
Stretching the Effectiveness of MLE from Accuracy to Bias for Pairwise Comparisons

Jingyan Wang Nihar B. Shah R. Ravi
Carnegie Mellon University

Abstract

A number of applications (e.g., AI bot tournaments, sports, peer grading, crowdsourcing) use pairwise comparison data and the Bradley-Terry-Luce (BTL) model to evaluate a given collection of items (e.g., bots, teams, students, search results). Past work has shown that under the BTL model, the widely-used maximum-likelihood estimator (MLE) is minimax-optimal in estimating the item parameters, in terms of the mean squared error. However, another important desideratum for designing estimators is fairness. In this work, we consider one specific type of fairness, which is the notion of bias in statistics. We show that the MLE incurs a suboptimal rate in terms of bias. We then propose a simple modification to the MLE, which “stretches” the bounding box of the maximum-likelihood optimizer by a small constant factor from the underlying ground truth domain. We show that this simple modification leads to an improved rate in bias, while maintaining minimax-optimality in the mean squared error. In this manner, our proposed class of estimators provably improves fairness in the sense of bias without loss in accuracy.

1 Introduction

A number of applications involve data in the form of pairwise comparisons among a collection of items, and entail an evaluation of the individual items from this data. An application gaining increasing popularity is competition between pairs of AI bots (e.g., Ontan

et al., 2013). Here a number of AI bots compete with each other in pairwise matchups for a certain task, where each bot plays every other bot a certain number of times in a *round robin* fashion, with the goal of evaluating the quality of each bot. A second example is the evaluation of self-play of AI algorithms in their training phase (Silver et al., 2017), where again, different copies of an AI bot play against each other a number of times. Applications involving humans include sports and online games such as the English Premier League of football (Király and Qian, 2017; SinceAWin.com, 2019) (unofficial ratings), and official world rankings for chess such as FIDE (International Chess Federation, 2017) and USCF (Glickman and Doan, 2017) ratings.

A common method of evaluating the items based on pairwise comparisons is to assume that the probability of an item beating another equals the logistic function of the difference in the true quality of the two items, and then infer the true quality from the observed outcomes of the comparisons (e.g., the Elo rating system). Various applications employ such an approach to rating from pairwise comparisons, with some modifications tailored to that specific application. Our goal is not to study the application-specific versions, but the foundational underpinnings of such rating systems.

In this paper, we study the pairwise-comparison model that underlies (Glickman and Jones, 1999; Aldous, 2017) these rating systems, namely the Bradley-Terry-Luce (BTL) model (Bradley and Terry, 1952; Luce, 1959). The BTL model assumes that each item is associated to an unknown real-valued parameter representing the quality of that item, and assumes that the probability of an item beating another is the logistic function applied to the difference of the parameters of these two items. The BTL model is also employed in the applications of peer grading (Shah et al., 2013; Lamon et al., 2016) (where the grades of the students are set as the BTL parameters to be estimated), crowdsourcing (Chen et al., 2016; Ponce-López et al., 2016), and understanding consumer choice in marketing (Green et al., 1981).

1.1 BTL model and maximum likelihood estimation

Now we present a formal definition of the BTL model. Let $d \geq 2$ denote the number of items. The d items are associated to an unknown parameter vector $\theta^* \in \mathbb{R}^d$ whose i^{th} entry represents the underlying quality of item $i \in [d]$. When any item $i \in [d]$ is compared with any item $j \in [d]$ in the BTL model, then item i beats item j with probability

$$\frac{1}{1 + e^{-(\theta_i^* - \theta_j^*)}}, \quad (1)$$

independent of all other comparisons. The probability of item j beating i is one minus the expression (1) above. We consider the “league format” (Aldous, 2017) of comparisons where every pair of items is compared k times.

We follow the usual assumption (Hajek et al., 2014; Shah et al., 2016) under the BTL model that the true parameter vector θ^* lies in the set Θ_B parameterized by a known *constant* $B > 0$, defined as:

$$\Theta_B = \{\theta \in \mathbb{R}^d \mid \|\theta\|_\infty \leq B \text{ and } \sum_{i=1}^d \theta_i = 0\}. \quad (2)$$

The first constraint requires that the magnitude of the parameters is bounded by some constant B . We call this constraint the “box constraint”. It has been shown that a box constraint is necessary, because otherwise the estimation error can diverge to infinity (Shah et al., 2016, Appendix G). The second constraint requires that the parameters sum to 0. This is without loss of generality due to the shift-invariance property of the BTL model.

A large amount of both theoretical (Hunter, 2004; Hajek et al., 2014; Szörényi et al., 2015; Negahban et al., 2016; Shah et al., 2016) and applied (Stigler, 1994; Sham and Curtis, 1995; Chen et al., 2016; Ponce-López et al., 2016) literature focuses on the goal of estimating the parameter vector θ^* of the BTL model. A standard and widely-studied estimator is the maximum-likelihood estimator (MLE):

$$\hat{\theta}^{(B)} = \underset{\theta \in \Theta_B}{\operatorname{argmin}} \ell(\theta), \quad (3)$$

where ℓ is the negative log-likelihood function. Letting W_{ij} denote a random variable representing the number of times that item $i \in [d]$ beats item $j \in [d]$, the log-likelihood function ℓ is given by:

$$\ell(\theta) := \ell(\{W_{ij}\}; \theta) = - \sum_{1 \leq i < j \leq d} \left[W_{ij} \log \left(\frac{1}{1 + e^{-(\theta_i - \theta_j)}} \right) + W_{ji} \log \left(\frac{1}{1 + e^{-(\theta_j - \theta_i)}} \right) \right].$$

1.2 Metrics

Accuracy. A common metric used in the literature on estimating the BTL model is the *accuracy* of the estimate, measured in terms of the mean squared error. Formally, the (worst-case) accuracy of any estimator $\hat{\theta}$ is defined as:

$$\alpha(\hat{\theta}) := \sup_{\theta^* \in \Theta_B} \mathbb{E}[\|\hat{\theta} - \theta^*\|_2^2]. \quad (4)$$

Importantly, past work (Hajek et al., 2014; Shah et al., 2016) has shown that the MLE (3) has the appealing property of being minimax-optimal in terms of the accuracy defined in (4).

