
Fairness in Ranking under Uncertainty

Ashudeep Singh
Cornell University
Ithaca, NY 14850
ashudeep@cs.cornell.edu

David Kempe
University of Southern California
Los Angeles, CA 90089
david.m.kempe@gmail.com

Thorsten Joachims
Cornell University
Ithaca, NY 14850
tj@cs.cornell.edu

Abstract

Fairness has emerged as an important consideration in algorithmic decision making. Unfairness occurs when an agent with higher merit obtains a worse outcome than an agent with lower merit. Our central point is that a primary cause of unfairness is uncertainty. A principal or algorithm making decisions never has access to the agents' true merit, and instead uses proxy features that only imperfectly predict merit (e.g., GPA, star ratings, recommendation letters). None of these ever fully capture an agent's merit; yet existing approaches have mostly been defining fairness notions directly based on observed features and outcomes.

Our primary point is that it is more principled to acknowledge and model the uncertainty explicitly. The role of observed features is to give rise to a posterior distribution of the agents' merits. We use this viewpoint to define a notion of approximate fairness in ranking. We call an algorithm ϕ -fair (for $\phi \in [0, 1]$) if it has the following property for all agents x and all k : if agent x is among the top k agents with respect to *merit* with probability at least ρ (according to the posterior merit distribution), then the algorithm places the agent among the top k agents in its *ranking* with probability at least $\phi\rho$.

We show how to compute rankings that optimally trade off approximate fairness against utility to the principal. In addition to the theoretical characterization, we present an empirical analysis of the potential impact of the approach in simulation studies. For real-world validation, we applied the approach in the context of a paper recommendation system that we built and fielded at the KDD 2020 conference.

1 Introduction

Fairness is an important consideration in decision-making, in particular when a limited resource must be allocated among multiple agents by a principal (or decision maker). A widely accepted tenet of fairness is that if an agent B does not have stronger merits for the resource than A, then B should not get more of the resource than A. Depending on the context, *merit* could be a qualification (e.g., job performance), a need (e.g., disaster relief), or some other measure of eligibility.

The motivation for our work is that uncertainty about merits is a primary reason that a principal's allocations can violate this tenet and thereby lead to unfair outcomes. Were agents' merits fully observable, it would be both fair and in the principal's best interest to rank agents by their merit. However, actual merits are practically always unobservable. Consider the following standard algorithmic decision making environments: (1) An e-commerce or recommender platform (the principal) displays items (the agents) in response to a user query. An item's merit is the utility the user would derive from it, whereas the platform can only observe numerical ratings, text reviews, the user's past history, and similar features. (2) A job recommendation site or employer (the principal) wants to recommend/hire one or more applicants (the agents). The merit of an applicant is her (future) performance on the job over a period of time, whereas the site or employer can only observe (past) grades, test scores, recommendation letters, performance in interviews, and similar assessments.

In both of these examples — and essentially all others in which algorithms are called upon to make allocation decisions between agents — uncertainty about merit is unavoidable, and arises from multiple sources: (1) the training data of a machine learning algorithm is a random sample, (2) the features themselves often come from a random process, and (3) the merit itself may only be revealed in the future after a random process (e.g., whether an item is sold or an employee performs well). Given that decisions will be made in the presence of uncertainty, it is important to define the notion of *fairness* under uncertainty. Extending the aforementioned tenet that “if agent B has less merit than A, then B should not be treated better than A,” we state the following generalization to uncertainty about merit, first for just two agents:

Axiom 1. If A has merit greater than or equal to B with probability at least ρ , then a fair policy should treat A at least as well as B with probability at least ρ .

This being an axiom, we cannot offer a mathematical justification. It captures an inherent sense of fairness in the absence of enough information, and it converges to the conventional tenet as uncertainty is reduced. In particular, consider the following situation: two items A, B with 10 reviews each have average star ratings of 3.9 and 3.8, respectively; or two job applicants A, B have GPAs of 3.9 and 3.8. While this constitutes some (weak) evidence that A may have more merit than B, this evidence leaves substantial uncertainty. The posterior merit distributions based on the available information should reflect this uncertainty by having non-trivial variance; our axiom then implies that A and B must be treated similarly to achieve fairness. In particular, it would be highly unfair to *deterministically* rank A ahead of B (or vice versa). Our main point is that this uncertainty, rather than the specific numerical values of 3.9 and 3.8, is the reason why a mechanism should treat A and B similarly.

1.1 Our Contributions

We study fairness in the presence of uncertainty specifically for the generalization where the principal must rank n items. Our main contribution is the fairness framework, giving definitions of fairness in ranking in the presence of uncertainty. This framework, including extensions to approximate notions of fairness, is presented and discussed in depth in § 2. We believe that uncertainty of merit is one of the most important sources of unfairness, and modeling it explicitly and axiomatically is key to addressing it.

Next, in § 3, we present algorithms for a principal to achieve (approximately) fair ranking distributions. A simple algorithm the principal may use to achieve approximate fairness is to mix between an optimal (unfair) ranking and (perfectly fair) Thompson sampling. We show that this policy is not optimal for the principal’s utility, and we present an efficient LP-based algorithm that achieves an optimal ranking distribution for the principal, subject to an approximate fairness constraint.

We next explore empirically to what extent a focus on fairness towards the agents reduces the principal’s utility. We do so with two extensive sets of experiments: one described in § 4 on existing data, and one described in § 5 “in the wild.” In the first set of experiments, we consider movie recommendations based on the standard MovieLens dataset and investigate to what extent fairness towards movies would result in lower utility for users of the system. The second experiment was carried out at the 2020 ACM SIGKDD Conference on Knowledge Discovery and Data Mining, where we implemented and fielded a paper recommendation system. Half of the conference attendees using the system received rankings that were modified to ensure greater fairness towards papers, and we report on various metrics that capture the level of engagement of conference participants based on which group they were assigned to.

The upshot of our experiments and theoretical analysis is that in the settings we have studied, high levels of fairness can be achieved at a small loss in utility for the principal and the system’s users.

Due to space constraints, essentially all proofs, as well as a detailed discussion of related work, are given in the supplementary material. Alternatively, a reader may want to read the full version (Singh et al., 2021).

2 Ranking with Uncertain Merits

We are interested in ranking policies for a principal (the ranking system, such as an e-commerce platform or a job portal in our earlier examples) whose goal is to rank a set $\mathcal{X}$ of n agents (such as

products or applicants). The principal observes some evidence for the merit of the agents, and must produce a distribution over rankings trading off fairness to the agents against the principal’s utility. For the agents, a higher rank is always more desirable than a lower rank.

2.1 Rankings and Ranking Distributions

We use $\Sigma(\mathcal{X})$ to denote the set of all $n!$ rankings, and $\Pi(\mathcal{X})$ for the set of all distributions over $\Sigma(\mathcal{X})$. We express a ranking $\sigma \in \Sigma(\mathcal{X})$ in terms of the agents assigned to given positions, i.e., $\sigma(k)$ is the agent in position k . A ranking distribution $\pi \in \Pi(\mathcal{X})$ can be represented by the $n!$ probabilities $\pi(\sigma)$ of the rankings $\sigma \in \Sigma(\mathcal{X})$. However, all the information relevant for our purposes can be represented more compactly using the *Marginal Rank Distribution*: we write $p_{x,k}^{(\pi)} = \sum_{\sigma: \sigma(k)=x} \pi(\sigma)$ for the probability under π that agent $x \in \mathcal{X}$ is in position k in the ranking. We let $\mathcal{P}^{(\pi)} = (p_{x,k}^{(\pi)})_{x,k}$ denote the $n \times n$ matrix of all marginal rank probabilities.

The matrix $\mathcal{P}^{(\pi)}$ is doubly stochastic, i.e., the sum of each row and column is 1. While π uniquely defines $\mathcal{P}^{(\pi)}$, the converse mapping may not be unique. However, given a doubly stochastic matrix $\mathcal{P}$, the Birkhoff-von Neumann decomposition (Birkhoff, 1946) can be used to compute *some* distribution π consistent with $\mathcal{P}$, i.e., $\mathcal{P}^{(\pi)} = \mathcal{P}$; any consistent distribution π will suffice for our purposes.

2.2 Merit, Uncertainty, and Fairness

The principal must determine a distribution over rankings of the agents. This distribution will be based on some evidence for the agents’ merits. This evidence could take the form of star ratings and reviews of products (combined with the site visitor’s partially known preferences), or GPA, test scores, and recommendation letters of an applicant. Our main departure from past work on individual fairness (following (Dwork et al., 2012)) is that we do not view this evidence as having inherent meaning; rather, its sole role is to induce a posterior joint distribution over the agents’ merits.

The merit of agent x is $v_x \in \mathbb{R}$, and we write $\mathbf{v} = (v_x)_{x \in \mathcal{X}}$ for the vector of all agents’ merits. Based on all observed evidence, the principal can infer a distribution Γ over agents’ merits using any suitable Bayesian inference procedure. Since the particular Bayesian model depends on the application, for our purposes, we merely assume that a posterior distribution Γ was inferred using best practices and that ideally, this model is open to verification and audit.

We write $\Gamma(\mathbf{v})$ for the probability of merits $\mathbf{v}$ under Γ . We emphasize that the distribution will typically not be independent over entries of $\mathbf{v}$ — for example, students’ merit conditioned on observed grades will be correlated via common grade inflation if they took the same class. To avoid awkward tie-breaking issues, we assume that $v_x \neq v_y$ for all distinct $x, y \in \mathcal{X}$ and all $\mathbf{v}$ in the support of Γ . This side-steps having to define the notion of top- k lists with ties, and comes at little cost in expressivity, as any tie-breaking would typically be encoded in slight perturbations to the v_x anyway.

We write $\mathcal{M}_{x,k}^{(\mathbf{v})}$ for the event that under $\mathbf{v}$, agent x is among the top k agents with respect to merit, i.e., that $|\{x' \mid v_{x'} > v_x\}| < k$. We now come to our key definition of approximate fairness.

Definition 2.1 (Approximately Fair Ranking Distribution). A ranking distribution π is ϕ -fair iff

$$\sum_{k'=1}^k p_{x,k'}^{(\pi)} \geq \phi \cdot \mathbb{P}_{\mathbf{v} \sim \Gamma} [\mathcal{M}_{x,k}^{(\mathbf{v})}] \quad (1)$$

for all agents x and positions k . That is, the ranking distribution π ranks x at position k or above with at least a ϕ fraction of the probability that x is actually among the top k agents according to Γ . Furthermore, π is *fair* iff it is 1-fair.

The reason for defining ϕ -approximately fair ranking distributions (rather than just fair distributions) is that fairness typically comes at a cost to the principal (such as lower expected clickthrough or lower expected performance of recommended employees). For example, if the v_x are probabilities that a user will purchase products on an e-commerce site, then deterministically ranking by decreasing $\mathbb{E}_{\Gamma} [v_x]$ is the principal’s optimal ranking under common assumptions about user behavior; yet, being deterministic, it is highly unfair. Our definition of approximate fairness allows, e.g., a policymaker to choose a trade-off regarding how much fairness (with resulting expected utility loss) to require from the principal. Notice that for $\phi = 0$, the principal is unconstrained.

We remark that the merit values v_x only matter insofar as comparison is concerned; in other words, they are used ordinally, not cardinally. This is captured by the following proposition.

Proposition 2.1. Let $f : \mathbb{R} \rightarrow \mathbb{R}$ be any strictly increasing function. Let Γ' be the distribution that draws the vector $(f(v_x))_{x \in \mathcal{X}}$ with probability $\Gamma(\mathbf{v})$ for all $\mathbf{v}$; that is, it replaces each entry v_x with $f(v_x)$. Then, a ranking distribution π is ϕ -fair with respect to Γ if and only if it is ϕ -fair with respect to Γ' .

