

Prophet Inequalities Beyond Utilitarian Social Welfare

Daniel Halpern¹, Abhiram Manohara², and Alexandros Psomas²

¹Google Research

²Purdue University

Abstract

In the classical i.i.d. prophet-inequality problem, a single item must be allocated to one of n agents who arrive sequentially, with values drawn independently from a known distribution. When an agent arrives, their value is revealed, and the algorithm must immediately decide whether to allocate the item or continue. The usual objective is to maximize utilitarian welfare: the expected value of the recipient. Guarantees for this objective immediately extend to the problem of allocating m indivisible items to sequentially arriving agents with i.i.d. additive values. Utilitarian welfare, however, ignores how expected utility is distributed across agents. Motivated by a rich literature in fair division, we instead evaluate an online rule by its generalized p -mean welfare. This family includes utilitarian welfare at $p = 1$, Nash welfare (the geometric mean of utilities) at $p = 0$, and egalitarian welfare (the minimum utility) as p tends to $-\infty$.

When the number of items is large, we show that our many-item fair-division problem is captured exactly by a single-item prophet problem in which an online rule is evaluated by the p -mean of the agents' expected utilities. We then characterize this single-item problem; every online rule is weakly Pareto dominated by a quantile-threshold rule, and an optimal egalitarian rule equalizes the agents' expected utilities. We prove that, for every n , the online optimum is at least $\Gamma \approx 0.7059$ times the prophet's egalitarian welfare; by monotonicity of the generalized means, the same guarantee extends to every $p \leq 1$ in the single-item problem. We further prove that, for egalitarian welfare, the optimal ratio converges to Γ as $n \rightarrow \infty$. Thus, asymptotically, optimizing egalitarian rather than utilitarian welfare costs only about four percentage points, relative to the classical utilitarian ratio of 0.7451. Finally, returning to fair division, we prove that the large-item assumption is necessary: when $m = n$, the worst-case competitive ratio converges to zero as $n \rightarrow \infty$ for every $p \leq 0$.

1 Introduction

Consider the following online allocation problem. There are m indivisible items and n agents with additive preferences. Agents arrive sequentially; when an agent arrives, their values for the remaining items are revealed, and we must irrevocably choose a bundle for them. Item values are drawn i.i.d. from a known distribution. When our objective is to maximize utilitarian social welfare (the sum of agents' utilities) this problem decomposes item by item. Each item gives rise to a classical i.i.d. prophet-inequality problem: a single item must be allocated to one of n sequentially arriving agents, and the algorithm must decide whether to allocate after observing each value. Running a prophet algorithm independently on every item therefore gives the same competitive guarantee in the many-item problem. The optimal ratio for the single-item problem converges to approximately 0.7451 as the number of agents grows [HK82, Ker86, CFH⁺17].

Utilitarian welfare, however, ignores how utility is distributed across agents: an allocation may generate high total value, but leave some agents with very little. A rich literature on fair division

asks how to balance efficiency with meaningful utility guarantees for individual agents. A standard way to capture this tradeoff is generalized p -mean welfare [Mou03, BBKS20, BKM22]. This family includes utilitarian welfare at $p = 1$ (the average utility), *Nash* welfare (the geometric mean of agents' utilities) at $p = 0$, and egalitarian welfare (the minimum utility across agents) at $p = -\infty$; as p decreases, the objective places increasing weight on agents with lower utility.

The picture becomes particularly clean when the number of items is large. Fixing n and the value distribution, independently applying a single-item rule to every item makes each agent's utility per item converge to its expected utility under that rule as the number of items grows large. We show that, in this limit, our fair-division problem is captured exactly by a single-item prophet problem in which an online rule is evaluated by the p -mean of the agents' expected utilities.

For the prophet, fairness over expected utilities comes at no cost: allocating to an agent of maximum value and breaking ties uniformly maximizes total value while giving every agent the same expected utility. Online, however, these goals can conflict. Allocating more readily benefits early agents, while waiting preserves opportunities for later agents. Our main question is how much of the prophet's welfare an online rule can retain when the objective places greater weight on agents with lower expected utility.

