
Random Utility Theory for Social Choice

Hossein Azari Soufiani
SEAS, Harvard University
azari@fas.harvard.edu

David C. Parkes
SEAS, Harvard University
parkes@eecs.harvard.edu

Lirong Xia
SEAS, Harvard University
lxia@seas.harvard.edu

Abstract

Random utility theory models an agent’s preferences on alternatives by drawing a real-valued score on each alternative (typically independently) from a parameterized distribution, and then ranking the alternatives according to scores. A special case that has received significant attention is the Plackett-Luce model, for which fast inference methods for maximum likelihood estimators are available. This paper develops conditions on general random utility models that enable fast inference within a Bayesian framework through MC-EM, providing concave log-likelihood functions and bounded sets of global maxima solutions. Results on both real-world and simulated data provide support for the scalability of the approach and capability for model selection among general random utility models including Plackett-Luce.

1 Introduction

Problems of learning with rank-based error metrics [16] and the adoption of learning for the purpose of rank aggregation in social choice [7, 8, 23, 25, 29, 30] are gaining in prominence in recent years. In part, this is due to the explosion of socio-economic platforms, where opinions of users need to be aggregated; e.g., judges in crowd-sourcing contests, ranking of movies or user-generated content.

In the problem of social choice, users submit ordinal preferences consisting of partial or total ranks on the alternatives and a single rank order must be selected to be representative of the reports. Since Condorcet [6], one approach to this problem is to formulate social choice as the problem of estimating a true underlying world state (e.g., a true quality ranking of alternatives), where the individual reports are viewed as noisy data in regard to the true state. In this way, social choice can be framed as a problem of inference. In particular, Condorcet assumed the existence of a true *ranking* over alternatives, with a voter’s preference between any pair of alternatives a, b generated to agree with the true ranking with probability $p > 1/2$ and disagree otherwise. Condorcet proposed to choose as the outcome of social choice the ranking that maximizes the likelihood of observing the voters’ preferences. Later, Kemeny’s rule was shown to provide the maximum likelihood estimator (MLE) for this model [32].

But Condorcet’s probabilistic model assumes identical and independent distributions on pairwise comparisons. This ignores the strength in agents’ preferences (the same probability p is adopted for all pairwise comparisons), and allows for cyclic preferences. In addition, computing the winner through the Kemeny rule is Θ_2^P -complete [13]. To overcome the first criticism, a more recent literature adopts the *random utility model* (RUM) from economics [26]. Consider $\mathcal{C} = \{c_1, \dots, c_m\}$ alternatives. In RUM, there is a ground truth utility (or score) associated with each alternative. These are real-valued parameters, denoted by $\vec{\theta} = (\theta_1, \dots, \theta_m)$. Given this, an agent independently samples a random utility (X_j) for each alternative c_j with conditional distribution $\mu_j(\cdot|\theta_j)$. Usually θ_j is the mean of $\mu_j(\cdot|\theta_j)$.¹ Let π denote a permutation of $\{1, \dots, m\}$, which naturally corresponds to a linear order: $[c_{\pi(1)} \succ c_{\pi(2)} \succ \dots \succ c_{\pi(m)}]$. Slightly abusing notation, we also use π to denote

¹ $\mu_j(\cdot|\theta_j)$ might be parameterized by other parameters, for example variance.

this linear order. Random utility $(X_1, \dots, X_m)$ generates a distribution on preference orders, as

$$\Pr(\pi \mid \vec{\theta}) = \Pr(X_{\pi(1)} > X_{\pi(2)} > \dots > X_{\pi(m)}) \quad (1)$$

The generative process is illustrated in Figure 1.

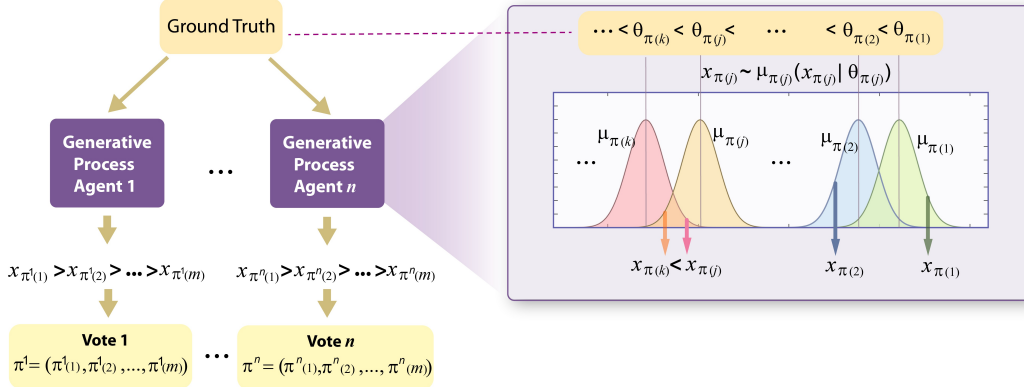


Figure 1: The generative process for RUMs.

Adopting RUMs rules out cyclic preferences, because each agent’s outcome corresponds to an order on real numbers, and it also captures the strength of preference, and thus overcomes the second criticism, by assigning a different parameter (θ_j) to each alternative.

