

Robust Online Selection with Uncertain Offer Acceptance

Sebastian Perez-Salazar,^{a,*} Mohit Singh,^b Alejandro Toriello^b

^aDepartment of Computational Applied Mathematics and Operations Research, Rice University, Houston, Texas 77005; ^bH. Milton Stewart School of Industrial and Systems Engineering, Georgia Institute of Technology, Atlanta, Georgia 30332

*Corresponding author

Contact: sperez@rice.edu,  <https://orcid.org/0000-0003-4534-7721> (SP-S); mohit.singh@isye.gatech.edu (MS); atoriello@isye.gatech.edu,

 <https://orcid.org/0000-0002-3147-0764> (AT)

Received: July 10, 2023
Revised: March 30, 2024
Accepted: June 28, 2024
Published Online in Articles in Advance:
August 29, 2024

MSC2020 Subject Classification: Primary:
90C39, 90C40, 68W25, 68W27

<https://doi.org/10.1287/moor.2023.0210>

Copyright: © 2024 INFORMS

Abstract. Online advertising has motivated interest in online selection problems. Displaying ads to the right users benefits both the platform (e.g., via pay-per-click) and the advertisers (by increasing their reach). In practice, not all users click on displayed ads, while the platform’s algorithm may miss the users most disposed to do so. This mismatch decreases the platform’s revenue and the advertiser’s chances to reach the right customers. With this motivation, we propose a secretary problem where a candidate may or may not accept an offer according to a known probability p . Because we do not know the top candidate willing to accept an offer, the goal is to maximize a robust objective defined as the minimum over integers k of the probability of choosing one of the top k candidates, given that one of these candidates will accept an offer. Using Markov decision process theory, we derive a linear program for this max-min objective whose solution encodes an optimal policy. The derivation may be of independent interest, as it is generalizable and can be used to obtain linear programs for many online selection models. We further relax this linear program into an infinite counterpart, which we use to provide bounds for the objective and closed-form policies. For $p \geq p^* \approx 0.6$, an optimal policy is a simple threshold rule that observes the first $p^{1/(1-p)}$ fraction of candidates and subsequently makes offers to the best candidate observed so far.

Funding: Financial support from the U.S. National Science Foundation [Grants CCF-2106444, CCF-1910423, and CMMI 1552479] is gratefully acknowledged.

Keywords: optimization • secretary problem • Markov decision processes

1. Introduction

The growth of online platforms has spurred renewed interest in online selection problems, auctions, and stopping problems (Alaei et al. [3], Devanur and Hayes [24], Edelman et al. [29], Lucier [55], Mehta et al. [58]). Online advertising has particularly benefited from developments in these areas. As an example, in 2005, Google reported about \$6 billion in revenue from advertising, roughly 98% of the company’s total revenue at that time; in 2020, Google’s revenue from advertising grew to almost \$147 billion. Targeting users is crucial for the success of online advertising. Studies suggest that targeted campaigns can double *click-through rates* in ads (Farahat and Bailey [31]), despite the fact that internet users have acquired skills to navigate the web while ignoring ads (Cho and Cheon [18], Drèze and Hussherr [26]). Therefore, it is natural to expect that not every displayed ad will be clicked on by a user, even if the user likes the product on the ad, whereas the platform and advertiser’s revenue depend on this event (Pujol et al. [61]). An ignored ad misses the opportunity of being displayed to another user willing to click on it and decreases the return on investment for the advertiser, especially in cases where the platform uses methods like pay-for-impression to charge the advertisers. At the same time, the ignored ad uses the space of another, possibly more suitable ad for that user. In this work, we take the perspective of a single ad, and we aim to understand the right time to begin displaying the ad to users as a function of the ad’s probability of being clicked.

We model the interaction between the platform and the users using a general online selection problem. We refer to it as the *secretary problem with uncertain acceptance* (SP-UA for short). Using the terminology of *candidate* and *decision maker*, the general interaction is as follows:

1. Similar to other secretary problems, a finite sequence of candidates of known length arrives online, in a random order. In our motivating application, candidates represent platform users.
2. Upon an arrival, the decision maker (DM) is able to assess the quality of a candidate compared with previously observed candidates and has to irrevocably decide whether to extend an offer to the candidate or move on to

the next candidate. This captures the online dilemma the platform faces: the decision of displaying an ad to a user is based solely on information obtained up to this point.

3. When the DM extends an offer, the candidate accepts with a known probability $p \in (0, 1]$, in which case the process ends, or turns down the offer, in which case the DM moves on to the next candidate. This models the users, who can click on the ad or ignore it.

4. The process continues until either a candidate accepts an offer or the DM has no more candidates to assess.

A DM that knows in advance that at least one of the top k candidates is willing to accept the offer would like to maximize the probability of making an offer to one of these candidates. In reality, the DM does not know k ; hence, the best she can do is maximize the minimum of all these scenario-based probabilities. We call the minimum of these scenario-based probabilities the *robust ratio* and our max-min objective the *optimal robust ratio* (see Subsection 1.2 for a formal description). Suppose that the DM implements a policy that guarantees a robust ratio $\gamma \in (0, 1]$. This implies the DM will succeed with probability at least γ in obtaining a top k candidate, in any scenario where a top k candidate is willing to accept the DM's offer. This is an *ex ante* guarantee when the DM knows the odds for each possible scenario, but the policy is independent of k and offers the same guarantee for any of these scenarios. Moreover, if the DM can assign a numerical valuation to the candidates, a policy with robust ratio γ can guarantee a factor at least γ of the optimal offline value. Tamaki [68] also studies the **SP-UA** and considers the objective of maximizing the probability of selecting the best candidate willing to accept the offer. Applying Tamaki's policy to value settings can also guarantee an approximation factor of the optimal offline cost; however, the policy with the optimal robust ratio attains the largest approximation factor of the optimal offline value among rank-based policies (see Proposition 1).

SP-UA captures the inherent unpredictability in online selection, as other secretary problems do, but also the uncertainty introduced by the possibility of candidates turning down offers. **SP-UA** is broadly applicable; the following are additional concrete examples.

1.1. Data-Driven Selection Problems

When selling an item in an auction, buyers' valuations are typically unknown beforehand. Assuming valuations follow a common distribution, the aim is to sell the item at the highest price possible; learning information about the distribution is crucial for this purpose. In particular auction settings, the auctioneer may be able to sequentially observe the valuations of potential buyers and can decide in an online manner whether to sell the item or continue observing valuations. Specifically, the auctioneer decides to consider the valuation of a customer with probability p , and, otherwise, the auctioneer moves on to see the next buyer's valuation. The auctioneer's actions can be interpreted as an *exploration-exploitation* process, which is often found in bandit problems and online learning (Cesa-Bianchi and Lugosi [15], Freund and Schapire [35], Hazan [42]). This setting is also closely related to data-driven online selection and the prophet inequality problem (Campbell and Samuels [14], Kaplan et al. [47], Kertz [48]); some of our results also apply in these models (see Section 6).

1.2. Human Resource Management

As its name suggests, the original motivation for the secretary problem is in hiring for a job vacancy. Screening resumes can be a time-consuming task that shifts resources away from the day-to-day job in Human Resources. Since the advent of the internet, several elements of the hiring process can be partially or completely automated; for example, multiple vendors offer automated resume screening (Salem and Gupta [64]), and machine learning algorithms can score and rank job applicants according to different criteria. Of course, a highly ranked applicant may nevertheless turn down a job offer. Although we consider the rank of a candidate as an absolute metric of their capacities, in reality, resume screening may suffer from different sources of bias (Salem and Gupta [63]), but addressing this goes beyond our scope. See also Smith [66], Tamaki [68], and Vanderbei [72] for classical treatments. Similar applications include apartment hunting (Bruss [12], Cowan and Zabczyk [23], Presman and Sonin [60]), among others.

1.3. Our Contributions

(1) We propose a framework and a robust metric to understand the interaction between a DM and competing candidates, when candidates can reject the DM's offer. (2) We state a linear program (LP) that computes the optimal robust ratio and the best strategy. We provide a general methodology to derive our LP, and this technique is generalizable to other online selection problems. (3) We provide bounds for the optimal robust ratio as a function of the probability of acceptance $p \in (0, 1]$. (4) We present a family of policies based on simple threshold rules; in particular, for $p \geq p^* \approx 0.594$, the optimal strategy is a simple threshold rule that skips the first $p^{1/(1-p)}$ fraction of candidates and then makes offers to the best candidate observed so far. We remark that as $p \rightarrow 1$, we recover the guarantees of

the standard secretary problem and its optimal threshold strategy. (5) Finally, for the setting where candidates also have nonnegative numerical values, we show that our solution is the optimal approximation among rank-based algorithms of the optimal offline value, where the benchmark knows the top candidate willing to accept the offer. The optimal approximation factor equals the optimal robust ratio.

1.4. Problem Formulation

A problem instance is given by a fixed probability $p \in (0, 1]$ and the number of candidates n . These are ranked by a total order, $1 \prec 2 \prec \dots \prec n$, with 1 being the best or highest-ranked candidate. The candidate sequence is given by a random permutation $\pi = (R_1, \dots, R_n)$ of $[n] \doteq \{1, 2, \dots, n\}$, where any permutation is equally likely. At time t , the DM observes the partial rank $r_t \in [t]$ of the t -th candidate in the sequence compared with the previous $t - 1$ candidates. The DM either makes an offer to the t -th candidate or moves on to the next candidate, without being able to make an offer to the t -th candidate ever again. If the t -th candidate receives an offer from the DM, she accepts the offer with probability p , in which case the process ends. Otherwise, if the candidate refuses the offer (with probability $1 - p$), the DM moves on to the next candidate and repeats the process until she has exhausted the sequence. A candidate with rank in $[k]$ is said to be a *top k candidate*. The goal is a policy that maximizes the probability of extending an offer to a highly ranked candidate that *will accept the offer*. However, because the DM does not know which candidate will accept the offer, the DM would like to be robust against any possible scenario. To measure the quality of a policy $\mathcal{P}$, we use the *robust ratio*

$$\gamma_{\mathcal{P}} = \gamma_{\mathcal{P}}(p) = \min_{k=1, \dots, n} \frac{\mathbf{P}(\mathcal{P} \text{ selects a top } k \text{ candidate, candidate accepts offer})}{\mathbf{P}(\text{At least one of the top } k \text{ candidates accepts offer})}. \quad (1)$$

The k -th term in the minimization operator, $\gamma_{\mathcal{P},k}(p)$, is the probability that policy $\mathcal{P}$ successfully selects a top k candidate, *given* that some top k candidate will accept the offer. Then, the *robust ratio* $\gamma_{\mathcal{P}} = \min_{k=1, \dots, n} \gamma_{\mathcal{P},k}(p)$ captures the situation where policy $\mathcal{P}$ has the worst possible performance over all such scenarios. When every candidate accepts an offer with certainty, $p = 1$, the robust ratio $\gamma_{\mathcal{P}}$ equals the probability of selecting the highest-ranked candidate, thus, we recover the standard secretary problem and $\gamma_{\mathcal{P}}(1) \approx 1/e$ for the optimal policy $\mathcal{P}$. The goal is to find a policy that maximizes this robust ratio, $\gamma_n^* \doteq \sup_{\mathcal{P}} \gamma_{\mathcal{P}}$. We say the policy $\mathcal{P}$ is *γ -robust* if $\gamma \leq \gamma_{\mathcal{P}}$.

1.4.1. The Robust Ratio and Related Objectives. The **SP-UA** has been studied before under different objectives. Smith [66] studied the **SP-UA** with the objective of maximizing the probability of selecting the top candidate and having that candidate accept the offer. This is the unconditional version of $\gamma_{\mathcal{P},k}$ for $k = 1$; however, the top candidate may not accept the offer, and the objective does not plan for this contingency. This is particularly inadequate when p is small, as in many of our motivating applications.

Tamaki [68] instead studied the **SP-UA** with the objective of maximizing the probability of choosing the top candidate willing to accept the offer. Despite being more realistic than Smith [66], this objective is often overly selective and may not make an offer, hoping to encounter a better candidate in the future. Our objective overcomes this selectiveness and makes an offer to a candidate as long as their rank is high compared with other candidates willing to accept the offer. A further distinction between Tamaki's objective and the robust ratio emerges when values are assigned to the candidates. In this case, a value-driven DM would like to maximize the value obtained from a candidate that accepts the offer. The robust ratio turns out to be the optimal approximation ratio of any rank-based algorithm in this setting (see Proposition 1). In Section 8, we provide extensive numerical experiments for the value version of the problem. Our policy consistently yields better results for small acceptance probabilities, $p < 0.2$, demonstrating its effectiveness compared with Tamaki [68].

1.5. Our Technical Contributions

Recent works have studied secretary models using linear programming methods (Buchbinder et al. [13], Chan et al. [16], Correa et al. [21], Dütting et al. [27]). We also give an LP formulation that computes the best robust ratio and the optimal policy for our model. Whereas these recent approaches derive an LP formulation using ad hoc arguments, our first contribution is to provide a general framework to obtain LP formulations that give optimal bounds and policies for different variants of the secretary problem. The framework is based on Markov decision process (MDP) theory (Altman [4], Puterman [62]). This is surprising because early literature on secretary problem used MDP techniques—for example, Dynkin [28] and Lindley [54]—though typically not LP formulations. In that sense, our results connect the early algorithms based on MDP methods with the recent literature based on LP methods. Specifically, we provide a mechanical way to obtain an LP using a simple MDP formulation (Section 4). Using this framework, we present a structural result that completely characterizes the space of policies for the **SP-UA**:

Theorem 1. Any policy $\mathcal{P}$ for the SP-UA can be represented as a vector in the set

$$\text{POL} = \left\{ (\mathbf{x}, \mathbf{y}) \geq 0 : x_{t,s} + y_{t,s} = \frac{1}{t} \sum_{\sigma=1}^{t-1} (y_{t-1,\sigma} + (1-p)x_{t-1,\sigma}), \forall t > 1, s \in [t], x_{1,1} + y_{1,1} = 1 \right\}.$$

Conversely, any vector $(\mathbf{x}, \mathbf{y}) \in \text{POL}$ represents a policy $\mathcal{P}$. The policy $\mathcal{P}$ makes an offer to the first candidate with probability $x_{1,1}$ and to the t -th candidate with probability $tx_{t,s}/(\sum_{\sigma=1}^{t-1} y_{t-1,\sigma} + (1-p)x_{t-1,\sigma})$ if the t -th candidate has partial rank $r_t = s$.

The variables $x_{t,s}$ represent the probability of reaching candidate t and making an offer to that candidate when that candidate has partial rank $s \in [t]$. Likewise, variables $y_{t,s}$ represent the probability of reaching candidate t and not making an offer when this candidate's partial rank is $s \in [t]$. We note that although the use of LP formulations in MDP is a somewhat standard technique—see, for example, Puterman [62]—the recent literature in secretary problems and related online selection models does not appear to make an explicit connection between LPs used in analysis and the underlying MDP formulation.

Problems solved via MDP can typically be formulated as reward models, where each action taken by the DM generates some immediate reward. Objectives in classical secretary problems fit in this framework, as the reward (e.g., the probability of selecting the top candidate) depends only on the current state (the number t of observed candidates so far and the current candidate's partial rank $r_t = s$), and on the DM's action (make an offer or not); see Section 4.1 for an example. Our robust objective, however, cannot be easily written as a reward depending only on $r_t = s$. Thus, we split the analysis into two stages. In the first stage, we deal with the space of policies and formulate an MDP for our model with a generic utility function. The feasible region of this MDP's LP formulation corresponds to POL and is independent of the utility function chosen; therefore, it characterizes all possible policies for the SP-UA. In the second stage, we use the structural result in Theorem 1 to obtain a linear program that finds the largest robust ratio.

Theorem 2. The best robust ratio γ_n^* for the SP-UA equals the optimal value of the linear program

$$\begin{aligned} \max_{\mathbf{x} \geq 0} \quad & \gamma \\ \text{s.t.} \quad & \\ (LP)_{n,p} \quad & x_{t,s} \leq \frac{1}{t} \left(1 - p \sum_{\tau=1}^{t-1} \sum_{\sigma=1}^{\tau} x_{\tau,\sigma} \right) \quad \forall t \in [n], s \in [t] \\ & \gamma \leq \frac{p}{1 - (1-p)^k} \sum_{t=1}^n \sum_{s=1}^t x_{t,s} \mathbf{P}(R_t \leq k | r_t = s) \quad \forall k \in [n], \end{aligned}$$

where $\mathbf{P}(R_t \leq k | r_t = s) = \sum_{i=s}^{k \wedge (n-t+s)} \binom{i-1}{s-1} \binom{n-i}{t-s} / \binom{n}{t}$ is the probability the t -th candidate is ranked in the top k given that her partial rank is s .

Moreover, given an optimal solution $(\mathbf{x}^*, \gamma_n^*)$ of $(LP)_{n,p}$, the (randomized) policy $\mathcal{P}^*$ that at state (t, s) makes an offer with probability $tx_{t,s}^*/(1 - p \sum_{\tau=1}^{t-1} \sum_{\sigma=1}^{\tau} x_{\tau,\sigma}^*)$ is γ_n^* -robust.

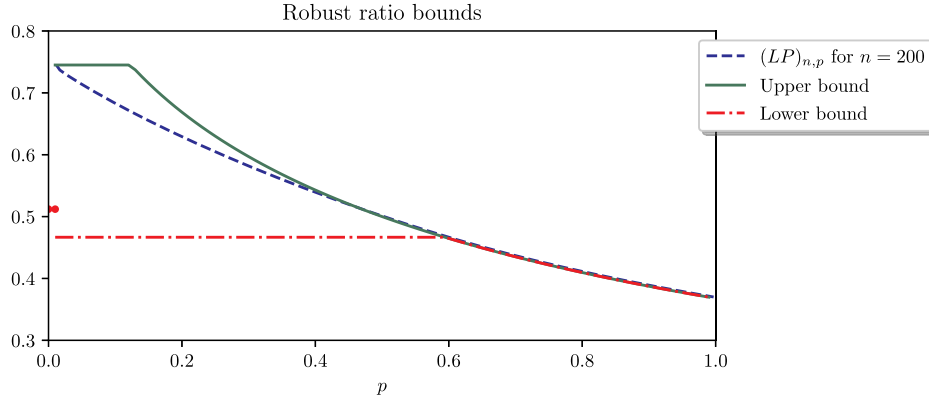
We show that γ_p can be written as the minimum of n linear functions on the $\mathbf{x}$ variables in POL, where these variables correspond to a policy's probability of making an offer in a given state. Thus, our problem can be written as the maximum of a concave piecewise linear function over POL, which we linearize with the variable γ . By projecting the feasible region onto the $(\mathbf{x}, \gamma)$ variables, we obtain $(LP)_{n,p}$.

As a by-product of our analysis via MDP, we show that γ_n^* is nonincreasing in n for fixed $p \in (0, 1]$ (Lemma 1), and, thus, $\lim_{n \rightarrow \infty} \gamma_n^* = \gamma_\infty^*$ exists. We show that this limit corresponds to the optimal value of an infinite version of $(LP)_{n,p}$ from Theorem 2, where n tends to infinity and we replace sums at time t with integrals (see Section 5). This allows us to show upper and lower bounds for γ_n^* by analyzing γ_∞^* . Our first result in this vein gives upper bounds on γ_∞^* .

Theorem 3. For any $p \in (0, 1]$, $\gamma_\infty^*(p) \leq \min\{p^{p/(1-p)}, 1/\beta\}$, where $1/\beta \approx 0.745$ and β is the (unique) solution of the equation $\int_0^1 (y(1 - \log y) + \beta - 1)^{-1} dy = 1$.

To show $\gamma_\infty^* \leq p^{p/(1-p)}$, we relax all constraints in the robust ratio except $k=1$. This becomes the problem of maximizing the probability of hiring the top candidate, which has a known asymptotic solution of $p^{1/(1-p)}$ (Smith [66]). For $\gamma_\infty^*(p) \leq 1/\beta$, we show that any γ -robust ordinal algorithm can be used to construct an algorithm

Figure 1. (Color online) Bounds for γ_∞^* as a function of p . The solid line represents the theoretical upper bound given in Theorem 3. The dashed-dotted line corresponds to the theoretical lower bound given in Theorem 4; for p close to 0, the guarantee rises to 0.51. In dashed line we present numerical results by solving $(LP)_{n,p}$ for $n = 200$ candidates.



for independent and identically distributed (i.i.d.) prophet inequality problems with a multiplicative loss of $(1 + o(1))\gamma$ and an additional $o(1)$ additive error. Using a slight modification of the impossibility result by Hill and Kertz [43] for the i.i.d. prophet inequality, we conclude that γ_∞^* cannot be larger than $1/\beta$.

By constructing solutions of the infinite LP, we can provide lower bounds for γ_n^* . For $1/k \geq p > 1/(k+1)$ with integer k , the policy that skips the first $1/e$ fraction of candidates and then makes an offer to any top k candidate afterward obtains a robust ratio of at least $1/e$. The following result gives improved bounds for $\gamma_\infty^*(p)$.

Theorem 4. Let $p^* \approx 0.594$ be the solution of $p^{(2-p)/(1-p)} = (1-p)^2$. There is a solution of the infinite LP for $p \geq p^*$ that guarantees $\gamma_n^* \geq \gamma_\infty^*(p) = p^{p/(1-p)}$. For $p \leq p^*$, we have $\gamma_\infty^*(p) \geq (p^*)^{p^*/(1-p^*)} \approx 0.466$. Moreover, for $p \rightarrow 0$, we obtain $\gamma_\infty^*(p) \geq 0.51$.

To prove this result, we use the following general procedure to construct feasible solutions for the infinite LP. For any numbers $0 < t_1 \leq t_2 \leq \dots \leq t_k \leq \dots \leq 1$, there is a policy that makes offers to any candidate with partial rank $r_t \in [k]$ when a fraction t_k of the total number of candidates has been observed (Proposition 2). For $p \geq p^*$, the policy corresponding to $t_1 = p^{1/(1-p)}$ and $t_2 = t_3 = \dots = 1$ has a robust ratio of at least $p^{p/(1-p)}$. For $p \leq p^*$, we show how to transform the solution for p^* into a solution for p with an objective value at least as good as the value $\gamma_\infty^*(p^*) = (p^*)^{p^*/(1-p^*)}$. For values of p close to 0, we construct a feasible solution of the infinite LP that guarantees $\gamma_\infty^*(p) \geq 0.51$.

Figure 1 depicts the various theoretical bounds we obtain. For reference, we also include numerical results for γ_n^* computed by solving $(LP)_{n,p}$ in Theorem 2 for $n = 200$ and with p ranging from $p = 10^{-2}$ to $p = 1$, with increments of 10^{-3} . Given that γ_n^* is nonincreasing in n , the numerical values obtained by solving $(LP)_{n,p}$ also provide an upper bound over γ_∞^* .

We follow this introduction with a brief literature review. In Section 3, we present preliminaries, including MDP notation and an alternative characterization of the robust ratio in terms of utility functions. In Section 4, we present the MDP framework and use it to prove Theorems 1 and 2. In Section 5, we introduce the infinite relaxation of (LP), then prove Theorem 3 in Section 6. In Section 7, we prove Theorem 4. In Section 8, we present a numerical comparison between the policies obtained by solving $(LP)_{n,p}$ and other benchmarks policies. We conclude in Section 9, and Appendices A–E include proofs and analysis omitted from the main article.

