilde{M} - M^* = \tilde{\Theta} J_K \tilde{\Theta}^\top - \Theta^* J_K \Theta^{*\top} = (\tilde{\Theta} - \Theta^*) J_K \tilde{\Theta}^\top + \Theta^* J_K (\tilde{\Theta} - \Theta^*)^\top$. Hence,

$$\|\tilde{M} - M^*\|_\infty \leq 2\|\tilde{\Theta} - \Theta^*\|_{2 \rightarrow \infty} \|J_K \tilde{\Theta}^\top\|_{2 \rightarrow \infty}.$$

Therefore, by (8) and $\|J_K \tilde{\Theta}^\top\|_{2 \rightarrow \infty} \leq C^*$, we have

$$\|\tilde{M} - M^*\|_\infty \lesssim n^{-1/4} p_n^{-3/4}$$

with probability converging to 1, and the proof of Theorem 11 is complete.

B.4 Proof of Theorem 12

Proof For ease of presentation, assume $\mathcal{S} = [n]$ and that the players are ordered such that $r_1(\Pi^*) \geq r_2(\Pi^*) \geq \dots \geq r_n(\Pi^*)$ without loss of generality. We write Δ_k for $\Delta_{k,\mathcal{S}}$ and $r_i(\Pi)$ for $r_{i,\mathcal{S}}(\Pi)$. It suffices to show that

$$\min_{i \in [k]} r_i(\tilde{\Pi}) > \max_{l \in \{k+1, \dots, n\}} r_l(\tilde{\Pi})$$

with probability converging to 1. For any $i \in [n]$, we have

$$|r_i(\tilde{\Pi}) - r_i(\Pi^*)| = \left| \frac{1}{n} \sum_{j=1}^n (\tilde{\Pi}_{ij} - \Pi_{ij}^*) \right| \leq \|\tilde{\Pi} - \Pi^*\|_{\max}.$$

Hence

$$\max_{i \in [n]} |r_i(\tilde{\Pi}) - r_i(\Pi^*)| \leq \|\tilde{\Pi} - \Pi^*\|_{\max}.$$

Consequently, for $i \in [k]$ and $l \in \{k+1, \dots, n\}$,

$$r_i(\tilde{\Pi}) \geq r_i(\Pi^*) - \|\tilde{\Pi} - \Pi^*\|_{\max}, \quad r_l(\tilde{\Pi}) \leq r_l(\Pi^*) + \|\tilde{\Pi} - \Pi^*\|_{\max}.$$

Recall that by (10), we have

$$\|\tilde{\Pi} - \Pi^*\|_\infty \leq \kappa_8 n^{-1/4} p_n^{-3/4}$$

with probability converging to 1. On this event, since $\Delta_k > 2\kappa_8 n^{-1/4} p_n^{-3/4}$, by assumption, we have then $\Delta_k > 2\|\tilde{\Pi} - \Pi^*\|_{\max}$.

Using $r_k(\Pi^*) = r_{k+1}(\Pi^*) + \Delta_k$, we obtain

$$\begin{aligned} \min_{i \in [k]} r_i(\tilde{\Pi}) &\geq r_k(\Pi^*) - \|\tilde{\Pi} - \Pi^*\|_{\max} \\ &= r_{k+1}(\Pi^*) + \Delta_k - \|\tilde{\Pi} - \Pi^*\|_{\max} \\ &> r_{k+1}(\Pi^*) + \|\tilde{\Pi} - \Pi^*\|_{\max} \\ &= \max_{l \in \{k+1, \dots, n\}} (r_l(\Pi^*) + \|\tilde{\Pi} - \Pi^*\|_{\max}) \\ &\geq \max_{l \in \{k+1, \dots, n\}} r_l(\tilde{\Pi}). \end{aligned}$$

Therefore, the top- k set based on $\tilde{\Pi}$ coincides with that based on Π^* with probability converging to 1, which completes the proof. $\blacksquare$

