
Optimal Selection Using Algorithmic Rankings with Side Information

Kate Donahue*
Massachusetts Institute of Technology

Nicole Immorlica
Microsoft Research

Brendan Lucier
Microsoft Research

Abstract

In this paper, we model an agent navigating a noisy ranking of candidates, each with different values. In addition to the ranking, the agent has access to a binary signal for each candidate about whether they are “free” or “busy”, where being busy is a signal both of increased candidate quality and decreased candidate availability. For example, in a job market, candidates might be busy if they are already employed. In this paper, we study the incentives and welfare of the three major actors - the firms selecting candidates, the company developing the ranking tool, and the candidates being ranked and society as a whole. First, we study the incentives of the firms, deriving the optimal strategy for selecting candidates, and studying when there are “benefits to congestion” where firms can benefit from free-riding on the hiring decisions of previous firms. Next, we study the welfare implications of this setting, showing that increasing the accuracy of the ranking tool can have paradoxical effects, such as reducing societal welfare (in terms of the total value of employed candidates) and increasing notions of unfairness among candidates. We conclude by discussing the implications our results have for how algorithmically-generated ranking systems should be constructed.

1 Introduction

Consider the following setting: you are a company, trying to find a candidate to fill a hiring slot you have available. You are aware that some candidates are better suited for your role than others - ideally, you would like to hire the best candidate available. To assist you with this goal, you might obtain a (noisy) ranking of the candidates: this could be produced by a human, or obtained from a commercial employment site, such as LinkedIn, Indeed, or other ranking tools. If the ranking is “good” (specifically, if the expected utility of candidates decreases as you navigate further down the list), the best action for the employer is always to pick the top-ranked candidate.

However, in many cases, employers have access to side information, besides the ranking itself. For example, you may be aware that a particular candidate is less likely to accept your offer, potentially because she is already employed. We call such a candidate “busy”, as compared to a “free” candidate without these competing sources of employment. If a “busy” candidate will *never* accept an offer, then the choice is clear: the best option is to hire the highest-ranked “free” candidate. However, if a busy candidate still has some chance of accepting an offer, the possible strategy space changes. In many markets there is an opposing force: it may be the case that high-value candidates are more likely to be “busy”, and thus being “busy” is a signal of candidate quality. For example, this could be

*Work partially completed while at Microsoft Research and Cornell University.

because high value (e.g. talented, highly skilled, or especially well-suited candidates) are more likely to have already been identified and hired by competing firms².

In this way, an employer faces a difficult question: given a noisy ranking of job candidates, along with information about which ones are free or busy, which candidate should she select? Should she pick the top-ranked busy candidate, following the old adage “if you need something done fast, ask a busy person”? Instead, should she aim to avoid the extra hassle of competing for a busy candidate and instead hire the top-ranked free candidate? Are there cases where she should do neither of these things, but pick some other candidate entirely? Moreover, a firm’s hiring strategy has impacts on the rest of society. Specifically, a firm’s choice to select a “busy” (already employed) candidate, rather than a “free” (unemployed) candidate has implications for the unemployment rate, as well as the ability of candidates to leverage competing offers. Additionally, firms’s hiring strategies could lead to disparities in selection rates among candidates with different employment statuses, leading to potential issues of fairness. Such questions have implications for regulation of algorithmic tools, especially the desired properties of ranking tools used in hiring settings.

While we will use the employer hiring job candidates as our motivating example, we note that this setting is quite broad and could model a range of other scenarios. For example, consider a hungry customer navigating a ranking of restaurants, each labeled as “crowded” or “not crowded”: a crowded restaurant is more likely to be of high quality, but will also involve a longer wait.

The rest of this paper proceeds as follows: in Section 2, we will present our model and assumptions, as well as a discussion as to when these assumptions are likely to hold. In Section 3, we discuss the connection of our paper to related work, including the extensive literature on herding and the Pandora’s Box problem. Next, Section 4 analyzes the optimal strategy that a firm navigating a ranking might take: should she pick the first busy candidate, the first free candidate, or some other candidate entirely? In Section 5 we study notions of social welfare, like how frequently firms preferentially select job candidates who are already employed (potentially modeling unemployment discrimination) or avoid doing so (potentially modeling collusion in job hiring). We also analyze incentives of job candidates and fairness notions, such as the gap in selection rate between candidates who are free and those who are busy. Throughout, we consider how increasing the accuracy of the ranking tool can affect the strategy, welfare, and fairness results in previous sections. Surprisingly, we show that increased accuracy can *reduce* measures of social welfare and fairness. In Section 6 we conclude by discussing implications for policy and regulation, as well as potential extensions. Appendix A discusses extensions of our model and all proofs are deferred to Appendix B.

2 Model, notation, and assumptions

In this section, we present our model, notation, and core assumptions.

2.1 Model and notation

First, we present our model. There are N candidates, of which the agent (alternatively, firm) wishes to select exactly one: for example, to interview a job candidate, or pick a restaurant for dinner. The candidates have different true values to the agent, where v_i denotes the value of candidate i . The goal of the agent is to maximize their expected utility. In order to help with this decision, the agent has access to a ranking tool which produces noisy permutations over the candidates: $\sigma \sim \mathcal{P}$, where σ_i denotes the index of the candidate ranked i th in permutation σ with value v_{σ_i} . Throughout, we will require that the expected value of candidates decreases as you go further down the ranking: that is, $\mathbb{E}_{\sigma \sim \mathcal{P}}[v_{\sigma_i}]$ is decreasing in i .

In addition to the ranking, the agent has access to a single bit of side information about each candidate: specifically, a status vector s , where $s_i = 1$ if the candidate ranked in position i is free, and $s_i = 0$ if the candidate is busy. Picking a busy candidate incurs a penalty: specifically, if a candidate has value v_i if picked when they are free, then they have value $\gamma \cdot v_i$ if they are selected when they are busy, for $\gamma \in [0, 1]$. This γ parameter reflects the fact that a busy (e.g. already employed) job candidate is less

²Throughout this paper, we will make the assumption that some candidates are better fits for the job than others (e.g. higher or lower value). We recognize that this assumption drastically simplifies the multi-faced strengths and weaknesses each candidate may bring, and discuss this (and other) assumptions in more depth in Section 2.

likely to accept an offer, for example. We may also view γ as a firm-specific parameter reflecting, for example, the attractiveness of that firm to candidates.

One central assumption is that the firm must select a candidate based solely on the ranking, and cannot re-select another candidate later. This could be motivated by hiring in stages, where “selecting a candidate” corresponds to bringing a candidate on for an intensive onsite, which cannot easily be filled with another candidate if one declines an offer. Thus, we model strategy involved in this first stage of hiring.

We assume that higher-valued candidates are less likely to be available, in particular, that $v_i \geq v_j$ implies $p_i \leq p_j$, where p_i gives the probability that the i th candidate is free. We also assume that each candidate’s probability of being free or busy is independent of every other candidate. This could reflect, for example, the steady state of a hiring market where candidates may be given offers by multiple firms each considering different subsets of candidates. Throughout most of this paper, we will take these probabilities $\{p_i\}$ as fixed, but in Section 6 we will briefly discuss other settings, such as where candidates may lie about their free/busy status (for example, lying about being employed or unemployed so as to increase their chances of getting an interview).

2.2 Superstar setting

In most of this paper, we will focus on the **super-star setting**, with exactly one high-value candidate, and all other candidates with value 0: $v_1 \geq v_2 = v_3 \dots = 0$ and $p_1 \leq p_2 = p_3 \dots$. This models settings where the distribution of candidate quality is skewed: the highest-value individual may have much higher value than other candidates, who are all roughly comparable. In this setting, the space of permutations becomes much smaller, because the only distinguishing feature is the index of the high-value item. We will use the notation σ^i to denote the set of permutations where the high-value item is in the i th index, and in general we will assume that $P[\sigma^i] \geq P[\sigma^{i+1}]$ for all i (that is, the ranking is more likely to rank the high value item higher). This also leads to a natural notion of accuracy: an increase in the ratio of $\frac{P[\sigma^1]}{P[\sigma^i]}$ for all $i \neq 1$.

Definition 1. A ranking $P[\sigma^i]'$ is more accurate than $P[\sigma^i]$ if: $\frac{P[\sigma^1]'}{P[\sigma^i]'} \geq \frac{P[\sigma^1]}{P[\sigma^i]} \forall i > 1$ with the inequality strict in at least one index.

Lemma 2.1 notes that this definition of accuracy implies that the probability of the high value candidate being in the top-ranked position strictly increases with increased accuracy. Note that the converse is *not* true: increasing the accuracy of a ranking distribution does not necessarily decrease $P[\sigma^2]$, for example. For intuition, a more accurate distribution could increase low ranks ($P[\sigma^1]$, $P[\sigma^2]$, $P[\sigma^3]$...) and sharply decrease high ranks ($P[\sigma^{n-1}]$, $P[\sigma^n]$...), so long as $P[\sigma^1]/P[\sigma^i]$ strictly increases (see Appendix A for further discussion).

Lemma 2.1. If distribution $\{P[\sigma^i]'\}$ is more accurate than $\{P[\sigma^i]\}$, then $P[\sigma^1]' > P[\sigma^1]$.

In Appendix A we discuss how our results translate to settings beyond the superstar model.

2.3 Random Utility Model

Our main results will hold for arbitrary permutation distributions in the superstar model $\{P[\sigma^i]\}$. However, we will often find it helpful to use a specific model of permutations to illustrate our main results. In particular, the Random Utility Model is commonly used as a model of permutations. In Proposition 2.1, we show that for the Random Utility Model with Gumbel noise, we can derive closed-form solutions for the permutation distribution. The proof involves an application of the Gumbel trick. To our knowledge, this derivation is novel and would potentially be of broader interest. In Appendix A we discuss further extensions, including its connection to RUM with Gaussian noise.

Proposition 2.1. Consider the superstar setting (1 item of value $v_1 > v_2 = 0$, and $n - 1$ elements of item 0). If i.i.d. Gumbel noise (with scale parameter β) is added to each value, then the probability of the high value item being in index i is given by:

$$P[\sigma^i] = \frac{\exp(v_1/\beta)}{\exp(v_1/\beta) + n - 1} \prod_{k=1}^{i-1} \frac{1}{1 + \frac{\exp(v_1/\beta) - 1}{n - k}}$$

As a direct consequence, for any $i > 2$, the ratio $\frac{P[\sigma^1]}{P[\sigma^i]}$ is given by: $\prod_{k=1}^{i-1} \left(1 + \frac{\exp(v_1/\beta) - 1}{n - k}\right)$.

As a corollary, we can see immediately from the derivation of the RUM Proposition 2.1 that a decrease in noise in the Gumbel distribution immediately leads to an increase in accuracy (according to Definition 1).

Corollary. *For the RUM with Gumbel noise, increasing v_1 or decreasing β induces a distribution $\{P[\sigma^i]'\}$ that is more accurate than $\{P[\sigma^i]\}$ (according to Definition 1).*

2.4 Assumptions and limitations

In the previous section, we described our technical model and assumptions. Here, we will focus on the assumptions and limitations in greater detail. In Section 6 we discuss further changes to our model and the implications they would have for our results. For example, one central assumption is that each candidate has some true intrinsic value to the agent (decision-maker). While this is a very common assumption in theoretical models, we recognize that job candidates (as people) bring nuanced and multi-faceted qualities to their that cannot be fully captured by a single number.

Additionally, one core assumption is that the agent has access to truthful, unbiased signals of free/busy status s . This assumption could easily be violated if such information is only available through the candidates themselves. For example, if candidates are aware that firms are preferentially hiring busy candidates, then they might lie about being employed or having competing offers. While we do not explore such issues in this paper, Horton et al. [2023] explores this question in a dynamical pricing model, giving conditions where pricing levels induce truthful response.

3 Related work

Our model has connections to celebrated models of agents strategically picking items, such as the hiring problem and Pandora’s box problem Krengel and Sucheston [1977], Weitzman [1978]. However, our setting has key differences from these: foremost among them is the assumption that the ordering of elements is generated from some (external) ranking. In particular, this means that agents must examine items in the order they are presented in, rather than being able to re-order them, for example. Another difference is the signaling mechanism: for example, rather than each item having an independent, known distribution of value, we model the scenario where the free/busy signal is generated by items with different probabilities, given different values. This structure, in conjunction with the ranking, means that the free/busy status of item i can affect the expected value of item j - generally not captured in existing models. Finally, one assumption in our setting is that the agent must pick a single item and commit to it without being able to observe its true value. This captures the setting where it is difficult or impossible to evaluate the item without committing to it first - for example, in hiring, when a candidate’s true “value” may not become apparent until after weeks or months of work, or in selecting a restaurant for dinner, where the only way to truly evaluate it may be to sit down for dinner - upon which there is no need (or ability) to consume a second dinner at a different restaurant.

Additionally, our work has connections to long literature on herding (also known as information cascades) Banerjee [1992], Bikhchandani et al. [1992], Welch [1992], a pattern where it may be optimal for a decision-maker to follow actions taken by previous people, even ignoring their own information. Such patterns can lead to “cascades” where multiple sequential decision-makers each follow actions taken by previous decision-makers, rather than following their own information. In our example, when an agent chooses to pick a lower ranked busy candidate over a higher ranked free, this could be seen as herding, since the free/busy status of candidates could be seen as created by actions taken by previous decision-makers, whereas the ranking is assumed to be private (personal) information. To our knowledge, herding has not been extensively studied in ranking settings: our model interestingly allows for herding to be moderated by the strength of private information through the quality of the ranking and the index of the first busy candidate.

Our paper also has connections to the literature on algorithmic monoculture Kleinberg and Raghavan [2021], which studies the case when utility may be *decreased* when two agents both rely on the same algorithmic ranking, rather than their own (uncorrelated) ranking over candidates. Variants of this setting have been explored in later work, such as how algorithmic monoculture can lead to individuals experiencing correlated outcomes (“outcome homogenization”) Bommasani et al. [2022], or the implications of monoculture and of noisy matching in two-sided matching markets Peng and Garg [2023, 2024]. Common features to this work includes the model of multiple agents taking strategic

actions in relation to imperfect rankings, as well as the notion of a penalty for picking an item that another agent has picked. Key differences in our model include the fact that this penalty is softened - in particular, an agent can derive non-zero utility for picking an item that another has already chosen. This dramatically complicates the strategy space, as later sections will show. Additionally, in our setting we modify how multiple firms interact - in particular, we assume that they each have their own independent realization of the ranking, but that they may have access to a common ranking tool (similar accuracy rates). Instead, the monoculture (or herding) phenomenon relates to how agents choose to use the free/busy signal - selecting items that are “busy” could relate to an agent choosing to “free ride” off of the ranking another agent has chosen (we discuss this further in Section 6).

Model multiplicity refers to the case where there might exist multiple equally accurate models that nevertheless differ in the predictions for some nontrivial fraction of the input space. This phenomenon has connections both to algorithmic monoculture and to our setting, because it motivates how firms could have access to rankings with similar accuracy rates and yet distinct orderings over elements. The fairness, welfare, and strategic implications of model multiplicity have been explored in works such as Black et al. [2022], Jain et al. [2023], Cooper et al. [2023], Marx et al. [2020], Hsu and Calmon [2022]. Separately, multiple works have considered the question of fairness in ranking - Zehlike et al. [2021] offers a helpful survey. In particular, Singh et al. [2021] describes a notion of fairness relating to how often items are presented in the top k , given that they have some probability of truly being the best k items (relating to our notion of an item being “picked”), while Peng et al. [2023] studies the tradeoff between diversity and accuracy, showing that consumption constraints (e.g. only top item can be consumed) explains away a tension between these two goals.

Finally, there are multiple papers (empirical and theoretical) exploring models of firms using prediction tools to compete with each other. For example, Filippas et al. [2023] explicitly studies an empirical version of this where job candidates could pay to send a signal of availability to potential employers. Intriguingly, in this setting it seemed like such a signal did *not* lead to adverse selection, with most employers choosing a “first free” strategy, perhaps indicating a parameter regime where this strategy was typically dominant. As mentioned previously, Horton et al. [2023] explores this setting in a dynamical pricing model, giving conditions where pricing levels for such a signal induce a truthful response. Additionally, Jagadeesan et al. [2023] explores a setting where improved data representation (as modeled by reduced Bayes risk) can paradoxically increase the error that users experience when they choose between two competing firms. Castera et al. [2022] studies statistical discrimination in stable matchings where candidates have preferences over firms, who only have noisy access to signals about candidate quality.

4 Optimal Selection Strategies

In this section, we begin by evaluating the firm’s strategy: given access to an algorithmically-generated ranking over candidates and side information about their free-busy status, what is the optimal strategy?

4.1 Derivation of Optimal Strategies

First, we will show that (given a status vector s) there are only two strategies we need to consider within the superstar setting: picking the first free candidate, or the first busy candidate. This dramatically reduces the strategy space firms need to consider³.

Lemma 4.1. *In the superstar setting, given a realized status vector s , it is always optimal to either pick the first free item or the first busy item.*

Having established that a firm will always choose either the first free candidate or the first busy one, our analysis hinges on when firms will select each strategy. Theorem 1 below exactly characterizes this. Specifically, it shows that the strategy space depends on two terms. First, free-busy ratio $r = \frac{p_2/(1-p_2)}{p_1/(1-p_1)}$ measures how strongly status is correlated with value - effectively, this could be seen as a measure of how efficient the market is. Secondly, γ measures the penalty for selecting a busy candidate. If $r \cdot \gamma \leq 1$, firms hire the top-ranked *free* candidate (so long as one is within the top j^* indices, defined formally in Theorem 1), and if $r \cdot \gamma > 1$, firms hire the top-ranked *busy* candidate

³Note that Lemma 4.1 does not necessarily hold beyond the superstar setting: see Appendix A for further discussion.

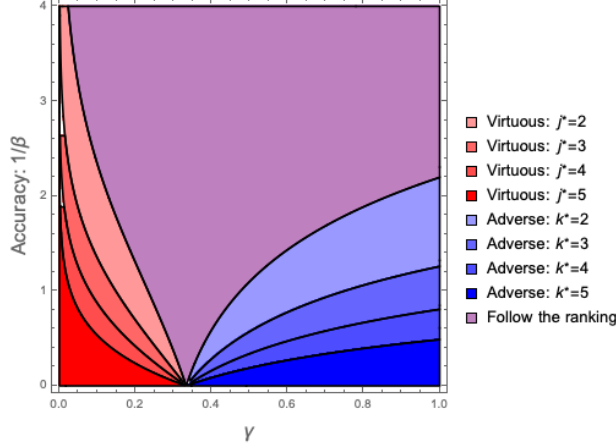


Figure 1: Illustration of optimal strategy for firm, given $n = 5$ candidates and ranking tool with accuracy governed by Random Utility Model with Gumbel noise (see Section 2.3). The x axis varies the γ busy penalty, while the y axis varies accuracy as parameterized by the $1/\beta$ Gumbel noise parameter (higher values increase accuracy). Shades of red indicate regions where virtuous selection is the optimal strategy and shades of blue indicate where adverse selection is (darker shades indicate regions where the firm has larger j^* or k^* and thus “hunts” further down the list to find a candidate with their preferred free/busy status). The purple region is where the optimal strategy is “follow the ranking”. This figure uses $r = 3$, so $1/r = 1/3$ is the critical point on the x-axis where firms switch from using virtuous selection to adverse selection.

(again, so long as they’re within the top k^* indices, defined formally in Theorem 1). For intuition, depending on $r \cdot \gamma$, firms have preferences for candidates who are either free or busy. However, the further down the ranking the firms search, the less likely it is that a given candidate will be high value. Therefore, the optimal strategy is for firms to pick the top-ranked candidate of their preferred status, so long as that candidate doesn’t come “too far” down the ranking. Effectively, Theorem 1 shows how firms combine information they have from the algorithmically generated ranking along with the free/busy status vector.

Theorem 1. [Optimal hiring for firms] If $r \cdot \gamma \leq 1$, then define j^* as the largest j such that $\frac{P[\sigma^1]}{P[\sigma^j]} < \frac{1}{r \cdot \gamma}$. Then, the optimal strategy for the firm is to use the j^* -**virtuous selection** strategy: select the highest-ranked free candidate within the top j^* . If all of the top j^* candidates are busy, hire the top-ranked candidate.

Conversely, if $r \cdot \gamma > 1$, define k^* as the largest k such that $\frac{P[\sigma^1]}{P[\sigma^k]} < r \cdot \gamma$. Then, the optimal strategy for the firm is to use the k^* -**adverse selection** strategy: select the highest-ranked busy candidate within the top k^* . If all of the top k^* candidates are free, hire the top-ranked candidate.

If $j^* = 1$ or $k^* = 1$, then optimal strategy for the firm is to use the **follow the ranking** strategy: always select the top-ranked candidate, regardless of whether they are free or busy.

Throughout the rest of this paper, we will assume that firms are always acting optimally, using the strategies defined in Theorem 1.

4.2 Accuracy of Ranking Tool

Next, we will discuss the implications of increasing accuracy on the firm’s optimal strategy. Figure 1 illustrates the strategy space for various different levels of accuracy and busy penalty γ . Note that this figure displays curves specifically for the RUM with Gumbel noise: this is purely to build intuition, as our theoretical results hold for arbitrary probability distributions. As Theorem 1 suggests, virtuous selection is only optimal for $\gamma < \frac{1}{r}$ and adverse selection is only optimal for $\gamma > \frac{1}{r}$.