Bias. Another important desideratum for designing and evaluating estimators is fairness. For example, in sports or online games, we do not want to assign scores in such a way that it systematically gives certain players higher scores than their true quality, but at the same time gives certain other players lower scores than their true quality. In this paper, we adopt the standard definition of *bias* in statistics as our notion of fairness. For any estimator, the bias incurred by this estimator on a parameter is defined as the difference between the expected value of the estimator and the true value of the parameter. Since our parameters are a vector, we consider the worst-case bias, that is, the maximum magnitude of the bias across all items. Formally, the worst-case bias of any estimator $\hat{\theta}$ is defined as:

$$\beta(\hat{\theta}) := \sup_{\theta^* \in \Theta_B} \|\mathbb{E}[\hat{\theta}] - \theta^*\|_\infty.$$

With this background, we now provide an overview of the contributions of this paper.

1.3 Contribution I: Performance of MLE

Our first contribution is to analyze the widely-used MLE (3) in terms of its bias. Let us begin with a visual illustration through simulation. Consider $d = 25$ items with parameter values equally spaced in the interval $[-1, 1]$ (here we have $B = 1$), where $k = 5$ pairwise comparisons are observed between each pair of items under the BTL model. We estimate the parameters using the MLE, and plot the bias on each item across 5000 iterations of the simulation in Fig. 1 (striped red). The MLE shows a systematic bias: it induces a negative bias (under-estimation) on the large positive parameters, and a positive bias (over-estimation) on the large negative parameters. In the applications of interest, the MLE thus systematically underestimates the abilities of the top players/students/items and overestimates the abilities of those at the bottom.

In this paper, we theoretically quantify the bias incurred by the MLE as follows.

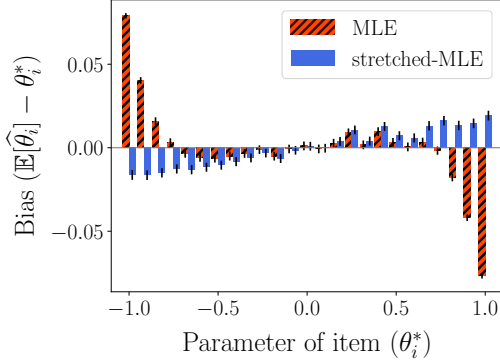


Figure 1: Biases on items of different parameters, induced by the MLE ($B = 1$) and our stretched-MLE ($A = 2$). Our estimator significantly reduces the maximum magnitude of the bias across the items. Note that this figure plots the bias including its sign: A positive bias means over-estimation of the parameter, and a negative bias means under-estimation of the parameter. Each bar is a mean over 5000 iterations.

Theorem 1.1 (MLE bias lower bound; Informal). *The MLE (3) incurs a bias $\beta(\hat{\theta}^{(B)})$ lower bounded as $\Omega(\frac{1}{\sqrt{dk}})$.*

As shown by our results to follow, this bias is suboptimal. Our proof for this result indicates that the bias is incurred because the MLE operates under the accurately specified model with the box constraint at B . That is, the MLE “clips” the estimate to lie within the set Θ_B . This issue is visible in the simulation of Fig. 1 where the bias is the largest when the true values of the parameters are near the boundaries $\pm B$. For example, consider a true parameter whose value equals B . The estimate of this parameter sometimes equals the largest allowed value B (due to the box constraint), and sometimes is smaller than B (due to the randomness of the data). Therefore, in expectation, the estimate of this parameter incurs a negative bias. An analogous argument explains the positive bias when the true parameter equals or is close to $-B$.

1.4 Contribution II: Proposed stretched estimator and its theoretical guarantees

Our goal is to design an estimator with a lower bias while maintaining high accuracy. Since the MLE (3) is already widely studied and used, it is also desirable from a practical and computational standpoint that the new estimator is a simple modification of the MLE (3). With this motivation in mind, an intuitive approach is to consider the MLE but without the box constraint “ $\|\theta\|_\infty \leq B$ ”. We call this estimator without the box constraint as the “unconstrained MLE” and denote it by $\hat{\theta}^{(\infty)}$, as removing the box constraint

Estimator	Bias	Mean squared error
Standard MLE $\hat{\theta}^{(B)}$	$\Omega(\frac{1}{\sqrt{dk}})$ (Thm. 2.1(a))	$\mathcal{O}(\frac{1}{k})$ minimax-optimal (Hajek et al., 2014; Shah et al., 2016)
Unconstrained MLE $\hat{\theta}^{(\infty)}$	Undefined	∞
Stretched MLE $\hat{\theta}^{(A)}$	$\tilde{\mathcal{O}}(\frac{1}{\sqrt{dk}})$ (Thm. 2.1(b))	$\mathcal{O}(\frac{1}{k})$ minimax-optimal (Thm. 2.2(b))

Table 1: Theoretical guarantees for the MLE $\hat{\theta}^{(B)}$, unconstrained MLE $\hat{\theta}^{(\infty)}$ and the proposed stretched-MLE $\hat{\theta}^{(A)}$ (with a constant A such that $A > B$). The proposed stretched-MLE achieves a better rate on bias, while retaining minimax optimality in terms of accuracy. Recall that d denotes the number of items and k denotes the number of comparisons per pair.

is equivalent to setting the box constraint to ∞ :

$$\hat{\theta}^{(\infty)} = \operatorname{argmin}_{\theta \in \Theta_\infty} \ell(\theta), \quad (5)$$

where $\Theta_\infty := \{\theta \in \mathbb{R}^d \mid \sum_{i=1}^d \theta_i = 0\}$. The unconstrained MLE $\hat{\theta}^{(\infty)}$ incurs an unbounded error in terms of accuracy. This is because with non-zero probability an item beats all others. Then the unconstrained MLE estimates the parameter of this item as ∞ , thereby inducing an unbounded mean squared error.