B.5 Proof of Proposition 8

Proof We consider the case where n is even; the proof for odd n is analogous. It is well known that M can be decomposed in the Murnaghan canonical form $M = QXQ^\top$ (Murnaghan and Wintner, 1931; Benner et al., 2000), where Q is orthogonal and X is block-diagonal of the form

$$X = \begin{pmatrix} 0 & \sigma_1 & 0 & 0 & \dots & 0 & 0 \\ -\sigma_1 & 0 & 0 & 0 & \dots & 0 & 0 \\ 0 & 0 & 0 & \sigma_2 & \dots & 0 & 0 \\ 0 & 0 & -\sigma_2 & 0 & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & 0 & 0 & \dots & 0 & \sigma_{n/2} \\ 0 & 0 & 0 & 0 & \dots & -\sigma_{n/2} & 0 \end{pmatrix},$$

where $\sigma_1, \dots, \sigma_{n/2}$ are the singular values of M . It can be verified that the projection operator $P_\tau(M)$ preserves the Murnaghan canonical form as $\mathcal{P}_\tau(M) = QYQ^\top$, where

$$Y = \begin{pmatrix} 0 & \max\{\sigma_1 - \lambda, 0\} & 0 & \dots & 0 \\ -\max\{\sigma_1 - \lambda, 0\} & 0 & 0 & \dots & 0 \\ 0 & 0 & \ddots & \dots & 0 \\ \vdots & \vdots & \vdots & 0 & \max\{\sigma_{n/2} - \lambda, 0\} \\ 0 & 0 & 0 & -\max\{\sigma_{n/2} - \lambda, 0\} & 0 \end{pmatrix}.$$

Hence we have $P_\tau(M) \in \text{Skew}_n$. ■

Appendix C. Additional Simulation Results

In this section, we present additional simulation results, including the computation time and out-of-sample predictive performances.

The computational time of the proposed method and the BT model is summarized in Table C. As expected, the BT model is consistently faster than the proposed method due to its simpler structure. While the computational time increases with n for both methods, it is also notable that the running time of the proposed method increases with the sparsity level and the latent dimension. This reflects an increase in computational complexity of the estimation problem when the underlying structure becomes more difficult to recover.

We also present results evaluating the out-of-sample predictive performances. Specifically, for each simulation replicate, after generating the training data and estimating the model parameters, we independently generate a test set from the same underlying comparison probability matrix Π^* . To mimic a realistic data-analysis scenario, we control the size of the test set to be approximately 25% of the training data.

In particular, for each pair (i, j) , we independently generate a potential number of test comparisons $\tilde{n}_{ij}$ using the same sampling mechanism as in the training phase. We then obtain the actual number of comparisons by generating $n_{ij}^{(\text{test})}$ following $\text{Bin}(\tilde{n}_{ij}, 0.25)$. The

k	n	Sparse		Less Sparse		Dense	
		proposed	BT	proposed	BT	proposed	BT
1	500	16.63	0.08	11.81	0.08	14.61	0.14
2	500	25.15	0.05	16.81	0.08	13.97	0.14
3	500	35.34	0.05	23.05	0.08	16.21	0.14
4	500	48.52	0.05	24.92	0.08	16.61	0.13
5	500	66.45	0.05	30.14	0.08	24.84	0.13
6	500	96.33	0.05	36.54	0.08	17.38	0.13
7	500	136.75	0.05	44.27	0.08	20.43	0.13
8	500	158.26	0.05	55.45	0.08	30.61	0.13
9	500	176.87	0.05	68.04	0.08	32.31	0.13
10	500	201.67	0.05	83.15	0.08	33.41	0.12
1	1000	125.86	0.22	76.32	0.25	86.25	0.69
2	1000	195.66	0.16	119.11	0.26	106.37	0.61
3	1000	287.85	0.16	141.20	0.26	121.19	0.60
4	1000	337.10	0.16	172.06	0.25	129.95	0.58
5	1000	497.85	0.16	218.85	0.26	123.24	0.60
6	1000	717.48	0.16	248.50	0.26	179.03	0.57
7	1000	954.86	0.16	256.35	0.25	180.23	0.60
8	1000	1199.15	0.16	325.61	0.25	176.08	0.56
9	1000	1477.26	0.15	405.13	0.26	204.43	0.57
10	1000	1666.11	0.15	488.08	0.25	212.17	0.56
1	1500	439.61	0.45	236.00	0.68	354.05	1.52
2	1500	658.29	0.40	368.61	0.67	391.12	1.55
3	1500	893.69	0.41	471.67	0.62	412.78	1.42
4	1500	1156.06	0.38	433.54	0.70	380.19	1.39
5	1500	1710.13	0.42	673.19	0.60	401.08	1.39
6	1500	2321.63	0.41	708.65	0.63	437.05	1.38
7	1500	3071.95	0.38	833.68	0.67	563.98	1.38
8	1500	3835.15	0.40	1022.79	0.62	661.77	1.35
9	1500	5032.98	0.41	1132.81	0.65	270.44	1.41
10	1500	5317.70	0.40	1388.54	0.64	709.70	1.36
1	2000	1054.74	0.86	548.95	1.09	725.27	2.89
2	2000	1383.09	0.83	874.44	1.13	885.54	3.09
3	2000	2046.04	0.82	1077.11	1.13	383.48	2.86
4	2000	2663.27	0.80	806.64	1.12	939.02	2.76
5	2000	4020.82	0.82	1455.05	1.13	934.34	2.80
6	2000	5335.95	0.83	1569.47	1.12	990.44	2.77
7	2000	7366.32	0.82	1957.48	1.11	1342.22	2.76
8	2000	10153.55	0.83	2086.00	1.14	1437.54	2.68
9	2000	12119.82	0.82	2410.81	1.11	1109.65	2.69
10	2000	14110.18	0.86	2873.36	1.10	1335.29	2.65