2. Related Work

2.1. Online Advertising and Online Selection

Online advertising has been extensively studied from the viewpoint of two-sided markets: advertisers and platform. There is extensive work in auction mechanisms to select ads (e.g. second-price auctions, the VCG mechanism, etc.), and the payment systems between platforms and advertisers (pay-per-click, pay-for-impression, etc.) (Devanur and Kakade [25], Edelman et al. [29], Fridgeirsdottir and Najafi-Asadolahi [36]); see also Choi et al. [19] for a review. On the other hand, works relating the platform, advertisers, and web users have been studied mainly from a learning perspective, to improve ad targeting (Devanur and Kakade [25], Farahat and Bailey [31], Hlynka and Sheahan [42]). In this work, we also aim to display an ad to a potentially interested user. Multiple online selection

problems have been proposed to display ads in online platforms—for example, packing models (Babaioff et al. [6], Korula and Pál [52]), secretary problems and auctions (Babaioff et al. [7]), prophet models (Alaei et al. [3]), and online models with “buyback” (Babaioff et al. [5]). In our setting, we add the possibility that a user ignores the ad; see, for example, Cho and Cheon [18] and Drèze and Hussherr [26]. Failure to click on ads has been considered in full-information models (Goyal and Udwan [38]); however, our setting considers only partial information, where the rank of an incoming customer can only be assessed relative to previously observed customers—a typical occurrence in many online applications. Our model is also disaggregated and looks at each ad individually. Our goal is to understand the right time to display an ad/make offers via the **SP-UA** and the robust ratio for each individual ad.

2.2. Online Algorithms and Arrival Models

Online algorithms have been extensively studied for adversarial arrivals (Borodin and El-Yaniv [10]). This *worst-case* viewpoint gives robust algorithms against any input sequence, which tend to be conservative. Conversely, some models assume distributional information about the inputs (Kertz [48], Kleywegt and Papastavrou [51], Lucier [55]). The random order model lies in between these two viewpoints, and perhaps the most studied example is the secretary problem (Dynkin [28], Gilbert and Mosteller [37], Lindley [54]). Random order models have also been applied in Adword problems (Devanur and Hayes [24]), online LPs (Agrawal et al. [2]), and online knapsacks (Babaioff et al. [6], Kesselheim et al. [49]), among others.

2.3. Secretary Problems

Martin Gardner popularized the secretary problem in his 1960 *Mathematical Games* column; for a historical review, see Ferguson et al. [33] and also the classical survey by Freeman [34]. For the classical secretary problem, the optimal strategy that observes the first n/e candidates and thereafter selects the best candidate was computed by Lindley [54] and Gilbert and Mosteller [37]. The model has been extensively studied in ordinal/ranked-based settings (Buchbinder et al. [13], Gusein-Zade [39], Lindley [54], Vanderbei [71]), as well as cardinal/value-based settings (Bateni et al. [8], Kleinberg [50]).

A large body of work has been dedicated to augment the secretary problem. Variations include cardinality constraints (Buchbinder et al. [13], Kleinberg [50], Vanderbei [71]), knapsack constraints (Babaioff et al. [6]), and matroid constraints (Feldman et al. [32], Lachish [53]), and Soto [67]. Model variants also incorporate different arrival processes, such as Markov chains (Hlynka and Sheahan [44]) and more general processes (Dütting et al. [27]). Closer to our problem are the data-driven variations of the model (Correa et al. [21], Correa et al. [22], Kaplan et al. [47]), where samples from the arriving candidates are provided to the decision maker. Our model can be interpreted as an online version of sampling, where a candidate rejecting the decision maker’s offer is tantamount to a sample. This also bears similarity to the exploration-exploitation paradigm often found in online learning and bandit problems (Cesa-Bianchi and Lugosi [15], Freund and Schapire [35], Hazan [42]).

2.4. Uncertain Availability in Secretary Problems

The **SP-UA** is studied by Smith [66] with the goal of selecting the top candidate— $k = 1$ in (1)—who gives an asymptotic probability of success of $p^{1/(1-p)}$. If the top candidate rejects the offer, this leads to zero value, which is perhaps excessively pessimistic in scenarios where other competent candidates could accept. Tamaki [68] considers maximizing the probability of selecting the top candidate *among* the candidates that will accept the offer. Although more realistic, this objective still gives zero value when the top candidate that accepts is missed because she arrives early in the sequence. In our approach, we make offers to candidates, even if we have already missed the top candidate that accepts the offer; this is also appealing in utility/value-based settings (see Proposition 1). We also further the understanding of the model and our objective by presenting closed-form solutions and bounds. See also Bruss [12], Presman and Sonin [60], and Cowan and Zabczyk [23].

2.5. Linear Programs in Online Selection

Linear programming has been used extensively in online selection (Agrawal et al. [2], Beyhaghi et al. [9], Epstein and Ma [30], Kesselheim et al. [49]). Typically, the LP is used as a structured bound over a benchmark that the algorithm designer can compare with. Our approach is different, as we provide an exact formulation of our robust objective. In secretary problems, early work used mostly MDPs (Lindley [54], Smith [66], Tamaki [68]), whereas LP formulations were recently introduced by Buchbinder et al. [13]; subsequently, multiple formulations have been used to solve variants of the secretary problem (Chan et al. [16], Correa et al. [21], Dütting et al. [27]). We extend this line of work and use an MDP to derive the exact polyhedron that encodes policies for the **SP-UA**; this helps explain why some LP formulations in secretary problems are exact (see Subsection 4.1). Jiang

et al. [46], Jiang et al. [45], and Perez-Salazar et al. [59] are closer to our work, as these studies characterize the optimal policies for prophet inequalities. Further connections between MDP and LPs in related models have been studied mostly in approximate regimes (Adelman [1], Torricco and Toriello [69], Torricco et al. [70]) and particularly in constrained MDPs (Altman [4], Haskell and Jain [40], Haskell and Jain [41]). To the best of the authors' knowledge, there was previously no explicit connection between the MDP formulation and the exact LP formulation in secretary problems.

3. Preliminaries

To discuss our model, we use standard MDP notation for secretary problems (Dynkin [28], Freeman [34], Lindley [54]). An instance is characterized by the number of candidates n and the probability $p \in (0, 1]$ that an offer is accepted. For $t \in [n]$ and $s \in [t]$, a *state* of the system is a pair (t, s) indicating that the candidate currently being evaluated is the t -th and the corresponding partial rank is $r_t = s$. To simplify notation, we add the states $(n+1, s)$, $s \in [n+1]$, and the state Θ as absorbing states where no decisions can be made. For $t < n$, transitions from a state (t, s) to a state $(t+1, \sigma)$ are determined by the random permutation $\pi = (R_1, \dots, R_n)$. We denote by $S_t \in \{(t, s)\}_{s \in [t]}$ the random variable indicating the state in the t -th stage. A simple calculation shows

$$\mathbf{P}(S_{t+1} = (t+1, \sigma) | S_t = (t, s)) = \mathbf{P}(r_{t+1} = \sigma | r_t = s) = \mathbf{P}(S_{t+1} = (t+1, \sigma)) = 1/(t+1),$$

for $t < n$, $s \in [t]$ and $\sigma \in [t+1]$. In other words, partial ranks at each stage are independent. For notational convenience, we assume the equality also holds for $t = n$. Let $\mathbf{A} = \{\text{offer}, \text{pass}\}$ be the set of *actions*. For $t \in [n]$, given a state (t, s) and an action $A_t = a \in \mathbf{A}$, the system transitions to a state S_{t+1} with the following probabilities:

$$P_{((t,s),a),(\tau,\sigma)} = \mathbf{P}(S_{t+1} = (\tau, \sigma) | S_t = (t, s), A_t = a) = \begin{cases} \frac{1-p}{t+1} & a = \text{offer}, \tau = t+1, \sigma \in [\tau] \\ p & a = \text{offer}, (\tau, \sigma) = \Theta \\ \frac{1}{t+1} & a = \text{pass}, \tau = t+1, \sigma \in [\tau]. \end{cases}$$

The randomness is over the permutation π and the random outcome of the t -th candidate's decision. We utilize states $(n+1, \sigma)$ as end states and the state Θ as the state indicating that an offer is accepted from the state S_t . A *policy* $\mathcal{P} : \{(t, s) : t \in [n], s \in [t]\} \rightarrow \mathbf{A}$ is a function that observes a state (t, s) and decides to extend an offer ($\mathcal{P}(t, s) = \text{offer}$) or move to the next candidate ($\mathcal{P}(t, s) = \text{pass}$). The policy specifies the actions of a decision maker at any point in time. The *initial state* is $S_1 = (1, 1)$, and the computation (of a policy) is a sequence of state and actions $(1, 1), a_1, (2, s_2), a_2, (3, s_3), \dots$ where the states transitions according to $P_{((t,s),a),(\tau,\sigma)}$ and $a_t = \mathcal{P}(t, s_t)$. Note that the computation always ends in a state $(n+1, \sigma)$ for some σ or the state Θ , either because the policy was able to go through all candidates or because some candidate t accepted an offer.

We say that a policy reaches stage t or reaches the t -th stage if the computation of a policy contains a state $s_t = (t, s)$ for some $s \in [t]$. We also refer to stages as *times*.

A randomized policy is a function $\mathcal{P} : \{(t, s) : t \in [n], s \in [t]\} \rightarrow \Delta_{\mathbf{A}}$, where $\Delta_{\mathbf{A}} = \{(q, 1-q) : q \in [0, 1]\}$ is the probability simplex over $\mathbf{A} = \{\text{offer}, \text{pass}\}$ and $\mathcal{P}(s_t) = (q_t, 1-q_t)$ means that $\mathcal{P}$ selects the **offer** action with probability q_t and otherwise selects **pass**.

We could also define policies that remember previously visited states and at state (t, s_t) make decisions based on the *history*, $(1, s_1), \dots, (t, s_t)$. However, MDP theory guarantees that it suffices to consider Markovian policies, which make decisions based only on (t, s_t) ; see Puterman [62].

We say that a policy $\mathcal{P}$ *collects* a candidate with rank k if the policy extends an offer to a candidate that has rank k and the candidate accepts the offer. Thus, our objective is to find a policy that solves

$$\begin{aligned} \gamma_n^* &= \max_{\mathcal{P}} \min_{k \in [n]} \frac{\mathbf{P}(\mathcal{P} \text{ collects a candidate with rank } \leq k)}{1 - (1-p)^k} \\ &= \max_{\mathcal{P}} \min_{k \in [n]} \mathbf{P}(\mathcal{P} \text{ collects a top } k \text{ candidate} | \text{a top } k \text{ candidate accepts}). \end{aligned}$$

The following result is an alternative characterization of γ_n^* based on utility functions. We use this result to relate **SP-UA** to the i.i.d. prophet inequality problem; the proof appears in Appendix A. Consider a nonzero utility function $U : [n] \rightarrow \mathbf{R}_+$ with $U_1 \geq U_2 \geq \dots \geq U_n \geq 0$ and any rank-based algorithm ALG for the **SP-UA**—that is, ALG only makes decisions based on the relative ranking of the values observed. In the value setting, if ALG collects a candidate with overall rank i , it obtains value U_i . We denote by $U(ALG)$ the value collected by such an algorithm.

Proposition 1. Let ALG be a γ -robust algorithm for **SP-UA**. For any $U : [n] \rightarrow \mathbf{R}_+$, we have $\mathbf{E}[U(ALG)] \geq \gamma \mathbf{E}[U(OPT)]$, where OPT is the maximum value obtained from candidates that accept. Moreover,

$$\gamma_n^* = \max_{ALG} \min \left\{ \frac{\mathbf{E}[U(ALG)]}{\mathbf{E}[U(OPT)]} : U : [n] \rightarrow \mathbf{R}_+, U \neq 0, U_1 \geq U_2 \geq \dots \geq U_n \geq 0 \right\}.$$

4. The LP Formulation

In this section, we present the proofs of Theorems 1 and 2. Our framework is based on MDP and can be used to derive similar LPs in the literature—for example, Buchbinder et al. [13], Chan et al. [16], and Correa et al. [21]. As a by-product, we also show that γ_n^* is a nonincreasing sequence in n (Lemma 1). For ease of explanation, we first present the framework for the classical secretary problem, then we sketch the approach for our model. Technical details are deferred to Appendix B.

4.1. Warm-Up: MDP to LP in the Classical Secretary Problem

We next show how to derive an LP for the classical secretary problem (Buchbinder et al. [13]) using an MDP framework. In this model, the goal is to maximize the probability of choosing the top candidate, and there is no offer uncertainty.

Theorem 5 (Buchbinder et al. [13]). *The maximum probability of choosing the top-ranked candidate in the classical secretary problem is given by*

$$\max \left\{ \sum_{t=1}^n \frac{t}{n} x_t : x_t \leq \frac{1}{t} \left(1 - \sum_{\tau=1}^{t-1} x_\tau \right), \forall t \in [n], x \geq 0 \right\}.$$

We show this as follows:

1. First, we formulate the secretary problem as a Markov decision process, where we aim to find the highest-ranked candidate. Let $v_{(t,s)}^*$ be the maximum probability of selecting the highest-ranked candidate in $t+1, \dots, n$, given that the current state is (t, s) . We define $v_{(n+1,s)}^* = 0$ for any s . The value v^* is called the *value function*, and it can be computed via the optimality equations (Puterman [62])

$$v_{(t,s)}^* = \max \left\{ \mathbf{P}(R_t = 1 | r_t = s), \frac{1}{t+1} \sum_{\sigma=1}^t v_{(t+1,\sigma)}^* \right\}. \quad (2)$$

The first term in the max operator corresponds to the expected value when the **offer** action is chosen in state (t, s) . The second corresponds to the expected value in stage $t+1$ when we decide to **pass** in (t, s) . Note that $\mathbf{P}(R_t = 1 | r_t = s) = t/n$ if $s = 1$ and $\mathbf{P}(R_t = 1 | r_t = s) = 0$ otherwise. The optimality Equations (2) can be solved via backward recursion, and $v_{(1,1)}^* \approx 1/e$ (for large n). An optimal policy can be obtained from the optimality equations by choosing at each state an action that attains the maximum, breaking ties arbitrarily.

2. Using a standard argument (Manne [56]), it follows that $v^* = (v_{(t,s)}^*)_{t,s}$ is an optimal solution of the linear program (D):

$$\begin{aligned} \min_{v \geq 0} \quad & v_{(1,1)} \\ (D) \quad & v_{(t,s)} \geq \mathbf{P}(R_t = 1 | r_t = s) \quad \forall t \leq n, \forall s \leq t, \end{aligned} \quad (3)$$

$$v_{(t,s)} \geq \frac{1}{1+t} \sum_{\sigma=1}^{t+1} v_{(t+1,\sigma)} \quad \forall t \leq n, s \leq t. \quad (4)$$

3. Taking the dual of (D), we obtain (P):

$$\begin{aligned} \max_{x, y \geq 0} \quad & \sum_{t=1}^n \frac{t}{n} x_{t,1} \\ (P) \quad & x_{1,1} + y_{1,1} \leq 1, \end{aligned} \quad (5)$$

$$x_{t,s} + y_{t,s} \leq \frac{1}{t} \sum_{\sigma=1}^{t-1} y_{t-1,\sigma} \quad \forall t \leq n, s \leq t. \quad (6)$$

Variables $x_{t,s}$ are associated with Constraints (3) and $y_{t,s}$ with Constraints (4). Take any solution $(\mathbf{x}, \mathbf{y})$ of the problem (P) and note that the objective does not depend on $\mathbf{y}$. Incrementing $\mathbf{y}$ to tighten all constraints does not alter the feasibility of the solution, and the objective does not change; thus, we can assume that all constraints are tight in (P) . Here, $x_{t,s}$ is the probability that a policy (determined by $\mathbf{x}, \mathbf{y}$) reaches the state (t, s) and makes an offer, whereas $y_{t,s}$ is the probability that the same policy reaches state (t, s) , but decides to not to make an offer.

4. Finally, projecting the feasible region of (P) onto the variables $(x_{t,1})_{t \in [n]}$ —for example, via Fourier-Motzkin elimination (see Schrijver [65] for a definition)—gives us Theorem 5. We skip this for brevity.

The same framework can be applied to obtain the linear program for the secretary problem with *rehire* (Buchbinder et al. [13]) and the formulation for the (J, K) -secretary problem (Buchbinder et al. [13], Chan et al. [16]). It can also be used to derive an alternative proof of the result by Smith [66]. Besides secretary problems, this approach using MDP has been applied for prophet problems (Jiang et al. [46]) and in online bipartite matching (Torricco and Toriello [69], Torricco et al. [70]).

4.2. Framework for the SP-UA

Next, we sketch the proof of Theorem 1 and use it to derive Theorem 2. Technical details are deferred to Appendix B.

In the classical secretary problem, the objective is to maximize the probability of choosing the top candidate, which we can write in the recursion of the value function v^* . For our model, the objective $\gamma_{\mathcal{P}}$ corresponds to a multiobjective criteria, and it is not clear a priori how to write the objective as a reward. We present a two-step approach: (1) First, we follow the previous subsection's argument to uncover the polyhedron of policies; and (2) second, we show that our objective function can be written in terms of variables in this polyhedron, and we maximize this objective in the polyhedron.

4.2.1. The Polyhedron of Policies via a Generic Utility Function. When we obtained the dual LP (P) (step 3 in the above framework), anything related to the objective of the MDP is moved to the objective value of the LP, while anything related to the actions of the MDP remained in Constraints (5) and (6). This suggests using a generic utility function to uncover the space of policies. Consider any vector $U : [n] \rightarrow \mathbf{R}_+$, and suppose that our objective is to maximize the utility collected, where choosing a candidate of rank i means obtaining $U_i \geq 0$ value. Let $v_{(t,s)}^*$ be the maximum value collected in times $t, t+1, \dots, n$, given that the current state is (t, s) , where $v_{(n+1,s)}^* = 0$. Then, the optimality equations yield

$$v_{(t,s)}^* = \max \left\{ pU_t(s) + (1-p) \frac{1}{t+1} \sum_{\sigma=1}^{t+1} v_{(t+1,\sigma)}^* + \frac{1}{t+1} \sum_{\sigma=1}^{t+1} v_{(t+1,\sigma)}^* \right\}, \quad (7)$$

where $U_t(s) = \sum_{i=1}^n U_i \mathbf{P}(R_t = i | r_t = s)$. The term in the left side of the max operator is the expected value obtained by an **offer** action, whereas the term in the right corresponds to the expected value of the **pass** action. Using an approach similar to the one used in steps 2 and 3 from the previous subsection, we can deduce that

$$\text{POL} = \left\{ (\mathbf{x}, \mathbf{y}) \geq 0 : x_{1,1} + y_{1,1} = 1, x_{t,s} + y_{t,s} = \frac{1}{t} \sum_{\sigma=1}^{t-1} (y_{t-1,\sigma} + (1-p)x_{t-1,\sigma}), \forall t > 1, s \in [t] \right\},$$

contains all policies (Theorem 1). A formal proof is presented in Appendix B.

4.2.2. The Linear Program. Next, we consider Theorem 2. Given a policy $\mathcal{P}$, we define $x_{t,s}$ to be the probability of reaching state (t, s) and making an offer to the candidate and $y_{t,s}$ to be the probability of reaching (t, s) and passing. Then, $(\mathbf{x}, \mathbf{y})$ belongs to POL. Moreover,

$$\mathbf{P}(\mathcal{P} \text{ collects a top } k \text{ candidate}) = p \sum_{t=1}^n \sum_{s=1}^t x_{t,s} \mathbf{P}(R_t \leq k | r_t = s). \quad (8)$$

Conversely, any point $(\mathbf{x}, \mathbf{y}) \in \text{POL}$ defines a policy $\mathcal{P}$: At state (t, s) , it extends an **offer** to the t -th candidate with probability $x_{1,1}$ if $t=1$, or probability $tx_{t,s}/(\sum_{\sigma=1}^{t-1} y_{t-1,\sigma} + (1-p)x_{t-1,\sigma})$ if $t > 1$. Also, $\mathcal{P}$ satisfies (8). Thus,

$$\begin{aligned} \gamma_n^* &= \max_{\mathcal{P}} \min_{k \in [n]} \frac{\mathbf{P}(\mathcal{P} \text{ collects a top } k \text{ candidate})}{1 - (1-p)^k} \\ &= \max_{(\mathbf{x}, \mathbf{y}) \in \text{POL}} \min_{k \in [n]} \frac{p \sum_{t=1}^n \sum_{s=1}^t x_{t,s} \mathbf{P}(R_t \leq k | r_t = s)}{1 - (1-p)^k} \\ &= \max \left\{ \gamma : (\mathbf{x}, \mathbf{y}) \in \text{POL}, \gamma \leq \frac{p \sum_{t=1}^n \sum_{s=1}^t x_{t,s} \mathbf{P}(R_t \leq k | r_t = s)}{1 - (1-p)^k}, \forall k \in [n] \right\}. \end{aligned} \quad (9)$$

Projecting the feasible region of (9) as in step 4 onto the (x, γ) -variables gives us Theorem 2. The details appear in Appendix B.

Our MDP framework also allows us to show the following monotonicity result.

Lemma 1. For a fixed $p \in (0, 1]$, we have $\gamma_n^* \geq \gamma_{n+1}^*$ for any $n \geq 1$.

We sketch the proof of this result here and defer the details to Appendix B.3. The dual of the LP (9) can be reformulated as

$$\begin{aligned} \min_{\substack{u: [n] \rightarrow \mathbf{R}_+ \\ \sum_i u_i \geq 1}} \quad & v_{(1,1)} \\ \text{(DLP) s.t.} \quad & v_{(t,s)} \geq \max \left\{ U_t(s) + \frac{1-p}{t+1} \sum_{\sigma=1}^{t-1} v_{(t+1,\sigma)}, \frac{1}{t+1} \sum_{\sigma=1}^{t-1} v_{(t+1,\sigma)} \right\} \quad \forall t \in [n], s \in [t] \\ & v_{(n+1,s)} = 0 \quad \forall s \in [n+1], \end{aligned}$$

where $U_t(s) = p \sum_{j=1}^n (\sum_{k \geq j} u_k / (1 - (1-p)^k)) \mathbf{P}(R_t = j | r_t = s)$. The variables $u_1, \dots, u_n$ correspond to the constraints involving γ in the LP (9). Note that (DLP) is the minimum value that an MDP can attain when the utility functions are given by $U_i = p \sum_{k \geq i} u_k / (1 - (1-p)^k)$. Taking any weighting $u: [n] \rightarrow \mathbf{R}_+$ with $\sum_i u_i \geq 1$, we extend it to $\hat{u}: [n+1] \rightarrow \mathbf{R}_+$ by setting $\hat{u}_{n+1} = 0$. We define accordingly $\hat{U}_i = p \sum_{k \geq i} \hat{u}_k / (1 - (1-p)^k)$ and note that $U_i = \hat{U}_i$ for $i \leq n$ and $\hat{U}_{n+1} = 0$. Using a coupling argument, from any policy for utilities $\hat{U}$ with $n+1$ candidates, we can construct a policy for utilities U , with n candidates, where both policies collect the same utility. Thus, the utility collected by the optimal policy for U upper bounds the utility collected by an optimal policy for $\hat{U}$. The conclusion follows because γ_{n+1}^* is a lower bound for the latter value.

Given that $\gamma_n^* \in [0, 1]$ and $(\gamma_n^*)_n$ is a monotone sequence in n , $\lim_{n \rightarrow \infty} \gamma_n^*$ must exist. In the next section, we show that the limit corresponds to the value of a continuous LP.