Figure 1 displays other phenomena: for example, increasing the accuracy of the ranking distribution always reduces how far down the ranking firms search (reducing j^* or k^*). Lemma 4.2 proves this property theoretically: increasing the accuracy of the ranking tool always (weakly) decreases how far down the list firms may hunt for candidates.

Lemma 4.2. *Holding γ constant, increasing the accuracy of the ranking tool always holds constant or decreases the furthest index a firm needs to consider (decreases j^* or k^* as defined in Theorem 1).*

Additionally, Figure 1 seems to show an asymmetry between virtuous selection and adverse selection: for a sufficiently accurate ranking, all firms with $\gamma > 1/r$ will use “follow the ranking”, but the same is not true for virtuous selection (observe how the red curves go towards infinity as γ goes to 0). Lemma 4.3 proves this theoretically. For intuition on this asymmetry, firms using virtuous selection are unable to pay the penalty of hiring a busy candidate, so may avoid busy candidates even if the ranking tool is very accurate, while firms doing adverse selection are simply hunting for the strongest signal and will follow that, whether it’s the ranking or the free/busy signal.

Lemma 4.3. *No matter how accurate a ranking is, there exists a firm whose optimal strategy is virtuous selection with $j^* = n$. More precisely, for all $\epsilon > 0$ such that $P[\sigma^1]/P[\sigma^n] = \epsilon$, there exists a busy penalty $\gamma < \frac{1}{r}$ such that $k^* = n$.*

For all firms with $\gamma \geq \frac{1}{r}$, if $P[\sigma^1]/P[\sigma^2] > r$, the optimal strategy is always “follow the ranking”.

Finally, we discuss the impact of increased accuracy on the ranking company itself. We may view the ranking company as wishing firms to preferentially select top-ranked candidates: this could be because of potential advertising revenue which is higher for top-ranked items, or because firms use click data to get feedback on the value of items. However, Lemma 4.4 shows that increasing the accuracy of the ranking tool can *decrease* the probability that a firm selects the top ranked candidate: intuitively, this can happen when a firm is using virtuous selection and the ranking tool becomes more likely to result in status vectors where the top-ranked candidate is busy.

Lemma 4.4. *Increasing the accuracy of the ranking tool can increase, decrease, or keep constant the probability that a firm selects the top ranked candidate.*

5 Welfare, Fairness, and Incentives

Section 4 described the optimal strategies that firms select under different parameter regimes: virtuous selection (preferentially select free candidates), adverse selection (preferentially select busy candidates), and follow the ranking (always hire the top ranked candidate). Additionally, these strategies have different impacts on the broader system, such as influences on welfare, fairness, and incentives, which we will analyze in this section.

5.1 Duplicative Effort and Benefits to Congestion

First, we will consider impacts that the firm’s hiring strategy has on overall welfare. For example, a firm using adverse selection will select busy candidates much more frequently than a firm using virtuous selection. We formalize this notion as *duplicative effort* in Definition 2. Intuitively, duplicative effort relates to the amount of welfare that a firm obtains by selecting candidates who are already busy. Duplicative effort could involve negative impacts: for example, it could be viewed as firms “poaching” candidates from other firms rather than hiring unemployed candidates, which may lead to a higher unemployment rate. However, duplicative effort could also be viewed positively: candidates who receive competing offers can use them to increase their compensation. Moreover, explicitly trying to reduce duplicative effort could run afoul of anti-trust laws.

Definition 2. *Duplicative effort is the amount of a firm’s utility that comes from giving offers who candidates who are already busy (e.g. employed), given by: $\gamma \cdot \mathbb{E}[v_{\rho_{i(s)}}] \cdot \mathbb{1}[s_{i(s)} = 0]$*

What is the impact of accuracy on duplicative effort? Lemma 5.1 shows that duplicative effort can *increase* or *decrease*. For intuition, increasing the accuracy of the ranking tool has multiple opposing effects: first, if the top-ranked candidate is free, a more accurate ranking tool can potentially overcome the negative signal of being free (unemployed), increasing the chance that the firm will select it (reducing duplicative effort). Secondly, if the top-ranked candidate is busy, a more accurate ranking tool increases the expected value of this candidate, and so a firm may decide that it’s now worth it to hire the busy candidate even with the γ penalty (increasing duplicative effort). Finally, increasing the accuracy changes the frequency with which different patterns of status vectors will occur (potentially increasing or decreasing duplicative effort).

Lemma 5.1. *There exist settings where increasing the accuracy of the ranking tool can increase or decrease duplicative effort.*

Specifically, in the example in Lemma 5.1, duplicative effort is increased for firms with $r \cdot \gamma < 1$ and decreased for firms with $r \cdot \gamma \geq 1$. If we view γ_i as an indication of how attractive a particular firm i is, then this result indicates that increasing the accuracy of the ranking tool may increase the number of competing offers that candidates receive from less attractive firms, while reducing the number from more attractive firms.

Next, we will consider another factor: whether firms benefit from the fact that other firms are also hiring. Specifically (given $p_1, p_2 < 1$), every candidate has some nonzero probability of being busy (e.g. employed). On the one hand, this reduces firms' utility because selecting a busy candidate incurs a γ penalty. On the other hand, because $p_1 < p_2$, seeing which candidates are free and busy gives firms valuable noisy information about their value, which they could use to guide hiring decisions. These two counterbalancing effects means that there are sometimes positive *benefits to congestion*: firms strictly benefit from the fact that some candidates are already employed (formally defined in Definition 3). This result could be seen in multiple ways, such as firms "free-riding" on the hiring decisions of other firms.

Definition 3. A firm's "benefit to congestion" is the difference between its utility with $p_1, p_2 < 1$ and its utility if $p_1 = p_2 = 1$ (every candidate is always free). If this quantity is positive, a firm has positive benefits to congestion and has higher utility when candidates have some nonzero probability of being busy.

Lemma 5.2 exactly specifies when a firm experiences positive benefits to congestion: specifically, when $r \cdot \gamma$ is high, and when the probability of the high value candidate being ranked first ($P[\sigma^1]$) isn't too high:

Lemma 5.2. A firm using virtuous selection or "follow the ranking" never experiences positive benefits to congestion.

A firm using adverse selection with $k^* \geq 2$ experiences positive benefits to congestion if and only if:
$$\frac{(1-p_1) \cdot \gamma}{p_1 \cdot (1-p_2^{k^*-1}) + (1-p_1) \cdot (1-\gamma)} \sum_{k=2}^{k^*} p_2^{k-1} P[\sigma^k] > P[\sigma^1]$$
 with k^* defined as in Theorem 1.

Finally, Lemma 5.3 shows that increasing the accuracy always *decreases* benefits to congestion.

Lemma 5.3. Increasing the accuracy of the ranking tool always reduces the firms' benefit to congestion.

5.2 Free-busy gap

In Section 5.1 we analyzed this setting from the perspective of the firm: the optimal selection strategy they should take and related factors like duplicative effort and benefits to congestion. In this section, we analyze the same setting from the perspective of job candidates. In particular, we will focus on the *free-busy gap*, which measures how more (or less) likely a free candidate is to be hired, holding value constant.

Definition 4. The free-busy gap is the difference in the probability that a free candidate is selected, as compared to the probability that a busy candidate is selected (holding value constant).

Again, the free-busy gap can be interpreted in multiple ways. If it's negative (busy candidates are more likely to be hired), then this could be viewed as unemployment discrimination: if it's positive, then it could be viewed as firms productively matching with available candidates. So far, we've assumed that job candidates are honest about their free-busy status: however, if the free-busy gap is very large in magnitude (either positive or negative), we could view this as a setting where job candidates may be strongly incentivized to lie in order to increase their chances of being selected.

As expected, if a firm is using virtuous selection (preferentially selecting free candidates), the free-busy gap will be positive, and negative if a firm is using adverse selection (as Lemma 5.4 formalizes). Part of this proof relies on Lemma B.3 which derives closed-form solutions for the free-busy gap (for conciseness, the statement and proof of this result is deferred to Appendix B).

Lemma 5.4. For all parameters, the free-busy gap is strictly positive if the firm is using virtuous selection, strictly negative if the firm is using adverse selection, and zero if the firm is using follow the ranking.

Next, we consider the impact of increased accuracy on the magnitude of the free-busy gap, which could be viewed unemployment discrimination or the incentive of candidates to lie about the employ-

ment status. In particular, Lemma 5.5 shows that there exist settings where increasing the accuracy could *increase* or *decrease* the magnitude of the free-busy gap. For intuition, consider the free-busy gap for the high value candidate when a firm is using virtuous selection. A free high value candidate is selected unless it's ranked second and the low value candidate is also free (event A), while a busy high value candidate would be selected only if it's ranked first and the low value candidate is also busy (event B). If $p_2 > 0.5$, increasing the accuracy of the ranking tool decreases the probability of event A more quickly than it increases the probability of event B, which means that the free-busy gap would increase. The full proof formalizes this intuition.

Lemma 5.5. *There exist settings where increasing the accuracy of the ranking tool can either increase or decrease the magnitude of the free-busy gap.*

However, this story becomes more complex if candidates are unsure about the strategy each firm is using. Recall from Section 4 that a firm's strategy depends on the quantity $r \cdot \gamma$, where γ is a penalty for picking a busy candidate. This penalty may be viewed as firm-specific and potentially unknown to candidates. In particular, candidates could have uncertainty in a single firm's parameters, or if multiple firms may simultaneously be using the same ranking tool (reflecting a diversity of preferences). In general, we would expect such uncertainty to *reduce* the free-busy gap.

6 Discussion

In this paper, we have proposed a simple model to study settings where strategic, self-interested agents use algorithmically generated rankings in the presence of side information, specifically a free-busy status vector for different candidates. Our results show that even this relatively simple model can have capture surprisingly complex phenomena. There are multiple fascinating extensions to our model. In particular, one interesting direction generalizing beyond the superstar model to include candidates with arbitrarily many values and probabilities of being available. We discuss some results in this direction in Appendix A: largely the same results hold, but with some surprising nuances. Additional extensions could consider other types of policy interventions, such as whether deliberately hiding the free-busy status of candidates could lead to better welfare or fairness guarantees. Finally, other extensions could consider cases when agents can deliberately mislead others, such as job candidates lying about their free-busy status.

Impacts on Policy and Regulation:

Strategy: Our results show that firms using algorithmic rankings can best-respond by strategically selecting candidates based on their free-busy status, so long as they are within the top k candidates the ranking returns. Increasing the accuracy of the tool can reduce incentives of firms to strategise, but can also decrease measures of social welfare like duplicative effort. These results generally assume that firms perfectly know certain parameters such as the r, γ (which could be viewed as measuring the efficiency of the market and the attractiveness of any given firm, respectively). In reality, firms only have access to estimates of these values, and in particular they may systematically mis-estimate these values, for example, assuming that their firm is more attractive than it really is. In this case, firms may use sub-optimal strategies, which may lead to increased levels of duplicative effort.

Fairness and incentives to lie: When firms select preferentially based on the free-busy (e.g. employment) status of candidates, this immediately leads to differences in selection rate based on employment status. This could immediately have fairness implications, but could also incentivize candidates to lie about their employment status: these incentives could be larger for high-value candidates or low-value candidates. We showed that these disparities can be reduced when firms have diverse preferences, but can surprisingly be *increased* when the accuracy of the ranking tool is increased.

Ranking company: Often, the incentives of ranking companies, their clients, and societal welfare are assumed to be at least weakly aligned: producing higher-quality ranking tools should increase the value to both companies and clients. However, our results show that their incentives may be misaligned: in particular, a more accurate ranking tool can reduce measures of social welfare and fairness, and can even reduce the welfare of the ranking company itself. Our results suggest that improving overall welfare may require more precise interventions than simply increasing the accuracy of tools.

Acknowledgments and Disclosure of Funding

We are grateful to Keegan Harris, John Horton, Meena Jagadeesen, Sloan Neitert, Michelle Si, Kiran Tomlinson, and Nicolas Wu for invaluable discussions.

References

- Abhijit V Banerjee. A simple model of herd behavior. *The quarterly journal of economics*, 107(3): 797–817, 1992.
- Sushil Bikhchandani, David Hirshleifer, and Ivo Welch. A theory of fads, fashion, custom, and cultural change as informational cascades. *Journal of political Economy*, 100(5):992–1026, 1992.
- Emily Black, Manish Raghavan, and Solon Barocas. Model multiplicity: Opportunities, concerns, and solutions. In *Proceedings of the 2022 ACM Conference on Fairness, Accountability, and Transparency*, FAccT ’22, page 850–863, New York, NY, USA, 2022. Association for Computing Machinery. ISBN 9781450393522. doi: 10.1145/3531146.3533149. URL <https://doi.org/10.1145/3531146.3533149>.
- Rishi Bommasani, Kathleen A. Creel, Ananya Kumar, Dan Jurafsky, and Percy S Liang. Picking on the same person: Does algorithmic monoculture lead to outcome homogenization? In S. Koyejo, S. Mohamed, A. Agarwal, D. Belgrave, K. Cho, and A. Oh, editors, *Advances in Neural Information Processing Systems*, volume 35, pages 3663–3678. Curran Associates, Inc., 2022. URL https://proceedings.neurips.cc/paper_files/paper/2022/file/17a234c91f746d9625a75cf8a8731ee2-Paper-Conference.pdf.
- Rémi Castera, Patrick Loiseau, and Bary S.R. Pradelski. Statistical discrimination in stable matchings. In *Proceedings of the 23rd ACM Conference on Economics and Computation*, EC ’22, page 373–374, New York, NY, USA, 2022. Association for Computing Machinery. ISBN 9781450391504. doi: 10.1145/3490486.3538364. URL <https://doi.org/10.1145/3490486.3538364>.
- A Feder Cooper, Katherine Lee, Solon Barocas, Christopher De Sa, Siddhartha Sen, and Baobao Zhang. Is my prediction arbitrary? measuring self-consistency in fair classification. *arXiv preprint arXiv:2301.11562*, 2023.
- Apostolos Filippas, John J Horton, and Diego Urraca. Advertising as coordination: Evidence from a field experiment. 2023.
- John Horton, Nicole Immorlica, Brendan Lucier, and Michelle Si. Costly signaling in online labor markets, 2023.
- Hsiang Hsu and Flavio Calmon. Rashomon capacity: A metric for predictive multiplicity in classification. In S. Koyejo, S. Mohamed, A. Agarwal, D. Belgrave, K. Cho, and A. Oh, editors, *Advances in Neural Information Processing Systems*, volume 35, pages 28988–29000. Curran Associates, Inc., 2022. URL https://proceedings.neurips.cc/paper_files/paper/2022/file/ba4caa85ecdcafbf9102ab8ec384182d-Paper-Conference.pdf.
- Meena Jagadeesan, Michael I. Jordan, Jacob Steinhardt, and Nika Haghtalab. Improved bayes risk can yield reduced social welfare under competition, 2023.
- Shomik Jain, Vinith Suriyakumar, Kathleen Creel, and Ashia Wilson. Algorithmic pluralism: A structural approach to equal opportunity, 2023.
- Jon Kleinberg and Manish Raghavan. Algorithmic monoculture and social welfare. *Proceedings of the National Academy of Sciences*, 118(22):e2018340118, 2021.
- Ulrich Krengel and Louis Sucheston. Semiamarts and finite values. 1977.
- Colin L Mallows. Non-null ranking models. i. *Biometrika*, 44(1/2):114–130, 1957.

- Charles Marx, Flavio Calmon, and Berk Ustun. Predictive multiplicity in classification. In Hal Daumé III and Aarti Singh, editors, *Proceedings of the 37th International Conference on Machine Learning*, volume 119 of *Proceedings of Machine Learning Research*, pages 6765–6774. PMLR, 13–18 Jul 2020. URL <https://proceedings.mlr.press/v119/marx20a.html>.
- Kenny Peng and Nikhil Garg. Monoculture in matching markets, 2023.
- Kenny Peng and Nikhil Garg. Wisdom and foolishness of noisy matching markets. *arXiv preprint arXiv:2402.16771*, 2024.
- Kenny Peng, Manish Raghavan, Emma Pierson, Jon Kleinberg, and Nikhil Garg. Reconciling the accuracy-diversity trade-off in recommendations, 2023.
- Ashudeep Singh, David Kempe, and Thorsten Joachims. Fairness in ranking under uncertainty. In M. Ranzato, A. Beygelzimer, Y. Dauphin, P.S. Liang, and J. Wortman Vaughan, editors, *Advances in Neural Information Processing Systems*, volume 34, pages 11896–11908. Curran Associates, Inc., 2021. URL https://proceedings.neurips.cc/paper_files/paper/2021/file/63c3ddcc7b23daa1e42dc41f9a44a873-Paper.pdf.
- Louis L Thurstone. A law of comparative judgment. *Psychological review*, 101(2):266, 1994.
- Martin Weitzman. *Optimal search for the best alternative*, volume 78. Department of Energy, 1978.
- Ivo Welch. Sequential sales, learning, and cascades. *The Journal of finance*, 47(2):695–732, 1992.
- Meike Zehlike, Ke Yang, and Julia Stoyanovich. Fairness in ranking: A survey. *arXiv preprint arXiv:2103.14000*, 2021.

A Ranking models

A.1 Beyond superstar

In this section, we further relax the model and consider cases where the N items can come in more than two types of values (and probabilities of being available). However, we will find that even relaxing the superstar setting slightly can cause the strategy space for the agent to become much more complex. Specifically, Lemma A.1 shows that it may be optimal to pick a strategy other than “first busy, first free”: for example, to pick the second-ranked free item rather than the top-ranked free item.

Lemma A.1. *There exists a probability distribution $\mathcal{P}$ such that for some realized status vector s , picking the second free item maximizes expected utility, even if $\mathcal{P}$ is descending in expected value.*

Proof sketch. We construct this example by creating a setting with exactly three items with $[v_1 = 1, v_2, v_3 = 0]$. We will create a permutation distribution with nonzero support on exactly two permutations:

$$\sigma^1 = [v_1, v_2, v_3] \quad \sigma^2 = [v_3, v_2, v_1]$$

where $P[\sigma^1] = 1 - \epsilon$, $P[\sigma^2] = \epsilon$. We will also set $p_1 < p_2 = p_3$, so items 2 and 3 have the same probability of being free, which is greater than for item 1.

We consider the status vector $[1, 1, 0]$, so the first two items are free, while the last is busy. At a high level, this status vector increases the posterior belief that the true permutation is σ^2 , with the high value ranked last. Carefully setting v_1, v_2, p_1, p_2 results in a posterior distribution where the expected value of the second item is higher than either the expected value of the first free item or first busy item. $\square$

Note that the ranking in Lemma A.1 did satisfy descending expected value. However, rankings were somewhat unusual - in particular, there were only two orderings with nonzero probability, even though, given 3 items, there are 6 possible permutations of each of them. A more natural ranking would probably have some nonzero weight over each of these, ideally with greater weight on rankings where more of the items are ordered “correctly” (that is, with $v_{\sigma_i} > v_{\sigma_j}$ for $i < j$).

Definition 5 exactly describes this intuition. Specifically, it defines an “inversion-monotone” probability distribution as one where, for any permutation σ with at least two items inverted (that is,

$v_{\sigma_i} < v_{\sigma_j}$ for $i < j$), there exists another permutation $\tilde{\sigma}$ that is identical to σ but with items in indices i, j flipped, and with probability at least as high as σ ($P[\tilde{\sigma}] \geq P[\sigma]$). Lemma A.2 proves that, if a probability distribution is inversion-monotone, we regain the same property we had in the superstar setting: the optimal strategy for the agent will always be to pick the first free item or first busy item.

Definition 5. A probability distribution over permutations $\mathcal{P}$ is called inversion-monotone if it satisfies the following condition: for any permutation σ with a pair of indices $i < j$ such that $v_{\sigma_i} > v_{\sigma_j}$, we construct a corresponding permutation $\tilde{\sigma}$ where $\tilde{\sigma} = \sigma$ except for indices i, j , which we have flipped: $v_{\tilde{\sigma}_i} = v_{\sigma_j}, v_{\tilde{\sigma}_j} = v_{\sigma_i}$. Then, we require:

$$P[\sigma] \geq P[\tilde{\sigma}]$$

Lemma A.2. Consider any vector of realized status vector s , with $s_i = s_j$, for some $i < j$. Then, if the probability distribution of permutations is inversion-monotone as in Definition 5, $\mathbb{E}[v_{\sigma_i} | a] \geq \mathbb{E}[v_{\sigma_j} | a]$. This implies that the optimal solution will always be to pick the first free item or the first busy item.

While Lemma A.2 shows that inversion-monotonicity will guarantee a more straightforward strategy space, it is natural to wonder how reasonable such a requirement is. In fact, it turns out that multiple commonly-used models of permutations already satisfy this property.

First, Lemma A.3 shows that the Mallows model is inversion-monotone. While helpful, this is not extremely surprising: the Mallows model Mallows [1957] is constructed such that the probability of a permutation σ occurring is directly tied to the number of pairs of items that are inverted, so this property follows naturally.

Lemma A.3. The Mallows model is inversion-monotone.