Table 2: Total computational time (in minutes) for the proposed method and the Bradley–Terry (BT) model across different sparsity levels (sparse, less sparse, dense), sample sizes ($n = 500, 1000, 1500, 2000$), and rank parameters $k = 1, \dots, 10$. Each value represents the total time required to complete 50 simulation runs.

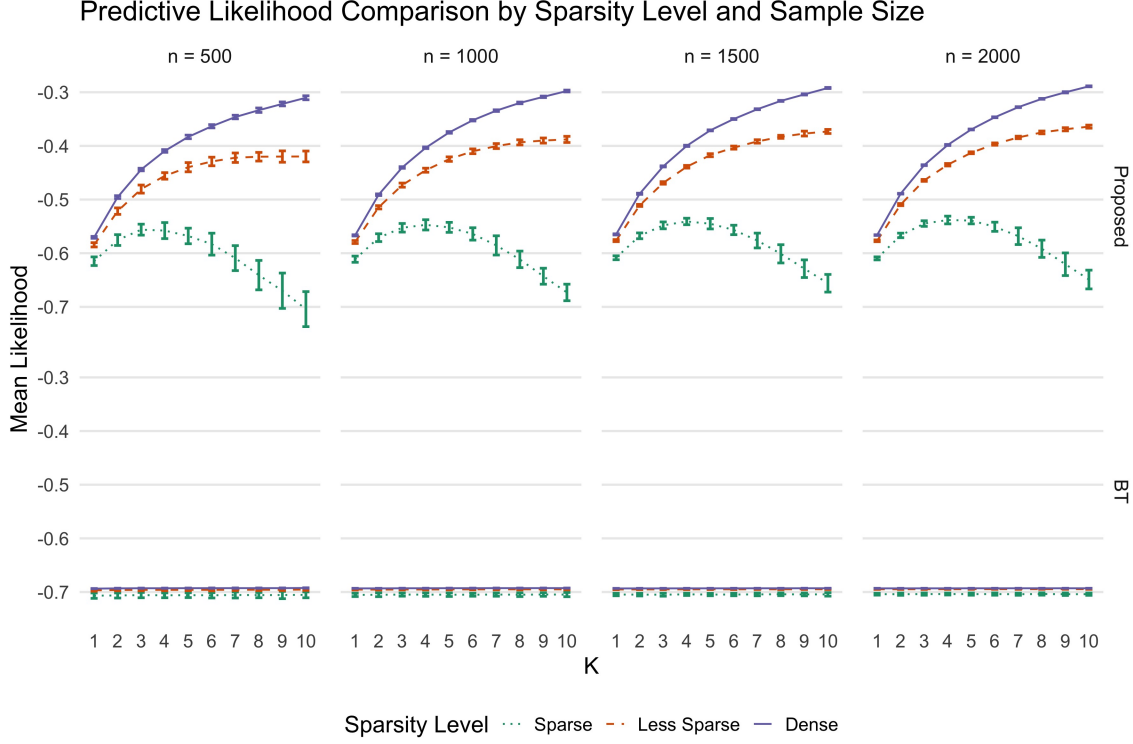


Figure 2: Comparison of predictive likelihood between the proposed method and the Bradley-Terry (BT) model across different sparsity levels (sparse, less sparse, dense). Each row corresponds to a method (top: proposed; bottom: BT), and each column corresponds to a sample size ($n = 500, 1000, 1500, 2000$). The x-axis represents the rank parameter k , while the y-axis shows the mean likelihood, computed as the average of the likelihood defined in (38). Error bars indicate ± 2 standard deviations across replicates.

comparison outcomes are then independently generated according to the true probability π_{ij}^* conditional on n_{ij}^{test} . We let $Y^{(\text{test})} = (y_{ij}^{(\text{test})})_{n \times n}$ denote the $n \times n$ matrix collecting all test comparison outcomes.

We compute the normalized predictive likelihood given by

$$L(Y^{(\text{test})} \mid \hat{\Pi}) = \frac{1}{\sum_{i=1}^n \sum_{j>i} n_{ij}} \sum_{i=1}^n \sum_{j>i} \left(y_{ij}^{(\text{test})} \log(\hat{\pi}_{ij}) + y_{ji}^{(\text{test})} \log(1 - \hat{\pi}_{ij}) \right), \quad (38)$$

and the average predictive likelihood across 50 simulations for each model is reported in Figure 2, considering different values of n , k , and sparsity levels.

References

- Jonathan Barzilai and Jonathan M Borwein. Two-point step size gradient methods. *IMA Journal of Numerical Analysis*, 8(1):141–148, 1988.
- Peter Benner, Ralph Byers, Heike Fassbender, Volker Mehrmann, and David Watkins. Cholesky-like factorizations of skew-symmetric matrices. *Electronic Transactions on Numerical Analysis*, 11:85–93, 2000.