5. The Continuous LP

In this section, we introduce the continuous linear program (CLP), and we show that its value γ_∞^* corresponds to the limit of γ_n^* when n tends to infinity. We also state Proposition 2, which allows us to construct feasible solutions of (CLP) using any set of times $0 < t_1 \leq t_2 \leq \dots \leq 1$. In the finite model, the solution constructed in this section has the natural interpretation of segmenting time: for a candidate arriving between times $t_i n$ and $t_{i+1} n$, we make an offer if the candidate has partial rank i or better. In the remainder of the section, *finite model* refers to the SP-UA with $n < \infty$ candidates, whereas the *infinite model* refers to SP-UA when $n \rightarrow \infty$.

We assume $p \in (0, 1]$ fixed. The continuous LP (CLP) is an infinite linear program with variables given by a function $\alpha: [0, 1] \times \mathbf{N} \rightarrow [0, 1]$ and a scalar $\gamma \geq 0$. Intuitively, if in the finite model we interpret $x_{t,s}$ as weights and the sums of $x_{t,s}$ over t as Riemann sums, then the limit of the finite model should have a robust ratio computed by the continuous LP (CLP)

$$\begin{aligned} \sup_{\substack{\alpha: [0,1] \times \mathbf{N} \rightarrow [0,1] \\ \gamma \geq 0}} \quad & \gamma \\ \text{s.t.} \quad & \text{(CLP)}_p \quad t\alpha(t,s) \leq 1 - p \int_0^t \sum_{\sigma \geq 1} \alpha(\tau, \sigma) d\tau \quad \forall t \in [0, 1], s \geq 1, \end{aligned} \quad (10)$$

$$\gamma \leq \frac{p \int_0^1 \sum_{s \geq 1} \alpha(t,s) \sum_{\ell=s}^k \binom{\ell-1}{s-1} t^s (1-t)^{\ell-s} dt}{(1 - (1-p)^k)} \quad \forall k \geq 1. \quad (11)$$

We denote by $\gamma_\infty^* = \gamma_\infty^*(p)$ the objective value of $(\text{CLP})_p$. The following result formalizes the fact that the value of the continuous LP $(\text{CLP})_p$ is, in fact, the robust ratio of the infinite model. The proof is similar to other continuous approximations (Chan et al. [16]); a small caveat in the proof is the restriction of the finite LP to the top $(\log n)/p$ candidates, as they carry most of the weight in the objective function. The proof is deferred to Appendix C.

Lemma 2. Let γ_n^* be the optimal robust ratio for n candidates, and let γ_∞^* be the value of the continuous LP $(\text{CLP})_p$. Then, $|\gamma_n^* - \gamma_\infty^*| \leq \mathcal{O}((\log n)^2 / (p\sqrt{n}))$.

The following proposition gives a recipe to find feasible solutions for $(CLP)_p$. We use it to construct lower bounds in the following sections.

Proposition 2. Consider $0 \leq t_1 \leq t_2 \leq \dots \leq 1$, and consider the function $\alpha : [0, 1] \times \mathbf{N} \rightarrow [0, 1]$ defined such that for $t \in [t_i, t_{i+1})$

$$\alpha(t, s) = \begin{cases} T_i / t^{i \cdot p + 1} & s \leq i \\ 0 & s > i, \end{cases}$$

where $T_i = (t_1 \dots t_i)^p$. Then, α satisfies Constraint (10).

Proof. We verify that Inequality (10) holds. We only need to verify it for $t \in [t_i, t_{i+1})$ with $i \geq 1$ because $\alpha(t, s) = 0$ for $t \in [0, t_1)$. We define $t_0 = 0$ and $T_0 = 0$. For $t \in [t_i, t_{i+1})$, we have

$$\begin{aligned} 1 - p \int_0^t \sum_{\sigma \geq 1} \alpha(\tau, \sigma) d\tau &= 1 - p \left(\sum_{j=1}^{i-1} \int_{t_j}^{t_{j+1}} j \frac{T_j}{\tau^{jp+1}} d\tau \right) - p \int_{t_i}^t i \frac{T_i}{\tau^{ip+1}} d\tau \\ &= 1 - p \sum_{j=1}^{i-1} j \cdot T_j \cdot \frac{1}{-pj} (t_{j+1}^{-jp} - t_j^{-jp}) - pi \cdot T_i \cdot \frac{1}{-ip} (t^{-ip} - t_i^{ip}) \\ &= 1 + \sum_{j=1}^{i-1} T_j (t_{j+1}^{-jp} - t_j^{-jp}) + T_i (t^{-ip} - t_i^{ip}) \\ &= 1 + \left(\sum_{j=1}^{i-1} T_{j+1} t_{j+1}^{-(j+1)p} - T_j t_j^{-jp} \right) + T_i (t^{-ip} - t_i^{ip}) \quad (\text{Since } T_j t_{j+1}^{-jp} = T_{j+1} t_{j+1}^{-(j+1)p}) \\ &= 1 + T_i t_i^{-ip} - T_1 t_1^{-p} + T_i (t^{-ip} - t_i^{-ip}) \\ &= T_i t^{-ip} \geq t\alpha(t, s), \end{aligned}$$

for any $s \geq 1$. This concludes the proof. $\square$

We use this result to show lower bounds for γ_∞^* . For instance, if $1/k \geq p > 1/(k+1)$ for some integer k , and we set $t_1 = 1/e$ and $t_2 = t_3 = \dots = 1$, we can show that $\gamma_\infty^*(p)$ is at least $1/e$. Thus, in combination with Lemma 1, we have that $\gamma_n^*(p) \geq 1/e$ for any n and $p > 0$; we skip this analysis for brevity. In Section 7, we use Proposition 2 to show exact solutions of γ_∞^* for large p .

6. Upper Bounds for the Continuous LP

We now consider upper bounds for (CLP) and prove Theorem 3, which states that $\gamma_\infty^*(p) \leq \min\{p^{p/(1-p)}, 1/\beta\}$, for any $p \in (0, 1]$, where $1/\beta \approx 0.745$ and β is the unique solution of $\int_0^1 (y(1 - \log y) + \beta - 1)^{-1} dy = 1$ (Kertz [48]).

We show that γ_∞^* is bounded by each term in the minimum operator. For the first bound, we have

$$\begin{aligned} \gamma_n^* &= \max_p \min_{k \in [n]} \frac{\mathbf{P}(\mathcal{P} \text{ collects a top } k \text{ candidate})}{1 - (1-p)^k} \\ &\leq \max_p \frac{\mathbf{P}(\mathcal{P} \text{ collects the top candidate})}{p}. \end{aligned}$$

The probability of collecting the highest candidate in **SP-UA** is shown by Smith [66] to be $p^{1/(1-p)} + o(1)$, where $o(1) \rightarrow 0$ as $n \rightarrow \infty$. Thus, by Lemma 1, we have

$$\gamma_\infty^*(p) \leq \gamma_n^*(p) \leq p^{p/(1-p)} + o(1)/p.$$

Taking the limit $n \rightarrow \infty$, we conclude $\gamma_\infty^*(p) \leq p^{p/(1-p)}$.

For the second bound, we use the following technical result; its proof is deferred to Appendix D, but we give a short explanation here. A γ -robust algorithm $\mathcal{A}$ for the **SP-UA**, in expectation, has pn candidates to choose from

and $(1-p)n$ candidates from which the algorithm can learn about candidate quality. We give an algorithm $\mathcal{A}'$ that solves the i.i.d. prophet inequality for any $m \approx pn$ i.i.d. random variables $X_1, \dots, X_m$, for m large. The algorithm $\mathcal{A}'$ runs a utility version of $\mathcal{A}$ in n values sampled from the distribution X_1 (see the discussion before Proposition 1), guaranteeing at least a factor γ of the maximum of $m \approx pn$ of these samples, which is the value of the prophet. $\mathcal{A}'$ is the capped utility version of $\mathcal{A}$, where no more than $m \approx pn$ offers can be made. Using concentration bounds, we show that the loss of these restrictions is minimal. Kaplan et al. [47] use a similar argument, with the difference that their sampling is a fixed fraction of the input and is done in advance, whereas in our case, the sampling is online and might deviate from the expectation, implying the need for concentration bounds. The following result summarizes the reduction and the upper bound.

Theorem 6. *Let $p \in (0, 1)$ and $\mathcal{A}$ be any algorithm that is γ -robust for the SP-UA for any n . Then, $\gamma \leq 1/\beta$, where $\beta \approx 1.34$ is the unique solution of the integral equation $\int_0^1 (y(1 - \log y) + (\beta - 1))^{-1} dr = 1$.*

The proof of Theorem 6 uses the reduction mentioned in the previous paragraph. The use of concentration bounds guarantees a $(1 - o(1))\gamma$ multiplicative approximation with $e^{-\Theta(n^2)}$ additive error for the i.i.d. prophet inequality problem. Specifically, for any distribution $\mathcal{D}$ over $[0, 1]$, we can guarantee

$$\mathbb{E}[\text{Val}(\mathcal{A}')] + e^{-\Theta(n^2)} \geq \gamma(1 - o(1))\mathbb{E}\left[\max_{i \leq m} X_i\right], \quad (12)$$

where $X_1, \dots, X_m$ distribute according to $\mathcal{D}$ and $\text{Val}(\mathcal{A}')$ is the value collected by $\mathcal{A}'$, the algorithm described in the previous paragraph, where no more than $m \approx pn$ offers are made. Here $o(1)$ is a term that can be chosen arbitrarily close to 0 for n large enough. Unfortunately, we cannot conclude that $\gamma \leq 1/\beta$ immediately from this bound because this bound only holds for multiplicative error in the i.i.d. prophet problem. We bypass this technical challenge as follows. A combination of results by Hill and Kertz [43] and Kertz [48] shows that, for any m and for $\varepsilon' > 0$ small enough, there is an i.i.d. instance $X_1, \dots, X_m$ with support in $[0, 1]$ such that

$$\mathbb{E}\left[\max_{i \leq m} X_i\right] \geq (a_m - \varepsilon') \sup\{\mathbb{E}[X_\tau] : \tau \in T_m\},$$

where T_m is the class of stopping times for $X_1, \dots, X_m$, and $a_m \rightarrow \beta$. Thus, using Inequality (12), we must have

$$\gamma(1 - o(1)) \leq \frac{1}{a_m - \varepsilon'} + \frac{e^{-\Theta(n^2)}}{\mathbb{E}[\max_{i \leq m} X_i]},$$

for $m \approx pn$. A slight reformulation of Hill and Kertz's result allows us to set $\varepsilon' = 1/m^3$ and $\mathbb{E}[\max_{i \leq m} X_i] \geq 1/m^3$ (see the discussion at the end of Appendix D). Thus, as $n \rightarrow \infty$, we have $m \rightarrow \infty$ and so $e^{-\Theta(n^2)}/\mathbb{E}[\max_{i \leq m} X_i] \rightarrow 0$. In the limit, we obtain $\gamma(1 - o(1)) \leq 1/\beta$, which implies our stated result.

An algorithm that solves $(LP)_{n,p}$ and implements the policy given by the solution is γ_∞^* -robust (Theorem 2 and the fact that $\gamma_n^* \geq \gamma_\infty^*$) for any n . Thus, by the previous analysis, we obtain $\gamma_\infty^* \leq 1/\beta \approx 0.745$.

7. Lower Bounds for the Continuous LP

In this section, we consider lower bounds for $(CLP)_p$ and prove Theorem 4. We first give optimal solutions of $(CLP)_p$ for large values of p . For $p \geq p^* \approx 0.594$, the optimal value of $(CLP)_p$ is $\gamma_\infty^*(p) = p^{p/(1-p)}$, and the optimal strategy is to observe $p^{1/(1-p)}$ fraction of the candidates and then make offers to the best observed candidate so far. We then show that for $p \leq p^*$, $\gamma_\infty^*(p) \geq (p^*)^{p^*/(1-p^*)} \approx 0.466$. At the end of the section, we show that $\gamma_\infty^*(p) \geq 0.51$ when $p \rightarrow 0$.

7.1. Exact Solution for Large p

We now show that for $p \geq p^*$, $\gamma_\infty^*(p) = p^{p/(1-p)}$, where $p^* \approx 0.594$ is the solution of $(1-p)^2 = p^{(2-p)/(1-p)}$. Thanks to the upper bound $\gamma_\infty^*(p) \leq p^{p/(1-p)}$ for any $p \in (0, 1]$, it is enough to exhibit feasible solutions (α, γ) of the continuous LP $(CLP)_p$ with $\gamma \geq p^{p/(1-p)}$.

Let $t_1 = p^{1/(1-p)}$, $t_2 = t_3 = \dots = 1$, and consider the function α defined by $t_1, t_2, \dots$ in Proposition 2. That is, for $t \in [0, p^{1/(1-p)})$, $\alpha(t, s) = 0$ for any $s \geq 1$ and for $t \in [p^{1/(1-p)}, 1]$, we have

$$\alpha(t, s) = \begin{cases} p^{p/(1-p)} / t^{1+p} & s = 1 \\ 0 & s > 1. \end{cases}$$

Let $\gamma \doteq \inf_{k \geq 1} p(1 - (1 - p)^k)^{-1} \int_{p^{1/(1-p)}}^1 \frac{p^{p/(1-p)}}{t^{1+p}} \sum_{\ell=1}^k t(1-t)^{\ell-1} dt$. Then, (α, γ) is feasible for the continuous LP $(CLP)_p$, and we aim to show that $\gamma \geq p^{p/(1-p)}$ when $p \geq p^*$. The result follows by the following lemma.

Lemma 3. For any $p \geq p^*$ and any $\ell \geq 0$, $\int_{p^{1/(1-p)}}^1 (1-t)^\ell t^{-p} dt \geq (1-p)^\ell$.

We defer the proof of this lemma to Appendix E. Now, we have

$$\begin{aligned} \gamma &= \inf_{k \geq 1} \frac{p}{1 - (1-p)^k} \int_{p^{1/(1-p)}}^1 \frac{p^{p/(1-p)}}{t^{1+p}} \sum_{\ell=1}^k t(1-t)^{\ell-1} dt \\ &= p^{p/(1-p)} \inf_{k \geq 1} \frac{\sum_{\ell=1}^k \int_{p^{1/(1-p)}}^1 t^{-p} (1-t)^{\ell-1} dt}{\sum_{\ell=1}^k (1-p)^{\ell-1}} \\ &\geq p^{p/(1-p)} \inf_{k \geq 1} \inf_{\ell \in [k]} \int_0^1 \frac{1}{t^p} \frac{(1-t)^{\ell-1}}{(1-p)^{\ell-1}} dt \geq p^{p/(1-p)}, \end{aligned}$$

where we use the known inequality $\sum_{\ell=1}^m a_\ell / \sum_{\ell=1}^m b_\ell \geq \min_{\ell \in [m]} a_\ell / b_\ell$ for $a_\ell, b_\ell > 0$, for any ℓ , and the lemma. This shows that $\gamma_\infty^* \geq p^{p/(1-p)}$ for $p \geq p^*$.

Remark 1. Our analysis is tight. For $k=2$, constraint

$$\frac{p}{1 - (1-p)^2} \int_{p^{1/(1-p)}}^1 \frac{p^{p/(1-p)}}{t^{1+p}} \sum_{\ell=1}^2 t(1-t)^{\ell-1} dt \geq p^{p/(1-p)},$$

holds if and only if $p \geq p^*$.

7.2. Lower Bounds for Small p

In this subsection, we present two lower bounds for $\gamma_\infty^*(p)$ when $p \leq p^*$, with $p^* \approx 0.594$ as obtained in the previous subsection. The first bound guarantees $\gamma_\infty^*(p) \geq (p^*)^{p^*/(1-p^*)} \approx 0.466$; the second guarantees $\gamma_\infty^*(p) \geq 0.51$ when p approaches 0. We present details for the first bound, as it includes a mechanism to transform the solution of $(CLP)_{p^*}$ into a solution of $(CLP)_p$ for $p \leq p^*$. We defer some details of the latter bound, as it uses a construction similar to Correa et al. [22] with a different limit argument.

Let $\varepsilon \in [0, 1]$ satisfy $p = (1 - \varepsilon)p^*$. For the argument, we take the solution α^* for $(CLP)_{p^*}$ that we obtained in the last subsection, and we construct a feasible solution for $(CLP)_p$ with objective value at least $(p^*)^{p^*/(1-p^*)}$. For simplicity, we denote $\tau^* = (p^*)^{1/(1-p^*)}$.

From the previous subsection, we know that the optimal solution α^* of $(CLP)_{p^*}$ has the following form. For $t \in [0, \tau^*)$, $\alpha^*(t, s) = 0$ for any s , whereas for $t \in [\tau^*, 1]$, we have

$$\alpha^*(t, s) = \begin{cases} (p^*)^{p^*/(1-p^*)} / t^{p^*+1} & s = 1 \\ 0 & s > 1. \end{cases}$$

For $(CLP)_p$, we construct a solution α as follows. Let $\alpha(t, s) = \varepsilon^{s-1} \alpha^*(t, 1)$ for any $t \in [0, 1]$ and $s \geq 1$; for example, $\alpha(t, 1) = \alpha^*(t, 1)$. If we interpret α^* as a policy, it only makes offers to the highest candidate observed. By contrast, in $(CLP)_p$, the policy implied by α makes offers to more candidates (after time τ^*), with a probability geometrically decreasing according to the relative ranking of the candidate.

Lemma 4. The solution α satisfies Constraints (10),

$$t\alpha(t, s) \leq 1 - p \int_0^t \sum_{\sigma \geq 1} \alpha(\tau, \sigma) d\tau,$$

for any $t \in [0, 1]$, $s \geq 1$.

Proof. Indeed,

$$\begin{aligned}
 1 - p \int_0^t \sum_{\sigma \geq 1} \alpha(\tau, \sigma) d\tau &= 1 - p^*(1 - \varepsilon) \int_0^t \sum_{\sigma \geq 1} \varepsilon^{\sigma-1} \alpha^*(\tau, 1) d\tau \\
 &= 1 - p^* \int_0^t \alpha^*(\tau, 1) d\tau && \left(\text{Since } \sum_{\sigma \geq 1} \varepsilon^{\sigma-1} = 1/(1 - \varepsilon) \right) \\
 &= 1 - p^* \int_0^t \sum_{\sigma \geq 1} \alpha^*(\tau, \sigma) d\tau && (\text{Since } \alpha^*(\tau, \sigma) = 0 \text{ for } \sigma > 1) \\
 &\geq t\alpha^*(t, 1). && (\text{By feasibility of } \alpha^*)
 \end{aligned}$$

Given that $\alpha(t, s) = \varepsilon^{s-1} \alpha^*(t, 1) \leq \alpha^*(t, 1)$, we conclude that α satisfies (10) for any t and s . $\square$

We now define $\gamma = \inf_{k \geq 1} p(1 - (1 - p)^k)^{-1} \int_0^1 \sum_{s \geq 1} \alpha(t, s) \sum_{\ell=s}^k \binom{\ell-1}{s-1} t^s (1-t)^{\ell-s} dt$. Using the claim, we know that (α, γ) is feasible for $(CLP)_p$ and need to verify that $\gamma \geq (p^*)^{p^*/(1-p^*)}$. Similar to the analysis in the previous section, the result follows by the following result.

Lemma 5. For any $\ell \geq 0$, $\int_{\tau^*}^1 (1 - (1 - \varepsilon)t)^\ell t^{-p^*} dt \geq (1 - p)^\ell$.

Before proving the claim, we establish the bound:

$$\begin{aligned}
 \gamma &= \inf_{k \geq 1} \frac{1}{\sum_{\ell=1}^k (1 - p)^{\ell-1}} \int_0^1 \sum_{s=1}^k \varepsilon^{s-1} \alpha^*(s, 1) \sum_{\ell=s}^k \binom{\ell-1}{s-1} t^s (1-t)^{\ell-s} dt && (\text{Using definition of } \alpha) \\
 &= (p^*)^{p^*/(1-p^*)} \inf_{k \geq 1} \frac{1}{\sum_{\ell=1}^k (1 - p)^{\ell-1}} \int_{\tau^*}^1 \frac{1}{t^{p^*+1}} \sum_{\ell=1}^k \sum_{s=1}^{\ell} \varepsilon^{s-1} \binom{\ell-1}{s-1} t^s (1-t)^{\ell-s} dt \\
 &&& (\text{Using the definition of } \alpha^* \text{ and changing order of summation}) \\
 &= (p^*)^{p^*/(1-p^*)} \inf_{k \geq 1} \frac{1}{\sum_{\ell=1}^k (1 - p)^{\ell-1}} \int_{\tau^*}^1 \frac{1}{t^{p^*}} \sum_{\ell=1}^k (1 - (1 - \varepsilon)t)^{\ell-1} dt && (\text{Using the binomial expansion}) \\
 &= (p^*)^{p^*/(1-p^*)} \inf_{k \geq 1} \frac{\sum_{\ell=1}^k \int_{\tau^*}^1 t^{-p^*} (1 - (1 - \varepsilon)t)^{\ell-1} dt}{\sum_{\ell=1}^k (1 - p)^{\ell-1}} \geq (p^*)^{p^*/(1-p^*)}.
 \end{aligned}$$

We again used the inequality $\sum_{\ell=1}^m a_\ell / \sum_{\ell=1}^m b_\ell \geq \min_{\ell \in [m]} a_\ell / b_\ell$ for $a_\ell, b_\ell > 0$, for any ℓ , and the claim.

Proof of Lemma 5. We have $1 - (1 - \varepsilon)t = (1 - \varepsilon)(1 - t) + \varepsilon$. Therefore,

$$\begin{aligned}
 \int_{\tau^*}^1 \frac{1}{t^{p^*}} (1 - (1 - \varepsilon)t)^\ell dt &= \int_{\tau^*}^1 \frac{1}{t^{p^*}} \sum_{j=0}^{\ell} \binom{\ell}{j} (1 - \varepsilon)^{\ell-j} (1 - t)^{\ell-j} \varepsilon^j dt && (\text{Binomial expansion}) \\
 &= \sum_{j=0}^{\ell} \binom{\ell}{j} (1 - \varepsilon)^{\ell-j} \varepsilon^j \int_{\tau^*}^1 \frac{1}{t^{p^*}} (1 - t)^{\ell-j} dt \\
 &\geq \sum_{j=0}^{\ell} \binom{\ell}{j} (1 - \varepsilon)^{\ell-j} \varepsilon^j (1 - p^*)^{\ell-j} dt && (\text{Using Lemma 3 for } p^*) \\
 &= (\varepsilon + (1 - \varepsilon)(1 - p^*))^\ell && (\text{Using binomial expansion}) \\
 &= (1 - (1 - \varepsilon)p^*)^\ell = (1 - p)^\ell,
 \end{aligned}$$

where we used $p = (1 - \varepsilon)p^*$. From this inequality the claim follows. $\square$

7.2.1. Improved Bound for p Close to 0. Now we present a better bound for $\gamma_\infty^*(p)$, for p close to 0, using an explicit construction of a solution α for $(CLP)_p$.

The following proposition gives a sufficient condition to ensure lower bounds over $\gamma_\infty^*(p)$.

Proposition 3. Let $0 \leq t_1 \leq t_2 \leq t_3 \leq \dots \leq 1$. If for all $k \geq 1$ we have

$$T_k \left(\frac{t_{k+1}^{1-kp} - t_k^{1-kp}}{1 - kp} \right) \geq \gamma p (1 - p)^{k-1},$$

where $T_k = (t_1 \dots t_k)^p$, then, $\gamma_\infty^*(p) \geq \gamma$.