While the Mallows model is frequently used as a model of permutations, it does have drawbacks. Specifically, one desirable property of rankings is that items with very different true values should be less likely to be inverted. For example, given $[v_1 = 10, v_2 = 9, v_3 = 0]$, we would expect items v_1, v_2 to be more frequently inverted than v_2, v_3 , even though their ordinal ranks differ by the same amount. This type of property cannot be expressed by the Mallows model - but it can be captured by the Random Utility Model Thurstone [1994]. In the Random Utility Model (RUM), while each item has some true value $\{v_i\}$, it is assumed that they are ranked by noised versions of these values, such as additive Gaussian noise ($\hat{v}_i \sim \mathcal{N}(v_i, \sigma^2)$). This automatically satisfies the property that items with more similar true values will be more likely to be swapped. Additionally, it seems likely that the Random Utility Model might better capture the performance of rankings produced by humans or algorithmic tools, which might have some sense of the “true value” of each item, but make small, independent errors in their estimation of these values. However, does the Random Utility Model satisfy inversion-monotonicity?

Theorem 2 answers this question in the affirmative. While we believe that this property may be of independent interest, to our knowledge, we are the first paper to prove such a property for the Random Utility model. The proof of Theorem 2 is largely proven by Lemma A.4, which is surprisingly subtle.

Theorem 2. The RUM with identical, symmetric, noise distributions across items is inversion-monotone.

Lemma A.4. Suppose we have two random variables given by $X_1 = \mu_1 + \epsilon_1, X_2 = \mu_2 + \epsilon_2$, for $\mu_1 > \mu_2$ and $\epsilon_1, \epsilon_2 \sim \mathcal{D}$ for a symmetric, single-peaked distribution $\mathcal{D}$. Then, if $|X_1 - X_2| = \Delta$,

$$P[X_1 > X_2 \mid |X_1 - X_2| = \Delta] > P[X_1 < X_2 \mid |X_1 - X_2| = \Delta]$$

These results tell us that, even in settings beyond the superstar model, for cases where permutations are generated by commonly-used models, the optimal strategy will still be to either pick the first free item or the first busy item.

A.2 General superstar setting: $v_2 > 0$

Next, we consider the general superstar setting, where $v_2 \geq 0$. Note that Lemma 4.1, which shows that either “first free” or “first busy” is optimal, still applies in this setting. However, the conditions for which of these strategies will be optimal is more complex. For intuition, when $v_2 = 0$, the

s	$\mathbb{E}[v_{\sigma_1}]$	$\mathbb{E}[v_{\sigma_2}]$	$\mathbb{E}[v_{\sigma_3}]$
[0, 0, 0]	0.461	0.431	0.408
[0, 0, 1]	0.502	0.463	0.517
[0, 1, 0]	0.517	0.523	0.443
[0, 1, 1]	0.611	0.534	0.527
[1, 0, 0]	0.533	0.495	0.458
[1, 0, 1]	0.553	0.594	0.533
[1, 1, 0]	0.563	0.548	0.578
[1, 1, 1]	0.706	0.663	0.63

Table 1: Expected utility for given different realized status vectors s . Permutations given by Random Utility Model with $v_1 = 1, v_2 = v_3 = 0.5, p_1 = 0.1, p_2 = p_3 = 0.53$ (giving $r \approx 10$), with noise given by normal distribution with standard deviation of 1.75, discounting factor of $\gamma = 0.65$. Expected value calculated based on $5 \cdot 10^5$ simulations. Within each row, **bold** numbers indicate the item with the highest expected value.

objective of the agent can be simplified to “find where the high-value item is likely to be” (moderated by the added penalty for picking a busy item). However, if $v_2 > 0$, the agent must trade off their desire to find the high-value item with their potential willingness to settle for the (nonzero) reward of the low-value item. In this section, we will demonstrate how this seemingly minor change has substantial implications for the agent’s best strategy.

In particular, we will show that the decision about whether to pick the first or j th item may depend on the status of items ranked $k > j$. This dependence arises because the status of any item can influence the posterior probability of where the high value item is most likely to be, which influences which strategy is best.

First, Table 1 shows this effect empirically for an example using the Random Utility Model (RUM) with $N = 3$ in the superstar setting, showing expected utility of each index given each of the possible status vectors. We can note two examples where this property is violated (highlighted in blue): first, consider the case where the status vector is given by $s = [0, 1, *]$ for $* \in \{0, 1\}$. If the last entry is 0, then picking “first free” is optimal, but if the last entry is 1, picking “first busy” is optimal. We can also consider the case where the status vector is $s = [1, 0, *]$. Here, if $* = 0$, “first free” is optimal, but if $* = 1$, “first busy” becomes optimal. For intuition, the reason why this happens is because when the last item is free, this increases the posterior belief that items that are busy are the high value item (and an opposite effect occurs when the last item is busy). This effect does *not* occur in the $v_2 = 0$ setting because the only permutations that gives the agent non-zero utility is when the high value item comes in index 1 or j - the status of items in other indices ends up being irrelevant.

We can encapsulate this informal reasoning in a theoretical condition for the optimal strategy for the agent, shown in Theorem 3. This extends the optimal strategy result in Theorem 1, which applied only in the $v_2 = 0$ setting. Note that Theorem 3 similarly has a pair of thresholds, which relate to the lowest index j with a different status from the first item ($s_j \neq s_1$), as well as the ratio of probabilities $r = \frac{p_2/(1-p_2)}{p_1/(1-p_1)}$ which indicates the strength of signal from an item being busy. Similar to Theorem 1, these thresholds are identical given the uniform distribution $P[\sigma^i] = 1/N$ and $j = 2$ (given $j > 2$, the thresholds are conditioning on different status vectors, and thus involve different signals about the presence of the high value item). However, the key difference from the $v_2 = 0$ setting is that the thresholds in Theorem 3 depend on the status of items $k > j$ - which directly violates the “irrelevance of lower ranked items” property we would wish holds.

Theorem 3. Consider the superstar setting with $v_2 > 0$. Then, the agent’s optimal decision is given by:

$$\begin{aligned}
 \text{if } s_{i < j} = 0, s_j = 1 & \begin{cases} \text{pick the first item} & \gamma > T_b^j \\ \text{pick the } j\text{th item} & \text{otherwise} \end{cases} \\
 \text{if } s_{i < j} = 1, s_j = 0 & \begin{cases} \text{pick the first item} & \gamma < T_f^j \\ \text{pick the } j\text{th item} & \text{otherwise} \end{cases}
 \end{aligned}$$

For:

$$T_b^j = \frac{(v_1 - v_2) \cdot P[\sigma^j] + v_2 \cdot (r \cdot P[\sigma^{i < j}] + P[\sigma^j] + C)}{(v_1 - v_2) \cdot r \cdot P[\sigma^1] + v_2 \cdot (r \cdot P[\sigma^{i < j}] + P[\sigma^j] + C)}$$

$$T_f^j = \frac{(v_1 - v_2) \cdot P[\sigma^1] + v_2 \cdot (P[\sigma^{i < j}] + r \cdot P[\sigma^j] + C)}{(v_1 - v_2) \cdot r \cdot P[\sigma^j] + v_2 \cdot (P[\sigma^{i < j}] + r \cdot P[\sigma^j] + C)}$$

where $r = \frac{p_2/(1-p_2)}{p_1/(1-p_1)}$ and $C = \sum_{i=j+1}^N r^{1-s_i^j} P[\sigma^i]$.

Theorem 3's dependence on the status of items $k > j$ has implications for the strategy of the agent - and for the best design of the ranking tool it uses. One potential hope might be that a "sufficiently accurate" ranking might avoid this reliance on the status of lower ranked items. However, this unfortunately is not the case - as Lemma A.5 shows, any ranking that has a nonzero chance of having the high-value item appear at some index $k > j$ could result in requiring the decision-maker to consider the status of items down to index k .

Lemma A.5. *For any $N > j \geq 2$, if there exists $k > j \geq 2$ such that $P[\sigma^k] = \epsilon > 0$, then there exists γ, s such that observing the status vector up to index k will change the choice of which item will be selected.*

However, in Lemma A.6 we provide tight bounds on the thresholds in Theorem 3 that the agent must use to make a decision. In particular, if the penalty γ falls outside of these bounds, then the agent can select make an optimal decision based solely on the status of the first j items. Because these bounds are tight, if γ falls within them, the agent would need to consider further items in order to ensure its decision is optimal.

Lemma A.6. *It is possible to construct tight bounds on the thresholds in Theorem 3 using the status of only the first $k \geq j$ items. Specifically, such bounds are given by the terms below (with $C_0 = \sum_{i=j+1}^k r^{1-s_i^j} P[\sigma^i]$ throughout):*

When $s_{i < j} = 0, s_j = 1$, we have:

$$B_{j,k}^1 \leq T_b^j \leq B_{j,k}^2$$

For:

$$B_{j,k}^1 = \frac{(v_1 - v_2) \cdot P[\sigma^j] + v_2 \cdot (r - 1) \cdot P[\sigma^{i < j}] + v_2 + v_2 \cdot (C_0 - P[\sigma^{j < i \leq k}])}{(v_1 - v_2) \cdot P[\sigma^1] + v_2 \cdot (r - 1) \cdot P[\sigma^{i < j}] + v_2 + v_2 \cdot (C_0 - P[\sigma^{j < i \leq k}])}$$

$$B_{j,k}^2 = \frac{(v_1 - v_2) \cdot P[\sigma^j] + v_2 \cdot (r - 1) \cdot (1 - P[\sigma^j]) + v_2 + v_2 \cdot (C_0 - r \cdot P[\sigma^{i > k}])}{(v_1 - v_2) \cdot r \cdot P[\sigma^1] + v_2 \cdot (r - 1) \cdot (1 - P[\sigma^j]) + v_2 + v_2 \cdot (C_0 - r \cdot P[\sigma^{i > k}])}$$

where the upper bound is given by the case where $s_\ell = 0 \forall \ell > k$, and the lower bound is given by $s_\ell = 1 \forall \ell > k$.

For the case when $s_{i < j} = 1, s_j = 0$, we have:

$$\begin{cases} F_{j,k}^1 \leq T_f^j \leq F_{j,k}^2 & P[\sigma^1] \leq r \cdot P[\sigma^j] \\ F_{j,k}^2 \leq T_f^j \leq F_{j,k}^1 & \text{otherwise} \end{cases}$$

for

$$F_{j,k}^1 = \frac{(v_1 - v_2) \cdot P[\sigma^1] + v_2 \cdot (P[\sigma^{i < j}] + r \cdot P[\sigma^j] + C_0 - P[\sigma^{i > k}])}{(v_1 - v_2) \cdot r \cdot P[\sigma^j] + v_2 \cdot (P[\sigma^{i < j}] + r \cdot P[\sigma^j] + C_0 - P[\sigma^{i > k}])}$$

$$F_{j,k}^2 = \frac{(v_1 - v_2) \cdot P[\sigma^1] + v_2 \cdot (1 + (r - 1) \cdot P[\sigma^{i > j}] + C_0 - r \cdot P[\sigma^{j < i \leq k}])}{(v_1 - v_2) \cdot r \cdot P[\sigma^j] + v_2 \cdot (1 + (r - 1) \cdot P[\sigma^{i > j}] + C_0 - r \cdot P[\sigma^{j < i \leq k}])}$$

where the C_j^2 bound occurs when $s_\ell = 0 \forall \ell > k$, and the C_j^1 bound occurs when $s_\ell = 1 \forall \ell > k$.

We can use this bound to reason about how far down the status vector an agent must look to ensure that they are making the correct decision, shown below.

Corollary. *An agent can ensure it is picking the best item by observing the first $k \geq j$ items, with k defined as follows: If the first item is busy, then*

$$\min_{k \geq j} \text{ s.t. } B_{j,k}^1 \geq \gamma \text{ or } B_{j,k}^2 \leq \gamma$$

If the first item is free, then

$$\min_{k \geq j} \text{ s.t. } \begin{cases} F_{j,k}^1 \geq \gamma \text{ or } F_{j,k}^2 \leq \gamma & P[\sigma^1] \leq r \cdot P[\sigma^j] \\ F_{j,k}^2 \geq \gamma \text{ or } F_{j,k}^1 \leq \gamma & \text{otherwise} \end{cases}$$

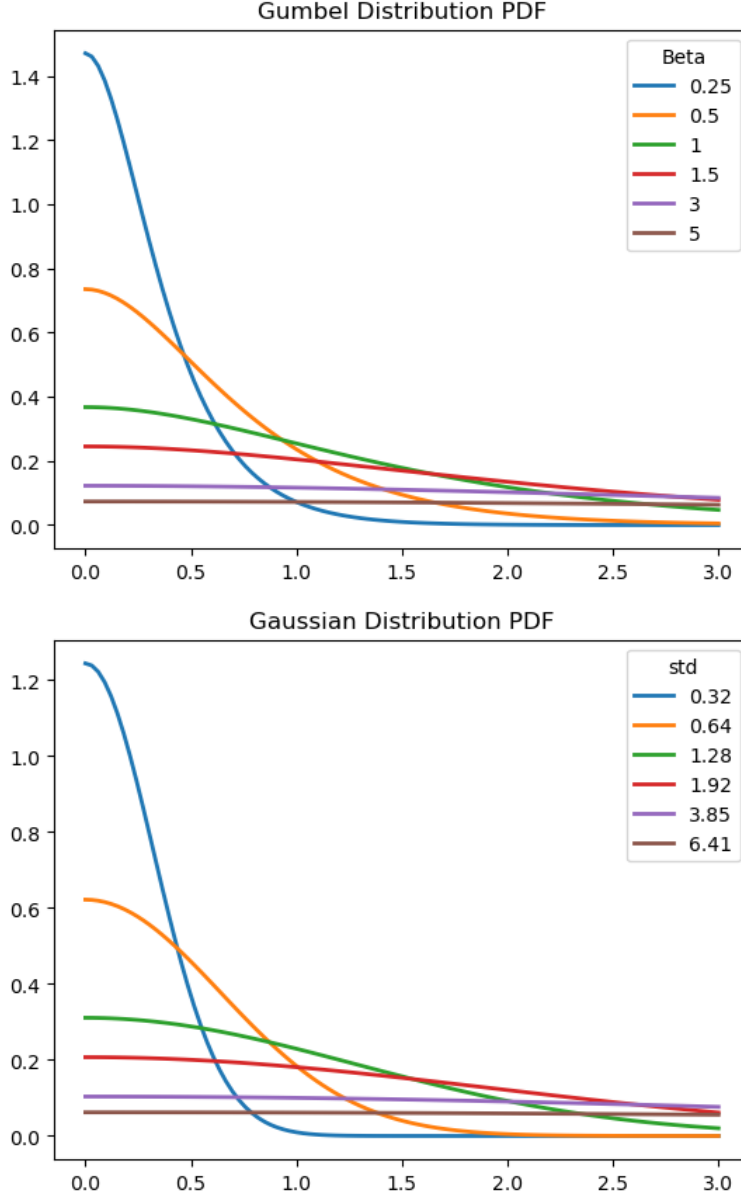


Figure 2: Figure of PDFs of Gumbel and Gaussian noise, respectively.

Taken together, in this section we have shown that even the superstar setting has the undesirable property of relevance of lower-ranked objects, meaning that the agent may need to consider items beyond j in order to be confident that it is picking the best item. However, we have also provided bounds that would allow the agent to limit how many items it has to consider.

A.3 Random Utility Model: Gumbel and Gaussian Noise

The most common noise distribution with the Random Utility Model is Gaussian noise, while our theoretical results (Proposition 2.1) uses Gumbel noise. In this section, we show that empirical results for Gaussian noise closely mimic theoretical results for Gumbel noise.

Figure 2 shows PDFs for Gumbel and Gaussian distributions, Figure 3 shows the $P[\sigma^i]$ values for the superstar model, and Figure 4 shows the ratio $P[\sigma^1]/P[\sigma^i]$, both for the Gumbel and Gaussian noise distribution. In all settings, the β and standard deviation values are set to be equal.

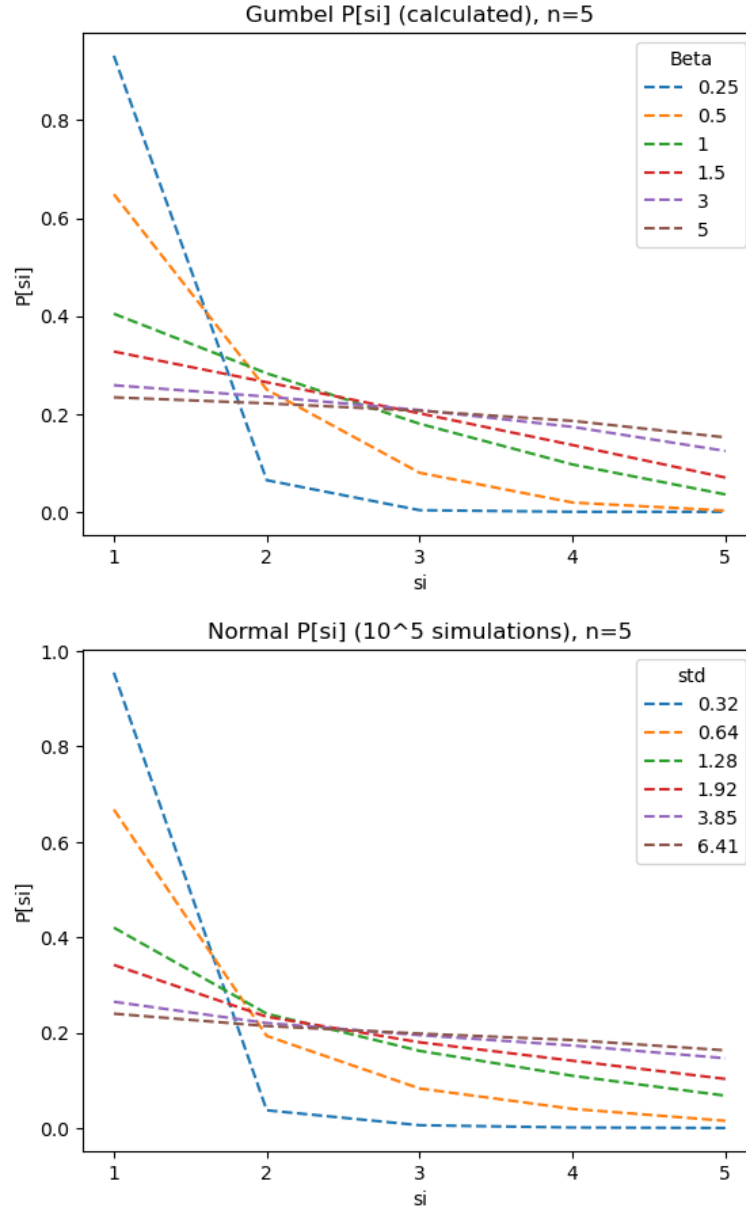


Figure 3: Figure of $P[\sigma^i]$ for RUM with Gumbel and Gaussian noise, respectively. Gumbel is calculated exactly while Gaussian is simulated numerically.

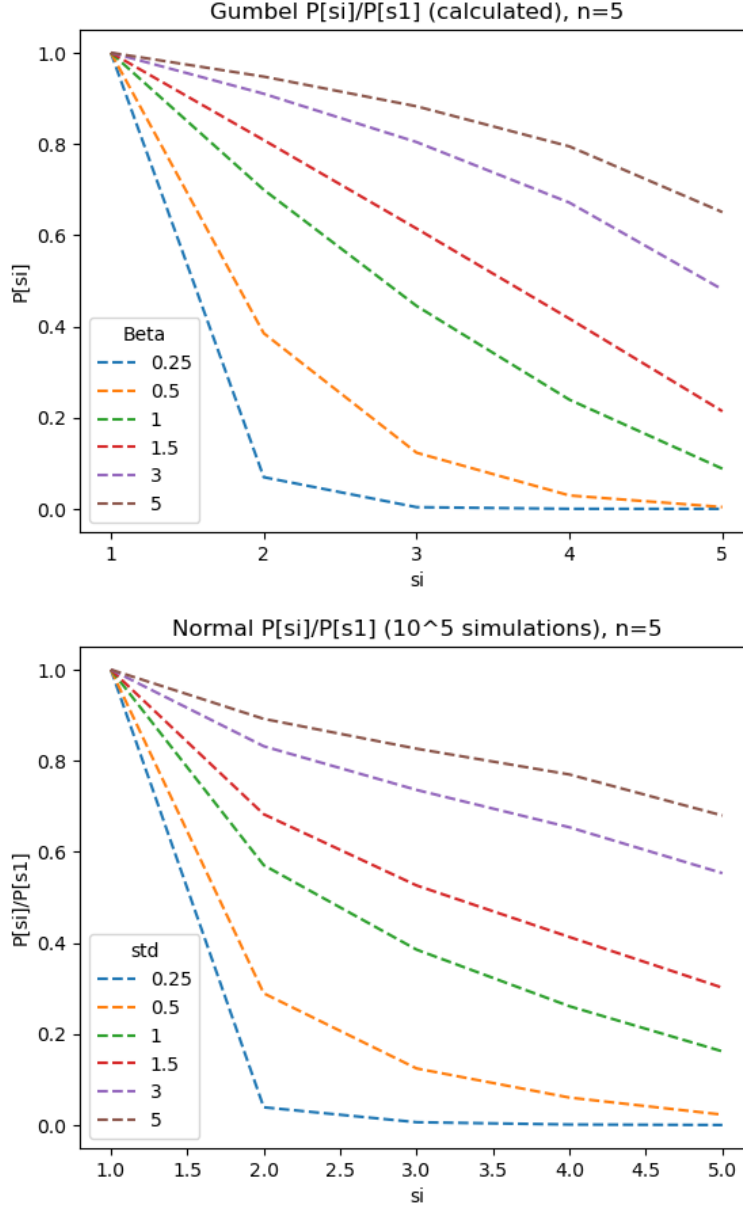


Figure 4: Figure of $\frac{P[\sigma^1]}{P[\sigma^i]}$ for RUM with Gumbel and Gaussian noise, respectively. Gumbel is calculated exactly while Gaussian is simulated numerically. Note that $P[\sigma^1]/P[\sigma^i]$ is strictly increasing as noise decreases for both Gumbel and Gaussian noise.