Proof. Because f is strictly increasing, $\mathbb{P}_{\mathbf{v} \sim \Gamma} [\mathcal{M}_{x,k}^{(\mathbf{v})}] = \mathbb{P}_{\mathbf{v} \sim \Gamma'} [\mathcal{M}_{x,k}^{(\mathbf{v})}]$ for all x and k . This immediately implies the claim by examining Definition 2.1. $\square$

Proposition 2.1 highlights a key aspect of our fairness definition: we avoid expressing any notion of fairness when “one agent has just ‘a little’ more merit than the other,” instead arguing that fairness is only truly violated when an agent with more merit is treated worse than one with less merit. In other words, fairness is inherently *ordinal* in our treatment. This viewpoint has implications for a principal seeking “high-risk high-reward” agents, which we discuss in more depth in § B.2.

2.3 The Principal’s Utility

The principal’s utility can be the profit of an e-commerce site or the satisfaction of its customers. We assume that the utility for a ranking σ with agent merits $\mathbf{v}$ takes the form $U(\sigma | \mathbf{v}) = \sum_{k=1}^n w_k v_{\sigma(k)}$, where w_k is the position weight for position k in the ranking, and we assume that the w_k are non-increasing, i.e., the principal derives the most utility from earlier positions of the ranking.

The assumption that the utility from each position is factorable (i.e., of the form $w_k \cdot v_{\sigma(k)}$) is quite standard in the literature (Järvelin and Kekäläinen, 2002; Taylor et al., 2008). The assumption that the utility is *linear* in $v_{\sigma(k)}$ is in fact not restrictive at all. To see this, assume that the principal’s utility were of the form $w_k \cdot f(v_{\sigma(k)})$ for some strictly increasing function f . By Proposition 2.1, the exact same fairness guarantees are achieved when the agents’ merits v_x are replaced with $f(v_x)$; doing so preserves fairness, while in fact making the principal’s utility linear in the merits. Some common examples of utility functions falling into this general framework are DCG with $w_k = 1/\log_2(1+k)$, Average Reciprocal Rank with $w_k = 1/k$, and Precision@K with $w_k = \mathbb{1}[k \leq K]/K$.

When the ranking and merits are drawn from distributions, the principal’s utility is the expected utility under both sources of randomness:

$$U(\pi | \Gamma) = \mathbb{E}_{\sigma \sim \pi, \mathbf{v} \sim \Gamma} [U(\sigma | \mathbf{v})]. \quad (2)$$

2.4 Discussion

Our definition is superficially similar to existing definitions of individual fairness (e.g., (Dwork et al., 2012; Joseph et al., 2016)), in that similar observable features often lead to similar outcomes. Importantly, though, it side-steps the need to define a similarity metric between agents in the feature space. Furthermore, it does not treat the observable attributes (such as star ratings) themselves as any notion of “merit.” Instead, our central point is that agents’ features should be viewed *solely* as noisy signals about the agents’ merits and that a comparison of their merits — and the principal’s uncertainty about the merits — should determine the agents’ relative ranking. That moves the key task of quantifying individual fairness from articulating which features should be considered relevant for similarity, to articulating what inferences can be drawn about merit from observed features.

Merit as an Abstraction Boundary between Data and Fairness. One may argue, rightfully, that from an operational perspective, our approach simply pushes the normative decisions into determining Γ . For example, if the distribution Γ were biased in favor of or against a particular group, then the decisions of a supposedly fair algorithm (with respect to Γ) would in fact be unfair to that group. However, our main point is that normative decisions *should* indeed be encoded in the distribution Γ . To appreciate the conceptual approach, first notice that any algorithm implicitly encodes normative decisions, merely in the outputs it produces, which will favor some agents over others. The key question is how these normative decisions are encoded, how they can be articulated, and whether they could possibly be audited. If they are encoded in ad hoc algorithmic choices, articulating and auditing them may be difficult.

As an example, consider an admissions officer at a university, who believes that the GPA or SAT scores of affluent applicants may be higher due to access to tutors, rather than true academic potential. One approach to compensate for this advantage could be to subtract some (wealth-dependent) amount from an applicant’s scores. A more principled approach — and the one we advocate — is for the admissions officer to explicitly express the possible distribution of merits given the test scores and wealth. The suitable notion of fairness is then derived by our framework, rather than as an ad hoc choice. A substantive discussion can then be had around the assumptions that go into the admission officer’s particular choice of distribution, whether a different distribution would be more suitable, etc.

In a sense, the notion of merit, and distributions thereof, serves as a clean abstraction boundary between available data, and the desired fairness and utility. We argue that frequently, the difficult question to address is not so much what is “fair,” but what the data truly reveal about an agent’s merit. The latter should be articulated by domain experts, whereas the role of computer science is to provide frameworks for deriving fair algorithms *given* the merit distributions, as well as statistical approaches that may guide the derivation of Γ from data.

Randomization, Fairness, and Single-Shot Scenarios. In the introduction, we discussed two possible applications in which fairness is desirable: ranking of products in online e-commerce, and ranking of job applicants. We note that these two settings differ along an important dimension: e-commerce sites typically display/rank the same set of products many times over a short period of time, and the stakes each time are fairly low. On the other hand, any one particular job is a one-shot setting with high stakes. Intuitively, it “feels” like the use of randomization as a means to achieve fairness is more natural in the former setting than the latter. We discuss this issue in more depth.

First, we consider two practical reasons for randomization being more natural in repeated low-stakes settings: (1) In a high-stakes situation, a principal may be less willing to trade off utility for fairness, and (2) If fairness is *required* of the principal (rather than the principal’s own goal), in a single-shot setting, it is much more difficult to verify that a decision was indeed made probabilistically; in contrast, for a repeated setting, statistical tools can be employed to keep the principal honest.

More fundamentally, the two settings differ in the point in time at which fairness is guaranteed. For concreteness, consider the simplest setting: two agents with identical posterior distributions vie for one position. In this case, a coin flip is *ex ante* fair: before the coin flip is realized, both agents have the same probability of being selected. However, it is not fair *ex post*: despite both having equal merit, one was selected, and the other was not. Contrast this with the alternative in which the same two agents compete multiple times, and a coin is flipped each time. Ex ante fairness is of course preserved, but even ex post, each agent was selected approximately the same number of times. This example may explain why randomization “feels” more fair for repeated than one-shot settings.

The fact that for a one-shot setting, a coin flip is only ex ante fair, however, does not obviate the need for making fair decisions in one-shot settings. There will be situations in which a principal is faced with multiple essentially indistinguishable agents and not enough positions for all of them.¹ While ex ante fairness may not be completely satisfactory, it still guarantees “more” fairness than arbitrary deterministic tie-breaking. Indeed, one may argue that the goal of many principals is not so much to make fair decisions as to make defensible ones. For example, if applicants are ranked strictly by GPA, choosing an applicant with GPA of 3.91 over one with GPA of 3.90 is essentially random tie-breaking, but with a rule that can be defended. Our point here is that randomization should be considered as a viable alternative, if the true goal is to achieve fairness.

Other Considerations. In extending the probabilistic fairness axiom from two to multiple agents in Equation (1), we chose to axiomatize fairness in terms of which position agents are assigned to. An equally valid generalization would have been to require for each pair x, y of agents that if x has more merit than y with probability at least ρ , then x must precede y in the ranking with probability at least $\phi \cdot \rho$. The main reason why we prefer Equation (1) is computational: the only linear programs we know for the alternative approach require variables for all rankings and are thus exponential (in n) in size. Exploring the alternative definition is an interesting direction for future work.

Due to space constraints, additional properties of our fairness definition are discussed in § B of the supplementary material. In particular, we discuss how fairness requirements give the principal

¹For example, in college admissions, qualified applicants typically outnumber available slots, and differences among the qualified applicants are frequently very small.

stronger incentives for obtaining more accurate estimates of agents' merits, and how to align the ordinal nature of our definition with a principal's interest in selecting high-risk high-reward agents.

3 Optimal and Fair Policies

For a distribution Γ over merits, let σ_Γ^* be the ranking which sorts the agents by expected merit, i.e., by non-increasing $\mathbb{E}_{v \sim \Gamma} [v_x]$. The following well-known proposition follows because the position weights w_k are non-increasing.

Proposition 3.1. σ_Γ^* is a utility-maximizing ranking policy for the principal.

If the principal's expected utility can be evaluated efficiently, computing σ_Γ^* only requires sorting the agents by utility, and thus takes time only $O(n \log n)$. While this policy conforms to the Probability Ranking Principle (Robertson, 1977), it violates Axiom 1 for ranking fairness when merits are uncertain. We define a natural solution for a 1-fair ranking distribution based on Thompson Sampling:

Definition 3.1 (Thompson Sampling Ranking Distribution). Define π_Γ^{TS} as follows: first, draw a vector of merits $v \sim \Gamma$, then rank the agents by decreasing merits in v .

That π_Γ^{TS} is 1-fair follows directly from the definition of fairness. By definition, it ranks each agent x in position k with exactly the probability that x has k -th highest merit.

Proposition 3.2. π_Γ^{TS} is a 1-fair ranking distribution.

Furthermore, computing π_Γ^{TS} only involves sampling from Γ and then sorting the agents by merit, so it can be efficiently performed in time $O(n \log n)$.

3.1 Trading Off Utility and Fairness

One straightforward way of trading off between the two objectives of fairness and principal's utility is to randomize between the two policies π_Γ^{TS} and π_Γ^* .

Definition 3.2 (OPT/TS-Mixing). The OPT/TS-Mixing ranking policy $\pi^{\text{Mix}, \phi}$ randomizes between π_Γ^{TS} and π_Γ^* with probabilities ϕ and $1 - \phi$, respectively.

This policy inherits a runtime of $O(n \log n)$ from π_Γ^{TS} and π_Γ^* . The following lemma gives guarantees for such randomization (but we will later see that this strategy is suboptimal).

Lemma 3.1. Consider two ranking policies π_1 and π_2 such that π_1 is ϕ_1 -fair and π_2 is ϕ_2 -fair. A policy that randomizes between π_1 and π_2 with probabilities q and $1 - q$, respectively, is at least $(q\phi_1 + (1 - q)\phi_2)$ -fair and obtains expected utility $qU(\pi_1 | \Gamma) + (1 - q)U(\pi_2 | \Gamma)$.

Corollary 3.1. The ranking policy $\pi^{\text{Mix}, \phi}$ is ϕ -fair.

By definition, $\pi^{\text{Mix}, \phi=0}$ has the highest utility among all 0-fair ranking policies. Furthermore, all 1-fair policies achieve the same utility since the fairness axiom for $\phi = 1$ completely determines the marginal rank probabilities.

Lemma 3.2. All 1-fair ranking policies have the same utility for the principal.

While $\pi^{\text{Mix}, \phi=0}$ and $\pi^{\text{Mix}, \phi=1}$ have the highest utility among 0-fair and 1-fair ranking policies, respectively, $\pi^{\text{Mix}, \phi}$ will typically not have maximum utility among all ϕ -fair ranking policies for other values of $\phi \in (0, 1)$. We illustrate this with a concrete example with $n = 3$ agents in § D.

3.2 Optimizing Utility for ϕ -Fair Rankings

We now formulate a linear program for computing the policy $\pi^{\text{LP}, \phi}$ that maximizes the principal's utility, subject to being ϕ -fair. The variables of the linear program are the marginal rank probabilities $p_{x,k}^{(\pi)}$ of the distribution π to be determined. Then, by Equation (2) and linearity of expectation, the principal's expected utility can be written as $U(\pi | \Gamma) = \sum_{x \in \mathcal{X}} \sum_k p_{x,k}^{(\pi)} \cdot \mathbb{E}_{v \sim \Gamma} [v_x] \cdot w_k$. We use this linear form of the utilities to write the optimization problem as the following LP with variables

$p_{x,k}$ (omitting π from the notation):

$$\begin{aligned}
& \text{Maximize} && \sum_x \sum_k p_{x,k} \cdot \mathbb{E}_{v \sim \Gamma} [v_x] \cdot w_k \\
& \text{subject to} && \sum_{k'=1}^k p_{x,k'} \geq \phi \cdot \mathbb{P}_{v \sim \Gamma} [\mathcal{M}_{x,k}^{(v)}] \quad \text{for all } x, k \\
& && \sum_{k=1}^n p_{x,k} = 1 \quad \text{for all } x \\
& && \sum_x p_{x,k} = 1 \quad \text{for all } k \\
& && 0 \leq p_{x,k} \leq 1 \quad \text{for all } x, k.
\end{aligned} \tag{3}$$

In the LP, the first set of constraints captures ϕ -approximate fairness for all agents and positions, while the remaining constraints ensure that the marginal probabilities form a doubly stochastic matrix.