Proof. Let $\alpha(t, s)$ be defined as follows. For $i \geq 1$, if $t \in [t_i, t_{i+1})$, then

$$\alpha(t, s) = \begin{cases} T_i / t^{ip+1} & s \leq i \\ 0 & s > i. \end{cases}$$

Then, by Proposition 2, we know that α defines a feasible solution to $(CLP)_p$. Now, using $\alpha(t, s) \geq \alpha(t, \ell)$ for any $s \leq \ell$, we obtain

$$\begin{aligned} p \int_0^1 \sum_{s \geq 1} \alpha(t, s) \sum_{\ell=s}^k \binom{\ell-1}{s-1} t^s (1-t)^{\ell-s} dt &\geq p \int_0^1 \sum_{\ell=1}^k \alpha(t, \ell) t dt \\ &= p \sum_{i=1}^{\infty} \min\{k, i\} T_i \left(\frac{t_{i+1}^{1-ip} - t_i^{1-ip}}{1 - ip} \right) \\ &= p \sum_{j=1}^k \sum_{i=j}^{\infty} T_i \left(\frac{t_{i+1}^{1-ip} - t_i^{1-ip}}{1 - ip} \right) \\ &\geq p \sum_{j=1}^k \sum_{i=j}^{\infty} \gamma p (1 - p)^{i-1} \\ &= \gamma p \sum_{j=1}^k (1 - p)^{j-1} = \gamma (1 - (1 - p)^k). \end{aligned}$$

This holds for any k ; thus, $\gamma_\infty^*(p) \geq \gamma$. $\square$

We now present an iterative method to generate a sequence $t_1, t_2, \dots$ as in Proposition 3. Fix $t_1 \in [0, 1]$ and $\gamma > 0$, and define $A_k = t_1(1 - p)^{-k} + \gamma p k$ for $k \geq 1$. Note that A_k is increasing in k . Define $t_2 = t_1(1 + \gamma p(1 - p)/t_1)^{1/(1-p)}$, and for $k \geq 2$, define t_{k+1} as follows:

$$\left(\frac{t_{k+1}}{t_k} \right)^{1-kp} = \frac{A_k(1 - p)}{A_{k-1}}. \quad (13)$$

Lemma 6. The sequence $t_1, \dots$ defined above satisfies $t_k \leq t_{k+1}$ for each $k \geq 1$. Moreover, for any $k \geq 1$,

$$T_k \left(\frac{t_{k+1}^{1-kp} - t_k^{1-kp}}{1 - kp} \right) = \gamma p (1 - p)^{k-1},$$

where $T_k = (t_1 \dots t_k)^p$.

Proof. The first part follows from the fact that $k < 1/p$ if and only if $A_k(1 - p) \geq A_{k-1}$. For the second part, let

$$B_k = T_k \left(\frac{t_{k+1}^{1-kp} - t_k^{1-kp}}{1 - kp} \right).$$

It is easy to verify that $B_1 = \gamma p$ using the definition of t_2 . Now, for $k \geq 2$,

$$\frac{B_{k+1}}{B_k} = t_{k+1}^p \left(\frac{t_{k+2}^{1-(k+1)p} - t_{k+1}^{1-(k+1)p}}{t_{k+1}^{1-kp} - t_k^{1-kp}} \right) \left(\frac{1 - kp}{1 - (k+1)p} \right).$$

Using Identity (13) and $A_j(1-p) - A_{j-1} = \gamma p(1-jp)$, we obtain

$$\frac{B_{k+1}}{B_k} = t_{k+1}^p \frac{A_{k-1}}{A_k} \frac{t_{k+1}^{1-(k+1)p}}{t_k^{1-kp}} = (1-p).$$

Then, inductively, we can show that $B_k = \gamma p(1-p)^{k-1}$. $\square$

Our goal is to give $t_1 \in [0, 1]$ and $\gamma \in [0, 1]$ as large as possible such that $\lim_k t_k \leq 1$, with t_k defined via (13).

Lemma 7. *We have*

$$\lim_{k \rightarrow \infty} \log t_k = \log t_1 + \int_0^\infty \frac{\gamma}{t_1 e^x + \gamma x} dx,$$

for $p \rightarrow 0$.

The proof of the lemma is technical and borrows a strategy used by Correa et al. [22]. The main insight is to apply logarithms to both sides of Identity (13) and find an expression for $\log t_{k+1}$ as a sum of terms only involving A_j , γ , and p . In the limit in k , and then in p , we can reinterpret the sums as Riemann sums that later translate into the integral term given in the lemma. We defer the details of this proof to Appendix E.

We find numerically that the combination of $t_1 \approx 0.228$ and $\gamma \approx 0.511$ ensures $\lim_{k \rightarrow \infty} \log t_k = 0$. Thus, $\gamma_\infty^*(p) \geq 0.51$ for p close to 0; note that t_1 remains bounded away from 0. This means that the policy given by the values $t_1, t_2, \dots$ spends the first t_1 fraction of time “exploring” before “exploiting,” and the fraction of exploration time remains a constant.

8. Computational Experiments

In this section, we aim to empirically test our policy; to do so, we focus on utility models. Recall from Proposition 1 that a γ -robust policy ensures at least γ fraction of the optimal offline utility, for any utility function that is consistent with the ranking—that is, $U_j < U_i$ if and only if $i < j$. This is advantageous for practical scenarios, where a candidate’s “value” may be unknown to the decision maker.

We evaluate the performance of two groups of solutions. The first group includes policies that are computed without the knowledge of any utility function:

- Robust policy (**Rob-Pol**(n, p)) corresponds to the optimal policy obtained by solving $(LP)_{n,p}$.
- Tamaki’s policy (**Tama-Pol**(n, p)), which maximizes the probability of selecting the best candidate willing to accept an offer. To be precise, Tamaki [68] studies two models of availability: MODEL 1, where the availability of the candidate is known after an offer has been made; and MODEL 2, where the availability of the candidate is known upon the candidate’s arrival. MODEL 2 has higher values and is computationally less expensive to compute; we use this policy. Note that in **SP-UA**, the expected value obtained by learning the availability of the candidate after making an offer is the same value obtained in the model that learns the availability upon arrival. Therefore, MODEL 2 is a better model to compare our solutions to than MODEL 1.

In the other group, we have policies that are computed with knowledge of the utility function.

- The expected optimal offline value ($\mathbf{E}[U(\text{OPT}(U, n, p))]$), which knows the outcome of the offers and the utility function. It can be computed via $\sum_{i=1}^n U_i p(1-p)^{i-1}$. For simplicity, we write OPT when the parameters are clear from the context.
- The optimal rank-based policy if the utility function is known in advance, (**Util-Pol**(U, n, p)), computed by solving the optimality equation

$$v_{(t,s)} = \max \left\{ U_t(s) + \frac{1-p}{t+1} \sum_{\sigma=1}^{t+1} v_{(t+1,\sigma)}, \frac{1}{t+1} \sum_{\sigma=1}^{t+1} v_{(t+1,\sigma)} \right\},$$

with boundary condition $v_{(n+1,\sigma)} = 0$ for any σ . We write **Util-Pol**(n, p) when U is clear from the context. We use a rank-based policy as opposed to a value-based policy for computational efficiency.

Note that $\mathbf{E}[U(\text{Rob-Pol})], \mathbf{E}[U(\text{Tama-Pol})] \leq \mathbf{E}[U(\text{Util-Pol})] \leq \mathbf{E}[U(\text{OPT})]$ and by Proposition 1, $\mathbf{E}[U(\text{Rob-Pol})] \geq \gamma_n^* \mathbf{E}[U(\mathcal{A})]$ for any $\mathcal{A}$ of the aforementioned policies.

We consider the following decreasing utility functions:

- **Top k candidates are valuable (top- k).** For $k \in [n]$, we consider utility functions of the form $U_i = 1 + \varepsilon^i$ for $i \in [k]$ and $U_i = \varepsilon^i$ for $i > k$ with $\varepsilon = 1/n$. Intuitively, we aim to capture the notion of an elite set of candidates, where

candidates outside the top k are not nearly as appealing to the decision maker. For instance, renowned brands like to target certain members of a population for their ads. We test $k = 1, 2, 3, 4$.

• **Power law population.** $U_i = i^{-1/(1+\delta)}$ for $i \geq 1$ and small $\delta > 0$. Empirical studies have shown that the distribution of individual performances in many areas follows a power law or Pareto distribution (Clauset et al. [20]). If we select a random person from $[n]$, the probability that this individual has a performance score of at least t is proportional to $t^{-(1+\delta)}$. We test $\delta \in \{10^{-2}, 10^{-1}, 2 \cdot 10^{-1}\}$.

We run experiments for $n=200$ candidates and range the probability of acceptance p from $p=10^{-2}$ to $p=9 \cdot 10^{-1}$.

8.1. Results for Top- k Utility Function

In this subsection, we present the results for utility function that has largest values in the top k candidates, where $k = 1, 2, 3, 4$. In Figure 2, we plot the ratio between the value collected by $\mathcal{A}$ and $\mathbf{E}[U(\text{OPT})]$, for $\mathcal{A}$ being **Util-Pol**, **Rob-Pol**, and **Tama-Pol**.

Naturally, of all sequential policies, **Util-Pol** attains the largest approximation factor of $\mathbf{E}[U(\text{OPT})]$. We observe empirically that **Rob-Pol** collects larger values than **Tama-Pol** for smaller values of k . Interestingly, we observe in the four experiments that the approximation factor for **Rob-Pol** is always better than **Tama-Pol** for small values of p . In other words, robustness helps online selection problems when the probability of acceptance is relatively low. In general, for this utility function, we observe in the experiments that **Rob-Pol** collects at least 50% of the optimal off-line value, except for the case $k = 1$. As n increases (not shown in the figures), we observe that the approximation factors of all three policies decrease; this is consistent with the fact that γ_n^* , the optimal robust ratio, is decreasing in n .

8.2. Results for Power-Law Utility Function

In this subsection, we present the result of our experiments for the power-law utility function $U_i = i^{-(1+\delta)}$ for $\delta = 10^{-2}, 10^{-1}$, and $2 \cdot 10^{-1}$. In Figure 3, we display the approximation factors of the three sequential policies.

Again, we note that **Util-Pol** collects the largest fraction of all sequential policies. We also observe a similar behavior as in the case of the top- k utility function. For small values of p , **Rob-Pol** empirically collects more value

Figure 2. (Color online) Approximation factors for the top k utility function, for $k = 1, 2, 3, 4$.

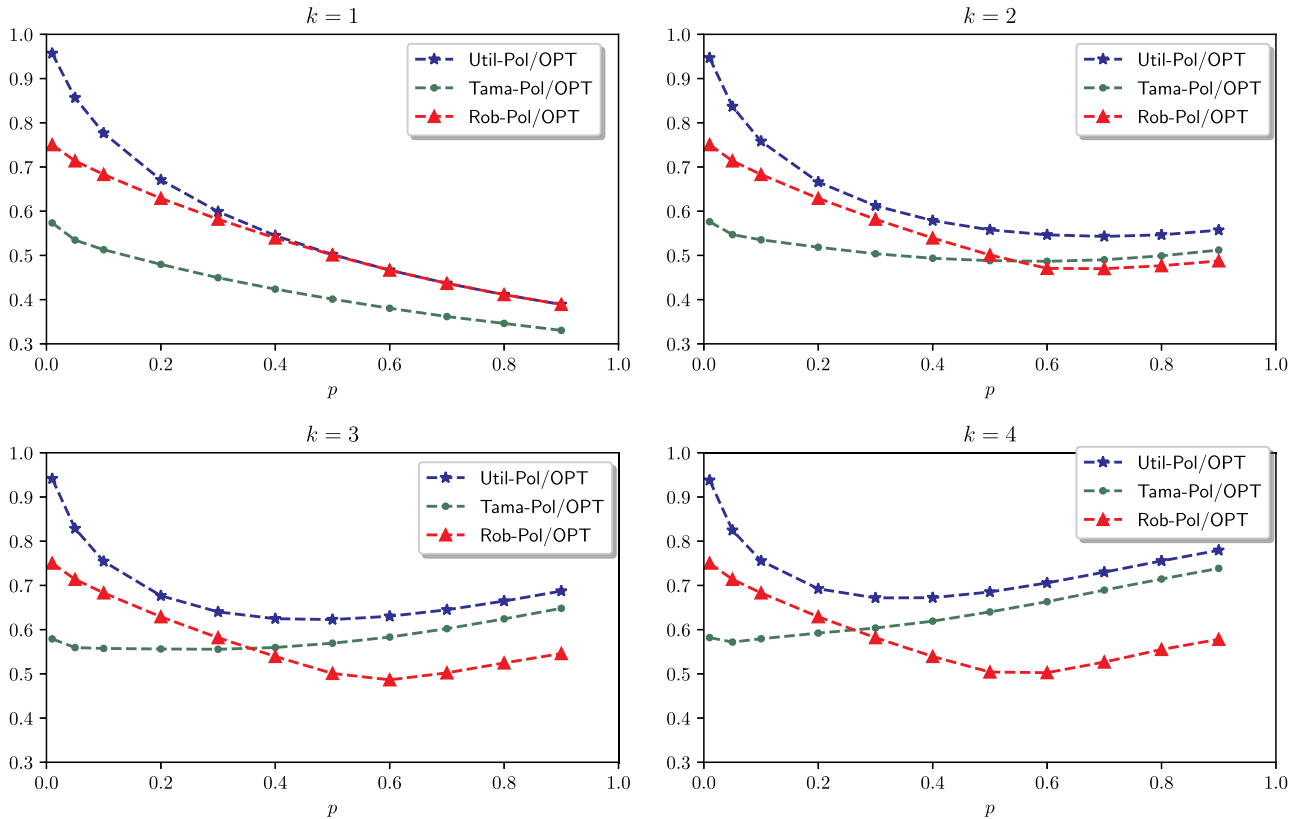
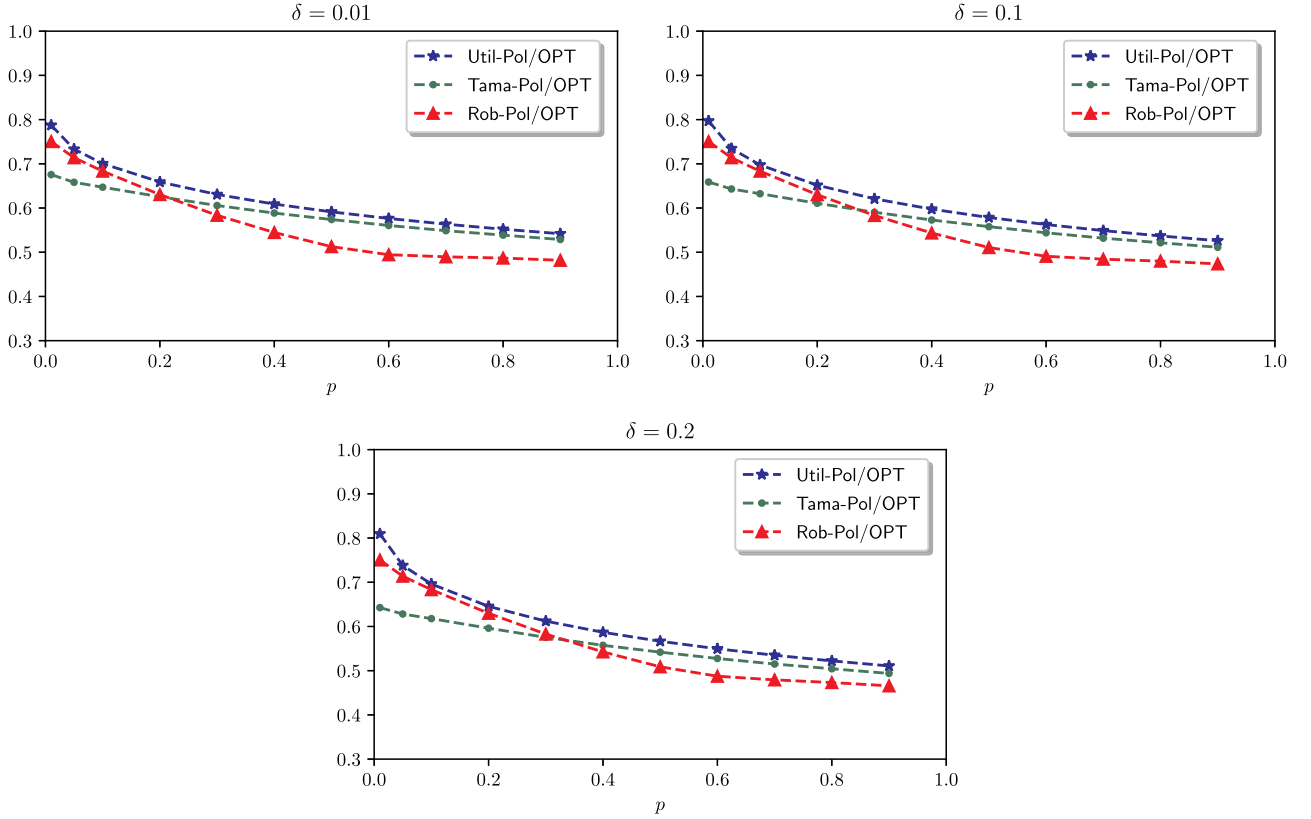


Figure 3. (Color online) Approximation factor for the power law utility function. The function has the form $U_i = i^{-(1+\delta)}$. Experiments are run for $\delta \in \{10^{-2}, 10^{-1}, 2 \cdot 10^{-1}\}$.



than **Tama-Pol**. As p increases, the largest valued candidate is more willing to accept an offer; hence, Tamaki's 1991 policy is able to capture that candidate.

In general, our experiments suggests that **Rob-Pol** is better than **Tama-Pol** for smaller values of p . This may be of interest in applications where the probability of acceptance is small, say, 20% or less. For instance, some sources state that click-through rates (the fraction of time that an ad is clicked on) are typically less than 1% (Farahat and Bailey [31]). Therefore, ad display policies based on **Rob-Pol** may be more appropriate than other alternatives.

9. Concluding Remarks

We have studied the **SP-UA**, which models an online selection problem where candidates can reject an offer. We introduced the robust ratio as a metric that tries to simultaneously maximize the probability of successfully selecting one of the best k candidates, given that at least one of these will accept an offer, for all values of k . This objective captures the worst-case scenario for an online policy against an offline adversary that knows in advance which candidates will accept an offer. We also demonstrated a connection between this robust ratio and online selection with utility functions. We presented a framework based on MDP theory to derive a linear program that computes the optimal robust ratio and its optimal policy. This framework can be generalized and used in other secretary problems (Section 4.1), for instance, by augmenting the state space. Furthermore, using the MDP framework, we were able to show that the robust ratio γ_n^* is a decreasing function in n . This enabled us to make connections between early works in secretary problems and recent advances. To study our LP, we allow the number of candidates to go to infinity and obtain a continuous LP. We provide bounds for this continuous LP and optimal solutions for large p .

We empirically observe that the robust ratio $\gamma_n^*(p)$ is convex and decreasing as a function of p , and, thus, we expect the same behavior from $\gamma_\infty(p)$, though this remains to be proved (see Figure 1). Based on numerical values obtained by solving $(LP)_{n,p}$, we conjecture that $\lim_{p \rightarrow 0} \gamma_\infty^*(p) = 1/\beta \approx 0.745$. This limit is also observed in a similar model Correa et al. [21], where a fraction of the input is given in advance to the decision maker as a sample. In our model, if we interpret the rejection from a candidate as a sample, then in the limit, both models might behave

similarly. Numerical comparisons between our policies and benchmarks suggest that our proposed policies perform especially well in situations where the probability of acceptance is small, say, less than 20%, as in the case of online advertisement.

A natural extension is the value-based model, where candidates reveal numerical values instead of partial rankings. Our algorithms are rank-based and guarantee an expected value at least a fraction $\gamma_n^*(p)$ of the optimal offline expected value (Proposition 1). Nonetheless, algorithms based on numerical values may attain higher expected values than the ones guaranteed by our algorithm. In fact, a threshold algorithm based on sampling may perhaps be enough to guarantee better values, although this requires an instance-dependent approach. The policies we consider are instance-agnostic, can be computed once, and used for any input sequence of values. In this value-based model, we would like to consider other arrivals processes. A popular arrival model is the adversarial arrival, where an adversary constructs values and the arrival order in response to the DM's algorithm. Unfortunately, a construction similar to the one in Marchetti-Spaccamela and Vercellis [57] for the online knapsack problem shows that it is impossible to attain a finite competitive ratio in an adversarial regime.

Customers belonging to different demographic groups may have different willingness to click on ads (Cheng and Cantú-Paz [17]). In this work, we considered a uniform probability of acceptance, and our techniques do not apply directly in the case of different probabilities. In ad display, one way to cope with different probabilities, depending on customers' demographic group, is the following. Upon observing a customer, a random variable (independent of the ranking of the candidate) signals the group of the customer. The probability of acceptance of a candidate depends on the candidate's group. Assuming independence between the rankings and the demographic group allows us to learn nothing about the global quality of the candidates beyond what we can learn from the partial rank. Using the framework presented in this work, with an augmented state space (time, partial rank, group type), we can write an LP that solves this problem exactly. Nevertheless, understanding the robust ratio in this new setting and providing a closed-form policy are still open questions.

Another interesting extension is the case of multiple selections. In practice, platforms can display the same ad to more than one user, and some job posts require more than one person for a position. In this setting, the robust ratio is less informative. If k is the number of possible selections, one possible objective is to maximize the number of top k candidates selected. We can apply the framework from this work to obtain an optimal LP. Although there is an optimal solution, no simple closed-form strategies have been found even for $p = 1$; see, for example, Buchbinder et al. [13]).

Appendix A. Missing Proofs from Section 3

Proof of Proposition 1. Let ALG be a γ -robust algorithm. Fix any algorithm ALG and any $U_1 \geq \dots \geq U_n \geq 0$. Let $\varepsilon > 0$ and let $\hat{U}_i = U_i + \varepsilon^i$. Thus, $\hat{U}_i > \hat{U}_{i+1}$, and so rank and utility are in one-to-one correspondence. Then,

$$\frac{\mathbb{E}[\hat{U}(ALG)]}{\mathbb{E}[\hat{U}(OPT)]} \geq \min_{x \in \{\hat{U}_1, \dots, \hat{U}_n\}} \frac{\mathbb{P}(\hat{U}(ALG) \geq x)}{\mathbb{P}(\hat{U}(OPT) \geq x)} = \min_{k \leq n} \frac{\mathbb{P}(ALG \text{ collects a top } k \text{ candidate})}{\mathbb{P}(A \text{ top } k \text{ candidate accepts})} \geq \gamma$$

where we used the fact that ALG is γ -robust. Notice that $\mathbb{E}[U(OPT)] \leq \mathbb{E}[\hat{U}(OPT)]$, and also $\mathbb{E}[\hat{U}(ALG)] \leq \mathbb{E}[U(ALG)] + \varepsilon$. Thus, doing $\varepsilon \rightarrow 0$ we obtain

$$\frac{\mathbb{E}[U(ALG)]}{\mathbb{E}[U(OPT)]} \geq \gamma, \tag{A.1}$$

for any nonzero vector U with $U_1 \geq U_2 \geq \dots \geq U_n \geq 0$. This finishes the first part. For the second part, let

$$\bar{\gamma}_n = \min_{ALG} \max \left\{ \frac{\mathbb{E}[U(ALG)]}{\mathbb{E}[U(OPT)]} : U : [n] \rightarrow \mathbf{R}_+, U_1 \geq \dots \geq U_n \right\}.$$

Note that the right-hand side (RHS) of Inequality (A.1) is independent of U , thus minimizing in U in the left-hand side (LHS) and, then, maximizing in ALG on both sides, we obtain $\bar{\gamma}_n \geq \gamma_n^*$.