A popular RUM is Plackett-Luce (P-L) [18, 21], where the random utility terms are generated according to Gumbel distributions with fixed shape parameter [2, 31]. For P-L, the likelihood function has a simple analytical solution, making MLE inference tractable. P-L has been extensively applied in econometrics [1, 19], and more recently in machine learning and information retrieval (see [16] for an overview). Efficient methods of EM inference [5, 14], and more recently expectation propagation [12], have been developed for P-L and its variants. In application to social choice, the P-L model has been used to analyze political elections [10]. EM algorithm has also been used to learn the *Mallows* model, which is closely related to the Condorcet’s probabilistic model [17].

Although P-L overcomes the two difficulties of the Condorcet-Kemeny approach, it is still quite restricted, by assuming that the random utility terms are distributed as Gumbel, with each alternative is characterized by one parameter, which is the mean of its corresponding distribution. In fact, little is known about inference in RUMs beyond P-L. Specifically, we are not aware of either an analytical solution or an efficient algorithm for MLE inference for one of the most natural models proposed by Thurstone [26], where each X_j is normally distributed.

1.1 Our Contributions

In this paper we focus on RUMs in which the random utilities are independently generated with respect to distributions in the *exponential family* (EF) [20]. This extends the P-L model, since the Gumbel distribution with fixed shape parameters belonging to the EF. Our main theoretical contributions are Theorem 1 and Theorem 2, which propose conditions such that the log-likelihood function is concave and the set of global maxima solutions is bounded for the *location family*, which are RUMs where the shape of each distribution μ_j is fixed and the only latent variables are the locations, i.e., the means of μ_j ’s. These results hold for existing special cases, such as the P-L model, and many other RUMs, for example the ones where each μ_j is chosen from Normal, Gumbel, Laplace and Cauchy.

We also propose a novel application of MC-EM. We treat the random utilities $(\vec{X})$ as latent variables, and adopt the Expectation Maximization (EM) method to estimate parameters $\vec{\theta}$. The E-step for this problem is not analytically tractable, and for this we adopt a Monte Carlo approximation. We establish through experiments that the Monte-Carlo error in the E-step is controllable and does not affect inference, as long as numerical parameterizations are chosen carefully. In addition, for the E-step we suggest a parallelization over the agents and alternatives and a Rao-Blackwellized method,

which further increases the scalability of our method. We generally assume that the data provides total orders on alternatives from voters, but comment on how to extend the method and theory to the case where the input preferences are *partial* orders.

We evaluate our approach on synthetic data as well as two real-world datasets, a public election dataset and one involving rank preferences on sushi. The experimental results suggest that the approach is scalable despite providing significantly improved modeling flexibility over existing approaches. For the two real-world datasets we have studied, we compare RUMs with normal distributions and P-L in terms of four criteria: log-likelihood, predictive log-likelihood, Akaike information criterion (AIC), and Bayesian information criterion (BIC). We observe that when the amount of data is not too small, RUMs with normal distributions fit better than P-L. Specifically, for the log-likelihood, predictive log-likelihood, and AIC criteria, RUMs with normal distributions outperform P-L with 95% confidence in both datasets.

2 RUMs and Exponential Families

In social choice, each agent $i \in \{1, \dots, n\}$ has a strict preference order on alternatives. This provides the data for an inferential approach to social choice. In particular, let $L(\mathcal{C})$ denote the set of all linear orders on $\mathcal{C}$. Then, a *preference-profile*, D , is a set of n preference orders, one from each agent, so that $D \in L(\mathcal{C})^n$. A *voting rule* r is a mapping that assigns to each preference-profile a set of winning rankings, $r : L(\mathcal{C})^n \mapsto (2^{L(\mathcal{C})} \setminus \emptyset)$. In particular, in the case of ties the set of winning rankings may include more than a singleton ranking.

In the maximum likelihood (MLE) approach to social choice, the preference profile is viewed as *data*, $D = \{\pi^1, \dots, \pi^n\}$. Given this, the probability (likelihood) of the data given ground truth $\vec{\theta}$ (and for a particular $\vec{\mu}$) is $\Pr(D \mid \vec{\theta}) = \prod_{i=1}^n \Pr(\pi^i \mid \vec{\theta})$, where,

$$P(\pi \mid \vec{\theta}) = \int_{x_{\pi(n)} = -\infty}^{\infty} \int_{x_{\pi(n-1)} = x_{\pi(n)}}^{\infty} \dots \int_{x_{\pi(1)} = x_{\pi(2)}}^{\infty} \mu_{\pi(n)}(x_{\pi(n)}) \dots \mu_{\pi(1)}(x_{\pi(1)}) dx_{\pi(1)} dx_{\pi(2)} \dots dx_{\pi(n)} \quad (2)$$

The MLE approach to social choice selects as the winning ranking that which corresponds to the $\vec{\theta}$ that maximizes $\Pr(D \mid \vec{\theta})$. In the case of multiple parameters that maximize the likelihood then the MLE approach returns a set of rankings, one ranking corresponding to each parameterization.

In this paper, we focus on probabilistic models where each μ_j belongs to the *exponential family* (EF). The density function for each μ in EF has the following format:

$$\Pr(X = x) = \mu(x) = e^{\eta(\theta)T(x) - A(\theta) + B(x)}, \quad (3)$$

where $\eta(\cdot)$ and $A(\cdot)$ are functions of θ , $B(\cdot)$ is a function of x , and $T(x)$ denotes the sufficient statistics for x , which could be multidimensional.