Consequently, in this work, we propose the following simple modification to the MLE which is a middle ground between the MLE and the unconstrained MLE. Specifically, we consider a “stretched-MLE” associated to a parameter A such that $A > B$. Given the parameter A , the stretched-MLE is identical to the MLE (3) but “stretches” the box constraint to A :

$$\hat{\theta}^{(A)} = \operatorname{argmin}_{\theta \in \Theta_A} \ell(\theta), \quad (6)$$

where $\Theta_A := \{\theta \in \mathbb{R}^d \mid \|\theta\|_\infty \leq A \text{ and } \sum_{i=1}^d \theta_i = 0\}$. That is, Θ_A simply replaces the box constraint $\|\theta\|_\infty \leq B$ in (2) by the “stretched” box constraint $\|\theta\|_\infty \leq A$.

The bias induced by the stretched-MLE (with $A = 2$) in the previous experiment is also shown in Fig. 1 (solid blue). Note that the maximum bias (at the leftmost item with the largest negative parameter, or the rightmost item with the largest positive parameter) is significantly reduced compared to the MLE. Moreover, the bias induced by the stretched-MLE looks qualitatively more evened out across the items.

Our second main theoretical result proves that the stretched-MLE indeed incurs a significantly lower bias.

Theorem 1.2 (Stretched-MLE bias upper bound; Informal). *When $B = 1$, the stretched-MLE (6) with $A = 2$ incurs a bias $\beta(\hat{\theta}^{(A)})$ upper bounded as $\tilde{\mathcal{O}}(\frac{1}{dk})$.*

Given the significant bias reduction by our estimator, a natural question is about the accuracy of the stretched-MLE, particularly given the unbounded error incurred by the unconstrained MLE. We prove that our stretched-MLE is able to maintain the same minimax-optimal rate on the mean squared error as the standard MLE.

Theorem 1.3 (Stretched-MLE accuracy upper bound; Informal). *When $B = 1$, the stretched-MLE (6) with $A = 2$ incurs a mean squared error $\alpha(\hat{\theta}^{(A)})$ upper bounded as $\mathcal{O}(\frac{1}{k})$, which is minimax-optimal.*

This result shows a **win-win** by our stretched-MLE: *reducing the bias while retaining the accuracy guarantee*. The comparison of the MLE and the stretched-MLE in terms of accuracy and bias is summarized in Table 1. Another attractive feature of our result is that the proposed stretched-MLE is a simple modification of the standard MLE, which can easily be incorporated in any existing implementation. While our modification to the estimator is simple to implement, our theoretical analyses and the proofs are non-trivial.

The MLE and the stretched-MLE also relate to the concepts of proper vs. improper learning in learning theory. The MLE can be considered a proper estimator that is required to output an estimate from true set Θ_B . The stretched-MLE is an improper estimator that outputs an estimate in the “stretched” set Θ_A , and the unconstrained MLE is a (fully) improper estimator that is allowed to output any arbitrary estimate.

1.5 Related work

The logistic nature (1) of the BTL model relates our work to studies of logistic regression (e.g., Portnoy, 1988; He and Shao, 2000; Fan et al., 2019), among which the work on high-dimensional logistic regression (Sur and Candès, 2019; Salehi et al., 2019; Zhao et al., 2020) is the most closely related to ours. The papers (Sur and Candès, 2019; Zhao et al., 2020) consider an unconstrained MLE in logistic regression, and shows its bias in the opposite direction as compared to our results on the standard MLE (constrained) in the BTL model. Specifically, the papers (Sur and Candès, 2019; Zhao et al., 2020) show that the large positive coefficients are overestimated, and the large negative coefficients are underestimated. There are several additional key differences between the results in Sur and Candès (2019) as compared to the present paper. The paper (Sur and Candès, 2019) studies the asymptotic

bias of the unconstrained MLE, showing that the unconstrained MLE is not consistent. On the other hand, we operate in a regime where the MLE is still consistent, and study finite-sample bounds. Moreover, the paper (Sur and Candès, 2019) assumes that the predictor variables are i.i.d. Gaussian. On the other hand, in the BTL model the probability that item i beats item j can be written as $\frac{1}{1+e^{-x_{ij}^T \theta^*}}$, where each predictor variable $x_{ij} \in \mathbb{R}^d$ has entry i equal to 1, entry j equal to -1 , and the remaining entries equal to 0. A subsequent paper (Salehi et al., 2019) considers a regularized MLE in logistic regression where the regularizer can reduce the asymptotic bias.

A common way to achieve bias reduction is to employ finite-sample correction, such as Jackknife (Que-nouille, 1949) and other methods (Cox and Snell, 1968; Anderson and Richardson, 1979; Firth, 1993) to the MLE (or other estimators). These methods operate in a low-dimensional regime (small d) where the MLE is asymptotically unbiased. Informally, these methods use a Taylor expansion, write the expression for the bias as an infinite sum, and modify the estimator in a variety of ways to eliminate the lower-order terms in this bias expression. However, since the expression is an infinite sum, eliminating the first term does not guarantee a low rate of the bias. Moreover, since the Taylor expansion terms in the infinite sum are implicit functions of θ^* , eliminating lower-order terms does not directly translate to explicit worst-case guarantees.