1.1 Our Contributions

For a distribution $\mathcal{D}$, let $\text{ALG}_n(p; \mathcal{D})$ denote the optimal online p -mean welfare. As shown in Section 2, the prophet's optimal value is the same for every $p \in [-\infty, 1]$; we denote this common value by $\text{OPT}_n(\mathcal{D})$. We study the worst-case competitive ratio

$$\text{CR}_n(p) = \inf_{\mathcal{D}} \frac{\text{ALG}_n(p; \mathcal{D})}{\text{OPT}_n(\mathcal{D})}.$$

Our main result gives a tight answer for egalitarian welfare as the number of agents grows. Let $\Gamma \approx 0.70591334$. We prove that for every $n \geq 2$, $\text{CR}_n(-\infty) \geq \Gamma$, and that this guarantee becomes tight as the number of agents grows: $\lim_{n \rightarrow \infty} \text{CR}_n(-\infty) = \Gamma$ (Theorem 3.1). Thus, even the least-favored agent can be guaranteed more than 70% of their share of the prophet's value. Compared with the classical asymptotic ratio of approximately 0.7451, the price of requiring an equal guarantee across arrival positions is only about four percentage points.

This result also gives guarantees across the entire spectrum of p -mean objectives. The generalized mean of a fixed utility profile is nondecreasing in p , while the offline benchmark (the prophet's performance) is the same for all $p \in [-\infty, 1]$. Hence $\text{CR}_n(p)$ is nondecreasing in p , and we can conclude that $\text{CR}_n(p) \geq \Gamma$ for every $n \geq 2$ and every $p \in [-\infty, 1]$. In fact, the same optimal egalitarian rule achieves this guarantee simultaneously for every $p \in [-\infty, 1]$. Together with the classical upper bound at $p = 1$, this gives, for every fixed $p \in [-\infty, 1]$,

$$0.7059 \approx \Gamma \leq \liminf_{n \rightarrow \infty} \text{CR}_n(p) \leq \limsup_{n \rightarrow \infty} \text{CR}_n(p) \leq \beta_{\text{iid}} \approx 0.7451,$$

where $\beta_{\text{iid}} = \lim_{n \rightarrow \infty} \text{CR}_n(1)$ is the classical i.i.d. prophet constant [HK82, Ker86, CFH⁺17].

Returning to the many-item fair division problem, we show that applying the p -mean to expected utilities gives an exact reduction to the single-item problem: for every m , the optimal online and offline welfares are m times their single-item counterparts, and hence the competitive ratios coincide (Proposition 4.1). When the p -mean is instead applied to realized (ex-post) utilities, the same connection holds asymptotically: for every fixed n , $p \in [-\infty, 1]$, and distribution $\mathcal{D}$, the competitive ratio converges to the corresponding single-item ratio as $m \rightarrow \infty$ (Proposition 4.3). Finally, we prove that this large-item assumption is necessary: when $m = n$, the worst-case ex-post ratio tends to zero for every $p \leq 0$ (Proposition 4.2).

1.2 Related Work

Our work lies at the intersection of prophet inequalities and fair division. The classical prophet inequality, originating with Krengel and Sucheston [KS77], guarantees one half of the prophet’s expected value for independent, nonnegative random variables; this guarantee can be achieved by a single-threshold rule [SC84]. The i.i.d. case admits a strictly better constant. Hill and Kertz characterized the finite-horizon extremal problem [HK82], Kertz identified its asymptotic value [Ker86], and Correa et al. gave a matching algorithm, establishing the tight asymptotic ratio of approximately 0.7451 [CFH⁺17]. This line of work asks how much total value an online rule can guarantee. We instead ask what guarantees are possible for the entire vector of expected utilities across arrival positions.

Our work is closest to Hill and Kennedy [HK90], who developed a framework for optimal stopping with generalized objective functions. In particular, they explicitly considered $\max_{\tau} \min_j \mathbb{E}[X_j \mathbf{1}\{\tau = j\}]$, exactly our egalitarian objective, for general sequences that need not be identically distributed. Their work establishes universal prophet bounds for this generalized objective, and the ordinary prophet inequality, applied to a weighted scalarization of the utility vector, gives a $1/2$ guarantee in the general independent setting. Our work focuses on what the additional i.i.d. structure buys: it leads to a strictly larger sharp asymptotic constant and gives guarantees for the entire family of p -mean objectives. Arsenis and Kleinberg [AK22] pursue a different notion of fairness, requiring identity-independent and time-independent treatment of candidates, and recover a $1/2$ prophet inequality even when both constraints are imposed. Thus, their fairness requirements constrain the stopping rule, whereas ours directly evaluates the resulting vector of agents’ utilities.