Example 1 (Plackett-Luce as an RUM [2]) In the RUM, let μ_j 's be Gumbel distributions. That is, for alternative $j \in \{1, \dots, m\}$ we have $\mu_j(x_j \mid \theta_j) = e^{-(x_j - \theta_j)} e^{-e^{-(x_j - \theta_j)}}$. Then, we have: $\Pr(\pi \mid \vec{\lambda}) = \Pr(x_{\pi(1)} > x_{\pi(2)} > \dots > x_{\pi(m)}) = \prod_{j=1}^m \frac{\lambda_{\pi(j)}}{\sum_{j'=j}^m \lambda_{\pi(j')}}$, where $\eta(\theta_j) = \lambda_j = e^{\theta_j}$, $T(x_j) = -e^{-x_j}$, $B(x_j) = -x_j$ and $A(\theta_j) = -\theta_j$. This gives us the Plackett-Luce model.

3 Global Optimality and Log-Concavity

In this section, we provide a condition on distributions that guarantees that the likelihood function (2) is log-concave in parameters $\vec{\theta}$. We also provide a condition under which the set of MLE solutions is bounded when any one latent parameter is fixed. Together, this guarantees the convergence of our MC-EM approach to a global mode with an accurate enough E-step.

We focus on the *location family*, which is a subset of RUMs where the shapes of all μ_j 's are fixed, and the only parameters are the means of the distributions. For the location family, we can write $X_j = \theta_j + \zeta_j$, where $X_j \sim \mu_j(\cdot \mid \theta_j)$ and $\zeta_j = X_j - \theta_j$ is a random variable whose mean is 0 and models an agent's *subjective noise*. The random variables ζ_j 's do not need to be identically distributed for all alternatives j ; e.g., they can be normal with different fixed variances. We focus on computing solutions ($\vec{\theta}$) to maximize the log-likelihood function,

$$l(\vec{\theta}; D) = \sum_{i=1}^n \log \Pr(\pi^i | \vec{\theta}) \quad (4)$$

Theorem 1 For the location family, if for every $j \leq m$ the probability density function for ζ_j is log-concave, then $l(\vec{\theta}; D)$ is concave.

Proof sketch: The theorem is proved by applying the following lemma, which is Theorem 9 in [22].

Lemma 1 Suppose $g_1(\vec{\theta}, \vec{\zeta}), \dots, g_R(\vec{\theta}, \vec{\zeta})$ are concave functions in $\mathbb{R}^{2m}$ where $\vec{\theta}$ is the vector of m parameters and $\vec{\zeta}$ is a vector of m real numbers that are generated according to a distribution whose pdf is logarithmic concave in $\mathbb{R}^m$. Then the following function is log-concave in $\mathbb{R}^m$.

$$L_i(\vec{\theta}, G) = \Pr(g_1(\vec{\theta}, \vec{\zeta}) \geq 0, \dots, g_R(\vec{\theta}, \vec{\zeta}) \geq 0), \quad \vec{\theta} \in \mathbb{R}^m \quad (5)$$

To apply Lemma 1, we define a set G^i of function g^i 's that is equivalent to an order π^i in the sense of inequalities implied by RUM for π^i and G^i (the joint probability in (5) for G^i to be the same as the probity of π^i in RUM with parameters $\vec{\theta}$). Suppose $g_r^i(\vec{\theta}, \vec{\zeta}) = \theta_{\pi^i(r)} + \zeta_{\pi^i(r)} - \theta_{\pi^i(r+1)} - \zeta_{\pi^i(r+1)}$ for $r = 1, \dots, m-1$. Then considering that the length of order π^i is $R+1$, we have:

$$L_i(\vec{\theta}, \pi^i) = L_i(\vec{\theta}, G^i) = \Pr(g_1^i(\vec{\theta}, \vec{\zeta}) \geq 0, \dots, g_R^i(\vec{\theta}, \vec{\zeta}) \geq 0), \quad \vec{\theta} \in \mathbb{R}^m \quad (6)$$

This is because $g_r^i(\vec{\theta}, \vec{\zeta}) \geq 0$ is equivalent to that in π^i alternative $\pi^i(r)$ is preferred to alternative $\pi^i(r+1)$ in the RUM sense.

To see how this extends to the case where preferences are specified as partial orders, we consider in particular an interpretation where an agent's report for the ranking of m_i alternatives implies that all other alternatives are worse for the agent, in some undefined order. Given this, define $g_r^i(\vec{\theta}, \vec{\zeta}) = \theta_{\pi^i(r)} + \zeta_{\pi^i(r)} - \theta_{\pi^i(r+1)} - \zeta_{\pi^i(r+1)}$ for $r = 1, \dots, m_i - 1$ and $g_r^i(\vec{\theta}, \vec{\zeta}) = \theta_{\pi^i(m_i)} + \zeta_{\pi^i(m_i)} - \theta_{\pi^i(r+1)} - \zeta_{\pi^i(r+1)}$ for $r = m_i, \dots, m-1$. Considering that $g_r^i(\cdot)$'s are linear (hence, concave) and using log concavity of the distributions of $\vec{\zeta}^i = (\zeta_1^i, \zeta_2^i, \dots, \zeta_m^i)$'s, we can apply Lemma 1 and prove log-concavity of the likelihood function. $\square$

It is not hard to verify that pdfs for normal and Gumbel are log-concave under reasonable conditions for their parameters, made explicit in the following corollary.