- Ernesto G Birgin, José Mario Martínez, and Marcos Raydan. Nonmonotone spectral projected gradient methods on convex sets. *SIAM Journal on Optimization*, 10(4):1196–1211, 2000.
- Ralph Allan Bradley and Milton E Terry. Rank analysis of incomplete block designs: I. the method of paired comparisons. *Biometrika*, 39(3/4):324–345, 1952.
- Tony Cai and Wen-Xin Zhou. A max-norm constrained minimization approach to 1-bit matrix completion. *Journal of Machine Learning Research*, 14(114):3619–3647, 2013.
- Tony Cai and Wen-Xin Zhou. Matrix completion via max-norm constrained optimization. *Electronic Journal of Statistics*, 10:1493–1525, 2016.
- Manuela Cattelan. Models for paired comparison data: A review with emphasis on dependent data. *Statistical Science*, 27(3):412–433, 2012.
- Manuela Cattelan, Cristiano Varin, and David Firth. Dynamic bradley–terry modelling of sports tournaments. *Journal of the Royal Statistical Society Series C: Applied Statistics*, 62(1):135–150, 2013.
- Sabyasachi Chatterjee and Sumit Mukherjee. Estimation in tournaments and graphs under monotonicity constraints. *IEEE Transactions on Information Theory*, 65(6):3525–3539, 2019.
- Sourav Chatterjee. Matrix estimation by universal singular value thresholding. *The Annals of Statistics*, 43(1):177–214, 2015.
- Pinhan Chen, Chao Gao, and Anderson Y Zhang. Partial recovery for top-k ranking: optimality of mle and suboptimality of the spectral method. *The Annals of Statistics*, 50(3):1618–1652, 2022a.
- Shuo Chen and Thorsten Joachims. Modeling intransitivity in matchup and comparison data. In *Proceedings of the Ninth ACM International Conference on Web Search and Data Mining*, pages 227–236, 2016.
- Xi Chen, Paul N Bennett, Kevyn Collins-Thompson, and Eric Horvitz. Pairwise ranking aggregation in a crowdsourced setting. In *Proceedings of the Sixth ACM International Conference on Web Search and Data Mining*, pages 193–202, 2013.
- Xi Chen, Kevin Jiao, and Qihang Lin. Bayesian decision process for cost-efficient dynamic ranking via crowdsourcing. *Journal of Machine Learning Research*, 17(216):1–40, 2016.