This objective-based view of fairness is motivated by a rich literature on fair division, where generalized p -mean welfare is used to interpolate between utilitarian, Nash, and egalitarian welfare [Mou03, BBKS20, BKM22]. Many works on online fair division fix the set of agents and allocate items that arrive online, focusing on various objectives, e.g., minimizing maximum envy [BKP⁺24, BHP25, HPVX25, MPV24, GPT21, HPPZ19, PSZ24, SZ25, ALMT25, KMNP26], maximizing Nash welfare [BGGJ22, GPK21, LGK22, YLK24, HLSW23] and egalitarian welfare [KS22, SHPK22], or, more generally, maximizing the generalized p -mean of utilities [HLSW25, BKM22, YKK26, BX23]. Our model reverses the direction of arrival: i.i.d. agents arrive against a fixed collection of items. Online arrival of agents has been primarily studied for (static) divisible items/resources [KPS14, LLL18, IMMP20, VPF22, SJBY23, BHS26]. The closest work to ours is that of Kulkarni et al. [KMS25], who study ex-post maximin-share guarantees for a fixed set of indivisible items when agents from a known set of additive valuation types arrive adversarially or stochastically. Finally, a related literature fixes the agents and studies repeated resource allocation under dynamic demands or resource availability [VFA⁺23, FAT24, BFT23, FBT25].

1.3 Organization.

Section 2 defines the single-item model. Section 3 states our main results for this model and characterizes the structure of optimal online rules. Section 4 develops the connection to online fair division. Finally, Section 5 presents the main proof ideas and proves our central result. The appendices contain the remaining proofs and technical details.

2 Model

Throughout the paper, all value distributions are supported on $\mathbb{R}_{\geq 0}$ and have finite, positive mean. Fix an integer $n \geq 2$, write $[n] := 1, \dots, n$, and let $\mathcal{D}$ be such a distribution with CDF F . The

agents' values $X_1, \dots, X_n$ are drawn independently from $\mathcal{D}$.

There is one item. The agents arrive in the order $1, \dots, n$. When agent i arrives, we observe X_i and must immediately decide whether to allocate the item to i or continue. The item may be allocated at most once.

Allocation Rules and Feasible Utilities. An *allocation rule* π is a random variable taking values in $[n] \cup \{\emptyset\}$, where $\pi = i$ means that agent i receives the item and $\pi = \emptyset$ means that the item is never allocated. A rule is *online* if its decision at step i may use only $X_1, \dots, X_i$ and internal randomness. An *offline* rule (or a *prophet*) may depend on the entire vector $(X_1, \dots, X_n)$.

Agent i 's expected utility under π is

$$v_i(\pi, \mathcal{D}) := \mathbb{E}[X_i \mathbf{1}\{\pi = i\}].$$

When $\mathcal{D}$ is clear from context, we abbreviate this to $v_i(\pi)$. The resulting utility profile is $\mathbf{v}(\pi, \mathcal{D}) = (v_1(\pi, \mathcal{D}), \dots, v_n(\pi, \mathcal{D}))$. Define

$$\begin{aligned} \mathcal{V}^{\text{on}}(\mathcal{D}) &:= \{\mathbf{v}(\pi, \mathcal{D}) : \pi \text{ is an online rule}\}, \\ \mathcal{V}^{\text{off}}(\mathcal{D}) &:= \{\mathbf{v}(\pi, \mathcal{D}) : \pi \text{ is an offline rule}\}. \end{aligned}$$

Social Welfare Objectives. For finite $p \leq 1$ with $p \neq 0$, define the generalized p -mean welfare of a utility profile $\mathbf{v}$ by

$$w_p(\mathbf{v}) := \left(\frac{1}{n} \sum_{i=1}^n v_i^p \right)^{1/p},$$

where this value is taken to be 0 for $p \leq 0$ when any $v_i = 0$. When $p = 1$, this is utilitarian social welfare, measured as average expected utility: $w_1(\mathbf{v}) = \frac{1}{n} \sum_{i=1}^n v_i$. It can also be continuously extended to $p = 0$ and $p = -\infty$. The limiting cases are Nash welfare (geometric mean), $w_0(\mathbf{v}) = (\prod_{i=1}^n v_i)^{1/n}$, and egalitarian welfare (minimum utility), $w_{-\infty}(\mathbf{v}) = \min_{i \in [n]} v_i$. For every $p \in [-\infty, 1]$, w_p is coordinatewise nondecreasing, concave, and positively homogeneous: $w_p(\lambda \mathbf{v}) = \lambda w_p(\mathbf{v})$ for $\lambda \geq 0$. Moreover, the power-mean inequality gives $w_{p_1}(\mathbf{v}) \geq w_{p_2}(\mathbf{v})$ whenever $p_1 \geq p_2$.