Corollary 1 For the location family where each ζ_j is a normal distribution with mean zero and with fixed variance, or Gumbel distribution with mean zeros and fixed shape parameter, $l(\vec{\theta}; D)$ is concave. Specifically, the log-likelihood function for P-L is concave.

The concavity of log-likelihood of P-L has been proved [9] using a different technique.

Using Fact 3.5. in [24], the set of global maxima solutions to the likelihood function, denoted by S_D , is convex since the likelihood function is log-concave. However, we also need that S_D is bounded, and would further like that it provides one unique order as the estimation for the ground truth.

For P-L, Ford, Jr. [9] proposed the following necessary and sufficient condition for the set of global maxima solutions to be bounded (more precisely, unique) when $\sum_{j=1}^m e^{\theta_j} = 1$.

Condition 1 Given the data D , in every partition of the alternatives $\mathcal{C}$ into two nonempty subsets $\mathcal{C}_1 \cup \mathcal{C}_2$, there exists $c_1 \in \mathcal{C}_1$ and $c_2 \in \mathcal{C}_2$ such that there is at least one ranking in D where $c_1 \succ c_2$.

We next show that Condition 1 is also a necessary and sufficient condition for the set of global maxima solutions S_D to be bounded in location families, when we set one of the values θ_j to be 0 (w.l.o.g., let $\theta_1 = 0$). If we do not bound any parameter, then S_D is unbounded, because for any $\vec{\theta}$, any D , and any number $s \in \mathbb{R}$, $l(\vec{\theta}; D) = l(\vec{\theta} + s; D)$.

Theorem 2 Suppose we fix $\theta_1 = 0$. Then, the set S_D of global maxima solutions to $l(\vec{\theta}; D)$ is bounded if and only if the data D satisfies Condition 1.

Proof sketch: If Condition 1 does not hold, then S_D is unbounded because the parameters for all alternatives in $\mathcal{C}_1$ can be increased simultaneously to improve the log-likelihood. For sufficiency, we first present the following lemma whose proof is omitted due to the space constraint.

Lemma 2 *If alternative j is preferred to alternative j' in at least in one ranking then the difference of their mean parameters $\theta_{j'} - \theta_j$ is bounded from above ($\exists Q$ where $\theta_{j'} - \theta_j < Q$) for all the $\vec{\theta}$ that maximize the likelihood function.*

Now consider a directed graph G_D , where the nodes are the alternatives, and there is an edge between c_j to $c_{j'}$ if in at least one ranking $c_j \succ c_{j'}$. By Condition 1, for any pair $j \neq j'$, there is a path from c_j to $c_{j'}$ (and conversely, a path from $c_{j'}$ to c_j). To see this, consider building a path between j and j' by starting from a partition with $\mathcal{C}_1 = \{j\}$ and following an edge from j to j_1 in the graph where j_1 is an alternatives in $\mathcal{C}_2$ for which there must be such an edge, by Condition 1. Consider the partition with $\mathcal{C}_1 = \{j, j_1\}$, and repeat until an edge can be followed to vertex $j' \in \mathcal{C}_2$. It follows from Lemma 2 that for any $\vec{\theta} \in S_D$ we have $|\theta_j - \theta_{j'}| < Qm$, using the telescopic sum of bounded values of the difference of mean parameters along the edges of the path, since the length of the path is no more than m (and tracing the path from j to j' and j' to j), meaning that S_D is bounded. $\square$

Now that we have the log concavity and bounded property, we need to declare conditions under which the bounded convex space of estimated parameters corresponds to a unique order. The next theorem provides a necessary and sufficient condition for all global maxima to correspond to the same order on alternatives. Suppose that we order the alternatives based on estimated θ 's (meaning that c_j is ranked higher than $c_{j'}$ iff $\theta_j > \theta_{j'}$).

Theorem 3 *The order over parameters is strict and is the same across all $\vec{\theta} \in S_D$ if, for all $\vec{\theta} \in S_D$ and all alternatives $j \neq j'$, $\theta_j \neq \theta_{j'}$.*

Proof: Suppose for the sake of contradiction there exist two maxima, $\vec{\theta}, \vec{\theta}^* \in S_D$ and a pair of alternatives $j \neq j'$ such that $\theta_j > \theta_{j'}$ and $\theta_{j'}^* > \theta_j^*$. Then, there exists an $\alpha < 1$ such that the j th and j' th components of $\alpha\vec{\theta} + (1 - \alpha)\vec{\theta}^*$ are equal, which contradicts the assumption. $\square$

Hence, if there is never a tie in the scores in any $\vec{\theta} \in S_D$, then any vector in S_D will reveal the unique order.