Returning to the pairwise-comparison setting, in addition to the mean squared error, past work has also considered accuracy in terms of the ℓ_1 norm error (Agarwal et al., 2018) and the ℓ_∞ norm error (Chen and Suh, 2015; Jang et al., 2017; Chen et al., 2019). The ℓ_∞ bound for a regularized MLE is analyzed in Chen et al. (2019). Our proof for bounding the bias of the standard MLE (unregularized) relies on a high-probability ℓ_∞ bound for the unconstrained MLE (unregularized). It is important to note that the bound for the regularized MLE from Chen et al. (2019) does not carry to the unregularized MLE, because the proof from Chen et al. (2019) relies on the strong convexity of the regularizer. On the other hand, *our intermediate result provides a partial answer to the open question in Chen et al. (2019) about the ℓ_∞ norm for the unregularized MLE* (Lemma A.5 in Appendix A): We establish an ℓ_∞ bound for the unregularized MLE when $p_{obs} = 1$, which has the same rate as that of the regularized MLE in (Chen et al., 2019).

Another common occurrence of bias is regression towards the mean (Stigler, 1997). It is the phenomenon that random variables taking large (or small) values in one measurement are likely to take more moderate

(closer to average) values in subsequent measurements. On the contrary, we consider items whose indices are fixed (and are not order statistics). For fixed indices, our results suggest that under the BTL model, the bias (under-estimation of large true values) is in the opposite direction as that in regression towards the mean (over-estimation of large observed values).

Finally, there is a vast literature on different notions of fairness in various problems such as classification and resource allocation (e.g. Barocas et al., 2019, Chapter 3; Marsh and Schilling, 1994 and references therein). While other notions of fairness are of interest for future research, in this work we focus on the classical statistical notion of bias to capture our intuition of minimizing the maximum disparity in estimation.

2 Main results

In this section, we formally provide our main theoretical results on bias and on the mean squared error.

2.1 Bias

Recall that d denotes the number of items and k denotes the number of comparisons per pair of items. The true parameter vector is $\theta^* \in \Theta_B$ for some pre-specified constant $B > 0$. The following theorem provides bounds on the bias of the standard MLE $\hat{\theta}^{(B)}$ and that of our stretched-MLE $\hat{\theta}^{(A)}$ with parameter A . Specifically, it shows that if A is a finite constant strictly greater than B , then our stretched-MLE has a much smaller bias than the MLE when d and k are sufficiently large.

Theorem 2.1. (a) *There exists a constant $c > 0$ that depends only on the constant B , such that*

$$\beta(\hat{\theta}^{(B)}) \geq \frac{c}{\sqrt{dk}}, \quad (7a)$$

for all $d \geq d_0$ and all $k \geq k_0$, where d_0 and k_0 are constants that depend only on the constant B .

(b) *Let A be any finite constant such that $A > B$. There exists a constant $c > 0$ that depends only on the constants A and B , such that*

$$\beta(\hat{\theta}^{(A)}) \leq c \frac{\log d + \log k}{dk}, \quad (7b)$$

for all $d \geq d_0$ and all $k \geq k_0$, where d_0 and k_0 are constants that depend only on the constants A and B .

We note that in Theorem 2.1(b), we allow A to be any positive constant as long as $A > B$. Therefore, the difference between A and B can be any arbitrarily small constant. It is perhaps surprising that stretching the

box constraint only by a small constant yields such a significant improvement in the bias. We provide intuition behind this result in Section 2.1.1. The complete proof is provided in Appendix A.

2.1.1 Intuition for Theorem 2.1

In this section, we provide intuition why stretching the box constraint from B to A significantly reduces the bias. Specifically, we consider a simplified setting with $d = 2$ items. Due to the centering constraint, we have $\theta_2^* = -\theta_1^*$ for the true parameters, and we have $\hat{\theta}_2 = -\hat{\theta}_1$ for any estimator $\hat{\theta}$ that satisfies the centering constraint. Therefore, it suffices to focus only on item 1. Denote μ as the random variable representing the fraction of times that item 1 beats item 2, and denote the true probability that item 1 beats item 2 as $\mu^* := \frac{1}{1+e^{-(\theta_1^* - \theta_2^*)}}$. We consider the true parameter of item 1 as $\theta_1^* \in [-B, B]$. Then we have $\mu^* \in [\mu_-, \mu_+]$, where $\mu_- = \frac{1}{1+e^{2B}}$ and $\mu_+ = \frac{1}{1+e^{-2B}}$. The standard MLE $\hat{\theta}^{(B)}$, the stretched-MLE $\hat{\theta}^{(A)}$ and the unconstrained MLE $\hat{\theta}^{(\infty)}$ can be solved in closed form:

$$\begin{aligned} \hat{\theta}_1^{(B)}(\mu) &= \begin{cases} -B & \text{if } \mu \in [0, \mu_-] \\ -\frac{1}{2} \log \left(\frac{1}{\mu} - 1 \right) & \text{if } \mu \in (\mu_-, \mu_+) \\ B & \text{if } \mu \in [\mu_+, 1]. \end{cases} \\ \hat{\theta}_1^{(A)}(\mu) &= \begin{cases} -A & \text{if } \mu \in \left[0, \frac{1}{1+e^{2A}} \right] \\ -\frac{1}{2} \log \left(\frac{1}{\mu} - 1 \right) & \text{if } \mu \in \left(\frac{1}{1+e^{2A}}, \frac{1}{1+e^{-2A}} \right) \\ A & \text{if } \mu \in \left[\frac{1}{1+e^{-2A}}, 1 \right]. \end{cases} \\ \hat{\theta}_1^{(\infty)}(\mu) &= -\frac{1}{2} \log \left(\frac{1}{\mu} - 1 \right). \end{aligned}$$

See Fig. 2a for a comparison of these three estimators.

Now we consider the bias incurred by these three estimators. For intuition, let us consider the case $\theta_1^* = B$, which incurs the largest bias in our simulation of Fig. 1. If the observation μ were noiseless (and thus equals the true probability μ_+), then all three estimators would output the true parameter B . However, the observation μ is noisy, and only concentrates around μ_+ . To investigate how these three estimators behave differently under this noise, we zoom in to the region around $\mu = \mu_+$ indicated by the grey box in Fig. 2a. (Note that the observation μ can lie outside the grey box, but for intuition we ignore this low-probability event due to concentration.)