As a second step, the algorithm uses the Birkhoff-von Neumann (BvN) Decomposition of the matrix $\mathcal{P} = (p_{x,k})_{x,k}$ to explicitly obtain a distribution π over rankings such that π has marginals $p_{x,k}$. The Birkhoff-von Neumann Theorem (Birkhoff, 1946) states that the set of doubly stochastic matrices is the convex hull of the permutation matrices, which means that we can write $\mathcal{P} = \sum_{\sigma} q_{\sigma} \mathcal{P}^{(\sigma)}$, where $\mathcal{P}^{(\sigma)}$ is the binary permutation matrix corresponding to the deterministic ranking σ , and the q_{σ} form a probability distribution. It was already shown by Birkhoff (1946) how to find a polynomially sparse decomposition in polynomial time.

Having to solve a linear program obviously makes the computation of $\pi^{\text{LP}, \phi}$ less efficient. It is still efficient enough to be feasible for several hundred agents. An interesting direction for future work would be whether the specific LP can be solved more efficiently, either exactly or approximately, by using algorithms other than the standard ones (Ellipsoid or Interior Point Methods).

In order to solve the Linear Program (3), one needs to know $\mathbb{P}_{v \sim \Gamma} [\mathcal{M}_{x,k}^{(v)}]$ for all i and k . For some distributions Γ (e.g., Example 2), these quantities can be calculated in closed form. For others, they can be approximated using Monte Carlo sampling. Small approximation errors only have a small impact on the final solution as captured by Propositions C.1 and C.2 in the appendix.

4 Experimental Evaluation: MovieLens Dataset

To evaluate our approach in a recommendation setting with a realistic preference distribution, we designed the following experimental setup based on the MovieLens 100K (ML-100K) dataset. The dataset contains 100,000 ratings, by 600 users, on 9,000 movies belonging to 18 genres (Harper and Konstan, 2015). In our setup, for each user, the principal is a recommender system that has to generate a ranking of movies $\mathcal{S}_g$ for one of the genres g (e.g., Horror, Romance, Comedy) according to a notion of *merit* of the movies we define as follows.

Modeling the Merit Distribution. We define the (unknown) merit v_m of a movie m as the average rating of the movie across the user population² — this merit is unknown because most users have not seen/rated most movies. To be able to estimate this merit based on ratings in the ML-100K dataset, and to concretely define its underlying distribution and the corresponding fairness criteria, we define a generative model of user ratings. The model assumes that the rating of a movie $m \in \mathcal{S}_g$ is drawn from a multinomial distribution over $\{1, 2, 3, 4, 5\}$ with parameters $\theta_m = (\theta_{m,1}, \dots, \theta_{m,5})$. Assuming a Dirichlet prior, we can infer the posterior distribution of the ratings (and hence the merit) from the dataset in closed form, using merely the counts of each rating for each movie in the dataset. (See § E.1 for details.) Based on this definition of merit and its uncertainty, we define and compare different ranking policies for this experimental setup as follows.

Utility Maximizing Ranking (π^*): We use the DCG function (Burges et al., 2005) with position weights $w_k = 1/\log_2(1+k)$ as our utility measure. Since the weights are indeed strictly decreasing, as described in § 3, the optimal ranking policy π^* sorts the movies (for the particular query) by decreasing expected merit, which is the expected average rating $\bar{v}_m \triangleq \mathbb{E}_{\theta \sim \mathbb{P}[\theta_m | \mathcal{D}]} [v_m(\theta)]$ under the posterior Dirichlet distribution in our case; here, $v_m(\theta)$ is the average rating of movie m corresponding to the parameter vector θ sampled from the posterior.

²For a personalized ranking application, an alternative would be to choose each user’s (mostly unknown) rating as the merit criterion instead of the average rating across the user population.

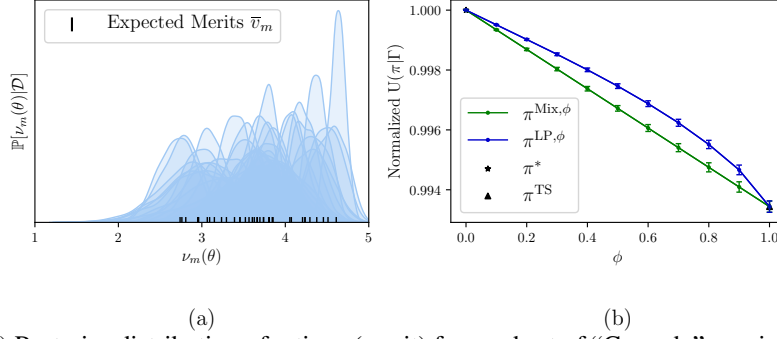


Figure 1: (a) Posterior distribution of ratings (merit) for a subset of “Comedy” movies, (b) Tradeoff between Utility and Fairness, as captured by ϕ .

Fully Fair Ranking Policy (π^{TS}): A fair ranking, in this case, ensures that, for all positions k , a movie is placed in the top k positions according to the posterior merit distribution. In this setup, a fully fair ranking policy π^{TS} is obtained by sampling the multinomial parameters θ_m for each movie $m \in \mathcal{S}_g$ and computing $v_m(\theta_m)$ to rank them:

$$\pi^{\text{TS}}(\mathcal{S}_g) \sim \text{argsort}_m v_m(\theta_m) \text{ s.t. } \theta_m \sim \mathbb{P}[\theta_m|\mathcal{D}].$$

OPT/TS-Mixing Ranking Policy ($\pi^{\text{Mix},\phi}$): The policies $\pi^{\text{Mix},\phi}$ randomize between the fully fair and utility-maximizing ranking policies with probabilities ϕ and $1 - \phi$, respectively.

LP Ranking Policy ($\pi^{\text{LP},\phi}$): The ϕ -fair policies $\pi^{\text{LP},\phi}$ require the principal to have access to the probabilities $\mathbb{P}_{v \sim \Gamma}[\mathcal{M}_{m,k}^{(v)}]$ which we estimate using $5 \cdot 10^4$ Monte Carlo samples, so that any estimation error becomes negligible.

Observations and Results. In the experiments presented, we used the ranking policies π^* , π^{TS} , $\pi^{\text{Mix},\phi}$ and $\pi^{\text{LP},\phi}$ to create separate rankings for each of the 18 genres. For each genre, the task is to rank a random subset of 40 movies from that genre. To get a posterior with an interesting degree of uncertainty, we take a 10% i.i.d. samples from $\mathcal{D}$ to infer the posterior for each movie. We observe that the results are qualitatively consistent across genres, and we thus focus on detailed results for the genre “Comedy” as a representative example. Its posterior merit distribution over a subset is visualized in Figure 1(a). Note that substantial overlap exists between the marginal merit distributions of the movies, indicating that as opposed to π^* (which sorts based on the expected merits), the policy π^{TS} will randomize over many different rankings.

Observation 1: We evaluate the cost of fairness to the principal in terms of loss in utility, as well as the ability of $\pi^{\text{LP},\phi}$ to minimize this cost for ϕ -fair rankings. Figure 1(b) shows this cost in terms of expected Normalized DCG (i.e., $\text{NDCG} = \text{DCG}/\max(\text{DCG})$ as in (Järvelin and Kekäläinen, 2002)). These results are averaged over 20 runs with different subsets of movies and different training samples. The leftmost end corresponds to the NDCG of π^* , while the rightmost point corresponds to the NDCG of the 1-fair policy π^{TS} .

The drop in NDCG is below one percentage point, which is consistent with the results for the other genres. We also conducted experiments with other values of s , data set sizes, and choices of w_k ; even under the most extreme conditions, the drop was at most 2 percent. While this rather small drop may be surprising at first, we point out that uncertainty in the estimates affects the utility of both π^* and π^{TS} . By industry standards, a 2% drop in NDCG is considered quite substantial; however, it is not catastrophic and hence bodes well for possible adoption.

Observation 2: Figure 1(b) also compares the trade-off in NDCG in response to the fairness approximation parameter ϕ for both $\pi^{\text{Mix},\phi}$ and $\pi^{\text{LP},\phi}$. We observe that the utility-optimal policy $\pi^{\text{LP},\phi}$ provides gains over $\pi^{\text{Mix},\phi}$, especially for large values of ϕ . Thus, using $\pi^{\text{LP},\phi}$ can further reduce the cost of fairness discussed above.

Observation 3: To provide intuition about the difference between $\pi^{\text{Mix},\phi}$ and $\pi^{\text{LP},\phi}$, Figure 4 in § E.3 visualizes the marginal rank distributions $p_{m,k}$, i.e., the probability that movie m is ranked

at position k . The key distinction is that, for intermediate values of ϕ , $\pi^{\text{LP},\phi}$ exploits a non-linear structure in the ranking distribution (to achieve a better trade-off) while $\pi^{\text{Mix},\phi}$ merely interpolates linearly between the solutions for $\phi = 0$ and $\phi = 1$.

Based on these observations, in general, the utility loss due to fairness is small, and can be further reduced by optimizing the ranking distribution with the LP-based approach. These results are based on the definition of merit as the average rating of movies over the entire user population. A more realistic setting would personalize rankings for each user, with merit defined as the expected relevance of a movie to the user. In our experiments, the results under such a setup were quite similar, and are hence omitted for brevity and clearer illustration.

5 Real-World Experiment: Paper Recommendation

To study the effect of deploying a fair ranking policy in a real ranking system, we built and fielded a paper recommendation system at the 2020 ACM SIGKDD Conference on Knowledge Discovery and Data Mining. The goal of the experiment is to understand the impact of fairness under real user behavior, as opposed to simulated user behavior that is subject to modeling assumptions. Specifically, we seek to answer two questions: (a) Does a fair ranking policy lead to a more equitable distribution of exposure among the papers? (b) Does fairness substantially reduce the utility of the system to the users?

The users of the paper recommendation system were the participants of the conference, which was held virtually in 2020. Signup and usage of the system was voluntary. Each user was recommended a personalized ranking of the papers published at the conference. This ranking was produced either by σ^* or by π^{TS} , and the assignment of users to treatment (π^{TS}) or control (σ^*) was randomized.

Modeling the Merit Distribution. The merit of a paper for a particular user is based on a relevance score $\mu_{u,i}$ that relates features of the user (e.g., bag-of-words representation of recent publications, co-authorship) to features of each conference paper (e.g., bag-of-words representation of paper, citations). Most prominently, the relevance score $\mu_{u,i}$ contains the TFIDF-weighted cosine similarity between the bag-of-words representations.

To model the uncertainty in $\mu_{u,i}$, we make the assumption that since all papers were accepted to the conference, they must have been deemed relevant to at least some fraction of the audience by the peer reviewers. Hence, papers with uniformly low relevance scores $\mu_{u,i}$ across users (e.g., ones introducing new research directions or bringing in novel techniques) must have a higher uncertainty in their relevance score. This assumption leads us to define the uncertainty as a normal distribution around the mean $\mu_{u,i}$ with standard deviation δ_i defined such that, for paper i , there exists at least one user who finds the paper highly relevant to their interests with high probability. (See § F.2 for details on how δ_i is calculated.)