To show the reverse inequality, fix $k \in [n]$ and let $\hat{U} : [n] \rightarrow \mathbf{R}_+$ given by $\hat{U}_i = 1$ for $i \leq k$ and $\hat{U}_i = 0$ for $i > k$. Then,

$$\frac{\mathbb{P}(ALG \text{ collects a top } k \text{ candidate})}{\mathbb{P}(A \text{ top } k \text{ candidate accepts})} = \frac{\mathbb{E}[\hat{U}(ALG)]}{\mathbb{E}[\hat{U}(OPT)]} \geq \min_{\substack{U : [n] \rightarrow \mathbf{R}_+ \\ U_1 \geq \dots \geq U_n}} \frac{\mathbb{E}[U(ALG)]}{\mathbb{E}[U(OPT)]}.$$

This bound holds for any k , thus minimizing over k and then maximizing over ALG on both sides, we obtain $\gamma_n^* \geq \bar{\gamma}_n$, which finishes the second part. $\square$

Appendix B. Missing Proofs from Section 4

Here, we present a detailed derivation of Theorem 1 and Theorem 2 by revisiting Section 4.

As stated in Section 4, we are going to proceed in two stages: (1) First, using a generic utility function, we uncover the space of policies POL. (2) Second, we show that our objective is a concave linear function of the variables of the space of policies that allows us to optimize it over POL.

B.1. Stage 1: The Space of Policies

Let $U: [n] \rightarrow \mathbf{R}_+$ be an arbitrary utility function, and suppose that a DM makes decisions based on partial rankings and her goal is to maximize the utility obtained, where she gets U_i if she is able to collect a candidate of ranking i . Let $v_{(t,s)}^*$ be the optimal value obtained by a DM in $\{t, t+1, \dots, n\}$ if she is currently observing the t -th candidate and this candidate has partial ranking $r_t = s$ —that is, the current state of the system is $s_t = (t, s)$. The function $v_{(t,s)}^*$ is called the *value function*. We define $v_{(n+1,\sigma)}^*$ for any $\sigma \in [n+1]$. Then, by optimality equations (Puterman [62]), we must have

$$v_{(t,s)}^* = \max \left\{ pU_t(s) + (1-p) \frac{1}{t+1} \sum_{\sigma=1}^{t+1} v_{(t+1,\sigma)}^*, \frac{1}{t+1} \sum_{\sigma=1}^{t+1} v_{(t+1,\sigma)}^* \right\}, \quad (\text{B.1})$$

with $v_{(n+1,s)}^* = 0$ for any $s \in [n+1]$. The first part in the max operator corresponds to the expected value obtained by selecting the current candidate, whereas the second part in the operator corresponds to the expected value collected by passing to the next candidate. Here, $U_t(s) = \sum_{i=1}^n U_i \mathbf{P}(R_t = i | r_t = s)$ is the expected value collected by the DM, given that the current candidate has partial ranking $r_t = s$ and accepts the offer. Although an optimal policy for this arbitrary utility function can be computed via the optimality equations, we are more interested in all the possible policies that can be obtained via these formulations. For this, we are going to use linear programming. This has been used in MDP theory (Altman [4], Puterman [62]) to study the space of policies.

The following proposition shows that the solution of the optimality Equations (B.1) solves the LP (D):

$$\begin{aligned} \max_{v \geq 0} \quad & v_{(1,1)} \\ (D) \quad & v_{(t,s)} \geq pU_t(s) + \frac{1-p}{t+1} \sum_{\sigma=1}^{t+1} v_{(t+1,\sigma)}, \quad \forall t \in [n], s \in [t], \end{aligned} \quad (\text{B.2})$$

$$v_{(t,s)} \geq \frac{1}{t+1} \sum_{\sigma=1}^{t+1} v_{(t+1,\sigma)}, \quad \forall t \in [n], s \in [t], \quad (\text{B.3})$$

which has as a dual the LP (P):

$$\begin{aligned} \min_{x, y \geq 0} \quad & \sum_{t=1}^n \sum_{s=1}^t U_t(s) x_{t,s} \\ (P) \quad & x_{1,1} + y_{1,1} \leq 1, \end{aligned} \quad (\text{B.4})$$

$$x_{t,s} + y_{t,s} \leq \frac{1}{t} \left(\sum_{\sigma=1}^{t-1} y_{t-1,\sigma} + (1-p)x_{t-1,\sigma} \right) \quad \forall t \in [n], s \in [t]. \quad (\text{B.5})$$

We denote by $v_{(D)}$ the value of the LP (D).

Proposition B.1. Let $v^* = (v_{(t,s)}^*)_{t,s}$ be a solution of (B.1); then, v^* is an optimal solution of the problem of (D) in Equations (B.2)–(B.5).

Proof. Given that v^* satisfies the optimality Equation (B.1), then it clearly satisfies Constraints (B.2) and (B.3). Thus, v^* is feasible and so $v_{(1,1)}^* \geq v_{(D)}$.

To show the optimality of v^* , we show that any solution $\bar{v}$ of the LP is an upper bound for the value function: $v^* \leq \bar{v}$. To show this, we proceed by backward induction in $t = n+1, n, \dots, 1$ and we prove that $v_{(t,s)}^* \leq \bar{v}_{(t,s)}$ for any $s \in [t]$.

We start with the case $t = n+1$. In this case, $v_{(n+1,s)}^* = 0$ for any s , and, because $\bar{v}_{(n+1,s)} \geq 0$ for any s , then the result follows.

Suppose the result is true for $t = \tau+1, \dots, n+1$ and let us show it for $t = \tau$. Using Constraints (B.2)–(B.3), we must have

$$\begin{aligned} \bar{v}_{(\tau,s)} &\geq \max \left\{ pU_\tau(s) + (1-p) \frac{1}{\tau+1} \sum_{\sigma=1}^{\tau+1} \bar{v}_{(\tau+1,\sigma)}, \frac{1}{\tau+1} \sum_{\sigma=1}^{\tau+1} \bar{v}_{(\tau+1,\sigma)} \right\} \\ &\geq \max \left\{ pU_\tau(s) + (1-p) \frac{1}{\tau+1} \sum_{\sigma=1}^{\tau+1} v_{(\tau+1,\sigma)}^*, \frac{1}{\tau+1} \sum_{\sigma=1}^{\tau+1} v_{(\tau+1,\sigma)}^* \right\} \quad (\text{backward induction}) \\ &= v_{(\tau,s)}^* \end{aligned}$$

where the last line follows by the Optimality Equations (B.1). Thus, $v_{(D)} = \bar{v}_{(1,1)} \geq v_{(1,1)}^*$. $\square$

The dual of the LP (D) is depicted in Equations (B.2)–(B.5) and named (P). The crucial fact to notice here is that the feasible region of (P) is oblivious of the utility function (or rewards) given initially to the MDP. This suggests that the region

$$\text{POL} = \left\{ (\mathbf{x}, \mathbf{y}) \geq 0 : x_{1,1} + y_{1,1} = 1, x_{t,s} + y_{t,s} = \frac{1}{t} \sum_{\sigma=1}^{t-1} y_{t-1,\sigma} + (1-p)x_{t-1,\sigma}, \forall t \in [n], s \in [t] \right\},$$

codifies all possible policies. The following two propositions formalize this.

Proposition B.2. For any policy $\mathcal{P}$ for the SP-UA, consider

$$x_{t,s} = \mathbf{P}(\mathcal{P} \text{ reaches state } (t,s), \text{ selects candidate})$$

and

$$y_{t,s} = \mathbf{P}(\mathcal{P} \text{ reaches state } (t,s), \text{ does not select candidate}).$$

Then, $(\mathbf{x}, \mathbf{y})$ belongs to POL.

Proof. Consider the event $D_t = \{t\text{-th candidate turns down offer}\}$. Then, $p = \mathbf{P}(D_t)$. Consider also the events

$$O_t = \{\mathcal{P} \text{ reaches } t\text{-th candidate and extends an offer}\},$$

and

$$\overline{O}_t = \{\mathcal{P} \text{ reaches } t\text{-th candidate and does not extend offer}\}.$$

Then, O_t and $\overline{O}_t$ are disjoint events and $O_t \cup \overline{O}_t$ equals the event of $\mathcal{P}$ reaching stage t . Thus

$$\mathbf{1}_{O_t \cap \{S_t=(t,s)\}} + \mathbf{1}_{\overline{O}_t \cap \{S_t=(t,s)\}} = \mathbf{1}_{\{\mathcal{P} \text{ reaches state } S_t=(t,s)\}}. \quad (\text{B.6})$$

Note that $x_{t,s} = \mathbf{P}(O_t \cap \{S_t = (t,s)\})$ and $y_{t,s} = \mathbf{P}(\overline{O}_t \cap \{S_t = (t,s)\})$. For $t=1$, then $S_1 = (1,1)$ and $\{\mathcal{P} \text{ reaches state } S_1 = (1,1)\}$ occurs with probability 1. Thus,

$$x_{1,1} + y_{1,1} = 1.$$

For $t > 1$, by the dynamics of the system, the only way that $\mathcal{P}$ reaches state t is by reaching stage $t-1$ and not extending an offer to the $t-1$ candidate or extending an offer, but this was turned down. Thus,

$$\mathbf{1}_{\{\mathcal{P} \text{ reaches state } S_t=(t,s)\}} = \sum_{\sigma=1}^{t-1} \mathbf{1}_{\overline{O}_{t-1} \cap \{S_{t-1}=(t-1,\sigma)\} \cap \{S_t=(t,s)\}} + \mathbf{1}_{O_{t-1} \cap \{S_{t-1}=(t-1,\sigma)\} \cap \overline{D}_{t-1} \cap \{S_t=(t,s)\}}. \quad (\text{B.7})$$

Note that

$$\begin{aligned} & \mathbf{P}(\overline{O}_{t-1} \cap \{S_{t-1} = (t-1, \sigma)\} \cap \{S_t = (t, s)\}) \\ &= \mathbf{P}(\overline{O}_{t-1} \cap \{S_{t-1} = (t-1, \sigma)\} | S_t = (t, s)) \mathbf{P}(S_t = (t, s)) \\ &= \mathbf{P}(\overline{O}_{t-1} \cap \{S_{t-1} = (t-1, \sigma)\}) \frac{1}{t} \\ &= \frac{1}{t} y_{t-1,\sigma}. \end{aligned}$$

Note that we use that $\mathcal{P}$'s action at stage $t-1$ only depends on S_{t-1} and not what is observed in the future. Likewise, we obtain

$$\begin{aligned} & \mathbf{P}(O_{t-1} \cap \{S_{t-1} = (t-1, \sigma)\} \cap \overline{D}_{t-1} \cap \{S_t = (t, s)\}) \\ &= \mathbf{P}(\overline{D}_{t-1}) \mathbf{P}(O_{t-1} \cap \{S_{t-1} = (t-1, \sigma)\} \cap \{S_t = (t, s)\}) \\ &= (1-p) \mathbf{P}(O_{t-1} \cap \{S_{t-1} = (t-1, \sigma)\} \cap \{S_t = (t, s)\}) \\ &= \frac{1-p}{t} x_{t-1,\sigma}. \end{aligned}$$

Using the equality between (B.6) and (B.7) and taking expectation, we obtain

$$x_{t,s} + y_{t,s} = \sum_{\sigma=1}^{t-1} \frac{1}{t} y_{t-1,\sigma} + \frac{1-p}{t} x_{t-1,\sigma},$$

which shows that $(\mathbf{x}, \mathbf{y}) \in \text{POL}$. $\square$

Conversely,

Proposition B.3. Let $(\mathbf{x}, \mathbf{y})$ be a point in POL. Consider the (randomized) policy $\mathcal{P}$ that in state (t, s) makes an offer to the candidate t with probability $x_{t,1}$ if $t = 1$ and

$$\frac{tx_{t,s}}{\sum_{\sigma=1}^{t-1} y_{t-1,\sigma} + (1-p)x_{t-1,\sigma}},$$

if $t > 1$. Then, $\mathcal{P}$ is a policy for **SP-UA** such that $x_{t,s} = \mathbf{P}(\mathcal{P} \text{ reaches state } (t, s), \text{ selects candidate } t \text{ and } y_{t,s} = \mathbf{P}(\mathcal{P} \text{ reaches state } (t, s), \text{ does not select candidate } t) \text{ for any } t \in [n] \text{ and } s \in [t]$.

Proof. We use the same events $O_t, \bar{O}_t$ and D_t as defined in the previous proof. Thus, we need to show that $x_{t,s} = \mathbf{P}(O_t \cap \{S_t = (t, s)\})$ and $y_{t,s} = \mathbf{P}(\bar{O}_t \cap \{S_t = (t, s)\})$ are the right marginal probabilities. For this, it is enough to show that

$$\mathbf{P}(O_t | S_t = (t, s)) = tx_{t,s} \quad \text{and} \quad \mathbf{P}(\bar{O}_t | S_t = (t, s)) = ty_{t,s},$$

for any $t \in [n]$ and for any $s \in [t]$. We prove this by induction in $t \in [n]$. For $t = 1$, the result is true by definition of the acceptance probability and the fact that $x_{1,1} + y_{1,1} = 1$. Let us assume the result is true for $t - 1$, and let us show it for t . First, we have

$$\begin{aligned} \mathbf{P}(O_t | S_t = (t, s)) &= \mathbf{P}(\text{Reach } t, \mathcal{P}(S_t) = \text{offer} | S_t = (t, s)) \\ &= \mathbf{P}(\text{Reach } t | S_t = (t, s)) \mathbf{P}(\mathcal{P}(S_t) = \text{offer} | S_t = (t, s)) \\ &= \mathbf{P}(\text{Reach } t | S_t = (t, s)) \cdot \frac{tx_{t,s}}{\left(\sum_{\sigma=1}^{t-1} y_{t-1,\sigma} + (1-p)x_{t-1,\sigma}\right)}. \end{aligned}$$

Now, we have

$$\begin{aligned} &\mathbf{P}(\text{Reach } t | S_t = (t, s)) \\ &= \mathbf{P}((O_{t-1} \cap \bar{D}_{t-1}) \cup \bar{O}_{t-1} | S_t = (t, s)) \\ &= (1-p) \sum_{\sigma=1}^{t-1} \mathbf{P}(O_{t-1} | S_{t-1} = (t-1, \sigma), S_t = (t, s)) \mathbf{P}(S_{t-1} = (t-1, \sigma) | S_t = (t, s)) \\ &\quad + \sum_{\sigma=1}^{t-1} \mathbf{P}(\bar{O}_{t-1} | S_{t-1} = (t-1, \sigma), S_t = (t, s)) \mathbf{P}(S_{t-1} = (t-1, \sigma) | S_t = (t, s)) \\ &= (1-p) \sum_{\sigma=1}^{t-1} \mathbf{P}(O_{t-1} | S_{t-1} = (t-1, \sigma)) \frac{1}{t-1} \\ &\quad + \sum_{\sigma=1}^{t-1} \mathbf{P}(\bar{O}_{t-1} | S_{t-1} = (t-1, \sigma)) \frac{1}{t-1} \quad (\mathcal{P} \text{ only makes decisions at stage } t-1 \text{ based on } S_{t-1}) \\ &= (1-p) \sum_{\sigma=1}^{t-1} x_{t-1,\sigma} + y_{t-1,\sigma} \quad (\text{induction}). \end{aligned}$$

Note that we used

$$\begin{aligned} \mathbf{P}(S_{t-1} = (t-1, \sigma) | S_t = (t, s)) &= \frac{\mathbf{P}(S_t = (t, s) | S_{t-1} = (t-1, \sigma)) \mathbf{P}(S_{t-1} = (t-1, \sigma))}{\mathbf{P}(S_t = (t, s))} \\ &= \frac{1}{t-1}. \end{aligned}$$

Thus, the induction holds for $\mathbf{P}(O_t | S_t = (t, s)) = tx_{t,s}$. Similarly, for

$$\begin{aligned} \mathbf{P}(\bar{O}_t | S_t = (t, s)) &= \mathbf{P}(\text{Reach } t, \mathcal{P}(S_t) = \text{pass} | S_t = (t, s)) \\ &= \mathbf{P}(\text{Reach } t | S_t = (t, s)) \mathbf{P}(\mathcal{P}(S_t) = \text{pass} | S_t = (t, s)) \\ &= \left(\sum_{\sigma=1}^{t-1} y_{t-1,\sigma} + (1-p)x_{t-1,\sigma} \right) \left(1 - \frac{tx_{t,s}}{\left(\sum_{\sigma=1}^{t-1} y_{t-1,\sigma} + (1-p)x_{t-1,\sigma}\right)} \right) \\ &= ty_{t,s}, \end{aligned}$$

where we used the fact that $(\mathbf{x}, \mathbf{y}) \in \text{POL}$. $\square$

B.2. Stage 2: The Robust Objective

Proposition B.4. Let $\mathcal{P}$ be any policy for **SP-UA** and let $(\mathbf{x}, \mathbf{y}) \in \text{POL}$ be its corresponding vector as in Proposition B.2. Then, for any $k \in [n]$,

$$\mathbf{P}(\mathcal{P} \text{ collects a top } k \text{ candidate}) = \sum_{t=1}^n \sum_{s=1}^t p x_{t,s} \mathbf{P}(R_t \leq k | r_t = s).$$

Proof. We use the same events as in the proof of Proposition B.2. Then,

$$\begin{aligned} \mathbf{P}(\mathcal{P} \text{ collects a top } k \text{ candidate}) &= \sum_{t=1}^n \sum_{s=1}^t \mathbf{P}(O_t \cap D_t \cap \{S_t = (t, s)\} \cap \{R_t \leq k\}) \\ &= \sum_{t=1}^n \sum_{s=1}^t p x_{t,s} \mathbf{P}(R_t \leq k | O_t \cap D_t \cap \{S_t = (t, s)\}) \\ &= \sum_{t=1}^n \sum_{s=1}^t p x_{t,s} \mathbf{P}(R_t \leq k | r_t = s). \end{aligned}$$

Note that R_t only depends on S_t , and $S_t = (t, s)$ is equivalent to $r_t = s$. $\square$

We are ready to prove of Theorem 2.

Theorem B.1 (Theorem 2 Restated). The largest robust ratio γ_n^* corresponds to the optimal value of the LP

$$\begin{aligned} &\max_{\mathbf{x} \geq 0} \gamma \\ &\text{s.t.} \\ (LP)_{n,p} \quad &x_{t,s} \leq \frac{1}{t} \left(1 - p \sum_{\tau=1}^{t-1} \sum_{\sigma=1}^{\tau} x_{\tau,\sigma} \right) \quad \forall t \in [n], s \in [t] \\ &\gamma \leq \frac{p}{1 - (1-p)^k} \sum_{t=1}^n \sum_{s=1}^t x_{t,s} \mathbf{P}(R_t \leq k | r_t = s) \quad \forall k \in [n]. \end{aligned}$$

Moreover, given an optimal solution $(\mathbf{x}^*, \gamma_n^*)$ of $(LP)_{n,p}$, the (randomized) policy $\mathcal{P}^*$ that at state (t, s) makes an offer with probability $t x_{t,s}^* / (1 - p \sum_{\tau=1}^{t-1} \sum_{\sigma=1}^{\tau} x_{\tau,\sigma}^*)$ is γ_n^* -robust.

Proof. We have

$$\begin{aligned} \gamma_n^* &= \max_{\mathcal{P}} \min_{k \in [n]} \frac{\mathbf{P}(\mathcal{P} \text{ collects a candidate with rank } \leq k)}{1 - (1-p)^k} \\ &= \max_{(\mathbf{x}, \mathbf{y}) \in \text{POL}} \min_{k \in [n]} \frac{p \sum_{t=1}^n \sum_{s=1}^t x_{t,s} \mathbf{P}(R_t \leq k | r_t = s)}{1 - (1-p)^k} \quad (\text{Propositions B.2, B.3 and B.4}) \end{aligned}$$

Now note that the function $\gamma : (\mathbf{x}, \mathbf{y}) \mapsto \min_{k \in [n]} \frac{p \sum_{t=1}^n \sum_{s=1}^t x_{t,s} \mathbf{P}(R_t \leq k | r_t = s)}{1 - (1-p)^k}$ is constant in $\mathbf{y}$. Thus any point $(\mathbf{x}, \mathbf{y})$ satisfying Constraints (B.4) and (B.5) has an equivalent point in $(\mathbf{x}', \mathbf{y}') \in \text{POL}$ with $\mathbf{x}' = \mathbf{x}$, $\mathbf{y}' \geq \mathbf{y}$ so all constraints tighten and the objective of γ is the same for both points. Thus, γ_n^* equals the optimal value of the LP (P') :

$$\begin{aligned} &\max_{\mathbf{x} \geq 0} \gamma \\ &\text{s.t.} \\ (P') \quad &x_{1,1} + y_{1,1} \leq 1 \\ &x_{t,s} + y_{t,s} \leq \frac{1}{t} \left(\sum_{\sigma=1}^{t-1} y_{t-1,\sigma} + (1-p)x_{t-1,\sigma} \right) \quad \forall t \in [n], s \in [t] \\ &\gamma \leq \frac{p \sum_{t=1}^n \sum_{s=1}^t x_{t,s} \mathbf{P}(R_t \leq k | r_t = s)}{1 - (1-p)^k} \quad \forall k \in [n], \end{aligned}$$

where we linearized the objective with the variable γ . By projecting the feasible region of (P') onto the variables $(\mathbf{x}, \gamma)$, we obtain $(LP)_{n,p}$. This is a routine procedure that can be carried out using Fourier-Motzkin (Schrijver [65]), but we skip it here for brevity.