- Xi Chen, Yunxiao Chen, and Xiaoou Li. Asymptotically optimal sequential design for rank aggregation. *Mathematics of Operations Research*, 47(3):2310–2332, 2022b.
- Yunxiao Chen and Xiaoou Li. Determining the number of factors in high-dimensional generalized latent factor models. *Biometrika*, 109(3):769–782, 2022.
- Yunxiao Chen and Xiaoou Li. A note on entrywise consistency for mixed-data matrix completion. *Journal of Machine Learning Research*, 25(343):1–66, 2024.
- Yunxiao Chen, Xiaoou Li, and Siliang Zhang. Structured latent factor analysis for large-scale data: Identifiability, estimability, and their implications. *Journal of the American Statistical Association*, 115(532):1756–1770, 2020a.
- Yuxin Chen and Changho Suh. Spectral mle: Top-k rank aggregation from pairwise comparisons. In *International Conference on Machine Learning*, pages 371–380. PMLR, 2015.
- Yuxin Chen, Jianqing Fan, Cong Ma, and Kaizheng Wang. Spectral method and regularized mle are both optimal for top-k ranking. *The Annals of Statistics*, 47(4):2204, 2019.
- Yuxin Chen, Yuejie Chi, Jianqing Fan, Cong Ma, and Yuling Yan. Noisy matrix completion: Understanding statistical guarantees for convex relaxation via nonconvex optimization. *SIAM Journal on Optimization*, 30(4):3098–3121, 2020b.
- Paul F Christiano, Jan Leike, Tom Brown, Miljan Martic, Shane Legg, and Dario Amodei. Deep reinforcement learning from human preferences. *Advances in Neural Information Processing Systems*, 30, 2017.
- Mark A Davenport, Yaniv Plan, Ewout Van Den Berg, and Mary Wootters. 1-bit matrix completion. *Information and Inference: A Journal of the IMA*, 3(3):189–223, 2014.
- Paul Erdős and Alfréd Rényi. On the evolution of random graphs. *A Magyar Tudományos Akadémia Matematikai Kutató Intézetének Közleményei*, 5:17–61, 1960.
- Jianqing Fan, Jikai Hou, and Mengxin Yu. Uncertainty quantification of mle for entity ranking with covariates. *Journal of Machine Learning Research*, 25(358):1–83, 2024.
- Ruijian Han, Rougang Ye, Chunxi Tan, and Kani Chen. Asymptotic theory of sparse bradley–terry model. *The Annals of Applied Probability*, 30(5):2491–2515, 2020.
- Ruijian Han, Yiming Xu, and Kani Chen. A general pairwise comparison model for extremely sparse networks. *Journal of the American Statistical Association*, 118(544):2422–2432, 2023.
- Reinhard Heckel, Nihar B Shah, Kannan Ramchandran, and Martin J Wainwright. Active ranking from pairwise comparisons and when parametric assumptions do not help. *The Annals of Statistics*, 47(6):3099–3126, 2019.
- Roger A Horn and Charles R Johnson. *Matrix Analysis*. Cambridge University Press, 2013.
- Michel Ledoux and Michel Talagrand. *Probability in Banach Spaces: Isoperimetry and Processes*, volume 23. Springer Science & Business Media, 1991.