Ranking Policies. Users in the control group $\mathcal{U}_{\pi^*}$ received rankings in decreasing order of $\mu_{u,i}$. Users in the treatment group $\mathcal{U}_{\pi^{\text{TS}}}$ received rankings from the fair policy that sampled scores from the uncertainty distribution, $\hat{\mu}_{u,i} \sim \mathcal{N}(\mu_{u,i}, \delta_i)$, and ranked the papers by decreasing $\hat{\mu}_{u,i}$.

Results and Observations. We first analyze if the fair policy provided more equitable exposure to the papers. In this real-world evaluation, exposure is not equivalent to rank, but depends on whether users actually browsed to a given position in the ranking. Users could browse their ranking in pages of 5 recommendations each; we count a paper as *exposed* if the user scrolled to its page.

Observation 1: Figure 2 compares the histograms of exposure of the papers in the treatment and control groups. Under the fair policy, the number of papers in the lowest tier of exposure is roughly halved compared to the control condition. This verifies that the fair policy does have an impact on exposure in a real-world setting, and it aligns with our motivation that a fair ranking policy distributes exposure more equally among the ranked agents. This is also supported by comparing the Gini inequality coefficient (Gini, 1936) for the two distributions: $G(\pi^*) = 0.3302$, while $G(\pi^{\text{TS}}) = 0.2996$ (where a smaller coefficient means less inequality in the distribution).

Observation 2: To evaluate the impact of fairness on user engagement, we analyze a range of engagement metrics as summarized in Table 1. While a total of 923 users signed up for the system ahead of the conference (and were randomized into treatment and control groups), 462 never came to

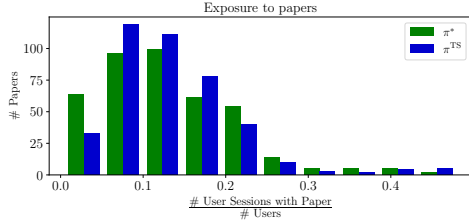


Figure 2: Distribution of the exposure of papers in treatment and control.

	Number of Users with activity		Average Activity Per User	
	π^*	π^{TS}	π^*	π^{TS}
Total Number of users	213	248	-	-
Num. of pages examined	-	-	10.8075	10.7984
Read Abstract	92	101	3.7230	2.6774
Add to Calendar	51	50	1.4366	0.8508
Read PDF	40	52	0.5258	0.5323
Add Bookmark	16	13	0.3192	0.6129

Table 1: User engagement under the two conditions π^* and π^{TS} . None of the differences are statistically significant. (For user actions, this is specifically due to the small sample size).

the system after all. Of the users that came, 213 users were in $\mathcal{U}_{\pi^*}$, and 248 users were in $\mathcal{U}_{\pi^{TS}}$. Note that this difference is not caused by the treatment assignment, since users had no information about their assignment/ranking before entering the system. The first engagement metric we computed is the average number of pages that users viewed under both conditions. With roughly 10.8 pages (about 54 papers), engagement under both conditions was almost identical. Users also had other options to engage, but there is no clear difference between the conditions, either. On average, they read more paper abstracts and added more papers to their calendar under the control condition, but read more PDF and added more bookmarks under the treatment condition. However, none of these differences is significant at the 95% level for either a Mann-Whitney U test or a two-sample t-test. While the sample size is small, these findings align with the findings on the synthetic data, namely that fairness did not appear to place a large cost on the principal (here representing the users).

6 Conclusions and Future Work

We believe that the focus on uncertainty we proposed in this paper constitutes a principled approach to capturing the intuitive notion of fairness to agents with similar features: rather than focusing on the features themselves, the key insight is that the features’ similarity entails significant statistical uncertainty about which agent has more merit. Randomization provides a way to fight fire with fire, and axiomatize fairness in the presence of such uncertainty.

Our work raises a wealth of questions for future work. Perhaps most importantly, as discussed in § 2.4, to operationalize our proposed notion of fairness, it is important to derive principled merit distributions Γ based on the observed features. Our experiments were based on “reasonable” notions of merit distributions and concluded that fairness might not have to be very expensive to achieve for a principal. However, much more experimental work is needed to truly evaluate the impact of fair policies on the utility that is achieved. It would be particularly intriguing to investigate which types of real-world settings lend themselves to implementing fairness at little cost, and which force a steep trade-off between the two objectives.

Our work also raises several interesting theoretical questions. In § B.1, we show one setting in which forcing the principal to use a fair policy drastically increases the principal’s incentives to form a more accurate posterior Γ for a minority group. We did not prove a general result in this vein. We ask: will the incentives of a principal to learn a better posterior Γ always (weakly) increase if the principal is forced to be fairer? If true, this would provide a fascinating additional benefit of fairness requirements.

Another interesting question concerns the utility loss incurred by using the policy OPT/TS-Mixing. As shown in § D, OPT/TS-Mixing is in general not optimal. However, in the example from § D as well as in our experiments in § 4, the loss in utility was quite small. An interesting question would be to bound the worst-case loss in the utility of OPT/TS-Mixing, compared to the LP-based policy. In particular, this question is of interest due to the simplicity of the OPT/TS-Mixing policy; it does not require the computationally expensive solution of an LP or an explicit estimate of marginal rank probabilities under Γ .

Acknowledgements

This research was supported in part by NSF Award IIS-2008139 and IIS-1901168. We thank Aleksandra Korolova, Cris Moore, and Stephanie Wykstra for useful discussions and feedback. Any opinions, findings, and conclusions or recommendations expressed in this material are those of the author(s) and do not necessarily reflect the views of the National Science Foundation.

References

- Asudehy, A., Jagadishy, H., Stoyanovich, J., and Das, G. (2019). Designing fair ranking schemes. *SIGMOD*.
- Barocas, S. and Selbst, A. D. (2016). Big data’s disparate impact. *Cal. L. Rev.*
- Beutel, A., Chen, J., Doshi, T., Qian, H., Wei, L., Wu, Y., Heldt, L., Zhao, Z., Hong, L., Chi, E. H., et al. (2019). Fairness in recommendation ranking through pairwise comparisons. In *KDD*.
- Biega, A. J., Gummadi, K. P., and Weikum, G. (2018). Equity of attention: Amortizing individual fairness in rankings. In *SIGIR*.
- Birkhoff, G. (1946). Tres observaciones sobre el algebra lineal. *Univ. Nac. Tucuman*.
- Burges, C., Shaked, T., Renshaw, E., Lazier, A., Deeds, M., Hamilton, N., and Hullender, G. (2005). Learning to rank using gradient descent. In *ICML*, pages 89–96.
- Calders, T., Kamiran, F., and Pechenizkiy, M. (2009). Building classifiers with independency constraints. In *Data mining workshops, ICDMW*, pages 13–18.
- Calmon, F. P., Wei, D., Vinzamuri, B., Ramamurthy, K. N., and Varshney, K. R. (2017). Optimized pre-processing for discrimination prevention. In *NIPS*.
- Carbonell, J. and Goldstein, J. (1998). The use of MMR, diversity-based reranking for reordering documents and producing summaries. In *SIGIR*.
- Celis, L. E., Mehrotra, A., and Vishnoi, N. K. (2020). Interventions for ranking in the presence of implicit bias. In *Proceedings of the Conference on Fairness, Accountability, and Transparency*.
- Celis, L. E., Straszak, D., and Vishnoi, N. K. (2018). Ranking with fairness constraints. *ICALP*.
- Clarke, C. L., Kolla, M., Cormack, G. V., Vechtomova, O., Ashkan, A., Büttcher, S., and MacKinnon, I. (2008). Novelty and diversity in information retrieval evaluation. In *SIGIR*.
- Diaz, F., Mitra, B., Ekstrand, M. D., Biega, A. J., and Carterette, B. (2020). Evaluating stochastic rankings with expected exposure. In *CIKM*.
- Dvoretzky, A., Kiefer, J., and Wolfowitz, J. (1956). Asymptotic minimax character of the sample distribution function and of the classical multinomial estimator. *Annals of Mathematical Statistics*, 27(3):642–669.
- Dwork, C., Hardt, M., Pitassi, T., Reingold, O., and Zemel, R. (2012). Fairness through awareness. In *ITCS*. ACM.
- Dwork, C., Kim, M. P., Reingold, O., Rothblum, G. N., and Yona, G. (2019). Learning from outcomes: Evidence-based rankings. In *FOCS*.
- Ghosh, A., Dutt, R., and Wilson, C. (2021). When fair ranking meets uncertain inference. *arXiv preprint arXiv:2105.02091*.
- Gini, C. (1936). On the measure of concentration with special reference to income and statistics. *Colorado College Publication, General Series*.
- Hardt, M., Price, E., and Srebro, N. (2016). Equality of opportunity in supervised learning. In *NIPS*.
- Harper, F. M. and Konstan, J. A. (2015). The MovieLens datasets: History and context. *ACM TIIS*.

- Heidari, H. and Krause, A. (2018). Preventing disparate treatment in sequential decision making. In *IJCAI*.
- Järvelin, K. and Kekäläinen, J. (2002). Cumulated gain-based evaluation of ir techniques. *TOIS*.
- Joseph, M., Kearns, M., Morgenstern, J., and Roth, A. (2016). Fairness in learning: Classic and contextual bandits. In *NIPS*.
- Kallus, N. and Zhou, A. (2019). The fairness of risk scores beyond classification: Bipartite ranking and the XAUC metric. *NeurIPS*.
- Kearns, M., Roth, A., and Wu, Z. S. (2017). Meritocratic fairness for cross-population selection. In *ICML*.
- Kilbertus, N., Carulla, M. R., Parascandolo, G., Hardt, M., Janzing, D., and Schölkopf, B. (2017). Avoiding discrimination through causal reasoning. In *NIPS*.
- Kim, M. P., Ghorbani, A., and Zou, J. (2019). Multiaccuracy: Black-box post-processing for fairness in classification. In *AAAI Conference on AI, Ethics, and Society*.
- Kleinberg, J. and Raghavan, M. (2018). Selection problems in the presence of implicit bias. In *ITCS*.
- Kusner, M. J., Loftus, J., Russell, C., and Silva, R. (2017). Counterfactual fairness. In *NIPS*.
- Lahoti, P., Gummadi, K. P., and Weikum, G. (2019). Operationalizing individual fairness with pairwise fair representations. *VLDB Endowment*.
- Liu, Y., Radanovic, G., Dimitrakakis, C., Mandal, D., and Parkes, D. C. (2017). Calibrated fairness in bandits. *FATML*.
- Lum, K. and Johndrow, J. (2016). A statistical framework for fair predictive algorithms. *FATML*.
- Massart, P. (1990). The tight constant in the Dvoretzky-Kiefer-Wolfowitz inequality. *Annals of Probability*, 18(3):1269–1283.
- Mehrabi, N., Morstatter, F., Saxena, N., Lerman, K., and Galstyan, A. (2019). A survey on bias and fairness in machine learning. *arXiv preprint arXiv:1908.09635*.
- Mehrotra, A. and Celis, L. E. (2021). Mitigating bias in set selection with noisy protected attributes. In *Proceedings of the ACM Conference on Fairness, Accountability, and Transparency*.
- Mehrotra, R., McInerney, J., Bouchard, H., Lalmas, M., and Diaz, F. (2018). Towards a fair marketplace: Counterfactual evaluation of the trade-off between relevance, fairness & satisfaction in recommendation systems. In *CIKM*.
- Morik, M., Singh, A., Hong, J., and Joachims, T. (2020). Controlling fairness and bias in dynamic learning-to-rank. In *SIGIR*.
- Narasimhan, H., Cotter, A., Gupta, M., and Wang, S. (2020). Pairwise fairness for ranking and regression. In *AAAI*.
- Patil, V., Ghalme, G., Nair, V., and Narahari, Y. (2020). Achieving fairness in the stochastic multi-armed bandit problem. In *AAAI*.
- Penha, G. and Hauff, C. (2021). On the calibration and uncertainty of neural learning to rank models. *arXiv preprint arXiv:2101.04356*.
- Pleiss, G., Raghavan, M., Wu, F., Kleinberg, J., and Weinberger, K. Q. (2017). On fairness and calibration. *NIPS*.
- Prost, F., Awasthi, P., Blumm, N., Kumthekar, A., Potter, T., Wei, L., Wang, X., Chi, E. H., Chen, J., and Beutel, A. (2021). Measuring model fairness under noisy covariates: A theoretical perspective. *arXiv preprint arXiv:2105.09985*.
- Radlinski, F., Bennett, P. N., Carterette, B., and Joachims, T. (2009). Redundancy, diversity and interdependent document relevance. In *ACM SIGIR Forum*.