The behaviors of the three estimators in the grey box are shown in Fig. 2b, Fig. 2c and Fig. 2d, respectively. For each of these estimators, the blue dots on the x-axis denotes the noisy observation of μ across different iterations, and the blue dots on the estimator function denotes the corresponding noisy estimates. The ex-

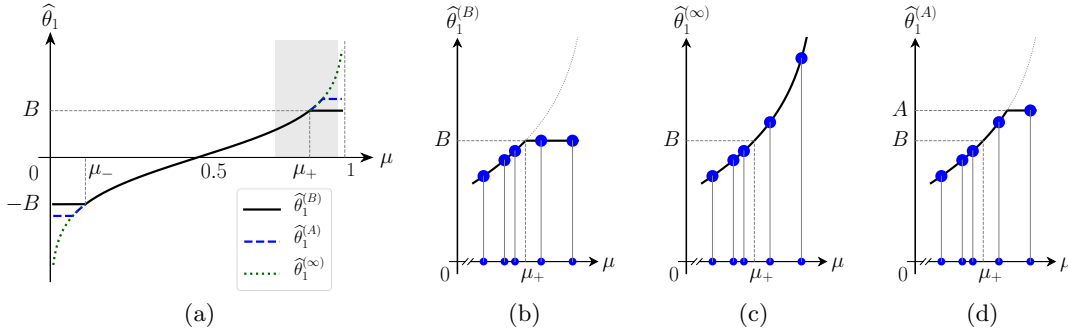


Figure 2: Intuition on the sources of bias. (a) The estimators standard MLE $\hat{\theta}^{(B)}$, stretched-MLE $\hat{\theta}^{(A)}$ and unconstrained MLE $\hat{\theta}^{(\infty)}$ (on item 1), as a function of μ when there are $d = 2$ items. We consider $\theta^* = [B, -B]$, under which the true probability that item 1 beats item 2 is μ_+ . We zoom in to the region around $\mu = \mu_+$ indicated by the grey box. (b) The standard MLE $\hat{\theta}^{(B)}$ incurs a negative bias, because the estimate is required to be at most B . (c) The unconstrained MLE $\hat{\theta}^{(\infty)}$ incurs a positive bias by Jensen’s inequality, because the estimator function is convex on $\mu \in (0.5, 1)$. (d) Our estimator balances out the negative bias and the positive bias.

pected value of the estimator is a mean over the blue dots on the estimator function. For the standard MLE $\hat{\theta}^{(B)}$ (Fig. 2b), the box constraint requires that the estimate shall never exceed B . We call this phenomenon the “clipping” effect, which introduces a negative bias. For the unconstrained MLE $\hat{\theta}^{(\infty)}$ (Fig. 2c), since the estimator function is convex, by Jensen’s inequality, the unconstrained MLE $\hat{\theta}^{(\infty)}$ introduces a positive bias. Our proposed stretched-MLE $\hat{\theta}^{(A)}$ (Fig. 2d) lies in the middle between the standard MLE and the unconstrained MLE. Therefore, the stretched-MLE balances out the negative bias from the “clipping” effect and the positive bias from the convexity of the estimator function, thereby yielding a smaller bias on the item parameter. In practice, one can numerically tune the parameter A to minimize the bias across all possible parameter vector $\theta^* \in \Theta_B$. Simulation results on different values of A are included in Section 3.

2.2 Accuracy

Given the result of Theorem 2.1 on the bias reduction of the estimator $\hat{\theta}^{(A)}$, we revisit the mean squared error. Past work (Hajek et al., 2014; Shah et al., 2016) has shown that the standard MLE $\hat{\theta}^{(B)}$ is minimax-optimal in terms of the mean squared error. The following theorem shows that this minimax-optimality also holds for our proposed stretched-MLE $\hat{\theta}^{(A)}$, where A is any constant such that $A > B$. The theorem statement and its proof follows Theorem 2 from (Shah et al., 2016), after some modification to accommodate the new bounding box parameter A .

Theorem 2.2. (a) [Theorem 2(a) from Shah et al. (2016)] *There exists a constant $c > 0$ that depends only on the constant B , such that any estimator*

$\hat{\theta}$ has a mean squared error lower bounded as

$$\alpha(\hat{\theta}) \geq \frac{c}{k}, \quad (8a)$$

for all $k \geq k_0$, where k_0 is a constant that depends only on the constant B .

(b) *Let A be any finite constant such that $A > B$. There exists a constant $c > 0$ that depends only on the constants A and B , such that*

$$\alpha(\hat{\theta}^{(A)}) \leq \frac{c}{k}. \quad (8b)$$

Theorem 2.2 shows that using the estimator $\hat{\theta}^{(A)}$ retains the minimax-optimality achieved by $\hat{\theta}^{(B)}$ in terms of the mean squared error. Combining Theorem 2.1 and Theorem 2.2 shows the Pareto improvement of our estimator $\hat{\theta}^{(A)}$: the estimator $\hat{\theta}^{(A)}$ decreases the rate of the bias, while still performing optimally on the mean squared error.

The proof of Theorem 2.2 closely mimics the proof of Theorem 2(b) from Shah et al. (2016), replacing the steps involving the domain Θ_B by the stretched domain Θ_A . The details are provided in Appendix B.