B Proofs from main body

Theorem 4. Consider sets of items with values v_i , and a Random Utility Model with i.i.d. Gumbel noise $G(\mu, \beta)$ is added to each value. Then, the probability that item i has the highest noised value $\hat{v}_i$ is given by:

$$\frac{\exp(v_i/\beta)}{\sum_{i \in [N]} \exp(v_i/\beta)}$$

Proof sketch. This proof is adapted from <https://homes.cs.washington.edu/~ewein/blog/2022/03/04/gumbel-max/>.

The probability of item i being ranked first can be written as:

$$P[I = i] = \int_{-\infty}^{\infty} f_i(m) \prod_{j \neq i} p(G_j < m) dm$$

Plugging in for the CDFs gives:

$$= \int_{-\infty}^{\infty} f_i(m) \exp\left(-\sum_{j \neq i} \exp(-(\theta_j - m)/\beta)\right) dm$$

The full proof shows that this can be reduced to:

$$\frac{\exp(v_i/\beta)}{\sum_{j \in [N]} \exp(v_j/\beta)}$$

□

Proposition 2.1. Consider the superstar setting (1 item of value $v_1 > v_2 = 0$, and $n - 1$ elements of item 0). If i.i.d. Gumbel noise (with scale parameter β) is added to each value, then the probability of the high value item being in index i is given by:

$$P[\sigma^i] = \frac{\exp(v_1/\beta)}{\exp(v_1/\beta) + n - 1} \prod_{k=1}^{i-1} \frac{1}{1 + \frac{\exp(v_1/\beta) - 1}{n - k}}$$

As a direct consequence, for any $i > 2$, the ratio $\frac{P[\sigma^1]}{P[\sigma^i]}$ is given by: $\prod_{k=1}^{i-1} \left(1 + \frac{\exp(v_1/\beta) - 1}{n - k}\right)$.

Proof. The $i = 1$ case is directly by the Gumbel trick. The $i = n$ case is the second most direct: the high value item v_1 is ranked last whenever:

$$\begin{aligned} P[I = i] &= \int_{-\infty}^{\infty} f_i(m) \prod_{j \neq i} p(\hat{v}_j > m) dm \\ &= \int_{-\infty}^{\infty} f_i(m) p(\hat{v}_j > m)^{n-1} dm \end{aligned}$$

Plugging in for the CDF gives us:

$$\int_{-\infty}^{\infty} f_1(m) (1 - \exp(-\exp(-(v_2 - m)/\beta)))^{n-1} dm$$

The binomial expansion tells us that:

$$(1 - x)^n = \sum_{k=0}^n \binom{n}{k} (-x)^k = 1 - n \cdot x + \binom{n}{2} x^2 - \binom{n}{3} x^3 + \dots$$

Applying the binomial expansion to our case gives us that:

$$(1 - \exp(-\exp(-(v_2 - m)/\beta)))^{n-1} = \sum_{k=0}^{n-1} \binom{n-1}{k} (-\exp(-\exp(-(v_2 - m)/\beta)))^k$$

$$= 1 - (n-1) \cdot \exp(-\exp(-(v_2 - m)/\beta)) + \binom{n-1}{2} \exp(-2 \exp(-(v_2 - m)/\beta)) - \dots$$

Plugging this into our integral gives us:

$$\int_{-\infty}^{\infty} f_1(m) \cdot \left(1 - (n-1) \cdot \exp(-\exp(-(v_2 - m)/\beta)) + \binom{n-1}{2} \exp(-2 \exp(-(v_2 - m)/\beta)) \dots \right) dm$$

We can break up this integral into multiple sums. The first is the most obvious:

$$\int_{-\infty}^{\infty} f_1(m) dm = 1$$

The second term is:

$$-(n-1) \cdot \int_{-\infty}^{\infty} f_1(m) \cdot \exp(-\exp(-(v_2 - m)/\beta)) dm$$

Looking at this, we can see this is exactly the probability that item 1 would be ranked first, we if had only two elements (v_1, v_2) . By the Gumbel trick (Theorem 4), we know that this is $\frac{\exp(v_1/\beta)}{\exp(v_1/\beta) + \exp(v_2/\beta)}$, multiplied by $-(n-1)$. Applying similar reasoning to other terms gives that the entire sum of integrals is given by:

$$\begin{aligned} & 1 - (n-1) \cdot \frac{\exp(v_1/\beta)}{\exp(v_1/\beta) + \exp(v_2/\beta)} + \binom{n-1}{2} \frac{\exp(v_1/\beta)}{\exp(v_1/\beta) + 2 \exp(v_2/\beta)} + \dots \\ &= \sum_{k=0}^{n-1} \frac{\exp(v_1/\beta)}{\exp(v_1/\beta) + k \cdot \exp(v_2/\beta)} \cdot (-1)^k \cdot \binom{n-1}{k} \end{aligned}$$

Finally, we consider the case with $i \in [2, n-1]$. This occurs when exactly $i-1$ low value items are ranked above the high value item and $n-i$ low value items are ranked below it. For any given i , there are $\binom{n-1}{i-1}$ ways that this could happen. Given a fixed set of $i-1$ low value items, the probability of this occurring is given by:

$$\int_{-\infty}^{\infty} f_i(m) P(\hat{v}_i > m)^{i-1} \cdot P(\hat{v}_i < m)^{n-i} dm$$

Substituting in for the CDFs gives us:

$$\int_{-\infty}^{\infty} f_1(m) \cdot (1 - \exp(-\exp(-(v_2 - m)/\beta)))^{i-1} \cdot (\exp(-(n-i) \exp(-(v_2 - m)/\beta))) dm$$

Expanding out using the binomial coefficient gives that:

$$\begin{aligned} (1 - \exp(-\exp(-(v_2 - m)/\beta)))^{i-1} &= 1 - (i-1) \cdot \exp(-\exp(-(v_2 - m)/\beta)) + \binom{i-1}{2} \exp(-2 \exp(-(v_2 - m)/\beta)) - \dots \\ &= \sum_{k=0}^{i-1} \binom{i-1}{k} (-1)^k \exp(-k \cdot \exp(v_2 - m)/\beta) \end{aligned}$$

Distributing out the terms gives:

$$\int_{-\infty}^{\infty} f_1(m) \cdot \sum_{k=0}^{i-1} \binom{i-1}{k} (-1)^k \exp(-(n-i+k) \cdot \exp(v_2 - m)/\beta) dm$$

This looks similar to the $i = n$ case, except that the first term is the probability of the element 1 being ranked first if we have $n-i$ low valued items, or

$$\frac{\exp(v_1/\beta)}{\exp(v_1/\beta) + (n-i) \cdot \exp(v_2/\beta)}$$

Adapting the results gives us that the overall sum is given by:

$$\binom{n-1}{i-1} \sum_{k=0}^{i-1} \frac{\exp(v_1/\beta)}{\exp(v_1/\beta) + (n-i+k) \cdot \exp(v_2/\beta)} \cdot (-1)^k \cdot \binom{i-1}{k}$$

where we have added the $\binom{n-1}{i-1}$ term to give us $P[\sigma^i]$.
We can strategically rewrite by pulling out a term in front:

$$= \exp(v_1/\beta) \binom{n-1}{i-1} \sum_{k=0}^{i-1} \frac{1}{\exp(v_1/\beta) + (n-i+k) \cdot \exp(v_2/\beta)} \cdot (-1)^k \cdot \binom{i-1}{k}$$

Setting $v_2 = 0$ gives:

$$= \exp(v_1/\beta) \binom{n-1}{i-1} \sum_{k=0}^{i-1} \frac{1}{\exp(v_1/\beta) + n - i + k} \cdot (-1)^k \cdot \binom{i-1}{k}$$

Applying the identity from Lemma B.1 gives with $x = \exp(v_1/\beta) + n$ gives:

$$\begin{aligned} & \exp(v_1/\beta) \cdot \binom{n-1}{i-1} \cdot \frac{(\exp(v_1/\beta) + n - i + 1)! \cdot (i-1)!}{(\exp(v_1/\beta) + n - 1)!} \\ &= \exp(v_1/\beta) \cdot \frac{(n-1)!}{(i-1)! \cdot (n-i)!} \cdot \frac{(\exp(v_1/\beta) + n - i + 1)! \cdot (i-1)!}{(\exp(v_1/\beta) + n - 1)!} \\ &= \exp(v_1/\beta) \cdot \frac{(n-1)!}{(n-i)!} \cdot \frac{(\exp(v_1/\beta) + n - i + 1)!}{(\exp(v_1/\beta) + n - 1)!} \\ &= \frac{\exp(v_1/\beta)}{\exp(v_1/\beta) + n - 1} \cdot \frac{\prod_{k=1}^{i-1} n - k}{\prod_{k=1}^{i-1} \exp(v_1/\beta) + n - k - 1} \\ &= \frac{\exp(v_1/\beta)}{\exp(v_1/\beta) + (n-1)} \cdot \prod_{k=1}^{i-1} \frac{1}{1 + \frac{\exp(v_1/\beta) - 1}{n-k}} \end{aligned}$$

as desired. $\square$

Lemma B.1. *The following relationship holds⁴:*

$$\sum_{k=0}^{i-1} \binom{i-1}{k} \cdot (-1)^k \cdot \frac{1}{x - i + k} = \frac{\Gamma(x-i) \cdot \Gamma(i)}{\Gamma(x)} = \frac{(x-i-1)! \cdot (i-1)!}{(x-1)!}$$

Proof. First, we note that

$$\sum_{k=0}^{i-1} \binom{i-1}{k} \cdot (-1)^k \cdot \frac{1}{x - i + k} = \sum_{k=0}^{i-1} \binom{i-1}{k} \cdot (-1)^k \int_0^1 z^{x-i+k-1} dz$$

Switching the order of the integral and summation gives:

$$= \int_0^1 \sum_{k=0}^{i-1} \binom{i-1}{k} \cdot (-1)^k \cdot z^{x-i+k-1} dz = \int_0^1 z^{x-i-1} \sum_{k=0}^{i-1} \binom{i-1}{k} \cdot (-z)^k dz$$

Applying the binomial formula gives:

$$= \int_0^1 z^{x-i-1} \cdot (1-z)^{i-1} = \text{Beta}(x-i, i) = \frac{\Gamma(x-i) \cdot \Gamma(i)}{\Gamma(x)} = \frac{(x-i-1)! \cdot (i-1)!}{(x-1)!}$$

where we have applied a relationship between the Beta and Gamma functions. $\square$

Lemma 2.1. *If distribution $\{P[\sigma^i]'\}$ is more accurate than $\{P[\sigma^i]\}$, then $P[\sigma^1]' > P[\sigma^1]$.*

Proof. We will prove this by contradiction: suppose there exists some distribution such that $\frac{P[\sigma^1]'}{P[\sigma^i]} \geq \frac{P[\sigma^1]'}{P[\sigma^j]'} \forall i > 1$ (with the inequality strict in at least one index), and yet $P[\sigma^1] \leq P[\sigma^1]'$. In order for the inequality to hold, we must have $P[\sigma^i] \leq P[\sigma^i]' \forall i > 1$.

We require the inequality to be strict in at least one index: call the set of such indices $\mathcal{J}$. For these indices, $\frac{P[\sigma^1]'}{P[\sigma^j]} > \frac{P[\sigma^1]'}{P[\sigma^j]'} \forall j \in \mathcal{J}$, we must have $P[\sigma^j] < P[\sigma^j]'$. Thus, for all indices $i \in [N]$, we have $P[\sigma^i] \leq P[\sigma^i]'$, with the inequality strict in at least some indices $j \in \mathcal{J}$. This implies that the total probability of $\{P[\sigma^i]'\}$ is greater than 1, which is a contradiction. $\square$

⁴The authors especially wish to thank Sloan Neiert for discussions on this identity.

Lemma 4.1. *In the superstar setting, given a realized status vector s , it is always optimal to either pick the first free item or the first busy item.*

Proof. In order to prove this, we will prove that the expected value of items given the status vector s is descending, or:

$$\mathbb{E}_{\sigma \sim \mathcal{P}}[v_{\sigma_i} | s] \geq \mathbb{E}_{\sigma \sim \mathcal{P}}[v_{\sigma_j} | s] \quad i < j$$

If this condition holds, then for any two i, j given $a_i = a_j$, we know that the lower ranked of the two has higher utility, and thus no optimal strategy could every pick the higher ranked item.

Next, we will show that the superstar setting always satisfies this property. The expected value of the i th entry is given by:

$$\mathbb{E}[v_{\sigma_i} | s] = \sum_{\sigma \sim \mathcal{P}} P[\sigma | s] \cdot v_{\sigma_i} = \sum_{k=1}^N \frac{P[s | \sigma^k] \cdot P[\sigma^k]}{P[s]} \cdot v_{\sigma_i^k}$$

where we have used that in the superstar setting, each of the permutations can be identified by the index of the high-value item. By identical reasoning, the expected value of the j th entry is given by:

$$\mathbb{E}[v_{\sigma_j} | s] = \sum_{\sigma \sim \mathcal{P}} P[\sigma | s] \cdot v_{\sigma_j} = \sum_{k=1}^N \frac{P[s | \sigma^k] \cdot P[\sigma^k]}{P[s]} \cdot v_{\sigma_j^k}$$

The condition we wish to show is:

$$\sum_{k=1}^N \frac{P[s | \sigma_k] \cdot P[\sigma_k]}{P[s]} \cdot v_{\sigma_i^k} \geq \sum_{k=1}^N \frac{P[s | \sigma_k] \cdot P[\sigma_k]}{P[s]} \cdot v_{\sigma_j^k}$$

Or:

$$\sum_{k=1}^N P[s | \sigma_k] \cdot P[\sigma_k] \cdot (v_{\sigma_i^k} - v_{\sigma_j^k}) \geq 0$$

Because we are in the superstar setting, we know that $v_{\sigma_i^k} = v_{\sigma_j^k}$ (both items are low value) unless $k \in \{i, j\}$. Dropping all but these values of k and using $v_1 \geq v_2$ for the high and low values respectively gives:

$$P[s | \sigma^i] \cdot P[\sigma^i] \cdot (v_1 - v_2) + P[s | \sigma^j] \cdot P[\sigma^j] \cdot (v_2 - v_1) \geq 0$$

which is positive whenever:

$$(v_1 - v_2) \cdot (P[s | \sigma^i] \cdot P[\sigma^i] - P[s | \sigma^j] \cdot P[\sigma^j]) \geq 0$$

By assumption, $v_1 \geq v_2$, and $s_i = s_j$. This latter assumption tells us that $P[s | \sigma^i] = P[s | \sigma^j]$: because the status in index i, j are identical, they are equally likely to occur if the high-value item is in index i or j . Finally, we require that the expected value of the distribution $\mathbb{E}\sigma \sim \mathcal{P}[v_{\sigma_i}]$ is always descending in i , which for the superstar setting implies that the high-value item is always (weakly) more likely to be at lower indices ($P[\sigma^i] \geq P[\sigma^j]$). Taken together, this implies that the above condition reduces to:

$$(v_1 - v_2) \cdot P[s | \sigma^i] \cdot (P[\sigma^i] - P[\sigma^j]) \geq 0$$

which always holds. $\square$

Theorem 1. *[Optimal hiring for firms] If $r \cdot \gamma \leq 1$, then define j^* as the largest j such that $\frac{P[\sigma^1]}{P[\sigma^j]} < \frac{1}{r \cdot \gamma}$. Then, the optimal strategy for the firm is to use the j^* -**virtuous selection** strategy: select the highest-ranked free candidate within the top j^* . If all of the top j^* candidates are busy, hire the top-ranked candidate.*

Conversely, if $r \cdot \gamma > 1$, define k^ as the largest k such that $\frac{P[\sigma^1]}{P[\sigma^k]} < r \cdot \gamma$. Then, the optimal strategy for the firm is to use the k^* -**adverse selection** strategy: select the highest-ranked busy candidate within the top k^* . If all of the top k^* candidates are free, hire the top-ranked candidate.*

If $j^ = 1$ or $k^* = 1$, then optimal strategy for the firm is to use the **follow the ranking** strategy: always select the top-ranked candidate, regardless of whether they are free or busy.*

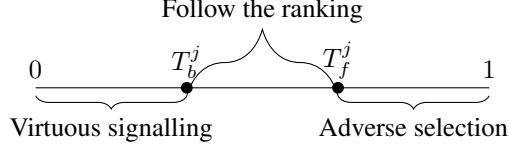


Figure 5: Axis shows γ going from 0 to 1, in the case that $T_b^j \leq T_f^j$. When $\gamma \leq T_b^j \leq T_f^j$, the optimal strategy is always to pick the first element that is free: virtuous signalling, modeling scenarios where picking a busy item has prohibitively high penalty. When $\gamma \geq T_f^j \geq T_b^j$, the optimal strategy is always to pick the first element that is *busy*: adverse selection, modeling scenarios where picking a busy item has small penalty, but higher utility due to the signal induced by being busy. Given $T_b^j \leq \gamma \leq T_f^j$, the optimal strategy is always pick the highest-ranked item: this is where the signal induced by the ranking is stronger than the signal from the free/busy status.

Proof. First, we use sub-lemma 4.1, which tells us that, for any given status vector, the optimal candidate to pick is always either the first free (highest-ranked free) or first busy candidate. Thus, the rest of our analysis simply needs to determine which of these two candidates is best in any given scenario.

Calculating thresholds:

First, we will show that the optimal strategy can be calculated as a function of index-specific thresholds T_b^j, T_f^j depending on whether the highest-ranked candidate is free or busy.

First, we will consider the case where first item is busy, j th item is the first that is free. We pick the 1st (busy) index over the j th (free) index when:

$$\mathbb{E}_{\sigma \sim \mathcal{P}}[v_{\sigma_1} \mid s] \cdot \gamma \geq \mathbb{E}_{\sigma \sim \mathcal{P}}[v_{\sigma_j} \mid s]$$

Note that in the superstar setting with $v_2 = 0$, we can write the expected value as:

$$\mathbb{E}[v_{\sigma_i} \mid s] = \sum_{\sigma \sim \mathcal{P}} P[\sigma \mid s] \cdot v_{\sigma_i} = \sum_{k=1}^N \frac{P[s \mid \sigma^k] \cdot P[\sigma^k]}{P[s]} \cdot v_{\sigma_i^k} = \frac{P[s \mid \sigma^i] \cdot P[\sigma^i]}{P[s]} \cdot v_1$$

where in the last step, we have used the property that $v_{\sigma_i^k} = 0$ unless $k = i$, in which case it equals v_1 . Our condition thus reduces to:

$$P[s \mid \sigma^1] \cdot P[\sigma^1] \cdot \gamma \geq P[s \mid \sigma^j] \cdot P[\sigma^j]$$

Next, we will reason about the relative probabilities of seeing the status vector s , given that the high value item is in the 1st or j th index. We can write:

$$P[s \mid \sigma^1] = (1 - p_1) \cdot (1 - p_2)^{j-2} \cdot p_2 \cdot \prod_{k=j+1}^N p_2^{s_k} \cdot (1 - p_2)^{1-s_k}$$

$$P[s \mid \sigma^j] = (1 - p_2)^{j-1} \cdot p_1 \cdot \prod_{k=j+1}^N p_2^{s_k} \cdot (1 - p_2)^{1-s_k}$$

Note that $P[s \mid \sigma^j] = P[s \mid \sigma^1] \cdot (1 - p_2) \cdot p_1 \cdot \frac{1}{(1-p_1) \cdot p_2} = P[s \mid \sigma^1] \cdot \frac{p_1/(1-p_1)}{p_2/(1-p_2)} = \frac{1}{r}$. Because $p_1 < p_2$, we have $r = \frac{r_1}{r_2} < 1$. We can rewrite the relevant condition as:

$$\gamma > \frac{r_1}{r_2} \cdot \frac{P[\sigma^j]}{P[\sigma^1]} = T_b^j$$

Next, we consider the case where the first item is free, the j th item is the first that is busy. Using similar reasoning to the first case, we pick the first (free) item whenever:

$$\mathbb{E}_{\sigma \sim \mathcal{P}}[v_{\sigma_1} \mid a] \geq \mathbb{E}_{\sigma \sim \mathcal{P}}[v_{\sigma_j} \mid a] \cdot \gamma$$

which simplifies to:

$$P[s \mid \sigma^1] \cdot P[\sigma^1] \geq P[s \mid \sigma^j] \cdot P[\sigma^j] \cdot \gamma$$

In this case, the status vector s is free for the first $j - 1$ values, meaning that:

$$P[a \mid \sigma^1] = p_1 \cdot p_2^{j-2} \cdot (1 - p_2) \cdot \prod_{k=j+1}^N p_2^{s_k} \cdot (1 - p_2)^{1-s_k}$$

$$P[a \mid \sigma^j] = p_2^{j-1} \cdot (1 - p_1) \cdot \prod_{k=j+1}^N p_2^{s_k} \cdot (1 - p_2)^{1-s_k}$$

This tells us that $P[s \mid \sigma^j] = P[s \mid \sigma^1] \cdot \frac{p_2 \cdot (1-p_1)}{p_1 \cdot (1-p_2)} = r$. We can rewrite the relevant condition as:

$$T_f^j = \frac{P[\sigma^1]}{P[\sigma^j]} \cdot \frac{1}{r} \geq \gamma$$

Calculating the optimal strategy:

Finally, we use these thresholds to calculate regimes for optimal strategies for firms:

We assume γ, r are fixed.