- Robertson, S. E. (1977). The probability ranking principle in IR. *Journal of documentation*.
- Schumann, C., Lang, Z., Mattei, N., and Dickerson, J. P. (2019). Group fairness in bandit arm selection. *arXiv preprint arXiv:1912.03802*.
- Singh, A. and Joachims, T. (2018). Fairness of exposure in rankings. In *KDD*.
- Singh, A. and Joachims, T. (2019). Policy learning for fairness in ranking. In *NeurIPS*.
- Singh, A., Kempe, D., and Joachims, T. (2021). Fairness in ranking under uncertainty. *arXiv preprint 2107.06720*.
- Soufiani, H. A., Parkes, D. C., and Xia, L. (2012). Random utility theory for social choice. In *NIPS*.
- Taylor, M., Guiver, J., Robertson, S., and Minka, T. (2008). Softrank: optimizing non-smooth rank metrics. In *WSDM*.
- Woodworth, B., Gunasekar, S., Ohannessian, M. I., and Srebro, N. (2017). Learning non-discriminatory predictors. *COLT*.
- Yang, K. and Stoyanovich, J. (2017). Measuring fairness in ranked outputs. *SSDBM*.
- Zafar, M. B., Valera, I., Gomez Rodriguez, M., and Gummadi, K. P. (2017). Fairness beyond disparate treatment & disparate impact: Learning classification without disparate mistreatment. In *WWW*.
- Zehlike, M., Bonchi, F., Castillo, C., Hajian, S., Megahed, M., and Baeza-Yates, R. (2017). Fa*ir: A fair top-k ranking algorithm. *CIKM*.
- Zehlike, M. and Castillo, C. (2020). Reducing disparate exposure in ranking: A learning to rank approach. In *The Web Conference*.
- Zemel, R., Wu, Y., Swersky, K., Pitassi, T., and Dwork, C. (2013). Learning fair representations. In *ICML*.
- Zliobaite, I. (2015). On the relation between accuracy and fairness in binary classification. *FATML Workshop at ICML*.

Supplementary Material: Fairness in Ranking under Uncertainty

A Related Work

As algorithmic techniques, especially machine learning, find widespread applications in decision making, there is notable interest in understanding its societal impacts. While algorithmic decisions can counteract existing biases by preventing human error and implicit bias, data-driven algorithms may also create new avenues for introducing unintended bias (Barocas and Selbst, 2016). There have been numerous attempts to define notions of fairness in the supervised learning setting, especially for binary classification and risk assessment (Calders et al., 2009; Zliobaite, 2015; Dwork et al., 2012; Hardt et al., 2016; Mehrabi et al., 2019). The group fairness perspective imposes constraints like demographic parity (Calders et al., 2009; Zliobaite, 2015) and equalized odds (Hardt et al., 2016). Follow-up work has proposed techniques for implementing fairness through pre-processing methods (Calmon et al., 2017; Lum and Johndrow, 2016), in process while learning the model (Zemel et al., 2013; Woodworth et al., 2017; Zafar et al., 2017) and post-processing methods (Hardt et al., 2016; Pleiss et al., 2017; Kim et al., 2019), in addition to causal approaches to fairness (Kilbertus et al., 2017; Kusner et al., 2017).

Individual fairness, on the other hand, is concerned with comparing the outcomes of agents directly, not in aggregate. Specifically, the individual fairness axiom states that two individuals similar with respect to a task should receive similar outcomes (Dwork et al., 2012). While the property of individual fairness is highly desirable, it is hard to define precisely; in particular, it is highly dependent on the definition of a suitable similarity notion. Although similar in spirit, our work sidesteps this need to define a similarity metric between agents in the feature space. Rather, we view an agent’s features solely as noisy signals about the agent’s merit and posit that a comparison of these merits — and the principal’s uncertainty about them — should determine the relative ranking. Individual fairness definitions have also been adopted in online learning settings such as stochastic multi-armed bandits (Patil et al., 2020; Heidari and Krause, 2018; Schumann et al., 2019; Celis et al., 2018), where the desired property is that a worse arm is never “favored” over a better arm despite the algorithm’s uncertainty over the true payoffs (Joseph et al., 2016), or a smooth fairness assumption that a pair of arms be selected with similar probability if they have a similar payoff distribution (Liu et al., 2017). While these definitions are derived from the same tenet of fairness as Axiom 1 for a pair of agents, we extend it to rankings, where n agents are compared at a time.

Rankings are a primary interface through which machine learning models support human decision making, ranging from recommendation and search in online systems to machine-learned assessments for college admissions and recruiting. One added difficulty with considering fairness in the context of rankings is that the decision for an agent (where to rank that agent) depends not only on their own merits, but on others’ merits as well (Dwork et al., 2019). The existing work can be roughly categorized into three groups: Composition-based, opportunity-based, and evidence-based notions of fairness. The notions of fairness based on the composition of the ranking operate along the lines of demographic parity (Zliobaite, 2015; Calders et al., 2009), proposing definitions and methods that minimize the difference in the (weighted) representation between groups in a prefix of the ranking (Yang and Stoyanovich, 2017; Celis et al., 2018; Asudehy et al., 2019; Zehlike et al., 2017; Mehrotra et al., 2018; Zehlike and Castillo, 2020). Other works argue against the winner-take-all allocation of economic opportunity (e.g., exposure, clickthrough, etc.) to the ranked agents or groups of agents, and that the allocation should be based on a notion of merit (Singh and Joachims, 2018; Biega et al., 2018; Diaz et al., 2020). Meanwhile, the metric-based notions equate a ranking with a set of pairwise comparisons, and define fairness notions based on parity of pairwise metrics within and across groups (Kallus and Zhou, 2019; Beutel et al., 2019; Narasimhan et al., 2020; Lahoti et al., 2019). Similar to pairwise accuracy definitions, evidence-based notions such as (Dwork et al., 2019) propose semantic notions such as *domination-compatibility* and *evidence-consistency*, based on relative ordering of subsets within the training data. Our fairness axiom combines the opportunity-based and evidence-based notions by stating that the economic opportunity allocated to the agents must be consistent with the existing evidence about their relative ordering.

Ranking has been widely studied in the field of Information Retrieval (IR), mostly in the context of optimizing user utility. The *Probability Ranking Principle (PRP)* (Robertson, 1977), a guiding principle for ranking in IR, states that user utility is optimal when documents (i.e., the agents) are ranked by expected values of their estimated relevance (merit) to the user. While this certainly holds

when the estimates are unbiased and devoid of uncertainty, we argue that it leads to unfair rankings for agents about whose merits the model might be uncertain. While the research on diversified rankings in IR appears related, in comparison to our work, the goal there is to maximize user utility alone by handling uncertainty about the user’s information needs (Radlinski et al., 2009) and to avoid redundancy in the ranking (Clarke et al., 2008; Carbonell and Goldstein, 1998). Besides ranking diversity, IR methods have dealt with uncertainty in relevance that comes via users’ implicit or explicit feedback (Penha and Hauff, 2021; Soufiani et al., 2012), as well as stochasticity arising from optimizing over probabilistic rankings instead of discrete combinatorial structures (Taylor et al., 2008; Burges et al., 2005). It is only recently that there has been an interest in developing evaluation metrics (Diaz et al., 2020) and learning algorithms (Singh and Joachims, 2019; Morik et al., 2020) that use stochastic ranking models to deal with unfair exposure.

Additional recent strands of work on fairness in selection problems focus on fairly selecting individuals distributed across different groups in the presence of group-based implicit bias (Kleinberg and Raghavan, 2018; Celis et al., 2020), noisy sensitive attributes (Mehrotra and Celis, 2021), or incomparable merits across different groups (Kearns et al., 2017). Kearns et al. (2017) present a way to fairly select k individuals distributed across d populations where each population can be sorted by merit without uncertainty but merit in one population cannot be directly compared to merit in another. Hence, they propose using the true CDF rank as a derived merit criterion that can be compared. There has also been recent interest in studying the effect of uncertainty regarding sensitive attributes, labels and other features used by the machine learning model on the accuracy-based fairness properties of the model (Ghosh et al., 2021; Prost et al., 2021). In contrast, our work takes a more fundamental approach to defining a merit-based notion of fairness arising due to the presence of uncertainty when estimating merits based on fully observed features and outcomes.

B Additional Model Discussion

B.1 Information Acquisition Incentives for the Principal

An additional benefit of requiring the use of fair ranking policies is that it makes the principal bear more of the cost of an inaccurate Γ , and thereby incentivizes the principal to improve the distribution Γ . To see this at a high level, notice that if Γ precisely revealed merits, then the optimal and fair policies would coincide. In the presence of uncertainty, an unrestricted principal will optimize utility, and in particular do better than a principal who is constrained to be (partially or completely) fair. Thus, a fair principal stands to gain more by obtaining perfect information. The following example shows that this difference can be substantial, i.e., the information acquisition incentives for a fair principal can be much higher.

Example 1. Consider again the case of a job portal. To keep the example simple, consider a scenario in which the portal tries to recommend exactly one candidate for a position.³ There are two groups of candidates, which we call majority and underrepresented minority (URM). The majority group contains exactly one candidate of merit 1, all others having merit 0; the URM group contains exactly one candidate of merit $1 + \epsilon$, all others having merit 0 as well. Due to past experience with the majority group, the portal’s distribution Γ over merits precisely pinpoints the meritorious majority candidate, but reveals no information about the meritorious URM candidate; that is, the distribution places equal probability on each of the URM candidates having merit $1 + \epsilon$.

A utility-maximizing portal will therefore go with “the known thing,” obtaining utility 1 from recommending the majority candidate. The loss in utility from ignoring the URM candidates is only ϵ . Now consider a portal required to be 1-fair. Because each of the URM candidates is the best candidate with probability $1/n$ (when there are n URM candidates), and the majority candidate is *known* to never be the best candidate, each URM candidate *must* be recommended with probability $1/n$. Here, the uncertainty about which URM candidate is meritorious will provide the portal with a utility that is only $(1+\epsilon)/n$.

In this example, fairness strengthens the incentive for the portal to acquire more information about the URM group; specifically, to learn to perfectly identify the meritorious candidate. Under full knowledge, the portal will now have utility $1 + \epsilon$ for both the fair and the utility-maximizing policy.

³This can be considered a ranking problem in which the first slot has $w_1 = 1$, while all other slots have weight 0.

For the utility-maximizing portal, this is the optimal choice; and for the fair strategy, it is perfectly fair to always select the (deterministically known) best candidate. Thus, a portal forced to use the fair strategy stands to increase its utility by a much larger amount; at least in this example, our definition of fairness splits the cost of a high-variance distribution Γ more evenly between the principal and the affected agents when compared to the utility-optimizing policy, where almost all the cost of uncertainty is borne by the agents in the URM group. This drastically increases the principal’s incentives for more accurate and equitable information gathering.

To what extent the insights from this example generalize to arbitrary settings (e.g., whether the principal *always* stands to gain more from additional information when forced to be fairer) is a fascinating direction for future research.