4 Monte Carlo EM for Parameter Estimation

In this section, we propose an MC-EM algorithm for MLE inference for RUMs where every μ_j belongs to the EF.²

The EM algorithm determines the MLE parameters $\vec{\theta}$ iteratively, and proceeds as follows. In each iteration $t + 1$, given parameters $\vec{\theta}^t$ from the previous iteration, the algorithm is composed of an E-step and an M-step. For the E-step, for any given $\vec{\theta} = (\theta_1, \dots, \theta_m)$, we compute the conditional expectation of the complete-data log-likelihood (latent variables $\vec{x}$ and data D), where the latent variables $\vec{x}$ are distributed according to data D and parameters $\vec{\theta}^t$ from the last iteration. For the M-step, we optimize $\vec{\theta}$ to maximize the expected log-likelihood computed in the E-step, and use it as the input $\vec{\theta}^{t+1}$ for the next iteration:

$$\begin{aligned} \text{E-Step : } Q(\vec{\theta}, \vec{\theta}^t) &= E_{\vec{X}} \left\{ \log \prod_{i=1}^n \Pr(\vec{x}^i, \pi^i | \vec{\theta}) \mid D, \vec{\theta}^t \right\} \\ \text{M-step : } \vec{\theta}^{t+1} &\in \arg \max_{\vec{\theta}} Q(\vec{\theta}, \vec{\theta}^t) \end{aligned}$$

4.1 Monte Carlo E-step by Gibbs sampler

The E-step can be simplified using (3) as follows:

$$\begin{aligned} E_{\vec{X}} \left\{ \log \prod_{i=1}^n \Pr(\vec{x}^i, \pi^i | \vec{\theta}) \mid D, \vec{\theta}^t \right\} &= E_{\vec{X}} \left\{ \log \prod_{i=1}^n \Pr(\vec{x}^i | \vec{\theta}) \Pr(\pi^i | \vec{x}^i) \mid D, \vec{\theta}^t \right\} \\ &= \sum_{i=1}^n \sum_{j=1}^m E_{X_j^i} \{ \log \mu_j(x_j^i | \theta_j) \mid \pi^i, \vec{\theta}^t \} = \sum_{i=1}^n \sum_{j=1}^m (\eta(\theta_j) E_{X_j^i} \{ T(x_j^i) \mid \pi^i, \vec{\theta}^t \} - A(\theta_j) + W, \end{aligned}$$

²Our algorithm can be naturally extended to compute a maximum *a posteriori* probability (MAP) estimate, when we have a prior over the parameters $\vec{\theta}$. Still, it seems hard to motivate the imposition of a prior on parameters in many social choice domains.

where $W = E_{X_j^i}\{B(x_j^i) \mid \pi^i, \bar{\theta}^t\}$ only depends on $\bar{\theta}^t$ and D (not on $\bar{\theta}^i$), which means that it can be treated as a constant in the M-step.

Hence, in the E-step we only need to compute $S_j^{i,t+1} = E_{X_j^i}\{T(x_j^i) \mid \pi^i, \bar{\theta}^t\}$ where $T(x_j^i)$ is the sufficient statistic for the parameter θ_j in the model. We are not aware of an analytical solution for $E_{X_j^i}\{T(x_j^i) \mid \pi^i, \bar{\theta}^t\}$. However, we can use a Monte Carlo approximation, which involves sampling $\bar{x}^i$ from the distribution $\Pr(\bar{x}^i \mid \pi^i, \bar{\theta}^t)$ using a Gibbs sampler, and then approximates $S_j^{i,t+1}$ by $\frac{1}{N} \sum_{k=1}^N T(x_j^{i,k})$ where N is the number of samples in the Gibbs sampler.

In each step of our Gibbs sampler for voter i , we randomly choose a position j in π^i and sample $x_{\pi^i(j)}^i$ according to a *TruncatedEF* distribution $\Pr(\cdot \mid x_{\pi^i(-j)}, \bar{\theta}^t, \pi^i)$, where $x_{\pi^i(-j)} = (x_{\pi^i(1)}, \dots, x_{\pi^i(j-1)}, x_{\pi^i(j+1)}, \dots, x_{\pi^i(m)})$. The *TruncatedEF* is obtained by truncating the tails of $\mu_{\pi^i(j)}(\cdot \mid \theta_{\pi^i(j)}^t)$ at $x_{\pi^i(j-1)}$ and $x_{\pi^i(j+1)}$, respectively. For example, a truncated normal distribution is illustrated in Figure 2.

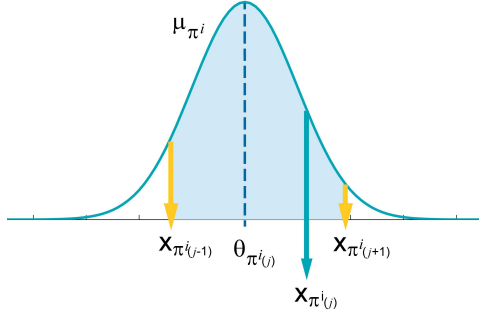


Figure 2: A truncated normal distribution.

Rao-Blackwellized: To further improve the Gibbs sampler, we use Rao-Blackwellized [4] estimation using $E\{T(x_j^{i,k}) \mid x_{-j}^{i,k}, \pi^i, \bar{\theta}^t\}$ instead of the sample $x_j^{i,k}$, where $x_{-j}^{i,k}$ is all of $\bar{x}^{i,k}$ except for $x_j^{i,k}$. Finally, we estimate $E\{T(x_j^{i,k}) \mid x_{-j}^{i,k}, \pi^i, \bar{\theta}^t\}$ in each step of the Gibbs sampler using M samples as $S_j^{i,t+1} \simeq \frac{1}{N} \sum_{k=1}^N E\{T(x_j^{i,k}) \mid x_{-j}^{i,k}, \pi^i, \bar{\theta}^t\} \simeq \frac{1}{NM} \sum_{k=1}^N \sum_{l=1}^M T(x_j^{i,l,k})$, where $x_j^{i,l,k} \sim \Pr(x_j^{i,l,k} \mid x_{-j}^{i,k}, \pi^i, \bar{\theta}^t)$. Rao-Blackwellization reduces the variance of the estimator because of conditioning and expectation in $E\{T(x_j^{i,k}) \mid x_{-j}^{i,k}, \pi^i, \bar{\theta}^t\}$.