Adverse selection:

$$\gamma > T_f^j = \frac{1}{r} \cdot \frac{P[\sigma^1]}{P[\sigma^j]} \geq T_b^j = \frac{1}{r} \cdot \frac{P[\sigma^j]}{P[\sigma^1]}$$

$$\gamma \cdot r > \frac{P[\sigma^1]}{P[\sigma^j]} \geq \frac{P[\sigma^j]}{P[\sigma^1]}$$

Can rewrite this as:

$$\gamma \cdot r \cdot P[\sigma^j] > P[\sigma^1] \text{ and } P[\sigma^1] > \frac{P[\sigma^j]}{r \cdot \gamma}$$

$$\gamma \cdot r \cdot P[\sigma^j] > P[\sigma^1] > \frac{P[\sigma^j]}{r \cdot \gamma}$$

Follow the ranking:

$$T_f^j = \frac{1}{r} \cdot \frac{P[\sigma^1]}{P[\sigma^j]} > \gamma \geq T_b^j = \frac{1}{r} \cdot \frac{P[\sigma^j]}{P[\sigma^1]}$$

Or, the first term gives you:

$$P[\sigma^1] > r \cdot \gamma \cdot P[\sigma^j]$$

The second term gives:

$$P[\sigma^1] > \frac{P[\sigma^j]}{r \cdot \gamma}$$

Taken together,

$$P[\sigma^1] > \max \left(r \cdot \gamma \cdot P[\sigma^j], \frac{P[\sigma^j]}{r \cdot \gamma} \right)$$

Virtuous selection:

$$T_f^j = \frac{1}{r} \cdot \frac{P[\sigma^1]}{P[\sigma^j]} \geq T_b^j = \frac{1}{r} \cdot \frac{P[\sigma^j]}{P[\sigma^1]} > \gamma$$

$$\frac{P[\sigma^1]}{P[\sigma^j]} \geq T_b^j = \frac{P[\sigma^j]}{P[\sigma^1]} > r \cdot \gamma$$

Or, the first term becomes:

$$P[\sigma^1] > r \cdot \gamma \cdot P[\sigma^j]$$

The second term becomes:

$$\frac{P[\sigma^j]}{r \cdot \gamma} > P[\sigma^1]$$

Or,

$$r \cdot \gamma \cdot P[\sigma^j] < P[\sigma^1] < \frac{P[\sigma^j]}{r \cdot \gamma}$$

Analysis:

We know that $P[\sigma^1] \geq P[\sigma^j] \forall j \geq 2$. Note that $r \geq 1$ and $\gamma < 1$, so $r \cdot \gamma$ could be either > 1 or < 1 . Which of this it is will affect which of these regimes the firm may go through.

First, let's consider the case where $r \cdot \gamma \leq 1$. Then, we know $r \cdot \gamma \cdot P[\sigma^j] \leq P[\sigma^j] \leq P[\sigma^1]$ holds automatically. We also know that $P[\sigma^1]$ could be less than $\frac{P[\sigma^j]}{r \cdot \gamma}$. However, note that we never have $\gamma \cdot r \cdot P[\sigma^j] > P[\sigma^1]$, so adverse selection is never optimal for this regime. The optimal strategy can be described as:

$$\begin{cases} P[\sigma^1] < \frac{P[\sigma^j]}{r \cdot \gamma} & \text{Virtuous selection} \\ P[\sigma^1] \geq \frac{P[\sigma^j]}{r \cdot \gamma} & \text{Follow the ranking} \end{cases}$$

So, as $P[\sigma^1]$ increases or $P[\sigma^j]$ decreases, we go from the virtuous selection regime (selecting a free candidate, whenever one is available) to the “follow the ranking” scheme, where the optimal strategy is to pick whichever is ranked first.

Next, let's consider the case where $r \cdot \gamma > 1$. Then, we know that we never have $P[\sigma^1] < \frac{P[\sigma^j]}{r \cdot \gamma}$ (because we know $\frac{P[\sigma^j]}{r \cdot \gamma} < P[\sigma^j] \leq P[\sigma^1]$), so virtuous selection is never possible. Instead, for low $P[\sigma^1] < \gamma \cdot r \cdot P[\sigma^j]$, we're in adverse selection, and as $P[\sigma^1]$ increases or $P[\sigma^j]$ decreases, we move to “follow the ranking”. The optimal strategy looks like:

$$\begin{cases} P[\sigma^1] < P[\sigma^j] \cdot r \cdot \gamma & \text{Adverse selection} \\ P[\sigma^1] \geq P[\sigma^j] \cdot r \cdot \gamma & \text{Follow the ranking} \end{cases}$$

□

Lemma 4.2. *Holding γ constant, increasing the accuracy of the ranking tool always holds constant or decreases the furthest index a firm needs to consider (decreases j^* or k^* as defined in Theorem 1).*

Proof. This can be seen almost immediately: from the strategy in Theorem 1, a firm's strategy depends on the ratio $\frac{P[\sigma^1]}{P[\sigma^i]}$ in comparison to a constant $r \cdot \gamma$ (or its inverse for adverse selection). As the accuracy increases, by Definition 1 the ratio $\frac{P[\sigma^1]}{P[\sigma^i]}$ increases. This means that fewer indices i will satisfy the $\frac{P[\sigma^1]}{P[\sigma^i]} < r \cdot \gamma$ criteria, so j^* or k^* will stay constant or decrease. □

Lemma 4.3. *No matter how accurate a ranking is, there exists a firm whose optimal strategy is virtuous selection with $j^* = n$. More precisely, for all $\epsilon > 0$ such that $P[\sigma^1]/P[\sigma^n] = \epsilon$, there exists a busy penalty $\gamma < \frac{1}{r}$ such that $k^* = n$.*

For all firms with $\gamma \geq \frac{1}{r}$, if $P[\sigma^1]/P[\sigma^2] > r$, the optimal strategy is always “follow the ranking”.

Proof. These results can be seen almost immediately from the derivation of the optimal strategies in Theorem 1.

First, assume $\frac{P[\sigma^1]}{P[\sigma^n]} = \epsilon > 0$. Then, set any $\gamma < \frac{\epsilon}{r}$. Note that $r \cdot \gamma = \epsilon < 1$, so the firm's optimal strategy is virtuous selection or “follow the ranking”. Because:

$$\frac{1}{r \cdot \gamma} > \epsilon = \frac{P[\sigma^1]}{P[\sigma^n]}$$

We know that the firm's optimal strategy is virtuous selection with $j^* = n$, as desired.

Next, we wish to show that if $\gamma > \frac{1}{r}$ given $\frac{P[\sigma^1]}{P[\sigma^2]}$, follow the ranking is always optimal. By assumption:

$$\frac{P[\sigma^1]}{P[\sigma^i]} \geq \frac{P[\sigma^1]}{P[\sigma^2]} > r \geq r \cdot \gamma$$

This condition means that the optimal strategy for firms is always “follow the ranking”. □

Lemma 4.4. *Increasing the accuracy of the ranking tool can increase, decrease, or keep constant the probability that a firm selects the top ranked candidate.*

Proof. Throughout, we'll focus on the $n = 2$ setting to give our existence examples. Given $n = 2$, there are exactly four possible status vectors: $[0, 0]$, $[0, 1]$, $[1, 0]$, $[1, 1]$. This proof will involve reasoning about the actions firms take when they observe different status vectors, and how likely these status vectors are to appear.

Less attractive firms: $r \cdot \gamma < 1$:

In this setting, if $\frac{P[\sigma^1]}{1-P[\sigma^1]} < (r \cdot \gamma)^{-1}$, firms are using virtuous selection with $j = 2$: they hire the top-ranked free candidate if one exists. Thus, the only time a firm in this setting hires the second-ranked candidate is if the realized status vector is $[0, 1]$. What happens as the accuracy of the ranking increases?

Case 1: If the accuracy increases such that $\frac{P[\sigma^1]'}{1-P[\sigma^1]'} > (r \cdot \gamma)^{-1}$, then the firm now always uses “follow the ranking” and selects the top-ranked candidate with probability 1.

Case 2: If the accuracy increases, but we still have $\frac{P[\sigma^1]'}{1-P[\sigma^1]'} < (r \cdot \gamma)^{-1}$, then the firm continues to hire the second-ranked candidate only if the realized status vector is $[0, 1]$. However, the probability of observing this status vector has *increased*, because a more accurate ranking tool more often places the high-value candidate first, and the high-value candidate is more likely to be unavailable. Thus, increasing the accuracy of the ranking tool can *decrease* the probability that the top-ranked candidate is selected. More formally,

$$\begin{aligned} P[s = [0, 1]] &= P[s = [0, 1] \mid \sigma^1] \cdot P[\sigma^1] + P[s = [0, 1] \mid \sigma^2] \cdot P[\sigma^2] \\ &= (1 - p_1) \cdot p_2 \cdot P[\sigma^1] + (1 - p_2) \cdot p_1 \cdot (1 - P[\sigma^1]) = P[\sigma^1] \cdot (p_2 - p_1) + (1 - p_2) \cdot p_1 \end{aligned}$$

Note that the coefficient on $P[\sigma^1]$ is positive, so the probability of the status vector $[0, 1]$ increases as $P[\sigma^1]$ increases (which decreases the probability that the top-ranked candidate is selected).

More attractive firms: $r \cdot \gamma > 1$:

The analysis in this setting is very similar to that of the above, except the result will be that increasing the accuracy always *increases* the chance that the top-ranked candidate will be selected. In this setting, if $\frac{P[\sigma^1]}{1-P[\sigma^1]} < r \cdot \gamma$, firms are using adverse selection with $j = 2$: they hire the top-ranked busy candidate if one exists. Thus, the only time a firm hires the second-ranked candidate is if the realized status vector is $[1, 0]$. What happens as the accuracy of the ranking increases?

Case 1: If the accuracy increases such that now $\frac{P[\sigma^1]'}{1-P[\sigma^1]'} > r \cdot \gamma$, then the firm always uses “follow the ranking” and selects the top-ranked candidate with probability 1.

Case 2: If the accuracy increases, but we still have $\frac{P[\sigma^1]'}{1-P[\sigma^1]'} < r \cdot \gamma$, then the firm continues to hire the second-ranked candidate only if the realized status vector is $[1, 0]$. This probability *decreases* because a more accurate ranking tool places the high-value candidate first more often, and the high value candidate is less likely to be free. More formally,

$$\begin{aligned} P[s = [1, 0]] &= P[s = [1, 0] \mid \sigma^1] \cdot P[\sigma^1] + P[s = [1, 0] \mid \sigma^2] \cdot P[\sigma^2] \\ &= p_1 \cdot (1 - p_2) \cdot P[\sigma^1] + p_2 \cdot (1 - p_1) \cdot (1 - P[\sigma^1]) = P[\sigma^1] \cdot (p_1 - p_2) + p_2 \cdot (1 - p_1) \end{aligned}$$

Note that the coefficient on $P[\sigma^1]$ is negative, so the probability of the status vector $[1, 0]$ *decreases* as $P[\sigma^1]$ increases (which increases the probability that the top-ranked candidate is selected). $\square$

Lemma 5.1. *There exist settings where increasing the accuracy of the ranking tool can increase or decrease duplicative effort.*

Proof. First, let's calculate the duplicative effort involved in each status vector. Recall that duplicative effort is the total value of selected busy candidates:

$$\gamma \cdot \mathbb{E}[v_{\rho_{i(s)}}] \cdot \mathbb{1}[s_{i(s)} = 0]$$

From prior analysis in Theorem 1, we know that in the superstar setting:

$$\mathbb{E}[v_{\rho_{i(s)}}] = \sum_{s \in S} \mathbb{E}[v_{\rho_{i(s)}} | s] \cdot P[s] = v_1 \cdot \gamma^{1-s_i} \cdot \sum_{s \in S} P[s | \sigma_{i(s)}] \cdot P[\sigma_{i(s)}]$$

Thus, duplicative effort can be written as:

$$\gamma \cdot v_1 \cdot \sum_{s \in S} P[s | \sigma_{i(s)}] \cdot P[\sigma_{i(s)}]$$

where $i(s)$ is the selected index, given the firm's strategy. Based on the results in Theorem 1, we know that the strategy of each firm depends on $r \cdot \gamma$.

Less attractive firms: $r \cdot \gamma < 1$:

First, we will analyze the case where $r \cdot \gamma < 1$, where a firm's strategy is virtuous selection for within the top k indices, where k is the largest index such that $P[\sigma^1] < \frac{P[\sigma^k]}{r \cdot \gamma}$.

Note that if $k = 1$, then the firm always selects the top-ranked candidate, and duplicative effort is exactly the governed by the probability that the top-ranked candidate is a busy high-value candidate:

$$v_1 \cdot \gamma \cdot (1 - p_1) \cdot P[\sigma^1]$$

In this case, increasing accuracy automatically increases duplicative effort, because by Lemma 2.1 it would increase $P[\sigma^1]$.

Next, we consider the case where $k > 1$. Then, the firm picks a busy candidate only if no free ones are available within the top k (and when they pick, they pick the top-ranked candidate). This gives duplicative effort equal to:

$$\gamma \cdot P[s_k | \sigma^1] \cdot P[\sigma^1] = v_1 \cdot \gamma \cdot (1 - p_1) \cdot (1 - p_2)^{k-1} \cdot P[\sigma^1]$$

How does this quantity change as the accuracy increases? Two effects can happen: first k can shrink (which would cause this quantity to increase), and secondly $P[\sigma^1]$ can increase (which would again cause this quantity to increase). Thus, for firms with $r \cdot \gamma < 1$, increasing the accuracy only increases duplicative effort.

More attractive firms: $r \cdot \gamma < 1$:

In this case, a firm's strategy is adverse selection within the first j indices, where j is the largest quantity such that $P[\sigma^1] < P[\sigma^j] \cdot r \cdot \gamma$. Note that if $j = 1$, then again the duplicative effort is given by:

$$v_1 \cdot \gamma \cdot (1 - p_1) \cdot P[\sigma^1]$$

Again, note that if the accuracy increases, then the duplicative effort also increases.

Next, we consider the case where $j > 1$. Here, the firm picks a busy candidate so long as one is available (that is, if one is available within the top j candidates). This gives a total duplicative effort of

$$\gamma \cdot v_1 \cdot \sum_{i=1}^j P[\sigma^i] \cdot (1 - p_1) \cdot p_2^{i-1} \\ P[\sigma^1] \cdot (1 - p_1) + P[\sigma^2] \cdot (1 - p_1) \cdot p_2 + \dots P[\sigma^j] \cdot (1 - p_1) \cdot p_2^{j-1}$$

In this case, increased accuracy has two opposing effects. First, it decreases j , which decreases duplicative effort. Secondly, it changes the values of $\{P[\sigma^i]\}$ terms: in general this *increases* the probability of the high value item being ranked highly (increasing $P[\sigma^i]$ for small i), but this change may not be monotonic.

We illustrate the fact that duplicative effort may increase or decrease by considering the $n = 2$ case, where duplicative effort with $j = 2$ reduces to:

$$P[\sigma^1] \cdot (1 - p_1) + (1 - P[\sigma^1]) \cdot (1 - p_1) \cdot p_2 = (1 - p_1) \cdot (P[\sigma^1] \cdot (1 - p_2) + p_2)$$

If the accuracy of the ranker increases ($P[\sigma^1]$ increases) but $j' = 2$ still, then duplicative effort increases: we can observe this by noting that the derivative of the term above is $(1 - p_1) \cdot (1 - p_2) > 0$. Intuitively, this means that the firm is using adverse selection still, but the top-ranked busy candidate is more likely to be a high value candidate, which means that selecting them has a greater cost than selecting a busy low value candidate.

However, if the accuracy of the ranker increases such that $j' = 1$ (the firm switches to “follow the ranking”), duplicative effort may decrease. This occurs exactly when:

$$(1 - p_1) \cdot (P[\sigma^1] \cdot (1 - p_2) + p_2) > P[\sigma^1]' \cdot (1 - p_1)$$

$$P[\sigma^1] \cdot (1 - p_2) + p_2 > P[\sigma^1]'$$

Thus, for cases where $P[\sigma^1]'$ is larger than $P[\sigma^1]$ but still lower than the quantity above, an increased accuracy can *decrease* duplicative effort.

For concreteness, set $p_1 = 0.5$ and $p_2 = 0.75$, which gives $r = \frac{p_2/(1-p_2)}{p_1/(1-p_1)} = 3$, and set $\gamma = 2/3$, which gives $r \cdot \gamma = 2$. Further set $P[\sigma^1] = 0.6$ and $P[\sigma^1]' = 0.8$.

Note that we have $\frac{P[\sigma^1]}{1-P[\sigma^1]} = 1.5 < 2 = \gamma \cdot r$, which means that the optimal strategy using $P[\sigma]$ is adverse selection with $k^* = 2$. However, we have $\frac{P[\sigma^1]'}{1-P[\sigma^1]'} = 4 > 2 = \gamma \cdot r$, so the optimal strategy using $P[\sigma^1]'$ is “follow the ranking”. Given the parameters, we have:

$$P[\sigma^1] \cdot (1 - p_2) + p_2 = 0.9 > P[\sigma^1]' = 0.8$$

as desired. $\square$

Lemma B.2. *A firm using virtuous selection or follow the ranking never experiences positive benefits to congestion.*

Proof. We can prove this result by reasoning about a) which index a firm selects when using virtuous selection or follow the ranking, and b) how the expected utility of the candidate in this index differs from the expected utility of the candidate ranked first. Note that without congestion, the optimal strategy is always to pick the top-ranked candidate, who will be free. Note that follow the ranking is identical to virtuous selection with $j^* = 1$ and can thus be analyzed using the same techniques.

First, consider the case where the status vector has the first candidate being free: then a firm using virtuous selection will pick this candidate and will obtain expected utility exactly identical to what it would obtain in the absence of congestion.

Second, consider the case where the status vector has all top j^* candidates being busy. Then, the firm using virtuous selection with parameter j^* will pick the top-ranked candidate (the same as it would do without congestion), but will pay a γ penalty, and thus will obtain strictly lower utility than without congestion.

Thirdly, consider the case where the status vector has the first candidate being busy and the first free candidate in index $i \in [2, j^*]$. Then, the firm using virtuous selection will select the free candidate in index i . Note that this has lower expected value (excluding the busy penalty) than the candidate ranked first: Because $P[\sigma^1] \geq P[\sigma^i]$, our prior is that the high value candidate is more likely to be ranked first. Because $p_1 < p_2$, having observed that the top-ranked candidate is busy and the i th ranked candidate is free only *increases* our posterior estimate of the top-ranked candidate’s value and decreases our posterior estimate of the i th candidate’s value. Without congestion, the firm would pick the 1st candidate (without a busy penalty γ) and thus would obtain strictly higher utility than the firm using virtuous selection.

Together, these three scenarios encompass all possible status vectors and show that a firm using virtuous selection would obtain equal or strictly lower utility among all of them, showing that there can never be positive benefits to congestion. $\square$

Lemma 5.2. *A firm using virtuous selection or “follow the ranking” never experiences positive benefits to congestion.*

A firm using adverse selection with $k^ \geq 2$ experiences positive benefits to congestion if and only if:*

$$\frac{(1-p_1) \cdot \gamma}{p_1 \cdot (1-p_2^{k^*-1}) + (1-p_1) \cdot (1-\gamma)} \sum_{k=2}^{k^*} p_2^{k-1} P[\sigma^k] > P[\sigma^1] \text{ with } k^* \text{ defined as in Theorem 1.}$$

Proof. First, Lemma B.2 shows that a firm using virtuous selection or follow the ranking never experiences positive benefits to congestion.

Next, we will consider the case where a firm is using adverse selection with $k^* \geq 2$. We will begin by building intuition by reasoning about the possible status vectors and then derive the formal conditions.

Case 1: If all candidates in the top k^* are free, then the firm using adverse selection will pick the top-ranked candidate. This is identical to what it would have picked without congestion (and has no busy penalty), and thus it obtains exactly identical utility.

Case 2: If the top-ranked candidate is busy, then the firm using adverse selection picks it. In the absence of congestion, the firm would still have picked the top-ranked candidate, but would have done so without have to pay a busy penalty. Thus, in this case, congestion hurts the firm.