B.2 Ordinal Merit and High-Risk High-Reward Agents

As we discussed earlier, Proposition 2.1 highlights the fact that our definition of fairness only considers ordinal properties, i.e., comparisons, of merit. This means that frequently selecting “moonshot” agents (those with very rare very high merit) would be considered unfair. We argue that this is not a drawback of our fairness definition; rather, if moonshot attempts are worth supporting frequently, then the definition of merit should be altered to reflect this understanding. As a result, viewing the merit definition under the prism of our fairness definition helps reveal misalignments between stated merit and actual preferences.

For a concrete example, consider two agents: agent A has known merit 1, while agent B has merit $M \gg 1$ with probability 1% and 0 with probability 99%. When $M > 100$, agent B has larger expected merit, but regardless of whether $M > 100$ or $M \leq 100$, a fully fair principal cannot select B with probability more than 1%. One may consider this a shortcoming of our model: it would prevent, for instance, a funding agency (which tries to be fair to research grant PIs) from focusing on high-risk high-reward research. We argue that the shortcoming will typically not be in the fairness definition, but in the chosen definition of merit. For concreteness, suppose that the status quo is to evaluate merit as the total number of citations which the funded work attracts during the next century.⁴ Also, for simplicity, suppose that “high-reward” research is research that attracts more than 100,000 citations over the next century. If we consider one unit of merit as 1000 citations, and assume that the typical research grant results in work attracting about that many citations, then the funding agency faces the problem from the previous paragraph, and will not be able to support PI B with probability more than 1%. This goes against the express preference of many funding agencies for high-risk high-reward work.

However, if one truly believes that high-reward work is fundamentally different (e.g., it will change the world), then this difference should be explicitly modeled in the notion of merit. For example, rather than “number of citations,” an alternative notion of merit would be “probability that the number of citations exceeds 100,000.” This approach would allow the agency to select PIs based on the posterior probability (based on observed attributes, such as track record and the proposal) of producing such high-impact work. Of course, in reality, different aspects of merit can be combined to define a more accurate notion of merit that reflects what society values as true merit of research.

The restrictions imposed on a principal by our framework will and should force the principal to articulate actual merit of agents carefully, rather than adding ad hoc objectives. Once merit has been clearly defined, we anticipate that the conflict between fairness and societal objectives will be significantly reduced.

C Omitted Proofs

Here, we provide proofs omitted in § 3. The results are restated here for convenience. Proposition 3.1 is standard, and we only include a proof for completeness.

Proposition 3.1 σ_Γ^* is a utility-maximizing ranking policy for the principal, even over randomized policies.

⁴This measure is chosen for simplicity of discussion, not to actually endorse this metric.

Proof. Let π be a randomized policy for the principal. We will use a standard exchange argument to show that making π more similar to σ_Γ^* can only increase the principal's utility. Recall that by Equation (2), the principal's utility under π can be written as

$$U(\pi | \Gamma) = \sum_{x \in \mathcal{X}} \sum_k p_{x,k}^{(\pi)} \cdot \mathbb{E}_{\mathbf{v} \sim \Gamma} [v_x] \cdot w_k.$$

Assume that π does not sort x by non-increasing $\mathbb{E}_{\mathbf{v} \sim \Gamma} [v_x]$. Then, there exist two positions $j < k$ and two agents x, y such that $\mathbb{E}_{\mathbf{v} \sim \Gamma} [v_x] > \mathbb{E}_{\mathbf{v} \sim \Gamma} [v_y]$, and $p_{x,k}^{(\pi)} > 0$ and $p_{y,j}^{(\pi)} > 0$. Let $\epsilon = \min(p_{x,k}^{(\pi)}, p_{y,j}^{(\pi)}) > 0$, and consider the modified policy which subtracts ϵ from $p_{x,k}^{(\pi)}$ and $p_{y,j}^{(\pi)}$ and adds ϵ to $p_{x,j}^{(\pi)}$ and $p_{y,k}^{(\pi)}$. This changes the expected utility of the policy by

$$\begin{aligned} & \epsilon \cdot (\mathbb{E}_{\mathbf{v} \sim \Gamma} [v_x] \cdot w_j + \mathbb{E}_{\mathbf{v} \sim \Gamma} [v_y] \cdot w_k - \mathbb{E}_{\mathbf{v} \sim \Gamma} [v_x] \cdot w_k - \mathbb{E}_{\mathbf{v} \sim \Gamma} [v_y] \cdot w_j) \\ &= \epsilon \cdot (w_j - w_k) \cdot (\mathbb{E}_{\mathbf{v} \sim \Gamma} [v_x] - \mathbb{E}_{\mathbf{v} \sim \Gamma} [v_y]) \geq 0. \end{aligned}$$

By repeating this type of update, the policy eventually becomes fully sorted, weakly increasing the utility with every step. Thus, the optimal policy must be sorted by $\mathbb{E}_{\mathbf{v} \sim \Gamma} [v_x]$. $\square$

Lemma 3.1 *Consider two ranking policies π_1 and π_2 such that π_1 is ϕ_1 -fair and π_2 is ϕ_2 -fair. A policy that randomizes between π_1 and π_2 with probabilities q and $1 - q$, respectively, is at least $(q\phi_1 + (1 - q)\phi_2)$ -fair and obtains expected utility $qU(\pi_1 | \Gamma) + (1 - q)U(\pi_2 | \Gamma)$.*

Proof. Both the utility and fairness proofs are straightforward. The proof of fairness decomposes the probability of agent i being in position k under the mixing policy into the two constituent parts, then pulls terms through the sum. The proof of the utility bound uses Equation (2) and linearity of expectations. We now give details of the proofs.

We write π_{Mix} for the policy that randomizes between π_1 and π_2 with probabilities q and $1 - q$, respectively. Using Equation (2), we can express the utility of π_{Mix} as

$$\begin{aligned} U(\pi_{\text{Mix}} | \Gamma) &= \mathbb{E}_{\sigma \sim \pi_{\text{Mix}}, \mathbf{v} \sim \Gamma} [U(\sigma | \mathbf{v})] \\ &= \mathbb{E}_{\mathbf{v} \sim \Gamma} \left[\sum_{\sigma} \pi_{\text{Mix}}(\sigma) \cdot U(\sigma | \mathbf{v}) \right] \\ &= \mathbb{E}_{\mathbf{v} \sim \Gamma} \left[\sum_{\sigma} (q \cdot \pi_1(\sigma) + (1 - q) \cdot \pi_2(\sigma)) \cdot U(\sigma | \mathbf{v}) \right] \\ &= q \cdot \mathbb{E}_{\mathbf{v} \sim \Gamma} \left[\sum_{\sigma} \pi_1(\sigma) \cdot U(\sigma | \mathbf{v}) \right] + (1 - q) \cdot \mathbb{E}_{\mathbf{v} \sim \Gamma} \left[\sum_{\sigma} \pi_2(\sigma) \cdot U(\sigma | \mathbf{v}) \right] \\ &= qU(\pi_1 | \Gamma) + (1 - q)U(\pi_2 | \Gamma). \end{aligned}$$

Similarly, we prove that π is at least $(q\phi_1 + (1 - q)\phi_2)$ -fair if π_1 and π_2 are ϕ_1 - and ϕ_2 -fair, respectively:

$$\begin{aligned} \sum_{k'=1}^k p_{x,k'}^{(\pi)} &= \sum_{k'=1}^k q \cdot p_{x,k'}^{(\pi_1)} + (1 - q) \cdot p_{x,k'}^{(\pi_2)} \\ &= q \cdot \sum_{k'=1}^k p_{x,k'}^{(\pi_1)} + (1 - q) \cdot \sum_{k'=1}^k p_{x,k'}^{(\pi_2)} \\ &\geq q\phi_1 \cdot \mathbb{P}_{\mathbf{v} \sim \Gamma} [\mathcal{M}_{x,k}^{(\mathbf{v})}] + (1 - q)\phi_2 \cdot \mathbb{P}_{\mathbf{v} \sim \Gamma} [\mathcal{M}_{x,k}^{(\mathbf{v})}] \\ &= (q \cdot \phi_1 + (1 - q) \cdot \phi_2) \cdot \mathbb{P}_{\mathbf{v} \sim \Gamma} [\mathcal{M}_{x,k}^{(\mathbf{v})}], \end{aligned}$$

where the inequality used that π_1 is ϕ_1 -fair and π_2 is ϕ_2 -fair. Hence, we have proved that π is $(q \cdot \phi_1 + (1 - q) \cdot \phi_2)$ -fair under Γ . $\square$

Lemma 3.2 *All 1-fair ranking policies have the same utility for the principal.*

Proof. Let π be a 1-fair ranking policy. By Equation (1), π must satisfy the following constraints:

$$\sum_{k'=1}^k p_{x,k'}^{(\pi)} \geq \mathbb{P}_{v \sim \Gamma} [\mathcal{M}_{x,k}^{(v)}] \quad \text{for all } x \text{ and } k. \quad (4)$$

Summing over all x (for any fixed k), both the left-hand side and right-hand side sum to k ; for the left-hand side, this is the expected number of agents placed in the top k positions by π , while for the right-hand side, it is the expected number of agents among the top k in merit. Because the weak inequality (4) holds for all x and k , yet the sum over x is equal, *each* inequality must hold with equality:

$$\sum_{k'=1}^k p_{x,k'}^{(\pi)} = \mathbb{P}_{v \sim \Gamma} [\mathcal{M}_{x,k}^{(v)}] \quad \text{for all } x \text{ and } k.$$

This implies that

$$p_{x,k}^{(\pi)} = \mathbb{P}_{v \sim \Gamma} [\mathcal{M}_{x,k}^{(v)}] - \mathbb{P}_{v \sim \Gamma} [\mathcal{M}_{x,k-1}^{(v)}],$$

which is completely determined by Γ . Substituting these values of $p_{x,k}^{(\pi)}$ into the principal's utility, we see that it is independent of the specific 1-fair policy used. $\square$

Proposition C.1. Consider an algorithm that draws $m = \frac{(\kappa+1) \log(2n)}{2\epsilon^2}$ i.i.d. samples of the agents' joint merits from Γ , and then estimates each probability $\mathbb{P}_{v \sim \Gamma} [\mathcal{M}_{x,k}^{(v)}]$ by the empirical frequency with which x was in position k or higher. Then, with probability at least $1 - n^{-\kappa}$, all $\mathbb{P}_{v \sim \Gamma} [\mathcal{M}_{x,k}^{(v)}]$ are estimated with additive error at most $\pm\epsilon$.

Proof. Focus on one agent x , and write $q_k = \mathbb{P}_{v \sim \Gamma} [\mathcal{M}_{x,k}^{(v)}]$. Notice that the q_k form the CDF of the rank of x . Let $Z_{k,j} = 1$ iff x is among the top k agents (by merit) in the j^{th} of the m samples. Then, $\mathbb{P}[Z_{k,j} = 1] = q_k$, and the estimate $Z_k = \frac{1}{m} \cdot \sum_j Z_{k,j}$ is the average of m independent $\text{Bin}(q_k)$ random variables. By the DKW Inequality for the uniform convergence of the empirical CDF to the true CDF (Dvoretzky et al., 1956; Massart, 1990), we get that with probability at least $1 - 2 \exp(-2m\epsilon^2) \geq 1 - n^{-(\kappa+1)}$, all of the estimates Z_k are within $\pm\epsilon$ of the true values $\mathbb{P}_{v \sim \Gamma} [\mathcal{M}_{x,k}^{(v)}]$. A union bound over all n agents now completes the proof. $\square$

While the estimates may be off by additive ϵ terms, it is fairly easy to compensate for such errors at a small loss in fairness and utility, as follows:

Proposition C.2. For each x, k , let $q_{x,k}$ be an empirical estimate of $\mathbb{P}_{v \sim \Gamma} [\mathcal{M}_{x,k}^{(v)}]$ such that $|q_{x,k} - \mathbb{P}_{v \sim \Gamma} [\mathcal{M}_{x,k}^{(v)}]| \leq \epsilon$ and $\sum_x q_{x,k} = k$ for all k . Consider the solution to the LP (3) with fairness parameter ϕ , using⁵ $q'_{x,k} = \frac{k(q_{x,k} + \epsilon)}{k + n\epsilon}$ in place of the (unknown) $\mathbb{P}_{v \sim \Gamma} [\mathcal{M}_{x,k}^{(v)}]$. Then, the resulting sampling distribution is at least $(\frac{\phi}{1+n\epsilon})$ -fair, and guarantees the principal a utility within a factor $\frac{1}{1+n\epsilon}$ of the optimum ϕ -fair solution.