4.2 M-step

In the E-step we have (approximately) computed $S_j^{i,t+1}$. In the M-step we compute $\bar{\theta}^{t+1}$ to maximize $\sum_{i=1}^n \sum_{j=1}^m (\eta(\theta_j) E_{X_j^i}\{T(x_j^i) \mid \pi^i, \bar{\theta}^t\} - A(\theta_j) + E_{X_j^i}\{B(x_j^i) \mid \pi^i, \bar{\theta}^t\})$. Equivalently, we compute θ_j^{t+1} for each $j \leq m$ separately to maximize $\sum_{i=1}^n \{\eta(\theta_j) E_{X_j^i}\{T(x_j^i) \mid \pi^i, \bar{\theta}^t\} - A(\theta_j)\} = \eta(\theta_j) \sum_{i=1}^n S_j^{i,t+1} - nA(\theta_j)$. For the case of the normal distribution with fixed variance, where $\eta(\theta_j) = \theta_j$ and $A(\theta_j) = (\theta_j)^2$, we have $\theta_j^{t+1} = \frac{1}{n} \sum_{i=1}^n S_j^{i,t+1}$. The algorithm is illustrated in Figure 3.

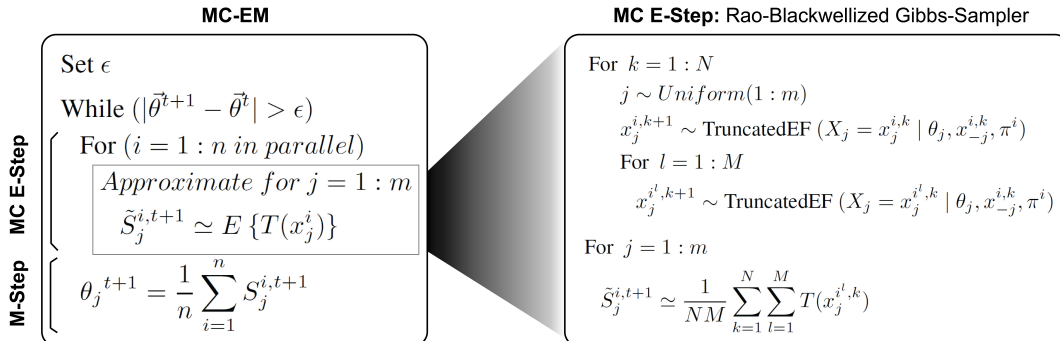


Figure 3: The MC-EM algorithm for normal distribution.

Theorem 1 and Theorem 2 guarantee the convergence of MC-EM for an exact E-step. In order to control the error of approximation in the MC-E step we can increase the number of samples with the iterations, in order to decrease the error in Monte Carlo step [28]. Details are omitted due to the space constraints and can be found in an extended version online.

5 Experimental Results

We evaluate the proposed MC-EM algorithm on synthetic data as well as two real world data sets, namely an election data set and a dataset representing preference orders on sushi. For simulated data we use the Kendall correlation [11] between two rank orders (typically between the true order and the method’s result) as a measure of performance.

5.1 Experiments for Synthetic Data

We first generate data from Normal models for the random utility terms, with means $\theta_j = j$ and equal variance for all terms, for different choices of variance ($Var = 2, 4$). We evaluate the performance of the method as the number of agents n varies. The results show that a limited number of iterations in the EM algorithm (at most 3), and samples $MN = 4000$ ($M=5, N=800$) are sufficient for inferring the order in most cases. The performance in terms of Kendall correlation for recovering ground truth improves for larger number of agents, which corresponds to more data. See Figure 4, which shows the asymptotic behavior of the maximum likelihood estimator in recovering the true parameters. Figure 4 left and middle panels show that the more the size of dataset the better the performance of the method. Moreover, for large variances in data generation, due to increasing noise in the data, the rate that performance gets better is slower than that for the case for smaller variances. Notice that the scales on the y-axis are different in the left and middle panels.

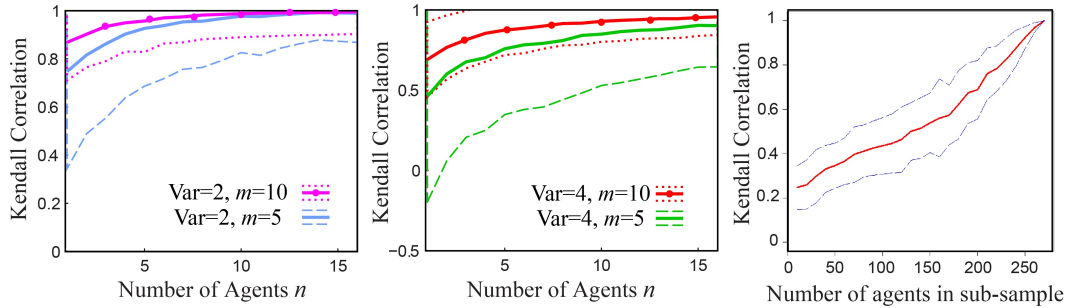


Figure 4: Left and middle panel: Performance for different number of agents n on synthetic data for $m = 5, 10$ and $Var = 2, 4$, with specifications $MN = 4000$, $EMiterations = 3$. Right panel: Performance given access to sub-samples of the data in the public election dataset, x-axis: size of sub-samples, y-axis: Kendall Correlation with the order obtained from the full data-set. Dashed lines are the 95% confidence intervals.