Case 3: If the top ranked candidate is free and the first busy candidate is in index $i \in [2, k^*]$, then the firm using adverse selection picks the candidate in index i , whereas it would have picked the top-ranked free candidate in the absence of congestion. Here: congestion could lead to higher or lower utility, depending on the strength of the signal from being busy as compared with the accuracy of the ranking tool.

The remainder of this proof will revolve around showing when the potential positive benefits of congestion from Case 3 outweigh the harms of congestion in Case 2 and Case 3.

Denote $i(s)$ as the index that is selected given a status vector s , given the optimal strategy. Note that this relies also on other parameters (γ, r) which we will omit for conciseness. We're interested in when the expected utility given the optimal strategy (and congestion) is higher than the expected utility without any congestion. If there's no congestion, then the best strategy is always to pick the first element, who is always free and has probability $P[\sigma^1]$ of being the high-valued candidate, which gives expected utility:

$$v_1 \cdot P[\sigma^1]$$

What is the utility of the optimal strategy?

For every status vector s , we need to calculate the optimal strategy, and the probability of that vector occurring.

We can write this expected utility as:

$$\sum_{s \in S} \mathbb{E}[v_{i(s)} \mid s] \cdot P[s]$$

Because $v_2 = 0$, we know that $v_{i(s)} = v_1$ if the high value item is in index $i(s)$ and 0 otherwise, which means we can simplify this down to:

$$\sum_{s \in S} v_1 \cdot \gamma^{1-s_i} \cdot P[\sigma_{i(s)} \mid s] \cdot P[s] = v_1 \cdot \gamma^{1-s_i} \cdot \sum_{s \in S} \frac{P[s \mid \sigma_{i(s)}] \cdot P[\sigma_{i(s)}]}{P[s]} \cdot P[s] = v_1 \cdot \gamma^{1-s_i} \cdot \sum_{s \in S} P[s \mid \sigma_{i(s)}] \cdot P[\sigma_{i(s)}]$$

where we have applied Bayes rule. Now, we analyze $i(s)$ for various values of s and how they affect the utility we get.

Note that the utility for no congestion given $v_2 = 0$ is given by:

$$v_1 \cdot \sum_{s \in S} P[s \mid \sigma^1] \cdot P[\sigma^1]$$

Next, we analyze the total difference in utility between the congested setting and uncongested setting for Case 2 and Case 3.

Case 2:

In this setting, the firm using adverse selection picks the top-ranked (busy) candidate, the same as it would do in the absence of congestion. However, it must pay a γ penalty for picking a busy candidate. The difference in utility is given by summing over all possible status vectors such that this is true, indexing by the location of the first candidate who is free:

$$\begin{aligned} & \sum_{s \mid s_i < j=0, s_j=1 \cup s_i=0 \forall i} P[s \mid \sigma^1] \cdot P[\sigma^1] \cdot (\gamma - 1) \\ &= P[\sigma^1] \cdot (\gamma - 1) \cdot \left(\sum_{j=2}^n (1 - p_1) \cdot (1 - p_2)^{j-2} \cdot p_2 + (1 - p_1) \cdot (1 - p_2)^{n-1} \right) \end{aligned}$$

where the last case is where the status vector has every agent being busy. We can simplify by collecting like terms:

$$(1 - p_1) \cdot P[\sigma^1] \cdot (\gamma - 1) \cdot \left(\sum_{j=2}^n (1 - p_2)^{j-2} \cdot p_2 + (1 - p_2)^{n-1} \right)$$

Note that we can further simplify this by using the sum of a geometric series:

$$\sum_{j=2}^n (1 - p_2)^{j-2} = \sum_{j=0}^{n-2} (1 - p_2)^j = \frac{1 - (1 - p_2)^{n-1}}{1 - (1 - p_2)} = \frac{1 - (1 - p_2)^{n-1}}{p_2}$$

which means that:

$$\sum_{j=2}^n (1 - p_2)^{j-2} \cdot p_2 + (1 - p_2)^{n-1} = \frac{1 - (1 - p_2)^{n-1}}{p_2} \cdot p_2 + (1 - p_2)^{n-1} = 1$$

So the overall term becomes:

$$(1 - p_1) \cdot P[\sigma^1] \cdot (\gamma - 1)$$

Case 3:

In this case, the firm using adverse selection picks the first busy candidate in index $j \geq 2$, while a candidate without congestion would pick the top-ranked candidate (who is free). Here, the difference in expected value is given by:

$$\sum_{s | s_i < j=1, s_j=0, j \leq k^*} P[s | \sigma^j] \cdot P[\sigma^j] \cdot \gamma - P[s | \sigma^1] \cdot P[\sigma^1]$$

We can again rewrite this by summing over the index of the first candidate that's busy:

$$= \sum_{j=2}^{k^*} p_2^{j-1} \cdot (1 - p_1) \cdot P[\sigma^j] \cdot \gamma - p_1 \cdot p_2^{j-2} \cdot (1 - p_2) \cdot P[\sigma^1]$$

We can simplify down some of the terms by noting:

$$\sum_{j=2}^{k^*} p_2^{j-2} = \sum_{j=0}^{k^*-2} p_2^j = \frac{1 - p_2^{k^*-1}}{1 - p_2}$$

So the term above becomes:

$$(1 - p_1) \cdot \gamma \sum_{j=2}^{k^*} p_2^{j-1} \cdot P[\sigma^j] - p_1 \cdot (1 - p_2^{k^*-1}) \cdot P[\sigma^1]$$

When is the total increase in utility from Case 3 larger than the total decrease in utility from Case 2?

Whenever:

$$(1 - p_1) \cdot \gamma \sum_{j=2}^{k^*} p_2^{j-1} \cdot P[\sigma^j] > p_1 \cdot (1 - p_2^{k^*-1}) \cdot P[\sigma^1] + (1 - p_1) \cdot P[\sigma^1] \cdot (1 - \gamma)$$

$$\frac{(1 - p_1) \cdot \gamma}{p_1 \cdot (1 - p_2^{k^*-1}) + (1 - p_1) \cdot (1 - \gamma)} \sum_{j=2}^{k^*} p_2^{j-1} P[\sigma^j] > P[\sigma^1]$$

as desired. □

Lemma 5.3. *Increasing the accuracy of the ranking tool always reduces the firms' benefit to congestion.*

Proof. This can be seen almost immediately from Lemma 5.2, which says that there are benefits to congestion if and only if:

$$\frac{(1 - p_1) \cdot \gamma}{p_1 \cdot (1 - p_2^{d-1}) + (1 - p_1) \cdot (1 - \gamma)} \sum_{j=2}^d p_2^{j-1} P[\sigma^j] > P[\sigma^1]$$

which can be rewritten as:

$$\frac{(1 - p_1) \cdot \gamma}{p_1 \cdot (1 - p_2^{d-1}) + (1 - p_1) \cdot (1 - \gamma)} \sum_{j=2}^d p_2^{j-1} \frac{P[\sigma^j]}{P[\sigma^1]} > 1$$

Increasing accuracy decreases $\frac{P[\sigma^j]}{P[\sigma^1]}$, which decreases the term on the lefthand side, reducing the benefits of congestion as desired. $\square$

Lemma 5.4. *For all parameters, the free-busy gap is strictly positive if the firm is using virtuous selection, strictly negative if the firm is using adverse selection, and zero if the firm is using follow the ranking.*

Proof. In proving this theorem, we will rely on sub-lemma B.3, which gives closed-form solutions for the free-busy gap (for conciseness, we will not restate these solutions here).

Note that the conditions in Lemma B.3 are equal to 0 if and only if $k^* = 1$ or $j^* = 1$ (the candidates are using “follow the ranking”).

Note that the free-busy gap is strictly positive when the firm is using virtuous selection (indicating that free candidates have a higher chance of being selected) and strictly negative when the firm is using adverse selection (indicating that busy candidates have a higher chance of being selected). $\square$

Lemma B.3. *The free-busy gaps are given by the following terms:*

High value candidate:

Virtuous selection:

$$P[\sigma^1] \cdot \left(1 - (1 - p_2)^{j^*-1}\right) + P[\sigma^2] \cdot (1 - p_2) + P[\sigma^3] \cdot (1 - p_2)^2 + \dots + P[\sigma^{j^*}] \cdot (1 - p_2)^{j^*-1}$$

Adverse selection:

$$-P[\sigma^1] \cdot \left(1 - p_2^{k^*-1}\right) - P[\sigma^2] \cdot p_2 - P[\sigma^3] \cdot p_2^2 - \dots - P[\sigma^i] \cdot p_2^{k^*-1}$$

Follow the ranking: 0.

Low value candidate:

Virtuous selection:

$$\frac{1 - P[\sigma^1]}{n - 1} \cdot (1 - \beta_1) + \frac{1 - P[\sigma^2]}{n - 1} \cdot \beta_2 + \dots + \frac{1 - P[\sigma^{j^*}]}{n - 1} \cdot \beta_{j^*}$$

where $\beta_\ell = P[\sigma^{i < \ell} \mid i \neq \ell] \cdot (1 - p_1) \cdot (1 - p_2)^{\ell-2} + (1 - P[\sigma^{i < \ell} \mid i \neq \ell]) \cdot (1 - p_2)^{\ell-1}$ for $\ell > 1$ and $\beta_1 = P[\sigma^{2 \leq i \leq j^*} \mid i \neq 1] \cdot (1 - p_1) \cdot (1 - p_2)^{j^*-2} + (1 - P[\sigma^{2 \leq i \leq j^*} \mid i \neq 1]) \cdot (1 - p_2)^{j^*-1}$.

Adverse selection:

$$-\frac{1 - P[\sigma^1]}{n - 1} \cdot (1 - \alpha_1) - \frac{1 - P[\sigma^2]}{n - 1} \cdot \alpha_2 - \dots - \frac{1 - P[\sigma^{k^*}]}{n - 1} \cdot \alpha_{k^*}$$

where $\alpha_\ell = P[\sigma^{i < \ell} \mid i \neq \ell] \cdot p_1 \cdot p_2^{\ell-2} + (1 - P[\sigma^{i < \ell} \mid i \neq \ell]) \cdot p_2^{\ell-1}$ for $\ell > 1$ and $\alpha_1 = P[\sigma^{2 \leq i \leq k^*} \mid i \neq 1] \cdot p_1 \cdot p_2^{k^*-2} + (1 - P[\sigma^{2 \leq i \leq k^*} \mid i \neq 1]) \cdot p_2^{k^*-1}$.

Proof. Here, we will derive the form of the free-busy gap. In each case, it will depend on the strategy of the firm. Throughout, we will calculate the probability of a candidate being hired, conditioned on them being free (alternatively busy). Note that this is different from the proportion of candidates hired who are free vs. busy.

Follow the ranking:

In this setting, a free candidate is hired if and only if they are the top ranked candidate. Similarly, a

busy candidate is hired if and only if they are the top ranked. Because free-busy status is independent of ranking, the probability of each is exactly equal.

Virtuous selection:

In this setting, for some $j^* > 1$, the firm's strategy is to hire the top-ranked free candidate within the top j^* indices. If all top j^* candidates are busy, then the firm hires the top candidate. Thus, a free candidate is selected if they are in the top j^* and every candidate ranked above them is busy. A busy candidate is selected only if they are ranked first and every other candidate in the top j^* is also busy. Next, we derive the form for these probabilities.

High value candidate:

The probability that a free high value candidate is selected is:

$$\begin{aligned} & \sum_{i=1}^{j^*} P[\text{Candidate is in index } i] \cdot P[\text{all candidates in lower indices are busy} \mid \text{Candidate in index } i] \\ &= P[\sigma^1] + P[\sigma^2] \cdot (1 - p_2) + P[\sigma^3] \cdot (1 - p_2)^2 + \dots P[\sigma^{j^*}] \cdot (1 - p_2)^{j^*-1} \end{aligned}$$

The probability that a busy high value candidate is selected is:

$$\begin{aligned} & P[\text{Candidate is in index 1}] \cdot P[\text{All candidates in lower indices are busy} \mid \text{Candidate in index 1}] \\ &= P[\sigma^1] \cdot (1 - p_2)^{j^*-1} \end{aligned}$$

Thus, the gap is given by:

$$P[\sigma^1] \cdot \left(1 - (1 - p_2)^{j^*-1}\right) + P[\sigma^2] \cdot (1 - p_2) + P[\sigma^3] \cdot (1 - p_2)^2 + \dots P[\sigma^{j^*}] \cdot (1 - p_2)^{j^*-1}$$

Low value candidates:

The probability that a free low value candidate is selected is:

$$\begin{aligned} & \sum_{i=1}^{j^*} P[\text{Candidate is in index } i] \cdot P[\text{all candidates in lower indices are busy} \mid \text{Candidate in index } i] \\ &= \frac{1 - P[\sigma^1]}{n - 1} + \frac{1 - P[\sigma^2]}{n - 1} \cdot \beta_2 + \dots \frac{1 - P[\sigma^{j^*}]}{n - 1} \cdot \beta_{j^*} \end{aligned}$$

where $\beta_\ell = P[\sigma^{i < \ell} \mid i \neq \ell] \cdot (1 - p_1) \cdot (1 - p_2)^{\ell-2} + (1 - P[\sigma^{i < \ell} \mid i \neq \ell]) \cdot (1 - p_2)^{\ell-1}$ where this term gives the probability of every other item being ranked above the candidate in index ℓ is busy, depending on whether the high value item is above or below index ℓ . Note that the $n - 1$ term comes because all low value candidates are interchangeable and have equal chance of being at any index (except for the one where the high value candidate is). The probability that a busy low value candidate is selected is:

$$\begin{aligned} & P[\text{Candidate is in index 1}] \cdot P[\text{All candidates in lower indices are busy} \mid \text{Candidate in index 1}] \\ &= \frac{1 - P[\sigma^1]}{n - 1} \cdot \beta_1 \end{aligned}$$

where $\beta_1 = P[\sigma^{2 \leq i \leq j^*} \mid i \neq 1] \cdot (1 - p_1) \cdot (1 - p_2)^{j^*-2} + (1 - P[\sigma^{2 \leq i \leq j^*} \mid i \neq 1]) \cdot (1 - p_2)^{j^*-1}$ gives the probability that all other candidates in the top j^* are busy.

The free-busy gap is given by:

$$\frac{1 - P[\sigma^1]}{n - 1} \cdot (1 - \beta_1) + \frac{1 - P[\sigma^2]}{n - 1} \cdot \beta_2 + \dots \frac{1 - P[\sigma^{j^*}]}{n - 1} \cdot \beta_{j^*}$$

Adverse selection:

In this setting, for some $k^* > 1$, the firm's strategy is to hire the top-ranked busy candidate within the top k^* indices. If all top k^* candidates are free, then the firm hires the top-ranked candidate. A busy candidate is selected if they are in the top k^* and every candidate above them is free. A free candidate is selected only if they are ranked first and every other candidate in the top k^* is also free.

High value candidate:

The probability that a busy high value candidate is selected is:

$$\begin{aligned} \sum_{i=1}^{j^*} P[\text{Candidate is in index } i] \cdot P[\text{all candidates in lower indices are free} \mid \text{Candidate in index } i] \\ = P[\sigma^1] + P[\sigma^2] \cdot p_2 + P[\sigma^3] \cdot p_2^2 + \dots P[\sigma^i] \cdot p_2^{k^*-1} \end{aligned}$$

The probability that a free candidate is selected is:

$$\begin{aligned} P[\text{Candidate is in index 1}] \cdot P[\text{All candidates in lower indices are free} \mid \text{Candidate in index 1}] \\ = P[\sigma^1] \cdot p_2^{k^*-1} \end{aligned}$$

Thus, the free-busy gap is given by:

$$-P[\sigma^1] \cdot (1 - p_2^{k^*-1}) - P[\sigma^2] \cdot p_2 - P[\sigma^3] \cdot p_2^2 - \dots - P[\sigma^i] \cdot p_2^{k^*-1}$$

Low value candidates:

The probability that a busy candidate is selected is given by:

$$\begin{aligned} \sum_{i=1}^{j^*} P[\text{Candidate is in index } i] \cdot P[\text{all candidates in lower indices are free} \mid \text{Candidate in index } i] \\ = \frac{1 - P[\sigma^1]}{n-1} + \frac{1 - P[\sigma^2]}{n-1} \cdot \alpha_2 + \dots \frac{1 - P[\sigma^{k^*}]}{n-1} \cdot \alpha_{k^*} \end{aligned}$$

where $\alpha_\ell = P[\sigma^{i < \ell} \mid i \neq \ell] \cdot p_1 \cdot p_2^{\ell-2} + (1 - P[\sigma^{i < \ell} \mid i \neq \ell]) \cdot p_2^{\ell-1}$ gives the probability of all candidates ranked above the candidate in index ℓ being free.

The probability that a free candidate is selected is given by:

$$\begin{aligned} P[\text{Candidate is in index 1}] \cdot P[\text{All candidates in lower indices are free} \mid \text{Candidate in index 1}] \\ = \frac{1 - P[\sigma^1]}{n-1} \cdot \alpha_1 \end{aligned}$$

where $\alpha_1 = P[\sigma^{2 \leq i \leq k^*} \mid i \neq 1] \cdot p_1 \cdot p_2^{k^*-2} + (1 - P[\sigma^{2 \leq i \leq k^*} \mid i \neq 1]) \cdot p_2^{k^*-1}$.

Thus, the free-busy gap is given by:

$$-\frac{1 - P[\sigma^1]}{n-1} \cdot (1 - \alpha_1) - \frac{1 - P[\sigma^2]}{n-1} \cdot \alpha_2 - \dots - \frac{1 - P[\sigma^{k^*}]}{n-1} \cdot \alpha_{k^*}$$

□

Lemma 5.5. *There exist settings where increasing the accuracy of the ranking tool can either increase or decrease the magnitude of the free-busy gap.*

Proof. First, we note that if a firm is within the “follow the ranking” strategy space, then increasing the accuracy of the ranking tool has *no* impact on the free-busy gap (it stays constant at 0).

We will again use the functional forms for the free-busy gap derived in Lemma B.3, simplifying the terms for $N = 2$. High value candidate:

Virtuous selection:

$$\begin{aligned} P[\sigma^1] \cdot (1 - (1 - p_2)) + (1 - P[\sigma^1]) \cdot (1 - p_2) \\ = 1 - P[\sigma^1] + p_2 \cdot (2 \cdot P[\sigma^1] - 1) \end{aligned}$$

Adverse selection :

$$\begin{aligned} -P[\sigma^1] \cdot (1 - p_2) - (1 - P[\sigma^1]) \cdot p_2 \\ = -p_2 \cdot (1 - 2 \cdot P[\sigma^1]) - P[\sigma^1] \end{aligned}$$

When does the *magnitude* of the free-busy gap increase with increased accuracy (here, increased $P[\sigma^1]$)? Taking the derivative of the free-busy gap under virtuous selection gives:

$$-1 + 2 \cdot p_2$$

Taking the derivative on the magnitude of the free-busy gap under adverse selection gives:

$$-2 \cdot p_2 + 1$$

Thus, if $p_2 > 0.5$, increasing the accuracy of the ranking tool increases the free-busy gap under virtuous selection and decreases it under adverse selection, and vice versa if $p_2 < 0.5$.

For intuition, under virtuous selection a busy candidate is only selected if it's ranked first and the other candidate happens to be busy, while a free candidate is always selected (unless it's ranked second and the other candidate is also free). If $p_2 > 0.5$, the low value candidate is more likely to be free than busy, so increasing the accuracy $P[\sigma^1]$ increases the probability of the free (high value) agent being selected, as compared to an alternative busy (high value) agent.

Similar reasoning follows for adverse selection and for the low value candidate's free-busy gap. $\square$

C Proofs from Appendix A

Lemma A.1. *There exists a probability distribution $\mathcal{P}$ such that for some realized status vector s , picking the second free item maximizes expected utility, even if $\mathcal{P}$ is descending in expected value.*

Proof. We will set the 3 items to have values $[v_1 = 1, v_2, v_3 = 0]$. We will create a permutation distribution with nonzero support on exactly two permutations:

$$\sigma^1 = [v_1, v_2, v_3] \quad \sigma^2 = [v_3, v_2, v_1]$$

where $P[\sigma^1] = 1 - \epsilon$, $P[\sigma^2] = \epsilon$. Before the item status vector is taken into account, the expected utility of picking item $i \in [1, 2, 3]$ is given by:

$$\mathbb{E}_{\sigma \sim \mathcal{P}}[v_{\sigma_1}] = v_1 \cdot (1 - \epsilon) \quad \mathbb{E}_{\sigma \sim \mathcal{P}}[v_{\sigma_2}] = v_2 \quad \mathbb{E}_{\sigma \sim \mathcal{P}}[v_{\sigma_3}] = \epsilon \cdot v_1$$

If we want the distribution to have descending expected value, we will set:

$$v_1 \cdot (1 - \epsilon) > v_2 > \epsilon \cdot v_1$$

We will next consider the probabilities of being free. For simplicity, we will set $p_1 < p_2 = p_3$, so items 2 and 3 have the same probability of being free, which is greater than for item 1.

Next, we will consider the case where we have realized status vector $[1, 1, 0]$. This increases our posterior belief that we are in σ^2 , which decreases the expected utility of picking the first item.