Proof. First, notice that by the assumption that the $q_{x,k}$ were good approximations for $\mathbb{P}_{v \sim \Gamma} [\mathcal{M}_{x,k}^{(v)}]$, we can bound that $q'_{x,k} \geq \frac{k}{k+n\epsilon} \cdot \mathbb{P}_{v \sim \Gamma} [\mathcal{M}_{x,k}^{(v)}]$.

⁵Notice that the $q'_{x,k}$ in fact satisfy that $\sum_x q'_{x,k} = \frac{k}{k+n\epsilon} \sum_x (q_{x,k} + \epsilon) = \frac{k}{k+n\epsilon} \cdot (k + n\epsilon) = k$, so they can be used as input to the LP.

Because the LP's solution $(p_{x,k})_{x,k}$ is ϕ -fair with respect to the $q'_{x,k}$, we get that

$$\sum_{k'=1}^k p_{x,k'} \geq \phi \cdot q'_{x,k} \geq \frac{k\phi}{k+n\epsilon} \cdot \mathbb{P}_{v \sim \Gamma} [\mathcal{M}_{x,k}^{(v)}] \geq \frac{\phi}{1+n\epsilon} \cdot \mathbb{P}_{v \sim \Gamma} [\mathcal{M}_{x,k}^{(v)}]$$

for all x, k ; thus, the solution is $(\frac{\phi}{1+n\epsilon})$ -fair.

Next, we analyze the principal's utility. Let $(p_{x,k}^*)_{x,k}$ be a ϕ -fair solution maximizing the principal's utility, and write $z_{x,k}^* = \sum_{k'=1}^k p_{x,k'}^*$ for the probability that agent x is ranked among the top k positions in the optimum solution. Now define $z'_{x,k} = \min(z_{x,k}^*, k - \phi \cdot \sum_{x' \neq x} q'_{x',k})$.

We will prove the following two facts: (1) The principal's utility under the probabilities $z'_{x,k}$ is not much smaller than under the original $z_{x,k}^*$, and (2) Every feasible solution $(p_{x,k})_{x,k}$ to the LP with fairness parameter ϕ and $q'_{x,k}$ satisfies $\sum_{k' \leq k} p_{x,k'} \geq z'_{x,k}$ for all x, k .

1. To show the first claim, we first use a standard way to rewrite the principal's objective in terms of the $z_{x,k}^*$ (or $z'_{x,k}$), using the definition $z_{x,0}^* := z'_{x,0} := 0$:

$$\begin{aligned} \sum_x \sum_{k=1}^n p_{x,k}^* \cdot \mathbb{E}_{v \sim \Gamma} [v_x] \cdot w_k &= \sum_x \mathbb{E}_{v \sim \Gamma} [v_x] \cdot \sum_{k=1}^n (z_{x,k}^* - z_{x,k-1}^*) \cdot w_k \\ &= \sum_x \mathbb{E}_{v \sim \Gamma} [v_x] \cdot \left(\sum_{k=1}^n z_{x,k}^* \cdot w_k - \sum_{k=0}^{n-1} z_{x,k}^* \cdot w_{k+1} \right) \\ &= \sum_x \mathbb{E}_{v \sim \Gamma} [v_x] \cdot \left(w_n + \sum_{k=1}^{n-1} z_{x,k}^* \cdot (w_k - w_{k+1}) \right). \quad (5) \end{aligned}$$

Because $z'_{x,k} \leq z_{x,k+1}^*$ for all x, k , writing $p'_{x,k} := z'_{x,k} - z'_{x,k-1}$, we can also express the principal's utility under $(z'_{x,k})_{x,k}$ in the same way, simply replacing the terms $z_{x,k}^*$ with $z'_{x,k}$ in (5). Note that the $p'_{x,k}$ do not form a valid solution to the LP, because the “probabilities” do not necessarily sum up to 1 each across agents or across positions. However, we are only using this “solution” to help with our bounds, and feasibility is not required.

We can write the principal's loss in utility going from $z_{x,k}^*$ to $z'_{x,k}$ as follows:

$$\begin{aligned} &\sum_x \mathbb{E}_{v \sim \Gamma} [v_x] \cdot \left(w_n + \sum_{k=1}^{n-1} z_{x,k}^* \cdot (w_k - w_{k+1}) \right) \\ &\quad - \sum_x \mathbb{E}_{v \sim \Gamma} [v_x] \cdot \left(w_n + \sum_{k=1}^{n-1} z'_{x,k} \cdot (w_k - w_{k+1}) \right) \\ &= \sum_x \mathbb{E}_{v \sim \Gamma} [v_x] \cdot \sum_{k=1}^{n-1} (z_{x,k}^* - z'_{x,k}) \cdot (w_k - w_{k+1}) \\ &= \sum_{k=1}^{n-1} (w_k - w_{k+1}) \cdot \sum_x \mathbb{E}_{v \sim \Gamma} [v_x] \cdot (z_{x,k}^* - z'_{x,k}). \quad (6) \end{aligned}$$

Notice that $w_k - w_{k+1} \geq 0$ for all k , and $\mathbb{E}_{v \sim \Gamma} [v_x] \geq 0$ for all x . To upper-bound the loss in utility, we therefore can apply bounds for each of the terms $z_{x,k}^* - z'_{x,k}$. Focus on one particular pair x, k . Notice that the LP constraints (specifically, the third constraint and the first constraint) imply that

$$z_{x,k}^* = k - \sum_{x' \neq x} z_{x',k}^* \leq k - \phi \cdot \sum_{x' \neq x} \mathbb{P}_{v \sim \Gamma} [\mathcal{M}_{x',k}^{(v)}] = k - \phi \cdot (k - \mathbb{P}_{v \sim \Gamma} [\mathcal{M}_{x,k}^{(v)}]).$$

If $z'_{x,k} < z^*_{x,k}$, then

$$z'_{x,k} = k - \phi \cdot \sum_{x' \neq x} q'_{x',k} = k - \phi \cdot \sum_{x' \neq x} \frac{k(q_{x,k} + \epsilon)}{k + n\epsilon} = k - \frac{k\phi}{k + n\epsilon} \cdot (k - q_{x,k} + (n-1)\epsilon).$$

Therefore, the difference is at most

$$\begin{aligned} z^*_{x,k} - z'_{x,k} &\leq \frac{k\phi}{k + n\epsilon} \cdot (k - q_{x,k} + (n-1)\epsilon) - \phi \cdot (k - \mathbb{P}_{\mathbf{v} \sim \Gamma} [\mathcal{M}_{x,k}^{(\mathbf{v})}]) \\ &= \frac{\phi}{k + n\epsilon} \cdot \left((k^2 - kq_{x,k} + k(n-1)\epsilon) \right. \\ &\quad \left. - (k^2 + kn\epsilon - (k + n\epsilon) \cdot \mathbb{P}_{\mathbf{v} \sim \Gamma} [\mathcal{M}_{x,k}^{(\mathbf{v})}]) \right) \\ &\stackrel{(*)}{\leq} \frac{\phi}{k + n\epsilon} \cdot \left((-k(\mathbb{P}_{\mathbf{v} \sim \Gamma} [\mathcal{M}_{x,k}^{(\mathbf{v})}] - \epsilon) - k\epsilon) + (k + n\epsilon) \cdot \mathbb{P}_{\mathbf{v} \sim \Gamma} [\mathcal{M}_{x,k}^{(\mathbf{v})}] \right) \\ &= \frac{\phi n\epsilon}{k + n\epsilon} \cdot \mathbb{P}_{\mathbf{v} \sim \Gamma} [\mathcal{M}_{x,k}^{(\mathbf{v})}] \\ &\stackrel{(**)}{\leq} \frac{n\epsilon}{k + n\epsilon} \cdot z^*_{x,k}. \end{aligned}$$

Here, the line labeled (*) used that the $q_{x,k}$ approximate the true probabilities $\mathbb{P}_{\mathbf{v} \sim \Gamma} [\mathcal{M}_{x,k}^{(\mathbf{v})}]$ to within additive error at most ϵ , and the line labeled (**) used that the $z^*_{x,k}$ formed a ϕ -fair solution.⁶

We now substitute this bound (for each k, x) into (6), obtaining that the principal's loss in utility is at most

$$\begin{aligned} &\sum_{k=1}^{n-1} (w_k - w_{k+1}) \cdot \sum_x \mathbb{E}_{\mathbf{v} \sim \Gamma} [v_x] \cdot \frac{n\epsilon}{k + n\epsilon} \cdot z^*_{x,k} \\ &\leq \frac{n\epsilon}{1 + n\epsilon} \sum_{k=1}^{n-1} (w_k - w_{k+1}) \cdot \sum_x \mathbb{E}_{\mathbf{v} \sim \Gamma} [v_x] \cdot z^*_{x,k}, \end{aligned}$$

which is exactly $\frac{n\epsilon}{1+n\epsilon}$ times the principal's utility under the solution $z^*_{x,k}$, i.e., the optimal utility. Thus, the utility obtained from using the approximate values is within at least a factor $1 - \frac{n\epsilon}{1+n\epsilon} = \frac{1}{1+n\epsilon}$ of optimal.

- Next, we show that every feasible solution $(p_{x,k})_{x,k}$ to the LP with fairness parameter ϕ and $q'_{x,k}$ satisfies $\sum_{k' \leq k} p_{x,k'} \geq z'_{x,k}$ for all x, k . In fact, we show that $\sum_{k' \leq k} p_{x,k'} \geq k - \phi \cdot \sum_{x' \neq x} q'_{x',k}$, which in turn is at least $z'_{x,k}$ by definition of $z'_{x,k}$.

To see this, note that for any feasible solution and for all x, k , the fairness constraint implies that $\sum_{k'=1}^k p_{x,k'} \geq \phi \cdot q'_{x,k}$ and furthermore, $\sum_x p_{x,k'} = 1$ for all k' . Therefore, for any fixed x, k ,

$$k = \sum_{x'} \sum_{k'=1}^k p_{x',k'} = \sum_{k'=1}^k p_{x,k'} + \sum_{x' \neq x} \sum_{k'=1}^k p_{x',k'} \geq \sum_{k'=1}^k p_{x,k'} + \sum_{x' \neq x} \phi \cdot q'_{x',k}.$$

Rearranging this inequality gives us the claimed bound.

Now consider the optimal solution $p_{x,k}$ (maximizing the principal's utility) with fairness parameter ϕ and estimated probabilities $q'_{x,k}$. For each x, k , define $z_{x,k} = \sum_{k'=1}^k p_{x,k'}$. Then, $z_{x,k} \geq z'_{x,k}$ for all x, k , and the utility under $(p_{x,k})_{x,k}$ is given by (5) (with $z_{x,k}$ in place of $z^*_{x,k}$). In particular, it is at least as large as under $(z'_{x,k})_{x,k}$, and thus within a factor of $\frac{1}{1+n\epsilon}$ of the optimum. $\square$

⁶For $\phi = 0$, the calculations do not apply, but in that case, the algorithm can completely ignore the estimated probabilities, and will obtain the optimum solution.

By Proposition C.2, if the principal wants to approximate fairness and utility to within a factor $1 - \epsilon$, it suffices to approximate the $\mathbb{P}_{v \sim \Gamma} [\mathcal{M}_{x,k}^{(v)}]$ to within an additive error of at most $\frac{\epsilon}{n(1-\epsilon)}$. In turn, by Proposition C.1, it is sufficient to draw $O(\frac{\kappa n^2 \log n}{2\epsilon^2})$ samples from Γ to achieve this approximation with probability at least $1 - n^{-\kappa}$; in particular, the number is polynomial in n and $1/\epsilon$.