For the second part, we can take an optimal solution $(\mathbf{x}^*, \gamma_n^*)$ and its corresponding point $(\mathbf{x}^*, \mathbf{y}^*) \in \text{POL}$. A routine calculation shows that $1 - p \sum_{\tau < t} \sum_{\sigma=1}^{\tau} x_{\tau,\sigma}^* = \sum_{\sigma=1}^{t-1} y_{t-1,\sigma}^* + (1-p)x_{t-1,\sigma}^*$. Thus, by Proposition B.3, we obtain the optimal policy. $\square$

B.3. Missing Proofs from Section 4: γ_n^* Is Decreasing in n

Proof of Lemma 1. We know that γ_n^* equals the value

$$\begin{aligned} & \min_{\substack{u: [n] \rightarrow \mathbf{R}_+ \\ \sum_i u_i \geq 1}} v_{(1,1)} \\ \text{(DLP)} \quad & \text{s.t.} \\ & v_{(t,s)} = \max \left\{ pU_t(s) + \frac{1-p}{t+1} \sum_{\sigma=1}^{t+1} v_{(t+1,\sigma)}, \frac{1}{t+1} \sum_{\sigma=1}^{t+1} v_{(t+1,\sigma)} \right\} \quad \forall t \in [n], s \in [t] \\ & v_{(n+1,s)} = 0 \quad \forall s \in [n+1], \end{aligned}$$

where $U_t(s) = \sum_{i=1}^n \frac{u_i}{1-(1-p)^i} \mathbf{P}(R_t \in [k] | r_t = s) = \sum_{j=1}^n \sum_{k \geq j} \left(\frac{u_k}{1-(1-p)^k} \right) \mathbf{P}(R_t = j | r_t = s)$. Thus, the utility collected by the policy if it collects a candidate with rank i is $U_i = \sum_{k \geq i} \frac{u_k}{1-(1-p)^k}$. Let $u: [n] \rightarrow [0, 1]$ such that $\sum_{i=1}^n u_i = 1$ and extend u to $\hat{u}: [n+1] \rightarrow [0, 1]$ by $\hat{u}_{n+1} = 0$ and define $\hat{U}_t(s)$ accordingly. Consider the optimal policy that solves the program

$$\hat{v}_{(t,s)} = \max \left\{ p\hat{U}_t(s) + \frac{1-p}{t+1} \sum_{\sigma=1}^{t+1} \hat{v}_{(t+1,\sigma)}, \frac{1}{t+1} \sum_{\sigma=1}^{t+1} \hat{v}_{(t+1,\sigma)} \right\}, \quad \forall t \in [n+1], s \in [t]$$

with the boundary condition $\hat{v}_{(n+2,s)} = 0$ for all $s \in [n+2]$. Call this policy $\hat{\mathcal{P}}$. Note that when policy $\hat{\mathcal{P}}$ collects a candidate with rank i , then it gets a utility of

$$\hat{U}_i = \sum_{k \geq i} \frac{\hat{u}_k}{1-(1-p)^k} = \sum_{k \geq i} \frac{u_k}{1-(1-p)^k} = U_i,$$

for $i \leq n$ and $\hat{U}_{n+1} = 0$. By the choice of $\hat{\mathcal{P}}$, the expected utility collected by $\hat{\mathcal{P}}$ is $\text{val}(\hat{\mathcal{P}}) = \hat{v}_{(1,1)}$. We can obtain a policy $\mathcal{P}$ for n elements out of $\hat{\mathcal{P}}$ by simulating an entry of $n+1$ elements as follows. Policy $\mathcal{P}$ randomly selects a time $t^* \in [n+1]$ and its availability b : we set $b=0$ (unavailable) with probability $1-p$ and $b=1$ (available) with probability p . Now, on an input of length n , the policy $\mathcal{P}$ will squeeze an item of rank $n+1$ in position t^* and it will run the policy $\hat{\mathcal{P}}$ in this input, simulating appropriately the new partial ranks. That is, before stage t^* , policy $\mathcal{P}$ behaves exactly as $\hat{\mathcal{P}}$ in the original input of $\mathcal{P}$. When the policy leaves the stage t^*-1 to transition to stage t^* , then the policy $\mathcal{P}$ simulates the simulated candidate t^* (with real rank $n+1$) that $\hat{\mathcal{P}}$ would have received and does the following: ignores the candidate and moves to stage t^* if the simulated candidate is unavailable ($b=0$) or if $\hat{\mathcal{P}}((t^*, t^*)) = \text{pass}$, whereas if $\hat{\mathcal{P}}((t^*, t^*)) = \text{offer}$ and the simulated candidate accepts ($b=1$), then the policy $\mathcal{P}$ accepts any candidate from that point on.

Coupling the input of length $n+1$ for $\hat{\mathcal{P}}$ and the input of length n with the random stage t^* for $\mathcal{P}$, we can see that the utilities collected by $\mathcal{P}$ and $\hat{\mathcal{P}}$ coincide. Thus, the optimal utility $v_{(1,1)}$ collected by a policy for n candidates and utilities given by $U: [n] \rightarrow \mathbf{R}_+$, hold $v_{(1,1)} \geq \hat{v}_{(1,1)}$. Given that $\hat{v}_{(1,1)} \geq \gamma_{n+1}^*$, by minimizing over u we obtain $\gamma_n^* \geq \gamma_{n+1}^*$. $\square$

Appendix C. Missing Proofs from Section 5

In this subsection, we show that $|\gamma_n^* - \gamma_\infty^*| \leq \mathcal{O}((\log n)^2 / \sqrt{n})$ for fixed p (Lemma 2). The proof is similar to other infinite approximations of finite models, and we require some preliminary results before showing the result. First, we introduce two relaxations, one for (LP) and one for (CLP). We show that the relaxations have values close to their nonrelaxed versions. After these preliminaries have been introduced, we present the proof of Lemma 2.

Consider the relaxation of (LP) $_{n,p}$ to the top q candidate constraints:

$$\begin{aligned} \gamma_{n,q}^* &= \max_{x \geq 0} \gamma \\ \text{(LP)}_{n,p,q} \quad & x_{t,s} \leq \frac{1}{t} \left(1 - p \sum_{\tau < t} \sum_{s'=1}^{\tau} x_{\tau,s'} \right) \quad \forall t, s, \end{aligned} \tag{C.1}$$

$$\gamma \leq \frac{p}{(1-(1-p)^q)} \sum_{t=1}^n \sum_{s=1}^t x_{t,s} \mathbf{P}(R_t \in [k] | r_t = s) \quad \forall k \in [q]. \tag{C.2}$$

Note that $\gamma_n^* \leq \gamma_{n,q}^*$ because (LP) $_{n,p,q}$ is a relaxation of (LP) $_{n,p}$. The following result gives a bound on γ_n^* compared with $\gamma_{n,q}^*$.

Proposition C.1. For any $q \in [n]$, $\gamma_n^* \geq (1 - (1-p)^q) \gamma_{n,q}^*$.

Proof. Let $(\mathbf{x}, \gamma_{n,q}^*)$ be an optimal solution of (LP) $_{n,p,q}$. Let $f_k = \frac{p}{1-(1-p)^q} \sum_{t=1}^n \sum_{s=1}^t x_{t,s} \mathbf{P}(R_t \in [k] | r_t = s)$. Then, $\gamma_{n,q}^* = \min_{k=1, \dots, q} f_k$. It is enough to show that $f_i \geq \left(\frac{1-(1-p)^q}{1-(1-p)^n} \right) f_q$ for $i \geq q$ because $\gamma_n^* \geq \min_{i=1, \dots, n} f_i \left(\frac{1-(1-p)^q}{1-(1-p)^n} \right) \min_{i=1, \dots, q} f_i \geq (1 - (1-p)^q) \gamma_{n,q}^*$.

For any j we have

$$f_{j+1} = \frac{p}{1 - (1-p)^{j+1}} \sum_{t=1}^n \sum_{s=1}^t x_{t,s} \mathbf{P}(R_t \in [j+1] | r_t = s) \geq \left(\frac{1 - (1-p)^j}{1 - (1-p)^{j+1}} \right) f_j.$$

Thus, iterating this for $j > q$, we get $f_j \geq \left(\frac{1 - (1-p)^q}{1 - (1-p)^j} \right) f_q$, and we obtain the desired result. $\square$

Likewise, we consider the relaxation of $(CLP)_p$ to the top q candidates:

$$\gamma_{\infty,q}^* = \max_{\substack{\alpha: [0,1] \times \mathbb{N} \rightarrow [0,1] \\ \gamma \geq 0}} \gamma$$

$$(CLP)_{p,q} \quad \alpha(t, s) \leq \frac{1}{t} \left(1 - p \int_0^t \sum_{\sigma \geq 1} \alpha(\tau, \sigma) d\tau \right), \quad \forall t \in [0, 1], s \geq 1, \quad (C.3)$$

$$\gamma \leq \frac{p \int_0^1 \sum_{s \geq 1} \alpha(t, s) \sum_{\ell=s}^k \binom{\ell-1}{s-1} t^s (1-t)^{\ell-s} dt}{(1 - (1-p)^k)}, \quad \forall k \in [q]. \quad (C.4)$$

We have $\gamma_{\infty,q}^* \geq \gamma_{\infty}^*$ and we have the approximate converse

Proposition C.2. For any $q \geq 1$, $\gamma_{\infty}^* \geq (1 - (1-p)^q) \gamma_{\infty,q}^*$.

Proof. Let $(\alpha, \gamma_{\infty,q}^*)$ be a feasible solution of $(CLP)_{p,q}$. Let

$$f_k = \frac{p}{1 - (1-p)^k} \int_0^1 \sum_{s \geq 1} \alpha(t, s) \sum_{\ell=s}^k \binom{\ell-1}{s-1} t^s (1-t)^{\ell-s} dt$$

Assume that $\gamma_{\infty,q}^* \leq \min_{k \leq q} f_k$. As in the previous proof, we aim to show that $f_i \geq (1 - (1-p)^q) f_q$ for $i \geq q$, because this will imply $\gamma_{\infty}^* \geq (1 - (1-p)^q) \gamma_{\infty,q}^*$ for any $(\alpha, \gamma_{\infty,q}^*)$ feasible for $(CLP)_{p,q}$.

Now, for any j , we have

$$\begin{aligned} f_{j+1} &= \frac{p}{1 - (1-p)^{j+1}} \int_0^1 \sum_{s \geq 1} \alpha(t, s) \sum_{\ell=s}^{j+1} \binom{\ell-1}{s-1} t^s (1-t)^{\ell-s} dt \\ &\geq \frac{p}{1 - (1-p)^{j+1}} \int_0^1 \sum_{s \geq 1} \alpha(t, s) \sum_{\ell=s}^j \binom{\ell-1}{s-1} t^s (1-t)^{\ell-s} dt \\ &= \left(\frac{1 - (1-p)^j}{1 - (1-p)^{j+1}} \right) f_j. \end{aligned}$$

Iterating the inequality until reaching q , we deduce that for any $j \geq q$, we have $f_j \geq \left(\frac{1 - (1-p)^q}{1 - (1-p)^j} \right) f_q$. From here, the result follows. $\square$

Remark C.1. If we set $q = (\log n)/p$, both results imply that $\gamma_n^* \geq (1 - 1/n) \gamma_{n,q}^*$ and $\gamma_{\infty}^* \geq (1 - 1/n) \gamma_{\infty,q}^*$. Thus, we lose at most $1/n$ by restricting the analysis to the top q candidates.

Proposition C.3. There is n_0 such that for $n \geq n_0$, for any t such that $\sqrt{n} \log n \leq t \leq n - \sqrt{n} \log n$, $\ell \leq \log(n)/p$ and $\ell \geq s$ it holds that for any $\tau \in [t/n, (t+1)/n]$ we have

$$1 - \frac{10}{p\sqrt{n}} \leq \frac{\binom{\ell-1}{s-1} \binom{n-\ell}{t-s} / \binom{n}{t}}{\binom{\ell-1}{s-1} \tau^s (1-\tau)^{\ell-s}} \leq 1 + \frac{10}{p\sqrt{n}}.$$

Proof. We only need to show that

$$1 - \frac{10}{p\sqrt{n}} \leq \frac{\binom{n-\ell}{t-s} / \binom{n}{t}}{\tau^s (1-\tau)^{\ell-s}} \leq 1 + \frac{10}{p\sqrt{n}}.$$

We have

$$\begin{aligned}
& \frac{\binom{n-\ell}{t-s}}{\binom{n}{t}} \\
&= \frac{t!}{(t-s)!} \frac{1}{n!} \frac{(n-\ell)!(n-t)!}{(n-t-(\ell-s))!} \\
&= \left(\frac{t \cdot (t-1) \cdots (t-s+1)}{n \cdot (n-1) \cdots (n-s+1)} \right) \left(\frac{(n-t)(n-t-1) \cdots (n-t-(\ell-s)+1)}{(n-s)(n-s-1) \cdots (n-s-(\ell-s)+1)} \right) \\
&= \underbrace{\left(\frac{t}{n} \right)^s \left(1 - \frac{t}{n} \right)^{\ell-s}}_A \underbrace{\left(\frac{1 \cdot (1 - \frac{1}{t}) \cdots (1 - \frac{s-1}{t})}{1 \cdot (1 - \frac{1}{n}) \cdots (1 - \frac{s-1}{n})} \right)}_B \underbrace{\left(\frac{1 \cdot (1 - \frac{1}{n-t}) \cdots (1 - \frac{(\ell-s)-1}{n-t})}{(1 - \frac{s}{n}) \cdot (1 - \frac{s+1}{n}) \cdots (1 - \frac{s+(\ell-s)-1}{n})} \right)}_C.
\end{aligned}$$

We bound terms A , B , and C separately. Given that $s \leq \ell$ and we are assuming that $\ell \leq (\log n)/p$ and $t \geq \sqrt{n} \log n$ for n large, then we will implicitly use that $s, \ell \leq \min\{t/2, n/2\}$.

Claim C.1. It holds $(1 - 4/(p\sqrt{n}))\tau^s(1 - \tau)^{\ell-s} \leq A = (t/n)^s(1 - t/n)^{\ell-s} \leq (1 + 4/(p\sqrt{n}))\tau^s(1 - \tau)^{\ell-s}$.

Proof. For the upper bound, we have

$$\begin{aligned}
\left(\frac{t}{n} \right)^s \left(1 - \frac{t}{n} \right)^{\ell-s} &\leq \tau^s \left(1 - \tau + \frac{1}{n} \right)^{\ell-s} && (\tau \in [t/n, (t+1)/n]) \\
&= \tau^s (1 - \tau)^{\ell-s} \left(1 + \frac{1}{(1 - \tau)n} \right)^{\ell-s} \\
&\leq \tau^s (1 - \tau)^{\ell-s} e^{(\ell-s)/((1-\tau)n)} \\
&\leq \tau^s (1 - \tau)^{\ell-s} e^{\ell/(n-t-1)} \\
&\leq \tau^s (1 - \tau)^{\ell-s} \left(1 + 2 \frac{\ell}{n-t-1} \right) && (\text{Using } e^x \leq 1 + 2x \text{ for } x \in [0, 1]).
\end{aligned}$$

The upper bound now follows by using the information over ℓ and t and that $(1 + 2\ell/(n-t-1)) \leq 1 + 2(\log n)p^{-1}(\sqrt{n} \log n - 1)^{-1} \leq 1 + 4/(p\sqrt{n})$ for n large.

For the lower bound, we have

$$\begin{aligned}
\left(\frac{t}{n} \right)^s \left(1 - \frac{t}{n} \right)^{\ell-s} &\geq \left(\tau - \frac{1}{n} \right)^s (1 - \tau)^{\ell-s} \\
&= \tau^s (1 - \tau)^{\ell-s} \left(1 - \frac{1}{\tau n} \right)^s \\
&\geq \tau^s (1 - \tau)^{\ell-s} e^{-\frac{s}{\tau n-1}} \\
&\geq \tau^s (1 - \tau)^{\ell-s} (1 - s/(\tau n - 1)).
\end{aligned}$$

Given that $s/(\tau n - 1) \leq \log n/(p(t-1)) \leq 2/(p\sqrt{n})$ for n large, the lower bound follows. $\square$

Claim C.2. We have $1 - 2s^2/t \leq B \leq 1 + 2s^2/n$.

Proof. For the upper bound, we upper bound the denominator

$$\begin{aligned}
\left(\frac{1 \cdot (1 - \frac{1}{t}) \cdots (1 - \frac{s-1}{t})}{1 \cdot (1 - \frac{1}{n}) \cdots (1 - \frac{s-1}{n})} \right) &\leq \frac{1}{1 \cdot (1 - \frac{1}{n}) \cdots (1 - \frac{s-1}{n})} \\
&\leq e^{\sum_{k=1}^{s-1} k/(n-k)} \\
&\leq e^{s^2/(n-s)} && (\text{Function } x \mapsto x/(n-x) \text{ is increasing}) \\
&\leq 1 + 2 \frac{s^2}{n}.
\end{aligned}$$

For the lower bound, we lower bound the numerator:

$$\begin{aligned} \left(\frac{1 \cdot \left(1 - \frac{1}{t}\right) \cdots \left(1 - \frac{(s-1)}{t}\right)}{1 \cdot \left(1 - \frac{1}{n}\right) \cdots \left(1 - \frac{(s-1)}{n}\right)} \right) &\geq 1 \cdot \left(1 - \frac{1}{t}\right) \cdots \left(1 - \frac{(s-1)}{t}\right) \\ &\geq e^{-\sum_{k=1}^{s-1} k/(t-k)} \quad (\text{Using } 1 - k/t = (1 + k/(t-k))^{-1} \geq e^{-k/(t-k)}) \\ &\geq 1 - \frac{s^2}{t-s} \geq 1 - 2\frac{s^2}{t}. \quad \square \end{aligned}$$

Claim C.3. We have $1 - 2\ell^2/(n-t) \leq C \leq 1 + 2\ell^2/n$.

Proof. Similar to the previous claim, we bound denominator for an upper bound and numerator for a lower bound.

$$\begin{aligned} \left(\frac{1 \cdot \left(1 - \frac{1}{n-t}\right) \cdots \left(1 - \frac{(\ell-s)-1}{n-t}\right)}{\left(1 - \frac{s}{n}\right) \cdot \left(1 - \frac{(s+1)}{n}\right) \cdots \left(1 - \frac{(s+(\ell-s)-1)}{n}\right)} \right) &\leq \frac{1}{\left(1 - \frac{s}{n}\right) \cdot \left(1 - \frac{(s+1)}{n}\right) \cdots \left(1 - \frac{(s+(\ell-s)-1)}{n}\right)} \\ &\leq e^{\sum_{k=0}^{\ell-s-1} (k+s)/(n-k)} \leq 1 + 2\frac{\ell^2}{n}, \end{aligned}$$

and

$$\begin{aligned} \left(\frac{1 \cdot \left(1 - \frac{1}{n-t}\right) \cdots \left(1 - \frac{(\ell-s)-1}{n-t}\right)}{\left(1 - \frac{s}{n}\right) \cdot \left(1 - \frac{(s+1)}{n}\right) \cdots \left(1 - \frac{(s+(\ell-s)-1)}{n}\right)} \right) &\geq 1 \cdot \left(1 - \frac{1}{n-t}\right) \cdots \left(1 - \frac{(\ell-s)-1}{n-t}\right) \\ &\geq e^{-\sum_{k=0}^{\ell-s-1} k/(n-t-k)} \geq 1 - 2\frac{\ell^2}{n-t}. \quad \square \end{aligned}$$

We can now upper bound ABC as

$$\begin{aligned} ABC &\leq \tau^s (1-\tau)^{\ell-s} \left(1 + \frac{4}{p\sqrt{n}}\right) \left(1 + 2\frac{s^2}{n}\right) \left(1 + 2\frac{\ell^2}{n}\right) \\ &\leq \tau^s (1-\tau)^{\ell-s} \left(1 + \frac{4}{p\sqrt{n}}\right) \left(1 + 2\frac{(\log n)^2}{p^2 n}\right)^2 \quad (\text{Using } t \leq n - \sqrt{n} \log n \text{ and } s \leq \ell \leq (\log n)/p) \\ &\leq \tau^s (1-\tau)^{\ell-s} \left(1 + \frac{4}{p\sqrt{n}}\right) \left(1 + 6\frac{(\log n)^2}{p^2 n}\right) \quad (\text{Using } (1+x)^2 \leq 1+3x \text{ if } x \in [0,1]) \\ &\leq \tau^s (1-\tau)^{\ell-s} \left(1 + \frac{10}{p\sqrt{n}}\right). \end{aligned}$$

Recall that we are assuming p constant and n large; thus, the dominating term is $1/\sqrt{n}$. Similarly, we can lower bound ABC as

$$\begin{aligned} ABC &\geq \tau^s (1-\tau)^{\ell-s} \left(1 - \frac{4}{p\sqrt{n}}\right) \left(1 - 2\frac{s^2}{t}\right) \left(1 - 2\frac{\ell^2}{n-t}\right) \\ &\geq \tau^s (1-\tau)^{\ell-s} \left(1 - \frac{4}{p\sqrt{n}}\right) \left(1 - 2\frac{(\log n)^2}{p^2 t}\right) \left(1 - 2\frac{(\log n)^2}{p^2 (n-t)}\right) \\ &\geq \tau^s (1-\tau)^{\ell-s} \left(1 - \frac{4}{p\sqrt{n}}\right) \left(1 - 2\frac{(\log n)^2}{p^2 \sqrt{n}}\right)^2 \geq \tau^s (1-\tau)^{\ell-s} \left(1 - \frac{10}{p\sqrt{n}}\right). \quad \square \end{aligned}$$

Proof of Lemma 2. We are going to show $|\gamma_n^* - \gamma_\infty^*| \leq \mathcal{O}((\log n)^2/\sqrt{n})$. As we can only guarantee good approximation of the binomial terms in Proposition C.3 for $\ell \leq (\log n)/p$, we need to restrict our analysis to $\gamma_{n,q}^*$ and $\gamma_{\infty,q}^*$ for $q = (\log n)/p$. This is enough because these values are within $1/n$ of γ_n^* and γ_∞^* , respectively, due to Propositions C.1 and C.2 (see Remark C.1).

Before proceeding, we give two technical results that allow us to control an error for values of t not considered by Proposition C.3. The deduction is a routine calculation and it is skipped for brevity.

Claim C.4. For any x feasible for Constraints (C.1) and such that $x_{t,s} = 0$ for $s > q$, we have for $k \leq q$

- $\sum_{t=1}^{\sqrt{n} \log n} \sum_{s=1}^t x_{t,s} \mathbf{P}(R_t \in [k] | r_t = s) \leq 10(\log n)^2/(p\sqrt{n})$.
- $\sum_{t=n-\sqrt{n} \log n}^n \sum_{s=1}^t x_{t,s} \mathbf{P}(R_t \in [k] | r_t = s) \leq 10(\log n)^2/(p\sqrt{n})$.

Claim C.5. For any α feasible for Constraints (C.3), we have for $k \leq q$

- $\int_{1-(\log n)/\sqrt{n}}^1 \sum_{s=1}^k \alpha(\tau, s) \sum_{\ell=s}^k \binom{\ell-1}{s-1} \tau^s (1-\tau)^{\ell-s} d\tau \leq (\log n)^2 / (p\sqrt{n})$.
- $\int_0^{(\log n)/\sqrt{n}} \sum_{s=1}^k \alpha(\tau, s) \sum_{\ell=s}^k \binom{\ell-1}{s-1} \tau^s (1-\tau)^{\ell-s} d\tau \leq (\log n)^2 / (p\sqrt{n})$.

First, we show that $\gamma_{n,q}^* \geq \gamma_{\infty}^* - 40(\log n)^2 / (p\sqrt{n})$. Let (α, γ) be a feasible solution of the continuous LP (CLP)_p. We construct a solution of the (LP)_{n,p,q} as follows. Define

$$x_{t,s} = \frac{t - (1-p)}{t} \int_{(t-1)/n}^{t/n} \alpha(\tau, s) d\tau \quad \forall t \in [n], \forall s \in [t],$$

and $\gamma_{n,q} = \min_{k \leq q} \frac{p}{1-(1-p)^k} \sum_{t=1}^n \sum_{s=1}^t x_{t,s} \mathbf{P}(R_t \in [k] | r_t = s)$. Let us show that $(\mathbf{x}, \gamma_{n,q})$ is feasible for (LP)_{n,p,q}—that is, it satisfies Constraints (C.1) and (C.2). First, for $\tau \in [(t-1)/n, t/n]$ we have

$$\begin{aligned} \tau \alpha(\tau, s) + p \int_{(t-1)/n}^{\tau} \alpha(\tau', s) d\tau' &\leq 1 - p \int_0^{\tau} \sum_{\sigma \geq 1} \alpha(\tau', \sigma) d\tau' + p \int_{(t-1)/n}^{\tau} \alpha(\tau', s) d\tau' \\ &\leq 1 - p \int_0^{(t-1)/n} \sum_{\sigma \geq 1} \alpha(\tau', \sigma) d\tau' \leq 1 - p \sum_{\tau'=1}^{t-1} \sum_{\sigma=1}^{\tau'} x_{\tau', \sigma}. \end{aligned}$$

We now integrate in $[(t-1)/n, t/n]$ on both sides of the inequality. After integration, the RHS equals $(1 - p \sum_{\tau=1}^t \sum_{\sigma=1}^{\tau} x_{\tau, \sigma})/n$. On the LHS, we obtain,

$$\begin{aligned} &\int_{(t-1)/n}^{t/n} \left(\tau \alpha(\tau, s) + p \int_{(t-1)/n}^{\tau} \alpha(\tau', s) d\tau' \right) d\tau \\ &= \frac{t}{n} \int_{(t-1)/n}^{t/n} \alpha(\tau, s) d\tau - (1-p) \int_{(t-1)/n}^{t/n} \left(\frac{t}{n} - \tau \right) \alpha(\tau, s) d\tau \\ &\geq \frac{(t - (1-p))}{n} \int_{(t-1)/n}^{t/n} \alpha(\tau, s) d\tau \quad (\text{Using } t/n - \tau \leq 1/n) \\ &= \frac{t}{n} x_{t,s}. \end{aligned}$$

Thus, Constraints (C.1) hold. By definition of $\gamma_{n,q}$, Constraints (C.2) also hold.