We will formalize this next. The posterior value for the first item is given by:

$$P[v_{\sigma_1} = v_1 \mid s] \cdot v_1 = \frac{P[s \mid \sigma^1] \cdot P[\sigma^1]}{P[s]} = v_1 \cdot \frac{p_1 \cdot p_2 \cdot (1 - p_2) \cdot (1 - \epsilon)}{P[s]}$$

The posterior value for the second item is always exactly v_2 , because both possible permutations involve having v_2 ranked second. However, we will find it useful to strategically rewrite this probability :

$$\begin{aligned} P[v_{\sigma_2} = v_2 \mid s] \cdot v_2 &= v_2 = v_2 \cdot \frac{P[s]}{P[s]} = v_2 \cdot \frac{P[s \mid \sigma^1] \cdot P[\sigma^1] + P[s \mid \sigma^2] \cdot P[\sigma^2]}{P[s]} \\ &= v_2 \cdot \frac{p_1 \cdot p_2 \cdot (1 - p_2) \cdot (1 - \epsilon) + p_2^2 \cdot (1 - p_1) \cdot \epsilon}{P[s]} \end{aligned}$$

The posterior value for the third item is given by:

$$P[v_{\sigma_3} = v_1 \mid s] \cdot v_1 \cdot \gamma = \gamma \cdot v_1 \cdot \frac{P[s \mid \sigma^2] \cdot P[\sigma^2]}{P[s]} = \gamma \cdot v_1 \cdot \frac{p_2^2 \cdot (1 - p_1) \cdot \epsilon}{P[s]}$$

Next, we will show that there exists parameters such that the optimal strategy (item with highest posterior utility) is the second (free) item, rather than the first free or first busy strategy. Specifically, we will show that this condition holds for $\epsilon = 0.1, v_1 = 1, v_2 = 2/3, v_3 = 0, p_1 = 0.1, p_2 = 0.4, \gamma \in [0, 1]$. Note that these parameters immediately satisfies $v_2 < (1 - \epsilon) \cdot v_1$, which implies that the prior ranking is descending in expected utility (as desired).

Note that each of these expected utility terms have a common denominator of $P[a]$, which we can drop. We know that picking the second (free) item has higher expected utility than picking the first (free) item whenever:

$$v_1 \cdot p_1 \cdot p_2 \cdot (1 - p_2) \cdot (1 - \epsilon) < v_2 \cdot (p_1 \cdot p_2 \cdot (1 - p_2) \cdot (1 - \epsilon) + p_2^2 \cdot (1 - p_1) \cdot \epsilon)$$

Dropping a common term of p_2 and distributing v_2 :

$$v_1 \cdot p_1 \cdot (1 - p_2) \cdot (1 - \epsilon) < v_2 \cdot p_1 \cdot (1 - p_2) \cdot (1 - \epsilon) + v_2 \cdot p_2 \cdot (1 - v_1) \cdot \epsilon$$

Substituting in for the given values sets:

$$1 \cdot 0.1 \cdot 0.6 \cdot 0.9 < \frac{2}{3} \cdot 0.1 \cdot 0.6 \cdot 0.9 + \frac{2}{3} \cdot 0.4 \cdot 0.9 \cdot 0.1$$

$$0.054 < 0.06$$

as desired.

We will additionally require that picking the third item (given s realized) has lower expected utility than picking the second item. This means we require:

$$v_2 \cdot (p_1 \cdot p_2 \cdot (1 - p_2) \cdot (1 - \epsilon) + p_2^2 \cdot (1 - p_1) \cdot \epsilon) > \gamma \cdot v_1 \cdot p_2^2 \cdot (1 - p_1) \cdot \epsilon$$

$$v_2 \cdot (p_1 \cdot (1 - p_2) \cdot (1 - \epsilon) + p_2 \cdot (1 - p_1) \cdot \epsilon) > \gamma \cdot v_1 \cdot p_2 \cdot (1 - p_1) \cdot \epsilon$$

Substituting in for given values:

$$v_2 \cdot (p_1 \cdot (1 - p_2) \cdot (1 - \epsilon) + p_2 \cdot (1 - p_1) \cdot \epsilon) > \gamma \cdot v_1 \cdot p_2 \cdot (1 - p_1) \cdot \epsilon$$

$$\frac{2}{3} \cdot 0.1 \cdot 0.6 \cdot 0.9 + \frac{2}{3} \cdot 0.4 \cdot 0.9 \cdot 0.1 > \gamma \cdot 1 \cdot 0.4 \cdot 0.9 \cdot 0.1$$

$$0.06 > 0.036 \cdot \gamma$$

which holds for any $\gamma \in [0, 1]$. □

Lemma A.2. Consider any vector of realized status vector s , with $s_i = s_j$, for some $i < j$. Then, if the probability distribution of permutations is inversion-monotone as in Definition 5, $\mathbb{E}[v_{\sigma_i} | a] \geq \mathbb{E}[v_{\sigma_j} | a]$. This implies that the optimal solution will always be to pick the first free item or the first busy item.

Proof. The expected value of the i th entry is given by:

$$\mathbb{E}[v_{\sigma_i} | s] = \sum_{\sigma \in \mathcal{P}} P[\sigma | s] \cdot v_{\sigma_i} = \sum_{\sigma \in \mathcal{P}} \frac{P[s | \sigma] \cdot P[\sigma]}{P[s]} \cdot v_{\sigma_i}$$

By identical reasoning, the expected value of the j th entry is given by:

$$\mathbb{E}[v_{\sigma_j} | s] = \sum_{\sigma \in \mathcal{P}} P[\sigma | s] \cdot v_{\sigma_j} = \sum_{\sigma \in \mathcal{P}} \frac{P[s | \sigma] \cdot P[\sigma]}{P[s]} \cdot v_{\sigma_j}$$

We wish to show that:

$$\sum_{\sigma \in \mathcal{P}} \frac{P[s | \sigma] \cdot P[\sigma]}{P[s]} \cdot v_{\sigma_i} \geq \sum_{\sigma \in \mathcal{P}} \frac{P[s | \sigma] \cdot P[\sigma]}{P[s]} \cdot v_{\sigma_j}$$

Consider some σ such that $v_{\sigma_i} < v_{\sigma_j}$. Then, we consider a unique mapping to some $\tilde{\sigma}$ such that $\sigma_k = \tilde{\sigma}_k$ except for $\tilde{\sigma}_i = \sigma_j, \tilde{\sigma}_j = \sigma_i$ (the i th and j th elements are swapped). Then, we will show that:

$$P[\sigma | s] \cdot v_{\sigma_i} + P[\tilde{\sigma} | s] \cdot v_{\tilde{\sigma}_i} > P[\sigma | s] \cdot v_{\sigma_j} + P[\tilde{\sigma} | s] \cdot v_{\tilde{\sigma}_j}$$

By construction, we have $\tilde{\sigma}_i = \sigma_j, \tilde{\sigma}_j = \sigma_i$, so we can rewrite this as:

$$P[\sigma | s] \cdot v_{\sigma_i} + P[\tilde{\sigma} | s] \cdot v_{\sigma_j} > P[\sigma | s] \cdot v_{\sigma_j} + P[\tilde{\sigma} | s] \cdot v_{\sigma_i}$$

$$(P[\tilde{\sigma} | s] - P[\sigma | s]) \cdot (v_{\sigma_j} - v_{\sigma_i}) > 0$$

Next, we can rewrite using Bayes rule:

$$\left(\frac{P[s | \tilde{\sigma}] \cdot P[\tilde{\sigma}]}{P[s]} - \frac{P[s | \sigma] \cdot P[\sigma]}{P[s]} \right) \cdot (v_{\sigma_j} - v_{\sigma_i}) > 0$$

Dropping the common denominator gives:

$$(P[s | \tilde{\sigma}] \cdot P[\tilde{\sigma}] - P[s | \sigma] \cdot P[\sigma]) \cdot (v_{\sigma_j} - v_{\sigma_i}) > 0$$

We will argue that $P[s | \tilde{\sigma}] = P[s | \sigma]$. Recall that $\tilde{\sigma}, \sigma$ are identical except for at entries i, j . Because $s_i = s_j$ by assumption, for the probability of observing s does not change if items are swapped between these two entries. Dropping this common term gives:

$$(P[\tilde{\sigma}] - P[\sigma]) \cdot (v_{\sigma_j} - v_{\sigma_i}) > 0$$

The second term is satisfied because $v_{\sigma_j} > v_{\sigma_i}$ by assumption, and the first term holds because $P[\tilde{\sigma}] > P[\sigma]$ by the requirement that the distribution is inversion-monotone as in Definition 5. $\square$

Lemma A.3. *The Mallows model is inversion-monotone.*

Proof. Consider some permutation σ , with $v_{\sigma_i} > v_{\sigma_j}$ for $i > j$. Recall that the probability of seeing a permutation is inversely proportional to the number of inversions (e.g. $k > l$ such that $v_{\sigma_k} < v_{\sigma_l}$).

If $i = j + 1$ then this is obvious: this only changes the number of inversions for i, j , and increases the number of inversions by exactly 1.

If $i > j + 1$, then there exists some set $S \in [i + 1, j - 1]$. Flipping $v_{\sigma_i}, v_{\sigma_j}$ also changes the number of inversions within this set (but maintains the number of inversions in elements outside of this set). We would like to show that if $v_{\sigma_i} > v_{\sigma_j}$, flipping the elements i, j only increases the number of inversions (and thus makes this alternative permutation less likely).

Within S , we can denote the set of elements that are *greater* and *less* than i and j respectively by:

$$\begin{aligned} G_i &= \{k \in S \mid v_{\sigma_k} > v_{\sigma_i}\} & G_j &= \{k \in S \mid v_{\sigma_k} > v_{\sigma_j}\} \\ L_i &= \{k \in S \mid v_{\sigma_k} < v_{\sigma_i}\} & L_j &= \{k \in S \mid v_{\sigma_k} < v_{\sigma_j}\} \end{aligned}$$

Note that $G_i \subseteq G_j$ and $L_j \subseteq L_i$: the set of elements that is greater than i is a subset of the set of elements that is greater than j , and the set of elements that is less than j is a subset of the elements that is smaller than i .

Note that within S in the original permutation π , the number of inversions is given by

$$|G_i| + |L_j|$$

the elements that are greater than v_{σ_i} and less than v_{σ_j} . After swapping the order of elements i, j , the number of inversions within S is given by:

$$|G_j| + |L_i|$$

the elements that are greater than π_j and less than π_i . By our prior reasoning:

$$|G_j| \geq |G_i| \quad |L_i| \geq |L_j|$$

and so:

$$|G_j| + |L_i| \geq |G_i| + |L_j|$$

We note that flipping i, j involves at least one more inversion (since now we have $v_{\sigma_j} < v_{\sigma_i}$) and so the swapping process strictly increased the total number of inversions. $\square$

Theorem 2. *The RUM with identical, symmetric, noise distributions across items is inversion-monotone.*

Proof. Consider any permutation σ with at least one inversion (a $v_{\sigma_i} < v_{\sigma_j}$ for $i < j$). Such a permutation can occur when the noised values of these items $\{\hat{v}_i\}$ fall within certain ranges. Consider the relevant pair of elements i, j . Lemma A.4 proves that for any pair of values $\hat{v}_i, \hat{v}_j$ with $|\hat{v}_i - \hat{v}_j| = \Delta$, given a RUM, it is more likely that we would have $\hat{v}_i < \hat{v}_j$ (ordered correctly) than $\hat{v}_i > \hat{v}_j$ (ordered incorrectly). Given values generated by a RUM, the value of every item $k \neq i, j$ is completely independent from the value of items i, j . Therefore, we know that there must exist a permutation $\tilde{\sigma}$ identical to σ except with items σ_i, σ_j flipped - and such a permutation must be strictly more likely than σ . $\square$

Lemma A.4. Suppose we have two random variables given by $X_1 = \mu_1 + \epsilon_1$, $X_2 = \mu_2 + \epsilon_2$, for $\mu_1 > \mu_2$ and $\epsilon_1, \epsilon_2 \sim \mathcal{D}$ for a symmetric, single-peaked distribution $\mathcal{D}$. Then, if $|X_1 - X_2| = \Delta$,

$$P[X_1 > X_2 \mid |X_1 - X_2| = \Delta] > P[X_1 < X_2 \mid |X_1 - X_2| = \Delta]$$

Proof. We will begin by analyzing the distribution $\mathcal{D}'$ induced by $X = X_1 - X_2$.

Symmetric:

First, we will show that it is symmetric around $\mu_1 - \mu_2$. In order to show this, we will show that:

$$P[X = \mu_1 - \mu_2 + \delta] = P[X = \mu_1 - \mu_2 - \delta]$$

for all $\delta > 0$. Suppose that we have some ϵ_1, ϵ_2 such that:

$$X = X_1 - X_2 = \mu_1 + \epsilon_1 - \mu_2 - \epsilon_2 = \mu_1 - \mu_2 + \delta$$

for $\delta = \epsilon_1 - \epsilon_2$. Then, we can show a mapping to an equally-likely event where $X = \mu_1 - \mu_2 - \delta$. Because ϵ_1, ϵ_2 are drawn from the same distribution $\mathcal{D}$, it is equally likely that we would have $\epsilon'_1 = \epsilon_2, \epsilon'_2 = \epsilon_1$. This would give us:

$$X = X_1 - X_2 = \mu_1 + \epsilon_2 - \mu_2 - \epsilon_1 = \mu_1 - \mu_2 + (\epsilon_2 - \epsilon_1) = \mu_1 - \mu_2 - (\epsilon_1 - \epsilon_2) = \mu_1 - \mu_2 - \delta$$

as desired.

Unimodal:

The next thing we would like to show is that $\mathcal{D}'$ is unimodal (strictly increasing and then decreasing).

We will look at the distribution $g(x)$ of $\mathcal{D}'$ directly. We'd like to show that this is decreasing for $x > \mu_1 - \mu_2$. What's the probability density of $g(\delta)$ for some $\delta = \epsilon_1 - \epsilon_2$? We can obtain this by integrating over all possible values of ϵ_1, ϵ_2 , along with the cdf $f(\cdot)$ for the noise distribution:

$$g(\delta) = \int_{-\infty}^{\infty} \mathbb{1}[\delta = \epsilon_1 - \epsilon_2] f(\epsilon_1) \cdot f(\epsilon_2) d\epsilon_1 d\epsilon_2$$

We can rewrite as a single integral over ϵ_1 and use $\epsilon_2 = \epsilon_1 - \delta$.

$$g(\delta) = \int_{-\infty}^{\infty} f(\epsilon_1) \cdot f(\epsilon_1 - \delta) d\epsilon_1$$

We would like to show that this is decreasing in δ for $\delta > 0$. We can take the derivative wrt δ , which gives us:

$$\frac{d}{d\delta} g(\delta) = \frac{d}{d\delta} \int_{-\infty}^{\infty} f(\epsilon_1) \cdot f(\epsilon_1 - \delta) d\epsilon_1$$

Integration and differentiation commutes, so we have:

$$= \int_{-\infty}^{\infty} f(\epsilon_1) \cdot \frac{d}{d\delta} f(\epsilon_1 - \delta) d\epsilon_1$$

by the chain rule:

$$= - \int_{-\infty}^{\infty} f(\epsilon_1) \cdot f'(\epsilon_1 - \delta) d\epsilon_1$$

Which means we want to show that:

$$\int_{-\infty}^{\infty} f(\epsilon_1) \cdot f'(\epsilon_1 - \delta) d\epsilon_1 > 0 \tag{1}$$

Let's consider a closely related term:

$$\int_{-\infty}^{\infty} f(\epsilon_1) \cdot f'(\epsilon_1) d\epsilon_1$$

We know that $f(\cdot)$ is symmetric around 0 and is increasing below 0 and decreasing above 0, which means that this above term equals 0. $f(\epsilon_1 - \delta)$ is shifted to the right: we will show that this suffices to show that Equation 1 is positive.

We can rewrite Equation 1 as:

$$\int_{-\infty}^{\infty} f(\epsilon + \delta) \cdot f'(\epsilon) d\epsilon$$

And what we want to show is

$$\int_{-\infty}^{\infty} f(\epsilon + \delta) \cdot f'(\epsilon) d\epsilon > \int_{-\infty}^{\infty} f(\epsilon) \cdot f'(\epsilon_1) d\epsilon$$

Or rewritten:

$$\int_{-\infty}^{\infty} (f(\epsilon + \delta) - f(\epsilon)) \cdot f'(\epsilon) d\epsilon > 0$$

First, we can divide this into two components based on whether ϵ is positive or negative:

$$\begin{aligned} & \int_{-\infty}^0 (f(\epsilon + \delta) - f(\epsilon)) \cdot f'(\epsilon) d\epsilon + \int_0^{\infty} (f(\epsilon + \delta) - f(\epsilon)) \cdot f'(\epsilon) d\epsilon \\ & \int_0^{\infty} (f(-\epsilon + \delta) - f(-\epsilon)) \cdot f'(-\epsilon) d\epsilon + \int_0^{\infty} (f(\epsilon + \delta) - f(\epsilon)) \cdot f'(\epsilon) d\epsilon \end{aligned}$$

Because f is symmetric by our prior reasoning, we have that $f(\epsilon) = f(-\epsilon)$ and $f'(\epsilon) = -f'(-\epsilon)$, which allows us to rewrite:

$$- \int_0^{\infty} (f(-\epsilon + \delta) - f(\epsilon)) \cdot f'(\epsilon) d\epsilon + \int_0^{\infty} (f(\epsilon + \delta) - f(\epsilon)) \cdot f'(\epsilon) d\epsilon$$

We can then combine these terms to give:

$$\begin{aligned} & \int_0^{\infty} f'(\epsilon) \cdot (f(\epsilon + \delta) - f(\epsilon) - f(-\epsilon + \delta) + f(\epsilon)) d\epsilon \\ & = \int_0^{\infty} f'(\epsilon) \cdot (f(\epsilon + \delta) - f(-\epsilon + \delta)) d\epsilon \end{aligned}$$

We know that $f'(\epsilon)$ is always negative for positive ϵ (by the assumption that $f(\cdot)$ is unimodal and centered at 0). In order for the total term to be positive, we need to show that the other term is also negative - that is, that

$$f(-\epsilon + \delta) > f(\epsilon + \delta) \quad \forall \epsilon > 0, \delta > 0$$

Again, we use that $f(\cdot)$ is symmetric and unimodal. This means that

$$f(\epsilon + \delta) = f(-\epsilon - \delta) < f(-\epsilon + \delta)$$

as desired.

Overall reasoning: Having shown that $\mathcal{D}'$ is unimodal and symmetric, we will now show that $P[X = \Delta] > P[X = -\Delta]$ for all $\Delta > 0$. We know that $\mathcal{D}'$ has a mean at $\mu_1 - \mu_2 > 0$.

If $0 < \Delta < \mu_1 - \mu_2$, then the event $X = \Delta$ occurs on the lefthand side of the unimodal distribution, as does the event $X = -\Delta$. Because $\mathcal{D}'$ is unimodal and strictly decreasing, this implies that $P[X = -\Delta] < P[X = \Delta]$.

If $0 < \mu_1 - \mu_2 \leq \Delta$, then $P[X = \Delta]$ occurs on the righthand side of the unimodal curve, at a distance $\Delta - (\mu_1 - \mu_2)$ above the peak. Then, reflecting across the axis of symmetry gives the point $P[X = (\mu_1 - \mu_2) - (\Delta - (\mu_1 - \mu_2))] = P[X = 2 \cdot (\mu_1 - \mu_2) - \Delta] = P[X = \Delta]$. This point is on the lefthand side of the unimodal curve, but because $\mu_1 - \mu_2 > 0$, this point also is above the point $P[X = -\Delta]$, and therefore we have:

$$P[X = \Delta] = P[X = 2 \cdot (\mu_1 - \mu_2) - \Delta] > P[X = -\Delta]$$

as desired. □

Theorem 3. Consider the superstar setting with $v_2 > 0$. Then, the agent's optimal decision is given by:

$$\text{if } s_{i < j} = 0, s_j = 1 \begin{cases} \text{pick the first item} & \gamma > T_b^j \\ \text{pick the } j\text{th item} & \text{otherwise} \end{cases}$$

$$\text{if } s_{i < j} = 1, s_j = 0 \begin{cases} \text{pick the first item} & \gamma < T_f^j \\ \text{pick the } j\text{th item} & \text{otherwise} \end{cases}$$

For:

$$T_b^j = \frac{(v_1 - v_2) \cdot P[\sigma^j] + v_2 \cdot (r \cdot P[\sigma^{i < j}] + P[\sigma^j] + C)}{(v_1 - v_2) \cdot r \cdot P[\sigma^1] + v_2 \cdot (r \cdot P[\sigma^{i < j}] + P[\sigma^j] + C)}$$

$$T_f^j = \frac{(v_1 - v_2) \cdot P[\sigma^1] + v_2 \cdot (P[\sigma^{i < j}] + r \cdot P[\sigma^j] + C)}{(v_1 - v_2) \cdot r \cdot P[\sigma^j] + v_2 \cdot (P[\sigma^{i < j}] + r \cdot P[\sigma^j] + C)}$$

where $r = \frac{p_2/(1-p_2)}{p_1/(1-p_1)}$ and $C = \sum_{i=j+1}^N r^{1-s_i} P[\sigma^i]$.