D Example for Suboptimality of $\pi^{\text{Mix},\phi}$

Here, we give an example showing that the policy $\pi^{\text{Mix},\phi}$ may not yield optimal utility for the principal among ϕ -fair policies. The example illustrates the types of tradeoffs to be considered for approximately fair solutions, and motivates the LP-based efficient algorithm in § 3.2.

Example 2. Consider $n = 3$ agents, namely a, b , and c . Under Γ , their merits $v_a = 1$, $v_b \sim \text{Bernoulli}(1/2)$, and $v_c \sim \text{Bernoulli}(1/2)$ are drawn independently.⁷ The position weights are $w_1 = 1$, $w_2 = 1$, and $w_3 = 0$.

Now, since $w_1 = w_2 = 1$ and agents b and c are i.i.d., any policy that always places agent a in positions 1 or 2 is optimal. In particular, this is true for the policy π^* which chooses uniformly at random from among $\sigma_1^* = \langle a, b, c \rangle$, $\sigma_2^* = \langle a, c, b \rangle$, $\sigma_3^* = \langle b, a, c \rangle$, and $\sigma_4^* = \langle c, a, b \rangle$.

For the specific distribution Γ , assuming uniformly random tie breaking, we can calculate the probabilities $\mathbb{P}_{v \sim \Gamma} [\mathcal{M}_{x,k}^{(v)}]$ in closed form:

$$\left(\mathbb{P}_{v \sim \Gamma} [\mathcal{M}_{x,k}^{(v)}] \right)_{x,k} = 1/24 \cdot \begin{pmatrix} 14 & 22 & 24 \\ 5 & 13 & 24 \\ 5 & 13 & 24 \end{pmatrix}.$$

The probability of a, b, c being placed in the top k positions by π^* can be calculated as follows:

$$\mathcal{P}^{(\pi^*)} = 1/24 \cdot \begin{pmatrix} 12 & 24 & 24 \\ 6 & 12 & 24 \\ 6 & 12 & 24 \end{pmatrix}.$$

In particular, this implies that π^* is ϕ -fair for every $\phi \leq 12/14 = 6/7$. This bound can be pushed up by slightly increasing the probability of ranking agent a at position 1 (hence increasing fairness to agent a in position 1 at the expense of agents b and c in positions 1–2). Figure 3 shows the principal’s optimal utility for different fairness parameters ϕ , derived from the LP (3). This optimal utility is contrasted with the utility of $\pi^{\text{Mix},\phi}$, which is the convex combination of the utilities of π^* and π^{TS} , by Lemma 3.1.

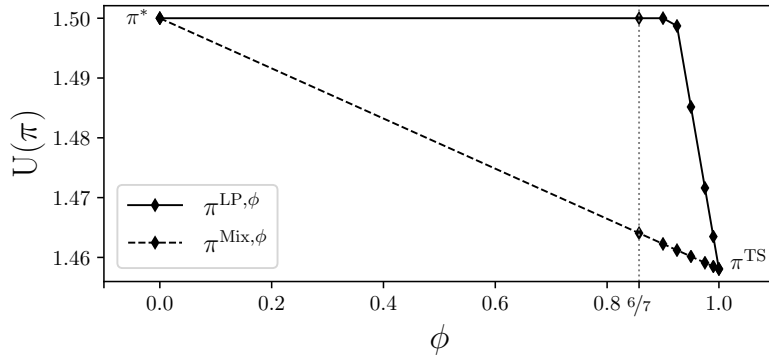


Figure 3: Utility of $\pi^{\text{Mix},\phi}$ and $\pi^{\text{LP},\phi}$ for Example 2 as one varies ϕ .

⁷Technically, this distribution violates the assumption of non-identical merit of agents under Γ . This is easily remedied by adding — say — i.i.d. $\mathcal{N}(0, \epsilon)$ Gaussian noise to all v_i , with very small ϵ . We omit this detail since it is immaterial and would unnecessarily overload notation.

E Details on the MovieLens Dataset Experiment

The MovieLens-100k dataset contains 100,000 ratings, by 600 users, on 9,000 movies belonging to 18 genres (Harper and Konstan, 2015). In our setup, for each user, the principal is a recommender system that has to generate a ranking of movies for one of the genres g (e.g., Horror, Romance, Comedy, etc.), according to a notion of *merit* of the movies we define as follows.

E.1 Experimental Setup

We assume that each rating of a movie $m \in \mathcal{S}_g$ is drawn from a multinomial distribution over $\{1, 2, 3, 4, 5\}$ with (unknown) parameters $\theta_m = (\theta_{m,1}, \dots, \theta_{m,5})$.

Prior: These parameters themselves follow a Dirichlet prior $\theta_m \sim \text{Dir}(\alpha)$ with known parameters $\alpha = (\alpha_1, \alpha_2, \alpha_3, \alpha_4, \alpha_5)$. We assume that the parameters of the Dirichlet prior are of the form $\alpha_r = s \cdot p_r$ where s is a scaling factor and $p_r = \mathbb{P}[\text{Rating} = r \mid \mathcal{D}]$ denotes the marginal probability of observing the rating r in the full MovieLens dataset.

The scaling factor s determines the weight of the prior compared to the observed data, since it acts as a pseudo-count in α' below. For the sake of simplicity, we use $s = 1$ in the following for all movies and genres.

Posterior: Since the Dirichlet distribution is the conjugate prior of the multinomial distribution, the posterior distribution based on the ratings observed in the dataset $\mathcal{D}$ is also a Dirichlet distribution, but with parameters $\alpha' = (\alpha + \mathbf{N}_m) = (\alpha_1 + N_{m,1}, \dots, \alpha_5 + N_{m,5})$ where $N_{m,r}$ is the number of ratings of r for the movie m in the dataset $\mathcal{D}$.

E.2 Expected Merit

The optimal ranking policy π^* sorts the movies (for the particular query) by decreasing expected merit, which is the expected average rating $\bar{v}_m$ under the posterior Dirichlet distribution, and can be computed in closed form as follows:

$$\bar{v}_m \triangleq \mathbb{E}_{\theta \sim \mathbb{P}[\theta_m \mid \mathcal{D}]} [v_m(\theta)] = \sum_{r=1}^5 r \cdot \frac{\alpha_r + N_{m,r}}{\sum_{r'} \alpha_{r'} + N_{m,r'}}. \quad (7)$$

E.3 Ranking Distribution Visualization

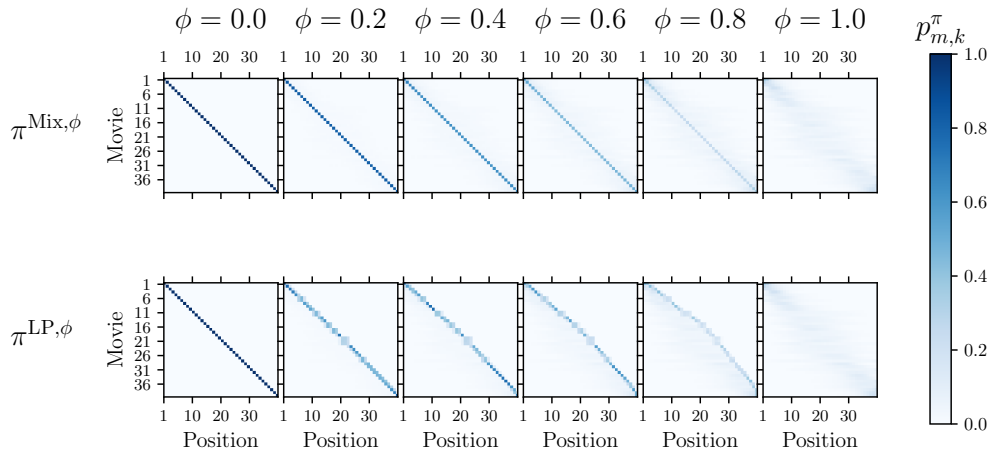


Figure 4: Comparison of marginal rank distribution matrices for $\pi^{\text{Mix},\phi}$ and $\pi^{\text{LP},\phi}$ on “Comedy” movies.

To provide intuition about the difference between the solutions of the LP and OPT/TS-Mixing for different values of ϕ , we visualize $\pi^{\text{Mix},\phi}$ and $\pi^{\text{LP},\phi}$ in Figure 4. The plotted matrices are the marginal

rank distributions: $p_{m,k}$ represents the probability that movie m is ranked at position k . Note that the distribution at $\phi = 0$ and $\phi = 1$ is identical for the two methods, as shown in our lemmas. The key distinction between the rank distributions for $\phi \in (0, 1)$ is that $\pi^{\text{LP},\phi}$ finds non-linear structure for intermediate values of ϕ , while $\pi^{\text{Mix},\phi}$ merely interpolates linearly between the solutions for $\phi = 0$ and $\phi = 1$.

F Details on the Real-World Experiment

As described in § 5, we designed a real-world experiment through a paper recommendation system where the users were the participants at the 2020 ACM SIGKDD Conference on Knowledge Discovery and Data Mining. Signup and usage of the system was voluntary.

Each user was recommended a personalized ranking of the papers published at the conference. This ranking was produced either by σ^* or by π^{TS} , and the assignment of users to treatment (π^{TS}) or control (σ^*) was randomized.

F.1 Users of the Paper Recommendation System

A total of 923 users signed up for the system ahead of the conference (and were randomized into treatment and control groups). Out of these 923 users, 462 did not use the system after all. Of the users that logged in at least once, 213 users were in $\mathcal{U}_{\pi^*}$, and 248 users were in $\mathcal{U}_{\pi^{\text{TS}}}$. Note that this difference is not caused by the treatment assignment, since users had no information about their assignment before entering the system. Users could either navigate through their recommendations by clicking next or previous buttons on their recommendation page, or had other options to engage with each paper such as reading the abstract, reading the PDF, adding the paper to their calendar, and adding a bookmark to the paper.

F.2 Modeling the Merit Distribution

The merit of a paper for a particular user is based on a relevance score $\mu_{u,i}$ that relates features of the user (e.g., bag-of-words representation of recent publications, co-authorship) to features of each conference paper (e.g., bag-of-words representation of paper, citations). Most prominently, the relevance score $\mu_{u,i}$ contains the TFIDF-weighted cosine similarity between the bag-of-words representations.

We model the uncertainty in $\mu_{u,i}$ with regard to the true relevance as follows. First, we observe that all papers were accepted to the conference and thus must have been deemed relevant to at least some fraction of the audience by the peer reviewers. This implies that papers with uniformly low $\mu_{u,i}$ across all/most participants are not irrelevant; we merely have high uncertainty as to which participants the papers are relevant to. For example, papers introducing new research directions or bringing in novel techniques may have uniformly low scores $\mu_{u,i}$ under the bag-of-words model that is less certain about who wants to read these papers compared to papers in established areas. To formalize uncertainty, we make the assumption that a paper’s relevance to a user follows a normal distribution centered at $\mu_{u,i}$, and with standard deviation equal to δ_i (dependent only on the paper, not the user) such that $\max_u \mu_{u,i} + \gamma \cdot \delta_i = 1 + \epsilon$. (For our experiments, we chose $\epsilon = 0.1$ and $\gamma = 2$.) This choice of δ_i ensures that there exists at least one user u such that the (sampled) relevance score $\hat{\mu}_{u,i}$ is greater than 1 with some significant probability; more specifically, we ensure that the probability of having relevance $1 + \epsilon$ is at least as large as that of exceeding the mean by two standard deviations. Furthermore, $\epsilon > 0$ ensures that all papers have a non-deterministic relevance distribution, even papers with $\max_u \mu_{u,i} = 1$.