3 Simulations

In this section, we explore our problem space and compare the standard MLE and our proposed stretched-MLE by simulations. Both estimators are solutions to convex optimization problems, so we use the CVXPY package for ease of implementation (for faster methods such as Minorization Maximization (MM), see Hunter, 2004; Hajek et al., 2014). In what follows, we set $B = 1$, and unless specified otherwise we set $A = 2$

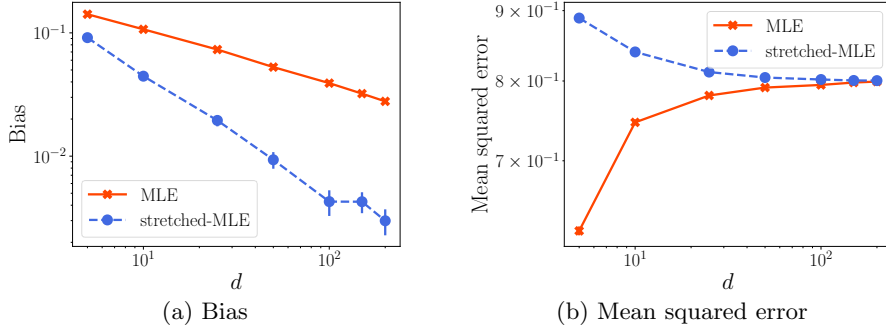


Figure 3: Performance of estimators for various values of d , with $k = 5$ and $A = 2$. Each point is a mean over 10000 iterations.

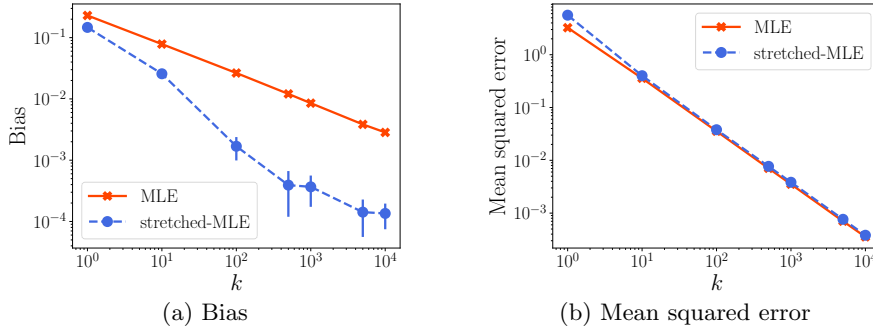


Figure 4: Performance of estimators for various values of k , with $d = 10$ and $A = 2$. Each point is a mean over 10000 iterations.

and $\theta^* = [1, -\frac{1}{d-1}, -\frac{1}{d-1}, \dots, -\frac{1}{d-1}]$. We also evaluate the performance of other values of θ^* subsequently. Error bars in all the plots represent the standard error of the mean.

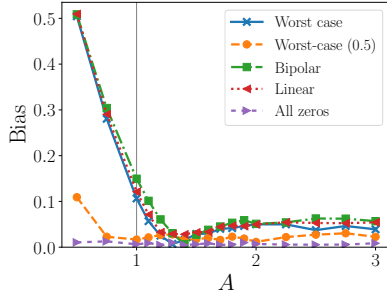
(i) **Dependence on d :** We vary the number of items d , while fixing $k = 5$. The results are shown in Fig. 3. Observe that the stretched-MLE has a significantly smaller bias, and performs on par with the MLE in terms of the mean squared error when d is large. Moreover, the simulations also suggest the rate of bias as of order $\frac{1}{\sqrt{d}}$ for the MLE and $\frac{1}{d}$ for the stretched-MLE, as predicted by our theoretical results.

(ii) **Dependence on k :** We vary the number of comparisons k per pair of items, while fixing $d = 10$. The results are shown in Fig. 4. As in the simulation (i) with varying d , we observe that the stretched-MLE has a significantly smaller bias, and performs on par with the MLE in terms of the mean squared error. Moreover, the simulations also suggest the rate of bias as of order $\frac{1}{\sqrt{k}}$ for the MLE and $\frac{1}{k}$ for the stretched-MLE, as predicted by our theoretical results.

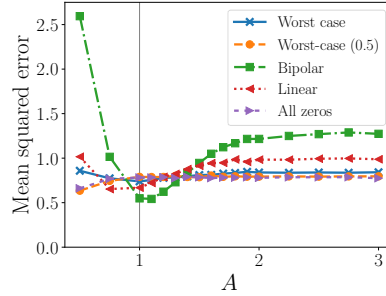
(iii) **Different settings of the true parameter θ^* :** Our theoretical result considers the worst-case bias and accuracy. In this simulation, we empirically compare the performance of the stretched-MLE under different settings of the true parameter vector θ^* (recall that setting $A = 1$ is equivalent to the standard MLE). Specifically, we consider the following values of θ^* :

- *Worst case:* $\theta^* = [1, -\frac{1}{d-1}, \dots, -\frac{1}{d-1}]$.
- *Worst case (0.5):* $\theta^* = [0.5, -\frac{0.5}{d-1}, \dots, -\frac{0.5}{d-1}]$.
- *Bipolar:* half of the values are 1, and the other half are -1 .
- *Linear:* the values are equally spaced in the interval $[-1, 1]$.
- *All zeros:* all parameters are 0.

We fix $d = 10$ and $k = 5$, varying $A \in [0.5, 3]$ under different settings of the true parameter vector θ^* . The results are shown in Fig. 5. Two high-level takeaways from the empirical evaluations are that the bias generally reduces with an increase in A till past B , and that the mean squared error remains relatively constant beyond $A = 1$ in the plotted range. In more detail, for the bias, we observe that the performance primarily depends on the largest magnitude of the items (that is, $\|\theta^*\|_\infty$). For the settings *worst case*, *bipolar* and *linear* (where $\|\theta^*\|_\infty = 1$), the bias keeps decreasing when A is past $B = 1$. For the setting *worst-case (0.5)* (where $\|\theta^*\|_\infty = 0.5$), the bias keeps decreasing when A is past 0.5. This makes sense since in this case we effectively have $B = 0.5$ (although the algorithm would not know this in practice). The bias for the setting *all zeros* stays small across values of A . For the mean squared error, the increase when A is past 1 is relatively small under most of the settings of the true parameter vector θ^* . The *bipolar* setting has the largest increase in the mean squared error. Under this setting, all parameters θ_i^* take values at the boundaries $\pm B$, and therefore the estimates of all parameters are