Online and Offline Optima. For each $p \in [-\infty, 1]$, define the optimal online value

$$\text{ALG}_n(p; \mathcal{D}) := \sup_{\mathbf{v} \in \mathcal{V}^{\text{on}}(\mathcal{D})} w_p(\mathbf{v}).$$

Observe that the offline optimum is the same for every $p \in [-\infty, 1]$: for any offline rule, the power-mean inequality gives

$$w_p(\mathbf{v}) \leq w_1(\mathbf{v}) = \frac{1}{n} \sum_{i=1}^n v_i \leq \frac{1}{n} \mathbb{E} \left[\max_{i \in [n]} X_i \right].$$

Conversely, allocating to an agent of maximum value and breaking ties uniformly at random gives every agent expected utility $\mathbb{E}[\max_{i \in [n]} X_i]/n$, achieving the previous upper bound. We denote this common offline value by

$$\text{OPT}_n(\mathcal{D}) := \sup_{\mathbf{v} \in \mathcal{V}^{\text{off}}(\mathcal{D})} w_p(\mathbf{v}) = \frac{1}{n} \mathbb{E} \left[\max_{i \in [n]} X_i \right] \quad \text{for every } p \in [-\infty, 1].$$

Thus the worst-case competitive ratio is

$$\text{CR}_n(p) := \inf_{\mathcal{D}} \frac{\text{ALG}_n(p; \mathcal{D})}{\text{OPT}_n(\mathcal{D})}.$$

Here and throughout, the infimum ranges over distributions on $\mathbb{R}_{\geq 0}$ with finite, positive mean. When $p = 1$, $\text{CR}_n(1)$ reduces to the classic i.i.d. prophet inequality, where it is known that $\lim_{n \rightarrow \infty} \text{CR}_n(1) \approx 0.7451$ [CFH⁺17].

Quantile Representation and Threshold Rules. It is convenient to parameterize values by their quantiles. Let $F^{-1}(q) = \inf\{t \in \mathbb{R} \mid F(t) \geq q\}$ be the quantile function of $\mathcal{D}$. We may assume without loss of generality that the online algorithm observes i.i.d. uniform quantiles $U_1, \dots, U_n \sim \text{Unif}[0, 1]$, where $X_i = F^{-1}(U_i)$ almost surely.¹

A *quantile-threshold rule* is specified by thresholds $q_1, \dots, q_{n-1} \in [0, 1]$. At step $i < n$, it allocates the item to agent i if the item is still available and $U_i \geq q_i$. If the item remains available at step n , it is allocated to agent n . For notational convenience, we may set $q_n = 0$, so that the same description applies to every step.

For a quantile-threshold rule with threshold vector $\mathbf{q} = (q_1, \dots, q_{n-1})$, let $r_i(\mathbf{q})$ denote the probability that the item reaches agent i . Then $r_1(\mathbf{q}) = 1$ and $r_{i+1}(\mathbf{q}) = r_i(\mathbf{q})q_i$. Equivalently, $r_i(\mathbf{q}) = \prod_{j=1}^{i-1} q_j$. When the threshold rule is clear from context, we abbreviate $r_i(\mathbf{q})$ to r_i .

For a distribution $\mathcal{D}$ with CDF F , define $S_{\mathcal{D}}(q) := \mathbb{E}[F^{-1}(U)\mathbf{1}\{U \geq q\}] = \int_q^1 F^{-1}(u) du$ where $U \sim \text{Unif}[0, 1]$. When $\mathcal{D}$ is clear from context, we abbreviate $S_{\mathcal{D}}$ to S . Conditional on the item reaching an agent, threshold q gives that agent expected utility $S(q)$ and leaves the item available for the next agent with probability q . Consequently, for a quantile-threshold rule π , $v_i(\pi, \mathcal{D}) = r_i S_{\mathcal{D}}(q_i)$.

3 Results

For $a, b > 0$, define

$$\rho(a, b) := \frac{1}{\left(1 - \frac{b}{a+b}e^{-a}\right) \left(1 + \frac{1 - e^{-b}}{a}\right)}.$$

and let $\Gamma := \inf_{a, b > 0} \rho(a, b)$. Numerical minimization gives $\Gamma \approx 0.70591334$, with a minimizer near $a \approx 0.81833$ and $b \approx 2.42237$. Our main theorem states that the worst-case competitive ratio for egalitarian welfare is bounded below by Γ for every n and converges to Γ as $n \rightarrow \infty$.

Theorem 3.1. *For every $n \geq 2$, $\text{CR}_n(-\infty) \geq \Gamma$. Moreover, $\lim_{n \rightarrow \infty} \text{CR}_n(-\infty) = \Gamma$.*

Section 5 is devoted to the proof of Theorem 3.1. It begins with a detailed overview of the main ideas, followed by proofs of the supporting lemmas and the theorem itself.