Proof. We will begin with the setting where $s_{i < j} = 0, s_j = 1$ (the highest-ranked item is busy). We pick the first item whenever:

$$\mathbb{E}_{\sigma \sim \mathcal{P}}[v_{\sigma_1} \mid s] \cdot \gamma \geq \mathbb{E}_{\sigma \sim \mathcal{P}}[v_{\sigma_j} \mid s]$$

Using similar analysis to the proof of Theorem 3, we can apply Bayes rule to rewrite this as:

$$\sum_{\ell=1}^N \frac{P[s \mid \sigma^\ell] \cdot P[\sigma^\ell]}{P[s]} \cdot v_{\sigma_1^\ell} \cdot \gamma \geq \sum_{\ell=1}^N \frac{P[s \mid \sigma^\ell] \cdot P[\sigma^\ell]}{P[s]} \cdot v_{\sigma_j^\ell}$$

$$\sum_{\ell=1}^N P[s \mid \sigma^\ell] \cdot P[\sigma^\ell] \cdot v_{\sigma_1^\ell} \cdot \gamma \geq \sum_{\ell=1}^N P[s \mid \sigma^\ell] \cdot P[\sigma^\ell] \cdot v_{\sigma_j^\ell}$$

Because we are in the superstar setting, we know that $v_{\sigma_i^\ell} = v_2$ unless $i = \ell$, so we can rewrite the condition as:

$$\gamma \cdot \left(P[s \mid \sigma^1] \cdot P[\sigma^1] \cdot v_1 + \sum_{i=2}^N P[s \mid \sigma^i] \cdot P[\sigma^i] \cdot v_2 \right) \geq P[s \mid \sigma^j] \cdot P[\sigma^j] \cdot v_1 + \sum_{i \neq j} P[s \mid \sigma^i] \cdot P[\sigma^i] \cdot v_2$$

Again using similar analysis to Theorem 3, we can reason about $P[s \mid \sigma^i]$, for a general i . Specifically, we can show that:

$$\begin{cases} P[s \mid \sigma^i] = P[s \mid \sigma^j] & s_i = s_j = 1 \\ P[s \mid \sigma^i] = r \cdot P[s \mid \sigma^j] & s_i = 0 \neq s_j = 1 \end{cases}$$

where the probability of observing a is *higher* if the high value item is in an index i with $s_i = 0$, because the high value item has $p_1 < p_2$ and is thus less likely to be free. Using this and the fact that $s_{i < j} = 0$, we can rewrite the conditions as:

$$\gamma \cdot P[s \mid \sigma^j] \left(r \cdot P[\sigma^1] \cdot v_1 + r \cdot P[\sigma^{2 \leq i < j}] \cdot v_2 + P[\sigma^j] \cdot v_2 + \sum_{i=j+1}^N r^{1-s_i} \cdot P[\sigma^i] \cdot v_2 \right)$$

$$\geq P[s \mid \sigma^j] \left(P[\sigma^j] \cdot v_1 + r \cdot P[\sigma^{1 \leq i < j}] \cdot v_2 + \sum_{i=j+1}^N r^{1-s_i} \cdot P[\sigma^i] \cdot v_2 \right)$$

which is equivalent to:

$$\gamma \geq \frac{v_1 \cdot P[\sigma^j] + v_2 \cdot (r \cdot P[\sigma^{1 \leq i < j}] + C)}{r \cdot P[\sigma^1] \cdot v_1 + v_2 \cdot (r \cdot P[\sigma^{2 \leq i < j}] + P[\sigma^j] + C)} = \frac{(v_1 - v_2) \cdot P[\sigma^j] + v_2 \cdot (r \cdot P[\sigma^{i < j}] + P[\sigma^j] + C)}{(v_1 - v_2) \cdot r \cdot P[\sigma^1] + v_2 \cdot (r \cdot P[\sigma^{i < j}] + P[\sigma^j] + C)}$$

as desired.

Next, we move to the setting where $s_{i < j} = 1, s_j = 0$ (the highest-ranking item is free). We pick the first item whenever:

$$\mathbb{E}_{\sigma \sim \mathcal{P}}[v_{\sigma_1} \mid s] \geq \mathbb{E}_{\sigma \sim \mathcal{P}}[v_{\sigma_j} \mid s] \cdot \gamma$$

By our prior analysis, we know this can be rewritten as:

$$P[s \mid \sigma^1] \cdot P[\sigma^1] \cdot v_1 + \sum_{i=2}^N P[s \mid \sigma^i] \cdot P[\sigma^i] \cdot v_2 \geq \gamma \cdot \left(P[s \mid \sigma^j] \cdot P[\sigma^j] \cdot v_1 + \sum_{i \neq j} P[s \mid \sigma^i] \cdot P[\sigma^i] \cdot v_2 \right)$$

Next, we again reason about $P[s \mid \sigma^i]$ for an arbitrary i . Note that the s status vector is different than in the previous part of this proof, so we will have a different correspondence:

$$\begin{cases} P[s \mid \sigma^i] = P[s \mid \sigma^1] & s_i = s_1 = 1 \\ P[s \mid \sigma^i] = r \cdot P[s \mid \sigma^1] & s_i = 0 \neq s_1 = 1 \end{cases}$$

This enables us to rewrite the condition as:

$$P[s \mid \sigma^1] \cdot (v_1 \cdot P[\sigma^1] + P[\sigma^{2 \leq i < j}] \cdot v_2 + r \cdot P[\sigma^j] \cdot v_2 + C \cdot v_2) \geq \gamma \cdot P[s \mid \sigma^1] \cdot (r \cdot P[\sigma^j] \cdot v_1 + v_2 \cdot P[\sigma^{1 \leq i < j}] + C \cdot v_2)$$

$$\gamma \leq \frac{v_1 \cdot P[\sigma^1] + P[\sigma^{2 \leq i < j}] \cdot v_2 + r \cdot P[\sigma^j] \cdot v_2 + C \cdot v_2}{r \cdot P[\sigma^j] \cdot v_1 + v_2 \cdot P[\sigma^{1 \leq i < j}] + C \cdot v_2} = \frac{(v_1 - v_2) \cdot P[\sigma^1] + v_2 \cdot (P[\sigma^{i < j}] + r \cdot P[\sigma^j] + C)}{(v_1 - v_2) \cdot r \cdot P[\sigma^j] + v_2 \cdot (P[\sigma^{i < j}] + r \cdot P[\sigma^j] + C)}$$

as desired. $\square$

Lemma A.5. *For any $N > j \geq 2$, if there exists $k > j \geq 2$ such that $P[\sigma^k] = \epsilon > 0$, then there exists γ, s such that observing the status vector up to index k will change the choice of which item will be selected.*

Proof. This can be seen directly by the bounds given in Lemma A.6. These bounds are tight - there exist realized status vectors s that result in the given upper and lower bounds on T_b^j, T_f^j . Additionally, the upper bound is always strictly higher than the lower bound whenever there is some positive probability $P[\sigma^k] \geq \epsilon$ for $k > j$ (and note, the decreasing expected value of the ranking implies $P[\sigma^l] \geq P[\sigma^k] = \epsilon$ for $l \in (j, k)$). Then, any γ value that falls within the upper and lower bounds on T_b^j, T_f^j is one where seeing the status vector up to index k would change the item the agent would choose. $\square$

Lemma A.6. *It is possible to construct tight bounds on the thresholds in Theorem 3 using the status of only the first $k \geq j$ items. Specifically, such bounds are given by the terms below (with $C_0 = \sum_{i=j+1}^k r^{1-s_i^j} P[\sigma^i]$ throughout):*

When $s_{i < j} = 0, s_j = 1$, we have:

$$B_{j,k}^1 \leq T_b^j \leq B_{j,k}^2$$

For:

$$B_{j,k}^1 = \frac{(v_1 - v_2) \cdot P[\sigma^j] + v_2 \cdot (r - 1) \cdot P[\sigma^{i < j}] + v_2 + v_2 \cdot (C_0 - P[\sigma^{j < i \leq k}])}{(v_1 - v_2) \cdot P[\sigma^1] + v_2 \cdot (r - 1) \cdot P[\sigma^{i < j}] + v_2 + v_2 \cdot (C_0 - P[\sigma^{j < i \leq k}])}$$

$$B_{j,k}^2 = \frac{(v_1 - v_2) \cdot P[\sigma^j] + v_2 \cdot (r - 1) \cdot (1 - P[\sigma^j]) + v_2 + v_2 \cdot (C_0 - r \cdot P[\sigma^{i > k}])}{(v_1 - v_2) \cdot r \cdot P[\sigma^1] + v_2 \cdot (r - 1) \cdot (1 - P[\sigma^j]) + v_2 + v_2 \cdot (C_0 - r \cdot P[\sigma^{i > k}])}$$

where the upper bound is given by the case where $s_\ell = 0 \forall \ell > k$, and the lower bound is given by $s_\ell = 1 \forall \ell > k$.

For the case when $s_{i < j} = 1, s_j = 0$, we have:

$$\begin{cases} F_{j,k}^1 \leq T_j^f \leq F_{j,k}^2 & P[\sigma^1] \leq r \cdot P[\sigma^j] \\ F_{j,k}^2 \leq T_j^f \leq F_{j,k}^1 & \text{otherwise} \end{cases}$$

for

$$F_{j,k}^1 = \frac{(v_1 - v_2) \cdot P[\sigma^1] + v_2 \cdot (P[\sigma^{i < j}] + r \cdot P[\sigma^j] + C_0 - P[\sigma^{i > k}])}{(v_1 - v_2) \cdot r \cdot P[\sigma^j] + v_2 \cdot (P[\sigma^{i < j}] + r \cdot P[\sigma^j] + C_0 - P[\sigma^{i > k}])}$$

$$F_{j,k}^2 = \frac{(v_1 - v_2) \cdot P[\sigma^1] + v_2 \cdot (1 + (r - 1) \cdot P[\sigma^{i > j}] + C_0 - r \cdot P[\sigma^{j < i \leq k}])}{(v_1 - v_2) \cdot r \cdot P[\sigma^j] + v_2 \cdot (1 + (r - 1) \cdot P[\sigma^{i > j}] + C_0 - r \cdot P[\sigma^{j < i \leq k}])}$$

where the C_j^2 bound occurs when $s_\ell = 0 \forall \ell > k$, and the C_j^1 bound occurs when $s_\ell = 1 \forall \ell > k$.

Proof. We will again begin with the $s_{i < j} = 0, s_j = 1$ case (where the highest-ranked item is busy). From Theorem 3 we know that the relevant threshold is

$$T_b^j = \frac{(v_1 - v_2) \cdot P[\sigma^j] + v_2 \cdot (r \cdot P[\sigma^{i < j}] + P[\sigma^j] + C_0 + C_1)}{(v_1 - v_2) \cdot r \cdot P[\sigma^1] + v_2 \cdot (r \cdot P[\sigma^{i < j}] + P[\sigma^j] + C_0 + C_1)}$$

where $r = \frac{p_2/(1-p_2)}{p_1/(1-p_1)} \geq 1$ and $C_0 = \sum_{i=j+1}^k r^{1-s_i^j} P[\sigma^i]$, $C_1 = \sum_{i=k+1}^N r^{1-s_i^j} P[\sigma^i]$. Note that we can immediately see that C_1 is bounded by two cases, where all items lower ranked than j are free or busy:

$$P[\sigma^{i > k}] \leq C_1 \leq r \cdot P[\sigma^{i > k}]$$

Next we will show that plugging in an upper bound on C_1 will always *upper* bound T_b^j . We note that for $c \geq 0$,

$$\frac{a}{b} \leq \frac{a+c}{b+c} \Leftrightarrow a \leq b$$

For T_b^j , this holds whenever:

$$\begin{aligned} (v_1 - v_2) \cdot P[\sigma^j] + v_2 \cdot (r \cdot P[\sigma^{i < j}] + P[\sigma^j] + C_0) &\leq (v_1 - v_2) \cdot r \cdot P[\sigma^1] + v_2 \cdot (r \cdot P[\sigma^{i < j}] + P[\sigma^j] + C_0) \\ (v_1 - v_2) \cdot P[\sigma^j] &\leq (v_1 - v_2) \cdot r \cdot P[\sigma^1] \end{aligned}$$

which is always satisfied.

Plugging in for the lower bound on C_1 gives a lower bound on T_b^j of: Note that we can rewrite:

$$\begin{aligned} r \cdot P[\sigma^{i < j}] + P[\sigma^j] + C_0 + P[\sigma^{i > k}] &= r \cdot P[\sigma^{i < j}] + P[\sigma^j] + C_0 + 1 - P[\sigma^{i \leq k}] \\ &= r \cdot P[\sigma^{i < j}] + P[\sigma^j] + C_0 + 1 - P[\sigma^{i < j}] - P[\sigma^j] - P[\sigma^{j < i \leq k}] = (r-1) \cdot P[\sigma^{i < j}] + 1 + C_0 - P[\sigma^{j < i \leq k}] \end{aligned}$$

which means that the lower bound can be written as:

$$\frac{(v_1 - v_2) \cdot P[\sigma^j] + v_2 \cdot (r-1) \cdot P[\sigma^{i < j}] + v_2 + v_2 \cdot (C_0 - P[\sigma^{j < i \leq k}])}{(v_1 - v_2) \cdot r \cdot P[\sigma^1] + v_2 \cdot (r-1) \cdot P[\sigma^{i < j}] + v_2 + v_2 \cdot (C_0 - P[\sigma^{j < i \leq k}])}$$

Similarly, plugging in for the upper bound on C_1 gives an upper bound on T_b^j of:

$$\frac{(v_1 - v_2) \cdot P[\sigma^j] + v_2 \cdot (r \cdot P[\sigma^{i < j}] + P[\sigma^j] + r \cdot P[\sigma^{i > k}] + C_0)}{(v_1 - v_2) \cdot r \cdot P[\sigma^1] + v_2 \cdot (r \cdot P[\sigma^{i < j}] + P[\sigma^j] + r \cdot P[\sigma^{i > k}] + C_0)}$$

Note that we can rewrite:

$$\begin{aligned} r \cdot P[\sigma^{i < j}] + P[\sigma^j] + r \cdot P[\sigma^{i > k}] + C_0 &= r \cdot P[\sigma^{i < j}] + P[\sigma^j] + r \cdot P[\sigma^{i > j}] - r \cdot P[\sigma^{j < i \leq k}] \\ &= r \cdot (1 - P[\sigma^j]) + P[\sigma^j] + C_0 - r \cdot P[\sigma^{j < i \leq k}] = (r-1) \cdot (1 - P[\sigma^j]) + 1 + C_0 - r \cdot P[\sigma^{j < i \leq k}] \end{aligned}$$

Substituting in to the full equation gives:

$$\frac{(v_1 - v_2) \cdot P[\sigma^j] + v_2 \cdot (r-1) \cdot (1 - P[\sigma^j]) + v_2 + v_2 \cdot (C_0 - r \cdot P[\sigma^{i > k}])}{(v_1 - v_2) \cdot r \cdot P[\sigma^1] + v_2 \cdot (r-1) \cdot (1 - P[\sigma^j]) + v_2 + v_2 \cdot (C_0 - r \cdot P[\sigma^{i > k}])}$$

Which gives:

$$\begin{aligned} &\frac{(v_1 - v_2) \cdot P[\sigma^j] + v_2 \cdot (r-1) \cdot P[\sigma^{i < j}] + v_2 + v_2 \cdot (C_0 - P[\sigma^{j < i \leq k}])}{(v_1 - v_2) \cdot r \cdot P[\sigma^1] + v_2 \cdot (r-1) \cdot P[\sigma^{i < j}] + v_2 + v_2 \cdot (C_0 - P[\sigma^{j < i \leq k}])} \leq T_b^j \\ &\leq \frac{(v_1 - v_2) \cdot P[\sigma^j] + v_2 \cdot (r-1) \cdot (1 - P[\sigma^j]) + v_2 + v_2 \cdot (C_0 - r \cdot P[\sigma^{i > k}])}{(v_1 - v_2) \cdot r \cdot P[\sigma^1] + v_2 \cdot (r-1) \cdot (1 - P[\sigma^j]) + v_2 + v_2 \cdot (C_0 - r \cdot P[\sigma^{i > k}])} \end{aligned}$$

as desired.

We can follow very similar analysis for the case where $s_{i < j} = 1, s_j = 0$ (the highest-ranked item is free). However, here upper bounding C can upper or lower bound T_f^j , depending on certain parameters. Specifically, upper bounding C upper bounds T_f^j exactly whenever:

$$(v_1 - v_2) \cdot P[\sigma^1] + v_2 \cdot (P[\sigma^{i < j}] + r \cdot P[\sigma^j] + C_0) \leq (v_1 - v_2) \cdot r \cdot P[\sigma^j] + v_2 \cdot (P[\sigma^{i < j}] + r \cdot P[\sigma^j] + C_0)$$

$$P[\sigma^1] \leq r \cdot P[\sigma^j]$$

This tells us that we must write the bounds as cases, given by:

$$\begin{cases} F_{j,k}^1 \leq T_j^f \leq F_{j,k}^2 & P[\sigma^1] \leq r \cdot P[\sigma^j] \\ F_{j,k}^2 \leq T_j^f \leq F_{j,k}^1 & \text{otherwise} \end{cases}$$

for

$$F_{j,k}^1 = \frac{(v_1 - v_2) \cdot P[\sigma^1] + v_2 \cdot (P[\sigma^{i < j}] + r \cdot P[\sigma^j] + C_0 + P[\sigma^{i > k}])}{(v_1 - v_2) \cdot r \cdot P[\sigma^j] + v_2 \cdot (P[\sigma^{i < j}] + r \cdot P[\sigma^j] + C_0 + P[\sigma^{i > k}])}$$

Note that we can rewrite:

$$\begin{aligned} P[\sigma^{i < j}] + r \cdot P[\sigma^j] + C_0 + P[\sigma^{i > k}] &= P[\sigma^{i < j}] + r \cdot P[\sigma^j] + C_0 + P[\sigma^{i > j}] - P[\sigma^{j < i \leq k}] \\ &= 1 - P[\sigma^j] + r \cdot P[\sigma^j] + C_0 - P[\sigma^{j < i \leq k}] = 1 + (r - 1) \cdot P[\sigma^j] + C_0 - P[\sigma^{j < i \leq k}] \end{aligned}$$

which gives:

$$F_{j,k}^1 = \frac{(v_1 - v_2) \cdot P[\sigma^1] + v_2 \cdot (1 + (r - 1) \cdot P[\sigma^j] + C_0 - P[\sigma^{j < i \leq k}])}{(v_1 - v_2) \cdot r \cdot P[\sigma^j] + v_2 \cdot (1 + (r - 1) \cdot P[\sigma^j] + C_0 - P[\sigma^{j < i \leq k}])}$$

For the other bound, we obtain:

$$F_{j,k}^2 = \frac{(v_1 - v_2) \cdot P[\sigma^1] + v_2 \cdot (P[\sigma^{i < j}] + r \cdot P[\sigma^j] + r \cdot P[\sigma^{i > k}] + C_0)}{(v_1 - v_2) \cdot r \cdot P[\sigma^j] + v_2 \cdot (P[\sigma^{i < j}] + r \cdot P[\sigma^j] + r \cdot P[\sigma^{i > k}] + C_0)}$$

Note that we can rewrite:

$$\begin{aligned} P[\sigma^{i < j}] + r \cdot P[\sigma^j] + r \cdot P[\sigma^{i > k}] + C_0 &= P[\sigma^{i < j}] + r \cdot P[\sigma^j] + r \cdot P[\sigma^{i > j}] + C_0 - r \cdot P[\sigma^{j < i \leq k}] \\ &= 1 - P[\sigma^{i > j}] + r \cdot P[\sigma^{i > j}] + C_0 - r \cdot P[\sigma^{j < i \leq k}] = 1 + (r - 1) \cdot P[\sigma^{i > j}] + C_0 - r \cdot P[\sigma^{j < i \leq k}] \end{aligned}$$

Which gives us:

$$F_{j,k}^2 = \frac{(v_1 - v_2) \cdot P[\sigma^1] + v_2 \cdot (1 + (r - 1) \cdot P[\sigma^{i > j}] + C_0 - r \cdot P[\sigma^{j < i \leq k}])}{(v_1 - v_2) \cdot r \cdot P[\sigma^j] + v_2 \cdot (1 + (r - 1) \cdot P[\sigma^{i > j}] + C_0 - r \cdot P[\sigma^{j < i \leq k}])}$$

□

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [\[Yes\]](#)

Justification: Each section backs up results within the abstract.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: Specific limitations section describes limitations.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory Assumptions and Proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [\[Yes\]](#)

Justification: Proofs given in appendix.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental Result Reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [NA]

Justification:

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [NA]

Justification:

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental Setting/Details

Question: Does the paper specify all the training and test details (e.g., data splits, hyper-parameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [NA]

Justification:

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment Statistical Significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [NA]

Justification:

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.

- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments Compute Resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [NA]

Justification:

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code Of Ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: Limitations of paper discussed, as well as ethical implications of paper.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader Impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [Yes]

Justification: Discussed in conclusion of paper.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to

generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.

- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: No models produced

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [NA]

Justification:

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. New Assets

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [NA]

Justification:

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. Crowdsourcing and Research with Human Subjects

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [NA]

Justification:

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. Institutional Review Board (IRB) Approvals or Equivalent for Research with Human Subjects

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

Justification:

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.