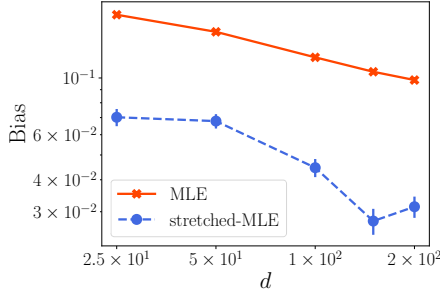


(a) Bias

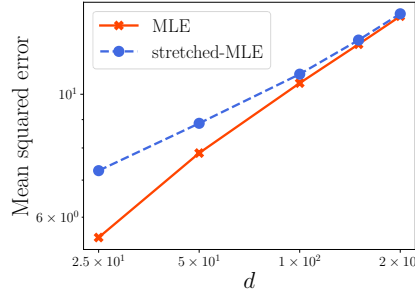


(b) Mean squared error

Figure 5: Performance of estimators for various values of A and various settings of θ^* , with $d = 10$ and $k = 5$. Setting $A = 1$ is equivalent to the standard MLE. Each point is a mean over 5000 iterations.



(a) Bias



(b) Mean squared error

Figure 6: Performance of estimators for various values of d under sparse observations, with $A = 2$. A number of $k = 5$ comparisons are observed between any pair independently with probability $p_{obs} = \frac{1}{\sqrt{d}}$ and none otherwise. Each point is a mean over 10000 iterations.

affected by the box constraint.

(iv) **Sparse observations:** So far we have considered a league format where k comparisons are observed between any pair of items. Now we consider a random-design setup, where k comparisons are observed between any pair of items independently with probability $p_{obs} \in (0, 1)$, and none otherwise (Negahban et al., 2016; Chen et al., 2019). In our simulations, we set $p_{obs} = \frac{1}{\sqrt{d}}$ and $k = 5$. We discard an iteration if the graph is not connected, since the problem is not identifiable under such a graph. The results are shown in Fig. 6. We observe that the stretched-MLE continues to outperform MLE in terms of bias, and perform on par in terms of the mean squared error.

4 Conclusions and discussions

In this work, we show that the widely-used MLE is suboptimal in terms of bias, and propose a class of estimators called the “stretched-MLE”, which provably reduces the bias while maintaining the minimax-optimality in terms of accuracy. These results on the performance of the MLE and the stretched-MLE are of both theoretical and practical interest. From the theoretical point of view, our analysis and proofs provide insights on the cause of the bias, explain why stretching the box alleviates this cause, and prove theoretical guarantees in bias reduction by stretching the box. Our results on the benefits of the stretched-MLE thus suggest theoreticians to consider the stretched-MLE

for analysis instead of the standard MLE.

From the practical point of view, when the constant B is unknown, practitioners often estimate the value of B by fitting the data or from past experience. Our results thus suggest that one should estimate B leniently, as an estimation smaller than or equal to the true B causes significant bias. Moreover, our proposed estimator is a simple modification to the MLE, which can be incorporated into any existing implementation at ease.

Our results lead to several open problems. First, it is of interest to extend our theoretical analysis to settings with sparse observations. For example, one may consider a random-design setup, where k comparisons are observed between any pair independently with probability p_{obs} and none otherwise (Negahban et al., 2016; Chen et al., 2019) (also see simulation (iv) in Section 3). In terms of the bias under this random-design setup, we think that the lower-bound for MLE and the upper-bound for our stretched-MLE depend on d and k also as $\Omega(\frac{1}{\sqrt{dk}})$ and $\tilde{O}(\frac{1}{dk})$ respectively; we think that the dependence of the stretched-MLE on p_{obs} is no worse than that of the standard MLE. Second, it is of interest to extend our results to other parametric models such as the Thurstone model (Thurstone, 1927), and other notions of fairness. Finally, in applications where the estimated parameters need to lie in a certain range (such as exam scores in between 0 and 100), it is of interest to consider how to map the estimates by the stretched-MLE to the required range.

Acknowledgements

The work of JW and NBS was supported in part by NSF grants 1755656 and 1763734. The work of RR was supported in part by U. S. Office of Naval Research award N00014-18-1-2099.