5.2 Experiments for Model Robustness

We apply our method to a public election dataset collected by Nicolaus Tideman [27], where the voters provided partial orders on candidates. A partial order includes comparisons among a subset of alternative, and the non-mentioned alternatives in the partial order are considered to be ranked lower than the lowest ranked alternative among mentioned alternatives.

The total number of votes are $n = 280$ and the number of alternatives $m = 15$. For the purpose of our experiments, we adopt the order on alternatives obtained by applying our method on the entire dataset as an assumed ground truth, since no ground truth is given as part of the data. After finding the ground truth by using all 280 votes (and adopting a normal model), we compare the performance of our approach as we vary the amount of data available. We evaluate the performance for sub-samples consisting of 10, 20, $\dots$, 280 of samples randomly chosen from the full dataset. For each sub-sample size, the experiment is repeated 200 times and we report the average performance and the variance. See the right panel in Figure 4. This experiment shows the robustness of the method, in the sense that the result of inference on a subset of the dataset shows consistent behavior with the case that the result on the full dataset. For example, the ranking obtained by using half of the data

can still achieve a fair estimate to the results with full data, with an average Kendall correlation of greater than 0.4.

5.3 Experiments for Model Fitness

In addition to a public election dataset, we have tested our algorithm on a sushi dataset, where 5000 users give rankings over 10 different kinds of sushi [15]. For each experiment we randomly choose $n \in \{10, 20, 30, 40, 50\}$ rankings, apply our MC-EM for RUMs with normal distributions where variances are also parameters.

In the former experiments, both the synthetic data generation and the model for election data, the variances were fixed to 1 and hence we had the theoretical guarantees for the convergence to global optimal solutions by Theorem 1 and Theorem 2. When we let the variances to be part of parametrization we lose the theoretical guarantees. However, the EM algorithm can still be applied, and since the variances are now parameters (rather than being fixed to 1), the model fits better in terms of log-likelihood.

For this reason, we adopt RUMs with normal distributions in which the variance is a parameter that is fit by EM along with the mean. We call this model a *normal model*. We compute the difference between the normal model and P-L in terms of four criteria: log-likelihood (LL), predictive log-likelihood (predictive LL), AIC, and BIC. For (predictive) log-likelihood, a positive value means that normal model fits better than P-L, whereas for AIC and BIC, a negative number means that normal model fits better than P-L. Predictive likelihood is different from likelihood in the sense that we compute the likelihood of the estimated parameters for a part of the data that is not used for parameter estimation.³ In particular, we compute predictive likelihood for a randomly chosen subset of 100 votes. The results and standard deviations for $n = 10, 50$ are summarized in Table 1.

	$n = 10$				$n = 50$			
Dataset	LL	Pred. LL	AIC	BIC	LL	Pred. LL	AIC	BIC
Sushi	8.8(4.2)	-56.1(89.5)	-7.6(8.4)	5.4(8.4)	22.6(6.3)	40.1(5.1)	-35.2(12.6)	-6.1(12.6)
Election	9.4(10.6)	91.3(103.8)	-8.8(21.2)	4.2(21.2)	44.8(15.8)	87.4(30.5)	-79.6(31.6)	-50.5(31.6)

Table 1: Model selection for the sushi dataset and election dataset. Cases where the normal model fits better than P-L statistically with 95% confidence are in bold.

When n is small ($n = 10$), the variance is high and we are unable to obtain statistically significant results in comparing fitness. When n is not too small ($n = 50$), RUMs with normal distributions fit better than P-L. Specifically, for log-likelihood, predictive log-likelihood, and AIC, RUMs with normal distributions outperform P-L with 95% confidence in both datasets.

5.4 Implementation and Run Time

The running time for our MC-EM algorithm scales linearly with number of agents on real world data (Election Data) with slope 13.3 second per agent on an Intel i5 2.70GHz PC. This is for 100 iterations of EM algorithm with Gibbs sampling number increasing with iterations as $2000 + 300 * \text{iteration steps}$.

Acknowledgments

This work is supported in part by NSF Grant No. CCF- 0915016. Lirong Xia is supported by NSF under Grant #1136996 to the Computing Research Association for the CIFellows Project. We thank Craig Boutilier, Jonathan Huang, Tyler Lu, Nicolaus Tideman, Paolo Viappiani, and anonymous NIPS-12 reviewers for helpful comments and suggestions, or help on the datasets.

References

- [1] Steven Berry, James Levinsohn, and Ariel Pakes. Automobile prices in market equilibrium. *Econometrica*, 63(4):841–890, 1995.
- [2] Henry David Block and Jacob Marschak. Random orderings and stochastic theories of responses. In *Contributions to Probability and Statistics*, pages 97–132, 1960.