Now, note that for $t \geq \sqrt{n} \log n$, we have

$$x_{t,s} \geq \left(1 - \frac{1}{\sqrt{n} \log n} \right) \int_{(t-1)/n}^{t/n} \alpha(\tau, s) d\tau.$$

Then,

$$\begin{aligned} \gamma_{n,q} &= \min_{k \leq q} \frac{p}{1 - (1-p)^k} \sum_{t=1}^n \sum_{s=1}^t x_{t,s} \sum_{\ell=s}^{k \wedge (n-t+s)} \frac{\binom{\ell-1}{s-1} \binom{n-\ell}{t-s}}{\binom{n}{t}} \quad (\text{Definition of } \mathbf{P}(R_t \in [k] | r_t = s)) \\ &\geq \min_{k \leq q} \frac{p}{1 - (1-p)^k} \sum_{t=\sqrt{n} \log n}^{n-\sqrt{n} \log n} \sum_{s=1}^t \int_{(t-1)/n}^{t/n} \alpha(\tau, s) d\tau \sum_{\ell=s}^{k \wedge (n-t+s)} \frac{\binom{\ell-1}{s-1} \binom{n-\ell}{t-s}}{\binom{n}{t}} \left(1 - \frac{1}{\sqrt{n} \log n} \right) \\ &\geq \min_{k \leq q} \frac{p}{1 - (1-p)^k} \sum_{t=\sqrt{n} \log n}^{n-\sqrt{n} \log n} \sum_{s=1}^k \int_{(t-1)/n}^{t/n} \alpha(\tau, s) \sum_{\ell=s}^k \binom{\ell-1}{s-1} \tau^s (1-\tau)^{\ell-s} d\tau \left(1 - \frac{20}{p\sqrt{n}} \right) \\ &\quad (\text{Since } n - t + s \geq \sqrt{n} \log n \geq k \text{ and } t \leq k \text{ and using Proposition C.3}) \\ &= \min_{k \leq q} \frac{p}{1 - (1-p)^k} \sum_{t=\sqrt{n} \log n}^{n-\sqrt{n} \log n} \int_{(t-1)/n}^{t/n} \sum_{\ell=1}^k \sum_{s=1}^{\ell} \alpha(\tau, s) \binom{\ell-1}{s-1} \tau^s (1-\tau)^{\ell-s} d\tau \left(1 - \frac{20}{p\sqrt{n}} \right) \\ &\geq \min_{k \leq q} \left(\frac{p}{1 - (1-p)^k} \int_0^1 \sum_{\ell=1}^k \sum_{s=1}^{\ell} \alpha(\tau, s) \binom{\ell-1}{s-1} \tau^s (1-\tau)^{\ell-s} d\tau - 2 \frac{(\log n)^2}{p\sqrt{n}} \right) \left(1 - \frac{20}{p\sqrt{n}} \right) \\ &\quad (\text{Claim C.5}) \\ &\geq \left(\gamma - 2 \frac{(\log n)^2}{p\sqrt{n}} \right) \left(1 - \frac{10}{p\sqrt{n}} \right) \geq \gamma - \frac{20}{p\sqrt{n}} - 2 \frac{\log n}{\sqrt{n}}. \end{aligned}$$

With this, we have proved that $(x, \gamma_{n,q})$ is a feasible solution of $(LP)_{n,p,q}$ with an objective value $\gamma_{n,q}$ at least $\gamma - 20/p\sqrt{n} - 2 \log n/\sqrt{n}$. Hence, the optimal value of $(LP)_{n,p,q}$ is at least $\gamma - 20/p\sqrt{n} - 2 \log n/\sqrt{n}$. Given that (α, γ) is any feasible solution of $(CLP)_{p,q}$ and $\gamma_{\infty,q}^* \geq \gamma_{\infty}^*$, we obtain $\gamma_{n,q}^* \geq \gamma_{\infty}^* - 40 \log(n)/(p\sqrt{n})$ for n large.

Now, we show that $\gamma_{\infty,q}^* \geq \gamma_{n,q}^* - 40(\log n)^2/(p\sqrt{n})$. Let $(x, \gamma_{n,q})$ be a solution of $(LP)_{n,p,q}$. Let us construct a solution of $(CLP)_{p,q}$. Note that we can assume $x_{t,s} = 0$ for $s > q$, as $(LP)_{n,p,q}$ does not improve its objective function by allocating any mass to these variables. Consider α defined as follows: for $\tau \in [0, 1]$, let

$$\alpha(\tau, s) = \begin{cases} nx_{t,s}(1 - \log(n)/\sqrt{n}) & t = \lceil \tau n \rceil \geq \sqrt{n}, s \leq \min\{t, \log n/p\} \\ 0 & t = \lceil \tau n \rceil < \sqrt{n} \text{ or } s > \min\{t, \log n/p\}. \end{cases}$$

Let $\gamma_{\infty,q} = \min_{k \leq q} \frac{p}{1 - (1-p)^k} \int_0^1 \sum_{s \geq 1} \alpha(\tau, s) \sum_{\ell=s}^k \binom{\ell-1}{s-1} \tau^s (1-\tau)^{\ell-s} d\tau$. We show first that $(\alpha, \gamma_{\infty,q})$ is feasible for $(CLP)_{p,q}$, and for this, it is enough to show that α holds Constraints (C.1). For $\tau < 1/\sqrt{n}$, we have $\alpha(\tau, s) = 0$ for any s ; thus, Constraint (C.3) is satisfied in this case. Let us verify that for $\tau \geq 1/\sqrt{n}$, the constraint also holds. Let $t = \lceil \tau n \rceil$ and $s \geq 1$. Then,

$$\begin{aligned} & \tau \alpha(\tau, s) + p \int_0^\tau \sum_{\sigma \geq 1} \alpha(\tau', \sigma) d\tau' \\ & \leq \tau \alpha(\tau, s) + p \sum_{t'=1}^{t-1} \int_{\frac{t'-1}{n}}^{\frac{t'}{n}} \sum_{s=1}^{t'} \alpha(\tau', s) d\tau' + p \int_{\frac{t-1}{n}}^{\frac{t}{n}} \sum_{\sigma=1}^{\frac{\log n}{p}} \alpha(\tau', \sigma) d\tau' \\ & \leq \left(1 - \frac{\log n}{\sqrt{n}}\right) \left(tx_{t,s} + p \sum_{t'=1}^{t-1} \sum_{s=1}^{t'} x_{t',s} + p \frac{\log n}{pt} \right) \quad \left(\text{Since } x_{t,s} \leq \frac{1}{t} \text{ always} \right) \\ & \leq \left(1 - \frac{\log n}{\sqrt{n}}\right) \left(1 + \frac{\log n}{t}\right) \\ & \leq \left(1 - \frac{\log n}{\sqrt{n}}\right) \left(1 + \frac{\log n}{\sqrt{n}}\right). \quad (t \geq \sqrt{n}) \end{aligned}$$

The last term is < 1 . Thus, $(\alpha, \gamma_{\infty,q})$ is feasible for $(CLP)_{p,q}$. Now,

$$\begin{aligned} \gamma_{\infty,q} &= \min_{k \leq q} \frac{p}{1 - (1-p)^k} \int_0^1 \sum_{s \geq 1} \alpha(\tau, s) \sum_{\ell=s}^k \binom{\ell-1}{s-1} \tau^s (1-\tau)^{\ell-s} d\tau \\ &\geq \min_{k \leq q} \frac{p}{1 - (1-p)^k} \sum_{t=\sqrt{n \log n}}^{n-\sqrt{n \log n}} \int_{\frac{t-1}{n}}^{\frac{t}{n}} \sum_{s \geq 1} \alpha(\tau, s) \sum_{\ell=s}^k \binom{\ell-1}{s-1} \tau^s (1-\tau)^{\ell-s} d\tau \\ &\geq \min_{k \leq q} \frac{p}{1 - (1-p)^k} \sum_{t=\sqrt{n \log n}}^{n-\sqrt{n \log n}} \int_{\frac{t-1}{n}}^{\frac{t}{n}} \sum_{s \geq 1} \alpha(\tau, s) \sum_{\ell=s}^k \frac{\binom{\ell-1}{s-1} \binom{n-\ell}{t-s}}{\binom{n}{t}} d\tau \left(1 - \frac{10}{p\sqrt{n}}\right) \quad (\text{Proposition C.3}) \\ &\geq \min_{k \leq q} \frac{p}{1 - (1-p)^k} \sum_{t=\sqrt{n \log n}}^{n-\sqrt{n \log n}} \sum_{s=1}^t x_{t,s} \sum_{\ell=s}^k \frac{\binom{\ell-1}{s-1} \binom{n-\ell}{t-s}}{\binom{n}{t}} d\tau \left(1 - \frac{\log n}{\sqrt{n}}\right) \left(1 - \frac{10}{p\sqrt{n}}\right) \\ &\geq \left(\min_{k \leq q} \frac{p}{1 - (1-p)^k} \sum_{t=1}^n \sum_{s=1}^t x_{t,s} \sum_{\ell=s}^k \frac{\binom{\ell-1}{s-1} \binom{n-\ell}{t-s}}{\binom{n}{t}} d\tau - 20 \frac{(\log n)^2}{p\sqrt{n}} \right) \left(1 - \frac{\log n}{\sqrt{n}}\right) \left(1 - \frac{10}{p\sqrt{n}}\right) \quad (\text{Claim 4.5}) \\ &\geq \gamma_{n,q} - 40 \frac{(\log n)^2}{p\sqrt{n}}. \end{aligned}$$

With this, we have formed a feasible solution $(\alpha, \gamma_{\infty,q})$ of $(CLP)_{p,q}$. Hence, $\gamma_{\infty,q}^* \geq \gamma_{n,q} - 40(\log n)^2/p\sqrt{n}$. Given that $(x, \gamma_{n,q})$ is any feasible solution of $(LP)_{n,p,q}$, we can optimize over $(x, \gamma_{n,q})$ and obtain the inequality $\gamma_{\infty,q}^* \geq \gamma_{n,q}^* - 40(\log n)^2/(p\sqrt{n})$.

Using Propositions C.1 and C.2, we can conclude that, for n large, $\gamma_{\infty}^* - 50(\log n)^2/(p\sqrt{n}) \leq \gamma_n^* \leq \gamma_{\infty}^* + 50(\log n)^2/(p\sqrt{n})$, where the additional constant factors appear as a by-product of choosing $q = \log n/p$ in both propositions. $\square$

Appendix D. Missing Proofs from Section 6

The following result is the reduction from SP-UA to i.i.d. prophet inequality problem

Lemma D.1. Then, there is an algorithm $\mathcal{A}'$ for the i.i.d. prophet inequality problem that for any $\varepsilon, \delta > 0$ satisfying

$$(1 + \varepsilon)p < 1, \quad n \geq \frac{2}{p\varepsilon^2} \log\left(\frac{2}{\delta}\right), \quad m = \lfloor (1 + \varepsilon)pn \rfloor,$$

ensures

$$\mathbf{E}[\text{Val}(\mathcal{A}')] + \delta \geq \gamma(1 - 4\varepsilon - \delta) \mathbf{E}\left[\max_{i \leq m} X_i\right],$$

for any $X_1, \dots, X_m$ sequence of i.i.d. random variables with support in $[0, 1]$, where $\text{Val}(\mathcal{A}')$ is the profit obtained by $\mathcal{A}'$ from the sequence of values $X_1, \dots, X_m$ in the prophet problem.

Proof. The input of the i.i.d. prophet inequality problem corresponds to a known distribution $\mathcal{D}$ with support in $[0, 1]$. The DM sequentially accesses at most m samples from $\mathcal{D}$, and, upon observing one of these values, she has to decide irrevocably whether to take it and stop the process or continue. We are going to use $\mathcal{A}$ to design a strategy for the prophet problem. We assume that the samples from $\mathcal{D}$ are all distinct. Indeed, we can add some small Gaussian noise to the distribution and consider a continuous distribution $\mathcal{D}'$ instead.

Note that $\mathcal{A}$ runs on an input of size n where a fraction p of the candidates accept an offer. We interpret $pn \approx m$ as the set of samples for the prophet inequality problem, while the remaining $(1 - p)n$ items are used as additional information for the algorithm. By concentration bounds, we are going to argue that we only need to run $\mathcal{A}$ in at most $(1 + \varepsilon)pn$ positive samples.

Formally, we proceed as follows. Fix n and $\varepsilon > 0$ and consider the algorithm $\mathcal{B}$ that receives an online input of n numbers $x_1, \dots, x_n$. The algorithm flips n coins with probability of heads p and marks item i as available if the corresponding coins turn out heads. Algorithm $\mathcal{B}$ feeds algorithm $\mathcal{A}$ with the partial rankings given by the ordering given by $x_1, \dots, x_n$. If $\mathcal{A}$ selects a candidate, but the candidate is marked as unavailable, then $\mathcal{B}$ moves to the next item. If $\mathcal{A}$ selects a candidate i and it is marked as available, then the process ends with $\mathcal{B}$ collecting the value x_i . Let us denote by $\text{Val}(\mathcal{B}, x_1, \dots, x_n)$ the value collected by $\mathcal{B}$ in the online input $x_1, \dots, x_n$. Then, we have the following claim.

Claim D.1. $\mathbf{E}_{X_1, \dots, X_n} [\text{Val}(\mathcal{B}, X_1, \dots, X_n)] \geq \gamma \mathbf{E}_{X_1, \dots, X_n} [\max_{i \in S} X_i]$, where S is the random set of items marked as available and $X_1, \dots, X_n$ are n i.i.d. random variables with common distribution $\mathcal{D}$.

Proof. Fix $x_1, \dots, x_n$ points in the support of $\mathcal{D}$. Then, a simple application of Proposition 1 shows

$$\frac{\mathbf{E}_{S, \pi} [\text{Val}(\mathcal{B}, x_{\pi(1)}, \dots, x_{\pi(n)})]}{\mathbf{E}_S [\max_{i \in S} X_i]} \geq \gamma,$$

Note that we need to feed $\mathcal{B}$ with all permutations of $x_1, \dots, x_n$ in order to obtain the guarantee of $\mathcal{A}$. From here, we obtain $\mathbf{E}_{S, \pi} [\text{Val}(\mathcal{B}, x_{\pi(1)}, \dots, x_{\pi(n)})] \geq \gamma \mathbf{E}_S [\max_{i \in S} X_i]$ and the conclusion follows by taking expectation in $X_1 = x_1, \dots, X_n = x_n$. $\square$

For ease of notation, we will refer by $\text{Val}(\cdot)$ to $\text{Val}(\cdot, X_1, \dots, X_n)$. We modify slightly $\mathcal{B}$. Consider $\mathcal{B}'$ that runs normally $\mathcal{B}$ if $|S| \leq (1 + \varepsilon)pn$ or simply return 0 value if $|S| > (1 + \varepsilon)pn$. Then, we have

Claim D.2. Let $\varepsilon, \delta > 0$. For $n \geq 2 \log(2/\delta)/(p\varepsilon^2)$, we have $\mathbf{E}_{X_1, \dots, X_n} [\max_{i \in S} X_i] \geq (1 - \delta) \mathbf{E}[\max_{i \leq \lfloor (1 - \varepsilon)pn \rfloor} X_i]$ and $\mathbf{E}[\text{Val}(\mathcal{B}')] + \delta \geq \mathbf{E}[\text{Val}(\mathcal{B})]$.

Proof. Using standard Chernoff concentration bounds (see, for instance, Boucheron et al. [11]), we get $\mathbf{P}_S(|S| - pn| \geq \varepsilon pn) \leq 2e^{-pn\varepsilon^2/2} = \delta$. Hence, for $n \geq 2 \log(2/\delta)/(p\varepsilon^2)$, we can guarantee that

$$\mathbf{E}_{X_1, \dots, X_n} \left[\max_{i \in S} X_i \right] \geq (1 - \delta) \mathbf{E} \left[\max_{i \leq \lfloor (1 - \varepsilon)pn \rfloor} X_i \right].$$

For the second part we have $\mathbf{E}[\text{Val}(\mathcal{B})] \leq \delta + \mathbf{E}[\text{Val}(\mathcal{B}) | |S| \leq (1 + \varepsilon)pn] = \delta + \mathbf{E}[\text{Val}(\mathcal{B}')]$. $\square$

Claim D.3. For any $\varepsilon > 0$, we have $\mathbf{E}[\max_{i \leq \lfloor (1 - \varepsilon)pn \rfloor} X_i] \geq (1 - \varepsilon)^2 \mathbf{E}[\max_{i \leq \lfloor (1 + \varepsilon)pn \rfloor} X_i]$.

Proof. Given that $\mathbf{P}(\max_{i \leq k} X_i \leq x) = \mathbf{P}(X_1 \leq x)^k$, then we have

$$\frac{\mathbf{E}[\max_{i \leq \lfloor (1 - \varepsilon)pn \rfloor} X_i]}{\mathbf{E}[\max_{i \leq \lfloor (1 + \varepsilon)pn \rfloor} X_i]} = \frac{\int_0^\infty (1 - \mathbf{P}(X_1 \leq x)^{pn(1 - \varepsilon)}) dx}{\int_0^\infty (1 - \mathbf{P}(X_1 \leq x)^{pn(1 + \varepsilon)}) dx} \geq \inf_{x \geq 0} \frac{1 - \mathbf{P}(X_1 \leq x)^{pn(1 - \varepsilon)}}{1 - \mathbf{P}(X_1 \leq x)^{pn(1 + \varepsilon)}} \geq \inf_{v \in [0, 1]} f(v),$$

where $f(v) = (1 - v^{1 - \varepsilon})/(1 - v^{1 + \varepsilon})$. Now, the conclusion follows by using the fact that the function f is nonincreasing and that $\inf_{v \in [0, 1]} f(v) = \lim_{v \rightarrow 1} f(v) = (1 - \varepsilon)/(1 + \varepsilon) \geq (1 - \varepsilon)^2$. $\square$

Putting together these two claims, we obtain an algorithm that checks at most $(1 + \varepsilon)pn$ items and guarantees

$$\gamma(1 - \varepsilon)^2(1 - \delta)\mathbf{E}\left[\max_{i \leq (1+\varepsilon)pn} X_i\right] \leq \mathbf{E}[\text{Val}(\mathcal{B}')] + \delta.$$

Now, fix $\varepsilon > 0$ small enough such that $(1 + \varepsilon)p < 1$. We know that the set $\{\lfloor (1 + \varepsilon)pn \rfloor\}_{n \geq 1}$ contains all nonnegative integers. Thus, for $n \geq 2 \log(2/\delta)/(p\varepsilon^2)$ algorithm $\mathcal{B}'$ in an input of length $m = \lfloor (1 + \varepsilon)pn \rfloor$ guarantees

$$\gamma(1 - \varepsilon)^2(1 - \delta)\mathbf{E}\left[\max_{i \leq m} X_i\right] \leq \mathbf{E}[\text{Val}(\mathcal{B}')] + \delta,$$

for any distribution $\mathcal{D}$ with support in $[0, 1]$. This finishes the proof. $\square$

The next result uses notation from the work by Hill and Kertz. For the details we refer the reader to the work, Hill and Kertz [43]. The result states that there is a hard instance for the i.i.d. prophet inequality problem where $\mathbf{E}[\max_{i \leq m} X_i]$ is away from 0 by a quantity at least $1/m^3$. The importance of this reformulation of the result by Hill and Kertz is that $e^{-\Theta(n)}/\mathbf{E}[\max_{i \leq m} X_i] \rightarrow 0$, which is what we needed to show that $\gamma \leq 1/\beta$. Recall that $\beta \approx 1.341$ is the unique solution of the integral equation $\int_0^1 (y(1 - \log y) + \beta - 1)^{-1} dy = 1$ (Kertz [48]).

Proposition D.1 (Reformulation of by Hill and Kertz [43, Proposition 4.4]). *Let a_m be the sequence constructed by Hill and Kertz—that is, such that $a_m \rightarrow \beta$, and for any sequence of i.i.d. random variables $X_1, \dots, X_m$ with support in $[0, 1]$ we have*

$$\mathbf{E}\left[\max_{i \leq m} X_i\right] \leq a_m \sup\{\mathbf{E}[X_t] : t \in T_m\},$$

where T_m is the set of stopping rules for $X_1, \dots, X_m$. Then, for m large enough, there is a sequence of i.i.d. random variables $\hat{X}_1, \dots, \hat{X}_m$ with support in $[0, 1]$ such that

- $\mathbf{E}[\max_{i \leq m} \hat{X}_i] \geq 1/m^3$, and
- $\mathbf{E}[\max_{i \leq m} \hat{X}_i] \geq (a_m - 1/m^3) \sup\{\mathbf{E}[\hat{X}_t] : t \in T_m\}$.

Proof. In Hill and Kertz [43, proposition 4.4], it is shown that that for any ε' sufficiently small, there is a random variable $\hat{X}$ with $\hat{p}_0 = \mathbf{P}(\hat{X} = 0)$, $\hat{p}_j = \mathbf{P}(\hat{X} = V_j(\hat{X}))$ for $j = 0, \dots, m-2$, $\mathbf{P}(\hat{X} = V_{m-1}(\hat{X})) = \hat{p}_{m-1} - \varepsilon'$ and $\mathbf{P}(\hat{X} = 1) = \varepsilon'$ such that $\mathbf{E}[\max_{i \leq m} \hat{X}_i] \geq (a_m - \varepsilon') \sup\{\mathbf{E}[\hat{X}_t] : t \in \hat{T}_m\}$, where $\hat{X}_1, \dots, \hat{X}_m$ are m independent copies of $\hat{X}$. Here, $V_j(\hat{X}) = \mathbf{E}[\hat{X} \wedge \mathbf{E}[V_{j-1}(\hat{X})]]$ corresponds to the optimal value computed via dynamic programming, and one can show that $\sup\{\mathbf{E}[\hat{X}_t] : t \in \hat{T}_m\} = V_m(\hat{X})$ (see Hill and Kertz [43, lemma 2.1]). We only need to show that we can choose $\varepsilon' = 1/m^3$. The probabilities $\hat{p}_0, \dots, \hat{p}_{m-1}$ are computed as follows: Let $\hat{s}_j = (\eta_{j,m}(\alpha_m))^{1/m}$ for $j = 1, \dots, n-2$, where $\alpha_m \in (0, 1)$ is the (unique) solution of $\eta_{m-1,m}(\alpha_m) = 1$, then $\hat{p}_0 = \hat{s}_0$, $\hat{p}_j = \hat{s}_j - \hat{s}_{j-1}$ for $j = 1, \dots, n-2$ and $\hat{p}_{n-1} = 1 - \hat{s}_{n-2}$. One can show that $\hat{s}_{m-2} = (1 - 1/m)^{1/(m-1)}(1 - \alpha_m/m)^{1/(m-1)}$ and α_m hold $1/(3e) \leq \alpha_m \leq 1/(e-1)$ (see Hill and Kertz [43, proposition 3.6]). For m large, we have

$$e^{-1/(m-1)} \leq \hat{s}_{m-2} \leq e^{-1/(3em^2)},$$

then $\hat{p}_{m-1} = 1 - \hat{s}_{m-2} \geq 1 - e^{-1/(m-1)} \geq 1/m^2$ for m large. Thus, we can set $\varepsilon' = 1/m^3$ and $\hat{p}_{m-1} - \varepsilon' > 0$, and the rest of the proof follows. Furthermore, $\mathbf{E}[\max_{i \leq m} X_i] \geq \varepsilon' \cdot 1 \geq 1/m^3$. $\square$

Appendix E. Missing Proofs from Section 7

Proof of Lemma 3. For $p \geq p^*$ and $\ell = 0, 1, \dots, 4$, we calculate tight lower bounds for the expression in the left-hand side of the inequality in the claim, and we show that these lower bounds are at least 1, with the lower bound attaining equality with 1 for $\ell = 1, 2$. For $\ell \geq 5$, we can generalize the previous bounds and show a universal lower bound of at least 1.