- Sze Ming Lee, Yunxiao Chen, and Tony Sit. A latent variable approach to learning high-dimensional multivariate longitudinal data. *Journal of the American Statistical Association*, pages 1–12, 2026.
- Yue Liu, Ethan X Fang, and Junwei Lu. Lagrangian inference for ranking problems. *Operations Research*, 71(1):202–223, 2023.
- Štefan Lyócsa and Tomáš Vřost. To bet or not to bet: a reality check for tennis betting market efficiency. *Applied Economics*, 50(20):2251–2272, 2018.
- Roberto Macrì Demartino, Leonardo Egidi, and Nicola Torelli. Alternative ranking measures to predict international football results. *Computational Statistics*, pages 1–19, 2024.
- Ian McHale and Alex Morton. A bradley-terry type model for forecasting tennis match results. *International Journal of Forecasting*, 27(2):619–630, 2011.
- FD Murnaghan and Aurel Wintner. A canonical form for real matrices under orthogonal transformations. *Proceedings of the National Academy of Sciences*, 17(7):417–420, 1931.
- Sahand Negahban, Sewoong Oh, and Devavrat Shah. Rank centrality: Ranking from pairwise comparisons. *Operations Research*, 65(1):266–287, 2017.
- Ivo FD Oliveira, Nir Ailon, and Ori Davidov. A new and flexible approach to the analysis of paired comparison data. *Journal of Machine Learning Research*, 19(60):1–29, 2018.
- Long Ouyang, Jeffrey Wu, Xu Jiang, Diogo Almeida, Carroll Wainwright, Pamela Mishkin, Chong Zhang, Sandhini Agarwal, Katarina Slama, Alex Ray, et al. Training language models to follow instructions with human feedback. *Advances in Neural Information Processing Systems*, 35, 2022.
- Philip Ramirez, J James Reade, and Carl Singleton. Betting on a buzz: Mispricing and inefficiency in online sportsbooks. *International Journal of Forecasting*, 39(3):1413–1423, 2023.
- Nihar B Shah and Martin J Wainwright. Simple, robust and optimal ranking from pairwise comparisons. *Journal of Machine Learning Research*, 18(199):1–38, 2018.
- Nihar B Shah, Sivaraman Balakrishnan, Joseph Bradley, Abhay Parekh, Kannan Ramchandran, and Martin J Wainwright. Estimation from pairwise comparisons: Sharp minimax bounds with topology dependence. *Journal of Machine Learning Research*, 17(58):1–47, 2016a.
- Nihar B Shah, Sivaraman Balakrishnan, Adityanand Guntuboyina, and Martin J Wainwright. Stochastically transitive models for pairwise comparisons: Statistical and computational issues. *IEEE Transactions on Information Theory*, 63(2):934–959, 2016b.
- Gordon Simons and Yi-Ching Yao. Asymptotics when the number of parameters tends to infinity in the Bradley-Terry model for paired comparisons. *The Annals of Statistics*, 27: 1041–1060, 1999.

- Harriet Spearing, Jonathan Tawn, David Irons, and Tim Paulden. Modeling intransitivity in pairwise comparisons with application to baseball data. *Journal of Computational and Graphical Statistics*, 32(4):1383–1392, 2023.
- Louis L Thurstone. The method of paired comparisons for social values. *The Journal of Abnormal and Social Psychology*, 21:384 – 400, 1927.
- Joel A Tropp et al. An introduction to matrix concentration inequalities. *Foundations and Trends® in Machine Learning*, 8(1-2):1–230, 2015.
- Alkeos Tsokos, Santhosh Narayanan, Ioannis Kosmidis, Gianluca Baio, Mihai Cucuringu, Gavin Whitaker, and Franz Király. Modeling outcomes of soccer matches. *Machine Learning*, 108:77–95, 2019.
- Ewout Van Den Berg and Michael P Friedlander. Probing the pareto frontier for basis pursuit solutions. *Siam Journal on Scientific Computing*, 31(2):890–912, 2008.
- Ting Yan, Yaning Yang, and Jinfeng Xu. Sparse paired comparisons in the bradley-terry model. *Statistica Sinica*, pages 1305–1318, 2012.
- Yi Yu, Tengyao Wang, and Richard J Samworth. A useful variant of the Davis–Kahan theorem for statisticians. *Biometrika*, 102(2):315–323, 2015.
- Banghua Zhu, Michael Jordan, and Jiantao Jiao. Principled reinforcement learning with human feedback from pairwise or k-wise comparisons. In *International Conference on Machine Learning*, pages 43037–43067. PMLR, 2023.