³The use of predictive likelihood allows us to evaluate the performance of the estimated parameters on the rest of the data, and is similar in this sense to cross validation for supervised learning.

- [3] James G. Booth and James P. Hobert. Maximizing Generalized Linear Mixed Model Likelihoods with an Automated Monte Carlo EM Algorithm. *JRSS. Series B*, 61(1):265–285, 1999.
- [4] Steve Brooks, Andrew Gelman, Galin Jones, and Xiao-Li Meng, editors. *Handbook of Markov Chain Monte Carlo*. Chapman and Hall/CRC, 2011.
- [5] Francois Caron and Arnaud Doucet. Efficient Bayesian Inference for Generalized Bradley-Terry Models. *Journal of Computational and Graphical Statistics*, 21(1):174–196, 2012.
- [6] Marquis de Condorcet. *Essai sur l'application de l'analyse à la probabilité des décisions rendues à la pluralité des voix*. Paris: L'Imprimerie Royale, 1785.
- [7] Vincent Conitzer, Matthew Rognlie, and Lirong Xia. Preference functions that score rankings and maximum likelihood estimation. In *Proc. IJCAI*, pages 109–115, 2009.
- [8] Vincent Conitzer and Tuomas Sandholm. Common voting rules as maximum likelihood estimators. In *Proc. UAI*, pages 145–152, 2005.
- [9] Lester R. Ford, Jr. Solution of a ranking problem from binary comparisons. *The American Mathematical Monthly*, 64(8):28–33, 1957.
- [10] Isobel Claire Gormley and Thomas Brendan Murphy. A grade of membership model for rank data. *Bayesian Analysis*, 4(2):265–296, 2009.
- [11] Przemyslaw Grzegorzewski. Kendall's correlation coefficient for vague preferences. *Soft Computing*, 13(11):1055–1061, 2009.
- [12] John Guiver and Edward Snelson. Bayesian inference for Plackett-Luce ranking models. In *Proc. ICML*, pages 377–384, 2009.
- [13] Edith Hemaspaandra, Holger Spakowski, and Jörg Vogel. The complexity of Kemeny elections. *Theoretical Computer Science*, 349(3):382–391, December 2005.
- [14] David R. Hunter. MM algorithms for generalized Bradley-Terry models. In *The Annals of Statistics*, volume 32, pages 384–406, 2004.
- [15] Toshihiro Kamishima. Nantonac collaborative filtering: Recommendation based on order responses. In *Proc. KDD*, pages 583–588, 2003.
- [16] Tie-Yan Liu. *Learning to Rank for Information Retrieval*. Springer, 2011.
- [17] Tyler Lu and Craig Boutilier. Learning mallows models with pairwise preferences. In *Proc. ICML*, pages 145–152, 2011.
- [18] R. Duncan Luce. *Individual Choice Behavior: A Theoretical Analysis*. Wiley, 1959.
- [19] Daniel McFadden. Conditional logit analysis of qualitative choice behavior. In *Frontiers of Econometrics*, pages 105–142, New York, NY, 1974. Academic Press.
- [20] Carl N. Morris. Natural Exponential Families with Quadratic Variance Functions. *Annals of Statistics*, 10(1):65–80, 1982.
- [21] R. L. Plackett. The analysis of permutations. *JRSS. Series C*, 24(2):193–202, 1975.
- [22] Andr s Pr kopa. Logarithmic concave measures and related topics. In *Stochastic Programming*, pages 63–82. Academic Press, 1980.
- [23] Ariel D. Procaccia, Sashank J. Reddi, and Nisarg Shah. A maximum likelihood approach for selecting sets of alternatives. In *Proc. UAI*, 2012.
- [24] Frank Proschan and Yung L. Tong. Chapter 29. log-concavity property of probability measures. *FSU technical report Number M-805*, pages 57–68, 1989.
- [25] Magnus Roos, J rg Rothe, and Bj rn Scheuermann. How to calibrate the scores of biased reviewers by quadratic programming. In *Proc. AAAI*, pages 255–260, 2011.
- [26] Louis Leon Thurstone. A law of comparative judgement. *Psychological Review*, 34(4):273–286, 1927.
- [27] Nicolaus Tideman. *Collective Decisions and Voting: The Potential for Public Choice*. Ashgate Publishing, 2006.
- [28] Greg C. G. Wei and Martin A. Tanner. A Monte Carlo Implementation of the EM Algorithm and the Poor Man's Data Augmentation Algorithms. *JASA*, 85(411):699–704, 1990.
- [29] Lirong Xia and Vincent Conitzer. A maximum likelihood approach towards aggregating partial orders. In *Proc. IJCAI*, pages 446–451, 2011.
- [30] Lirong Xia, Vincent Conitzer, and J r me Lang. Aggregating preferences in multi-issue domains by using maximum likelihood estimators. In *Proc. AAMAS*, pages 399–406, 2010.
- [31] John I. Jr. Yellott. The relationship between Luce's Choice Axiom, Thurstone's Theory of Comparative Judgment, and the double exponential distribution. *J. of Mathematical Psychology*, 15(2):109–144, 1977.
- [32] H. Peyton Young. Optimal voting rules. *Journal of Economic Perspectives*, 9(1):51–64, 1995.