- For $\ell = 0$, we have

$$\int_{p^{1/(1-p)}}^1 \frac{1}{t^p} dt = \frac{1}{1-p}(1-p) = 1 = (1-p)^0.$$

- For $\ell = 1$, we have

$$\int_{p^{1/(1-p)}}^1 \frac{(1-t)}{t^p} dt = 1 - \int_{p^{1/(1-p)}}^1 t^{1-p} dt = 1 - \frac{1}{2-p}(1-p^{(2-p)/(1-p)}).$$

The last value is at least $1 - p$ if and only if (iff) $p(2-p) \geq 1 - p^{(2-p)/(1-p)}$ iff $p^{(2-p)/(1-p)} \geq (1-p)^2$. The last inequality holds iff $p \geq p^* \approx 0.594134$, where p^* is computed numerically by solving $(1-p)^2 = p^{(2-p)/(1-p)}$.

• For $\ell = 2$, we use the approximation $p^{1/(1-p)} \leq (1+p)/(2e)$ that follows from the concavity of the function $p^{1/(1-p)}$ and the first-order approximation of the function at $p = 1$. With this, we can lower bound the integral

$$\begin{aligned}
 \int_{p^{1/(1-p)}}^1 \frac{(1-t)^2}{t^p} dt &\geq \int_{(1+p)/(2e)}^1 \frac{(1-t)^2}{t^p} dt \\
 &= \int_0^{1-(1+p)/(2e)} u^2 (1-u)^{-p} du && \text{(change of variable } u = 1-t) \\
 &\geq \int_0^{1-(1+p)/(2e)} u^2 \left(1 + pu + p(p+1)\frac{u^2}{2}\right) du \\
 &= \frac{1}{3} \left(1 - \frac{1+p}{2e}\right)^3 + \frac{p}{4} \left(1 - \frac{1+p}{2e}\right)^4 + \frac{p(p+1)}{10} \left(1 - \frac{1+p}{2e}\right)^5.
 \end{aligned}$$

(Using the series $(1-u)^{-p} = \sum_{k \geq 0} \binom{-p}{k} (-u)^k$)

By solving the polynomial, we see that the last expression is $\geq (1-p)^2$ if and only if $p \geq 0.585395$; thus, the inequality holds for $p \geq p^*$.

• For $\ell = 3, 4$, we can use a similar approach to get

$$\int_{p^{1/(1-p)}}^1 \frac{(1-t)^\ell}{t^p} dt \geq \frac{1}{\ell+1} \left(1 - \frac{1+p}{2e}\right)^{\ell+1} + \frac{p}{\ell+2} \left(1 - \frac{1+p}{2e}\right)^{\ell+2}.$$

The last expression is $\geq (1-p)^\ell$ for $\ell = 3, 4$ if and only if $p \geq 0.559826$.

• For $\ell \geq 5$, we have

$$\int_{p^{1/(1-p)}}^1 \frac{(1-t)^\ell}{t^p} dt \geq \frac{(1 - (1+p)/(2e))^{\ell+1}}{\ell+1}.$$

We show that $(1 - (1+p)/(2e))^{\ell+1}/(\ell+1) \geq (1-p)^\ell$. This is equivalent to

$$\left(\frac{1 - (1+p)/(2e)}{1-p}\right)^\ell \left(1 - \frac{1+p}{2e}\right) \geq \ell+1.$$

Note that the function $f(p) = (1 - (1+p)/(2e))/(1-p)$ is increasing because $f'(p) = (1 - 1/e)/(1-p)^2 > 0$. For $\bar{p} = (2e-1)/(4e-3) \approx 0.56351$, we have $f(\bar{p}) = 2 - 1/e$. Thus, for $p \geq p^* > \bar{p}$ and $\ell = 5$ we have $f(p)^5 (1 - (1+p)/(2e)) \geq (2 - 1/e)^5 (1 - 1/e) \geq 7.32 \geq 6$. By an inductive argument, we can show that $f(p)^\ell (1 - (1+p)/(2e)) \geq \ell+1$ for any $\ell \geq 5$, and this finishes the proof. $\square$

Proof of Lemma 7. During the proof, we assume that $1/p \notin \mathbb{N}$. This is an assumption that is easy to remove with a density argument. We divide the proof into a series of propositions and lemmas.

Taking logarithm on both sides of Identity (13), we obtain

$$\log t_{k+1} - \log t_k = \frac{1}{1-kp} \log \left(\frac{A_k(1-p)}{A_{k-1}} \right).$$

From here, we obtain

$$\begin{aligned}
 \log t_{\lfloor 1/p \rfloor + 1} - \log t_2 &= \sum_{j=2}^{\lfloor 1/p \rfloor} \frac{1}{1-jp} \log \left(\frac{A_j(1-p)}{A_{j-1}} \right) \\
 &= \sum_{j=2}^{\lfloor 1/p \rfloor} \frac{1}{1-jp} \int_{A_{j-1}}^{(1-p)A_j} \frac{1}{x} dx,
 \end{aligned}$$

and also

$$\begin{aligned}
 \log t_{k+1} - \log t_{\lfloor 1/p \rfloor + 1} &= \sum_{j=\lfloor 1/p \rfloor}^k \frac{1}{jp-1} \log \left(\frac{A_{j-1}}{A_j(1-p)} \right) \\
 &= \sum_{j=\lfloor 1/p \rfloor}^k \frac{1}{jp-1} \int_{A_j(1-p)}^{A_j} \frac{1}{x} dx.
 \end{aligned}$$

Proposition E.1. We have

1. For $k < 1/p$,

$$\frac{\gamma p(1 - kp)}{A_k(1 - p)} \leq \int_{A_{k-1}}^{A_k(1-p)} \frac{1}{x} dx \leq \frac{\gamma p(1 - kp)}{A_{k-1}}$$

2. For $k > 1/p$,

$$\frac{\gamma p(kp - 1)}{A_{k-1}} \leq \int_{A_k(1-p)}^{A_{k-1}} \frac{1}{x} dx \leq \frac{\gamma p(kp - 1)}{A_k(1 - p)}.$$

Proof. Both results follow by using the monotonicity of $1/x$ and that

$$A_k(1 - p) - A_{k-1} = t_1(1 - p)^{-k+1} + \gamma p k(1 - p) - t_1(1 - p)^{-k+1} - \gamma p(k - 1) = \gamma p(1 - kp). \quad \square$$

The next result shows bound over $\log t_{k+1}$. We use this result to interpret the bounds as Riemann sums.

Proposition E.2 (Bounds on $\log t_{k+1}$). For $k \geq 1/p$, we have

$$\sum_{j=2}^{k-1} \frac{\gamma p}{A_j} \leq \log t_{k+1} - \log t_2 \leq \sum_{j=1}^k \frac{\gamma p}{A_j} + \frac{p}{1-p} \sum_{j=\lfloor 1/p \rfloor + 1}^k \frac{\gamma p}{A_j}.$$

Proof. For the upper bound, we have

$$\begin{aligned} \log t_{k+1} &\leq \log t_2 + \sum_{j=2}^{\lfloor 1/p \rfloor} \frac{\gamma p}{A_{j-1}} + \frac{1}{1-p} \sum_{j=\lfloor 1/p \rfloor + 1}^k \frac{\gamma p}{A_j} \\ &\leq \log t_2 + \sum_{j=1}^k \frac{\gamma p}{A_j} + \frac{p}{1-p} \sum_{j=\lfloor 1/p \rfloor + 1}^k \frac{\gamma p}{A_j}. \end{aligned}$$

For the lower bound, we have

$$\log t_{k+1} \geq \log t_2 + \frac{1}{1-p} \sum_{j=2}^{\lfloor 1/p \rfloor} \frac{\gamma p}{A_j} + \sum_{j=\lfloor 1/p \rfloor + 1}^k \frac{\gamma p}{A_{j-1}} \geq \log t_2 + \sum_{j=2}^{k-1} \frac{\gamma p}{A_j}. \quad \square$$

For $p > 0$ but small enough, $t_1 e^{jp} + \gamma p j \leq A_j \leq t_1 e^{jp/(1-p)} + \gamma p j$. Using this in the bounds of the previous proposition, we obtain

$$\begin{aligned} \int_2^\infty \frac{\gamma p}{t_1 e^{xp/(1-p)} + \gamma p x} dx &\leq \lim_{k \rightarrow \infty} \log t_{k+1} - \log t_2 \\ &\leq \frac{\gamma p}{A_1} + \int_1^\infty \frac{\gamma p}{t_1 e^{xp} + \gamma p x} dx + \frac{p}{1-p} \int_{\lfloor 1/p \rfloor}^\infty \frac{\gamma p}{t_1 e^{px} + \gamma p x} dx. \end{aligned}$$

Note that

$$\frac{p}{1-p} \int_{\lfloor 1/p \rfloor}^k \frac{\gamma p}{t_1 e^{px} + \gamma p x} dx \leq \frac{p}{1-p} \int_1^\infty \frac{\gamma}{t_1} e^{-x} dx = \frac{p}{1-p} \frac{\gamma}{t_1} e^{-1}.$$

Then, taking $p \rightarrow 0$, we obtain

$$\int_0^\infty \frac{\gamma}{t_1 e^x + \gamma x} dx = \lim_{k \rightarrow \infty} \log t_k - \log t_1,$$

where we used that $t_2 = t_1(1 + \gamma p(1 - p)/t_1)^{1/(1-p)} \rightarrow t_1$ when $p \rightarrow 0$. This concludes the proof of Lemma 7. $\square$

References

- [1] Adelman D (2007) Dynamic bid prices in revenue management. *Oper. Res.* 55(4):647–661.
- [2] Agrawal S, Wang Z, Ye Y (2014) A dynamic near-optimal algorithm for online linear programming. *Oper. Res.* 62(4):876–890.
- [3] Alaei S, Hajiaghayi M, Liaghat V (2012) Online prophet-inequality matching with applications to ad allocation. *Proc. 13th ACM Conf. Electronic Commerce* (Association for Computing Machinery, New York), 18–35.
- [4] Altman E (1999) *Constrained Markov Decision Processes*, vol. 7 (CRC Press, Boca Raton, FL).
- [5] Babaioff M, Hartline J, Kleinberg R (2008) Selling banner ads: Online algorithms with buyback. *Fourth Workshop Ad Auctions (Chicago, IL)*.
- [6] Babaioff M, Immorlica N, Kempe D, Kleinberg R (2007) A knapsack secretary problem with applications. Charikar M, Jansen K, Reingold O, Rolim JDP, eds. *Approximation, Randomization, and Combinatorial Optimization. Algorithms and Techniques* (Springer, Berlin), 16–28.

- [7] Babaioff M, Immorlica N, Kempe D, Kleinberg R (2008) Online auctions and generalized secretary problems. *ACM SIGecom Exchanges* 7(2):1–11.
- [8] Bateni M, Hajiaghayi M, Zadimoghaddam M (2013) Submodular secretary problem and extensions. *ACM Trans. Algorithms* 9(4):1–23.
- [9] Beyhaghi H, Golrezaei N, Leme RP, Pál M, Sivan B (2021) Improved revenue bounds for posted-price and second-price mechanisms. *Oper. Res.* 69(6):1805–1822.
- [10] Borodin A, El-Yaniv R (2005) *Online Computation and Competitive Analysis* (Cambridge University Press, Cambridge, UK).
- [11] Boucheron S, Lugosi G, Massart P (2013) *Concentration Inequalities: A Nonasymptotic Theory of Independence* (Oxford University Press, Oxford, UK).
- [12] Bruss FT (1987) On an optimal selection problem of Cowan and Zabczyk. *J. Appl. Probab.* 24(4):918–928.
- [13] Buchbinder N, Jain K, Singh M (2014) Secretary problems via linear programming. *Math. Oper. Res.* 39(1):190–206.
- [14] Campbell G, Samuels SM (1981) Choosing the best of the current crop. *Adv. Appl. Probab.* 13(3):510–532.
- [15] Cesa-Bianchi N, Lugosi G (2006) *Prediction, Learning, and Games* (Cambridge University Press, Cambridge, UK).
- [16] Chan T-HH, Chen F, Jiang SH-C (2014) Revealing optimal thresholds for generalized secretary problem via continuous LP: Impacts on online K-item auction and bipartite K-matching with random arrival order. *Proc. Twenty-Sixth Annu. ACM-SIAM Sympos. Discrete Algorithms* (Society for Industrial and Applied Mathematics, Philadelphia), 1169–1188.
- [17] Cheng H, Cantú-Paz E (2010) Personalized click prediction in sponsored search. *Proc. Third ACM Internat. Conf. Web Search Data Mining* (Association for Computing Machinery, New York), 351–360.
- [18] Cho C-H, Cheon HJ (2004) Why do people avoid advertising on the Internet? *J. Advertising* 33(4):89–97.
- [19] Choi H, Mela CF, Balseiro SR, Leary A (2020) Online display advertising markets: A literature review and future directions. *Inform. Systems Res.* 31(2):556–575.
- [20] Clauset A, Shalizi CR, Newman MEJ (2009) Power-law distributions in empirical data. *SIAM Rev.* 51(4):661–703.
- [21] Correa J, Cristi A, Epstein B, Soto JA (2024) Sample-driven optimal stopping: From the secretary problem to the iid prophet inequality. *Math. Oper. Res.* 49(1):441–475.
- [22] Correa J, Cristi A, Feuilletoy L, Oosterwijk T, Tsigonias-Dimitriadis A (2021) The secretary problem with independent sampling. *Proc. 2021 ACM-SIAM Sympos. Discrete Algorithms SODA* (Society for Industrial and Applied Mathematics, Philadelphia), 2047–2058.
- [23] Cowan R, Zabczyk J (1979) An optimal selection problem associated with the Poisson process. *Theory Probab. Appl.* 23(3):584–592.
- [24] Devanur NR, Hayes TP (2009) The adwords problem: Online keyword matching with budgeted bidders under random permutations. *Proc. 10th ACM Conf. Electronic Commerce* (Association for Computing Machinery, New York), 71–78.
- [25] Devanur NR, Kakade SM (2009) The price of truthfulness for pay-per-click auctions. *Proc. 10th ACM Conf. Electronic Commerce* (Association for Computing Machinery, New York), 99–106.
- [26] Drèze X, Husherr F-X (2003) Internet advertising: Is anybody watching? *J. Interactive Marketing* 17(4):8–23.
- [27] Dütting P, Lattanzi S, Leme RP, Vassilvitskii S (2021) Secretaries with advice. *Proc. 22nd ACM Conf. Econom. Comput.* (Association for Computing Machinery, New York), 409–429.
- [28] Dynkin EB (1963) The optimum choice of the instant for stopping a Markov process. *Soviet Math.* 4:627–629.
- [29] Edelman B, Ostrovsky M, Schwarz M (2007) Internet advertising and the generalized second-price auction: Selling billions of dollars worth of keywords. *Amer. Econom. Rev.* 97(1):242–259.
- [30] Epstein B, Ma W (2024) Selection and ordering policies for hiring pipelines via linear programming. *Oper. Res.*, ePub ahead of print February 28, <https://doi.org/10.1287/opre.2023.0061>.
- [31] Farahat A, Bailey MC (2012) How effective is targeted advertising? *Proc. 21st Internat. Conf. World Wide Web* (Association for Computing Machinery, New York), 111–120.
- [32] Feldman M, Svensson O, Zenklus R (2014) A simple $o(\log \log(\text{rank}))$ -competitive algorithm for the matroid secretary problem. *Proc. Twenty-Sixth annual ACM-SIAM Sympos. Discrete Algorithms* (Society for Industrial and Applied Mathematics, Philadelphia), 1189–1201.
- [33] Ferguson TS (1989) Who solved the secretary problem? *Statist. Sci.* 4(3):282–289.
- [34] Freeman PR (1983) The secretary problem and its extensions: A review. *Internat. Statist. Rev.* 51(2):189–206.
- [35] Freund Y, Schapire RE (1999) Adaptive game playing using multiplicative weights. *Games Econom. Behav.* 29(1–2):79–103.
- [36] Friedgeirsdottir K, Najafi-Asadolahi S (2018) Cost-per-impression pricing for display advertising. *Oper. Res.* 66(3):653–672.
- [37] Gilbert JP, Mosteller F (1966) Recognizing the maximum of a sequence. *J. Amer. Statist. Assoc.* 61(313):35–73.
- [38] Goyal V, Udhwani R (2020) Online matching with stochastic rewards: Optimal competitive ratio via path based formulation. *Proc. 21st ACM Conf. Econom. Comput.* (Association for Computing Machinery, New York), 791.
- [39] Gusein-Zade SM (1966) The problem of choice and the optimal stopping rule for a sequence of independent trials. *Theory Probab. Appl.* 11(3):472–476.
- [40] Haskell WB, Jain R (2013) Stochastic dominance-constrained Markov decision processes. *SIAM J. Control Optim.* 51(1):273–303.
- [41] Haskell WB, Jain R (2015) A convex analytic approach to risk-aware Markov decision processes. *SIAM J. Control Optim.* 53(3):1569–1598.
- [42] Hazan E (2019) Introduction to online convex optimization. Preprint, submitted September 7, <https://arxiv.org/abs/1909.05207>.
- [43] Hill TP, Kertz RP (1982) Comparisons of stop rule and supremum expectations of iid random variables. *Ann. Probab.* 10(2):336–345.
- [44] Hlynka M, Sheahan JN (1988) The secretary problem for a random walk. *Stochastic Process. Appl.* 28(2):317–325.
- [45] Jiang J, Ma W, Zhang J (2023) Tightness without counterexamples: A new approach and new results for prophet inequalities. *Proc. 24th ACM Conf. Econom. Comput.* (Association for Computing Machinery, New York), 909.
- [46] Jiang J, Ma W, Zhang J (2024) Tight guarantees for multiunit prophet inequalities and online stochastic knapsack. *Oper. Res.*, ePub ahead of print June 3, <https://doi.org/10.1287/opre.2022.0309>.
- [47] Kaplan H, Naori D, Raz D (2020) Competitive analysis with a sample and the secretary problem. *Proc. Fourteenth Annu. ACM-SIAM Sympos. Discrete Algorithms* (Society for Industrial and Applied Mathematics, Philadelphia), 2082–2095.
- [48] Kertz RP (1986) Stop rule and supremum expectations of iid random variables: A complete comparison by conjugate duality. *J. Multivariate Anal.* 19(1):88–112.
- [49] Kesselheim T, Tönnis A, Radke K, Vöcking B (2014) Primal beats dual on online packing LPs in the random-order model. *Proc. Forty-Sixth Annu. ACM Sympos. Theory Comput.* (Association for Computing Machinery, New York), 303–312.

- [50] Kleinberg R (2005) A multiple-choice secretary algorithm with applications to online auctions. *Proc. Sixteenth Annu. ACM-SIAM Sympos. Discrete Algorithms* (Society for Industrial and Applied Mathematics, Philadelphia), 630–631.
- [51] Kleywegt AJ, Papastavrou JD (1998) The dynamic and stochastic knapsack problem. *Oper. Res.* 46(1):17–35.
- [52] Korula N, Pál M (2009) Algorithms for secretary problems on graphs and hypergraphs. *Internat. Colloquium Automata Languages Program.* (Springer, Berlin, Heidelberg), 508–520.
- [53] Lachish O (2014) $O(\log \log \text{rank})$ competitive ratio for the matroid secretary problem. *2014 IEEE 55th Annu. Sympos. Foundations Comput. Sci.* (IEEE, Piscataway, NJ), 326–335.
- [54] Lindley DV (1961) Dynamic programming and decision theory. *J. Roy. Statist. Soc. Ser. C Appl. Statist.* 10(1):39–51.
- [55] Lucier B (2017) An economic view of prophet inequalities. *ACM SIGecom Exchanges* 16(1):24–47.
- [56] Manne AS (1960) Linear programming and sequential decisions. *Management Sci.* 6:259–267.
- [57] Marchetti-Spaccamela A, Vercellis C (1995) Stochastic on-line knapsack problems. *Math. Program.* 68(1):73–104.
- [58] Mehta A, Waggoner B, Zadimoghaddam M (2014) Online stochastic matching with unequal probabilities. *Proc. Twenty-Sixth Annu. ACM-SIAM Sympos. Discrete Algorithms* (Society for Industrial and Applied Mathematics, Philadelphia), 1388–1404.
- [59] Perez-Salazar S, Singh M, Toriello A (2022) The iid prophet inequality with limited flexibility. Preprint, submitted October 11, <https://arxiv.org/abs/2210.05634>.
- [60] Presman EL, Mikhailovich Sonin I (1973) The best choice problem for a random number of objects. *Theory Probab. Appl.* 17(4):657–668.
- [61] Pujol E, Hohlfeld O, Feldmann A (2015) Annoyed users: Ads and ad-block usage in the wild. *Proc. 2015 Internet Measurement Conf.* (Association for Computing Machinery, New York), 93–106.
- [62] Puterman ML (2014) *Markov Decision Processes: Discrete Stochastic Dynamic Programming* (John Wiley & Sons, Hoboken, NJ).
- [63] Salem J, Gupta S (2019) Closing the gap: Group-aware parallelization for the secretary problem with biased evaluations. Preprint, submitted September 3, <https://dx.doi.org/10.2139/ssrn.3444283>.
- [64] Salem J, Gupta S (2023) Secretary problems with biased evaluations using partial ordinal information. *Management Sci.* 70(8):5337–5366.
- [65] Schrijver A (1998) *Theory of Linear and Integer Programming* (John Wiley & Sons, Hoboken, NJ).
- [66] Smith MH (1975) A secretary problem with uncertain employment. *J. Appl. Probab.* 12(3):620–624.
- [67] Soto JA (2013) Matroid secretary problem in the random-assignment model. *SIAM J. Comput.* 42(1):178–211.
- [68] Tamaki M (1991) A secretary problem with uncertain employment and best choice of available candidates. *Oper. Res.* 39(2):274–284.
- [69] Torrico A, Toriello A (2022) Dynamic relaxations for online bipartite matching. *INFORMS J. Comput.* 34(4):1871–1884.
- [70] Torrico A, Ahmed S, Toriello A (2018) A polyhedral approach to online bipartite matching. *Math. Program.* 172(1):443–465.
- [71] Vanderbei RJ (1980) The optimal choice of a subset of a population. *Math. Oper. Res.* 5(4):481–486.
- [72] Vanderbei RJ (2012) The postdoc variant of the secretary problem. Technical report, Princeton University, Princeton, NJ.