References

- Agarwal, A., Patil, P., and Agarwal, S. (2018). Accelerated spectral ranking. In *International Conference on Machine Learning*.
- Aldous, D. (2017). Elo ratings and the sports model: A neglected topic in applied probability? *Statistical Science*, 32(4):616–629.
- Anderson, J. A. and Richardson, S. C. (1979). Logistic discrimination and bias correction in maximum likelihood estimation. *Technometrics*, 21(1):71–78.
- Barocas, S., Hardt, M., and Narayanan, A. (2019). *Fairness and Machine Learning*. fairmlbook.org. <http://www.fairmlbook.org>.
- Bradley, R. A. and Terry, M. E. (1952). Rank analysis of incomplete block designs: I. the method of paired comparisons. *Biometrika*, 39(3/4):324–345.
- Chen, B., Escalera, S., Guyon, I., Ponce-López, V., Shah, N., and Simón, M. O. (2016). Overcoming calibration problems in pattern labeling with pairwise ratings: application to personality traits. In *European Conference on Computer Vision*.
- Chen, Y., Fan, J., Ma, C., and Wang, K. (2019). Spectral method and regularized MLE are both optimal for top-K ranking. *Ann. Statist.*, 47(4):2204–2235.
- Chen, Y. and Suh, C. (2015). Spectral MLE: top-K rank aggregation from pairwise comparisons. In *International Conference on Machine Learning*.
- Cox, D. R. and Snell, E. J. (1968). A general definition of residuals. *Journal of the Royal Statistical Society. Series B (Methodological)*, 30(2):248–275.
- Fan, Y., Demirkaya, E., and Lv, J. (2019). Nonuniformity of p-values can occur early in diverging dimensions. *Journal of Machine Learning Research*, 20(77):1–33.
- Firth, D. (1993). Bias reduction of maximum likelihood estimates. *Biometrika*, 80(1):27–38.
- Ford, Jr., L. R. (1957). Solution of a ranking problem from binary comparisons. *The American Mathematical Monthly*, 64(8P2):28–33.
- Gilbert, E. N. (1959). Random graphs. *The Annals of Mathematical Statistics*, 30(4):1141–1144.
- Glickman, M. E. and Doan, T. (2017). The US chess rating system. <http://www.glicko.net/ratings/rating.system.pdf> [Online; accessed May 21, 2019].
- Glickman, M. E. and Jones, A. C. (1999). Rating the chess rating system. *Chance*, 12:21–28.
- Green, P. E., Carroll, J. D., and DeSarbo, W. S. (1981). Estimating choice probabilities in multiattribute decision making. *Journal of Consumer Research*, 8(1):76–84.
- Hajek, B., Oh, S., and Xu, J. (2014). Minimax-optimal inference from partial rankings. In *Advances in Neural Information Processing Systems*.
- He, X. and Shao, Q.-M. (2000). On parameters of increasing dimensions. *Journal of Multivariate Analysis*, 73(1):120 – 135.
- Hunter, D. R. (2004). MM algorithms for generalized Bradley-Terry models. *Ann. Statist.*, 32(1):384–406.
- International Chess Federation (2017). FIDE rating regulations effective from 1 July 2017. <https://www.fide.com/fide/handbook.html?id=197&view=article> [Online; accessed May 21, 2019].
- Jang, M., Kim, S., Suh, C., and Oh, S. (2017). Optimal sample complexity of m-wise data for top-K ranking. In *Advances in Neural Information Processing Systems*.
- Király, F. J. and Qian, Z. (2017). Modelling competitive sports: Bradley-Terry-Élő models for supervised and on-line learning of paired competition outcomes. *preprint arXiv:1701.08055*.
- Lamon, A., Comroe, D., Fader, P., McCarthy, D., Ditto, R., and Huesman, D. (2016). Making WHOOPPEE: A collaborative approach to creating the modern student peer assessment ecosystem. In *EDUCAUSE*.
- Luce, R. D. (1959). *Individual Choice Behavior: A Theoretical Analysis*. Wiley, New York, NY, USA.
- Marsh, M. T. and Schilling, D. A. (1994). Equity measurement in facility location analysis: A review and framework. *European Journal of Operational Research*, 74(1):1 – 17.
- Negahban, S., Oh, S., and Shah, D. (2016). RankCentrality: Ranking from pair-wise comparisons. *Operations Research*, 65:266–287.
- Ontan, S., Synnaeve, G., Uriarte, A., Richoux, F., Churchill, D., and Preuss, M. (2013). A survey of real-time strategy game AI research and competition in StarCraft. *IEEE Transactions on Computational Intelligence and AI in Games*, 5(4):293–311.
- Ponce-López, V., Chen, B., Oliu, M., Corneanu, C., Clapés, A., Guyon, I., Baró, X., Escalante, H. J., and Escalera, S. (2016). ChaLearn LAP 2016: First round challenge on first impressions-dataset and results. In *European Conference on Computer Vision*.

- Portnoy, S. (1988). Asymptotic behavior of likelihood methods for exponential families when the number of parameters tends to infinity. *Ann. Statist.*, 16(1):356–366.
- Quenouille, M. H. (1949). Approximate tests of correlation in time-series. *Journal of the Royal Statistical Society. Series B (Methodological)*, 11(1):68–84.
- Rudin, W. (1976). *Principles of Mathematical Analysis*. McGraw-Hill.
- Salehi, F., Abbasi, E., and Hassibi, B. (2019). The impact of regularization on high-dimensional logistic regression. In *Advances in Neural Information Processing Systems 32*, pages 12005–12015. Curran Associates, Inc.
- Shah, N. B., Balakrishnan, S., Bradley, J., Parekh, A., Ramchandran, K., and Wainwright, M. J. (2016). Estimation from pairwise comparisons: Sharp minimax bounds with topology dependence. *Journal of Machine Learning Research*, 17(58):1–47.
- Shah, N. B., Bradley, J. K., Parekh, A., Wainwright, M., and Ramchandran, K. (2013). A case for ordinal peer-evaluation in MOOCs. In *NIPS Workshop on Data Driven Education*.
- Sham, P. C. and Curtis, D. (1995). An extended transmission/disequilibrium test (TDT) for multi-allele marker loci. *Annals of Human Genetics*, 59(3):323–336.
- Silver, D., Schrittwieser, J., Simonyan, K., Antonoglou, I., Huang, A., Guez, A., Hubert, T., Baker, L. R., Lai, M., Bolton, A., Chen, Y., Lillicrap, T. P., Hui, F. F. C., Sifre, L., van den Driessche, G., Graepel, T., and Hassabis, D. (2017). Mastering the game of Go without human knowledge. *Nature*, 550:354–359.
- SinceAWin.com (2019). Elo ratings - English Premier League. <https://sinceawin.com/data/elo/league/div/e0> [Online; accessed May 21, 2019].
- Stigler, S. M. (1994). Citation patterns in the journals of statistics and probability. *Statistical Science*, 9(1):94–108.
- Stigler, S. M. (1997). Regression towards the mean, historically considered. *Statistical methods in medical research*, 6(2):103–14.
- Sur, P. and Candès, E. J. (2019). A modern maximum-likelihood theory for high-dimensional logistic regression. *Proceedings of the National Academy of Sciences*, 116(29):14516–14525.
- Szörényi, B., Busa-Fekete, R., Paul, A., and Hüllermeier, E. (2015). Online rank elicitation for Plackett-Luce: A dueling bandits approach. In *Advances in Neural Information Processing Systems*.
- Thurstone, L. L. (1927). A law of comparative judgement. *Psychological Review*, 34:278–286.
- Zhao, Q., Sur, P., and Candès, E. J. (2020). The asymptotic distribution of the mle in high-dimensional logistic models: Arbitrary covariance. *preprint arXiv:2001.09351*.