Because $w_p(\mathbf{v})$ is nondecreasing in p , both $\text{ALG}_n(p; \mathcal{D})$ and $\text{CR}_n(p)$ are nondecreasing in p . The following corollary is immediate.

Corollary 3.2. *For every $n \geq 2$ and every $p \in [-\infty, 1]$, $\Gamma \leq \text{CR}_n(p)$, and*

$$0.7059 \approx \Gamma \leq \liminf_{n \rightarrow \infty} \text{CR}_n(p) \leq \limsup_{n \rightarrow \infty} \text{CR}_n(p) \leq \beta_{\text{iid}} \approx 0.7451.$$

¹This introduces no extra informational power: upon observing X_i , the algorithm can draw an independent uniform variable $R_i \sim \text{Unif}[0, 1]$ and set $U_i = F(X_i^-) + R_i(F(X_i) - F(X_i^-))$, where $F(t^-) := \lim_{s \rightarrow t^-} F(s)$. The resulting variables U_i are i.i.d. $\text{Unif}[0, 1]$ and satisfy $F^{-1}(U_i) = X_i$ almost surely.

Finally, it is worth noting that for two agents, the egalitarian competitive ratio coincides with the classical utilitarian ratio [HK82]. By monotonicity, the same is then true for every intermediate p -mean objective. However, this coincidence already disappears with three agents. The proof, given in Appendix A, combines a direct two-agent analysis with a three-agent instantiation of the limiting hard distribution from Theorem 3.1.

Proposition 3.3. *For every $p \in [-\infty, 1]$, $\text{CR}_2(p) = \frac{2+\sqrt{2}}{4} \approx 0.8536$. By contrast, $\text{CR}_3(-\infty) \leq \frac{128}{159} \approx 0.805 < 0.818 \approx \text{CR}_3(1)$.*

3.1 Structure of Optimal Online Rules

Our proof relies on two structural results about optimal rules. The first, proved in Appendix B, shows that quantile-threshold rules dominate arbitrary online rules and characterize the Pareto frontier of online expected-utility profiles. This is analogous to the utilitarian setting [HK82].

Proposition 3.4. *For every distribution $\mathcal{D}$ with CDF F , the following hold:*

1. *For every online rule π , there is a threshold rule π' such that $v_i(\pi', \mathcal{D}) \geq v_i(\pi, \mathcal{D})$ for all $i \in [n]$.*
2. *For every threshold rule π parameterized by $q_1, \dots, q_{n-1}$ such that $q_i \geq F(0) = \Pr[X = 0]$ for all $i \in [n-1]$, $\mathbf{v}(\pi, \mathcal{D})$ is Pareto optimal, i.e., there is no $\mathbf{u} \in \mathcal{V}^{\text{on}}(\mathcal{D})$ such that $u_i \geq v_i(\pi, \mathcal{D})$ for all $i \in [n]$, with at least one inequality strict.*

Second, as in the utilitarian setting, an optimal rule can be constructed by working backward from the final agent. For utilitarian welfare, the final agent receives the item whenever it remains available, and agent i receives it whenever their value exceeds the expected value obtained by continuing with agents $i+1, \dots, n$. Our egalitarian construction is analogous: we choose the threshold for agent i so that, conditional on the item reaching that agent, their expected utility equals that of the final agent. Proceeding backward in this way produces an optimal rule that equalizes the agents' expected utilities.

Proposition 3.5. *Fix a distribution $\mathcal{D}$ with CDF F and let $\mu = \mathbb{E}[X]$. There is a sequence of thresholds $q_1, \dots, q_{n-1} \in (F(0), 1)$ satisfying*

$$\frac{S(q_i)}{q_i} = \mu \prod_{j=i+1}^{n-1} q_j,$$

where the empty product is 1. Let π^ be the corresponding quantile-threshold rule, and define its reach probabilities by $r_1 = 1$ and $r_{i+1} = r_i q_i$. Then π^* is optimal for egalitarian welfare and gives every agent expected utility $v_i(\pi^*, \mathcal{D}) = \mu r_n$. Consequently, $\text{ALG}_n(-\infty; \mathcal{D}) = \mu r_n$.*

We prove Proposition 3.5 at the beginning of Section 5, since its recurrence is a key ingredient in the proof of our main result.

4 Connection to Online Fair Division

The fairness objective in our prophet problem measures how expected utility is distributed across agents. Here, we show how the same objective arises in an online fair-division problem with many indivisible items and sequentially arriving agents.

There are n agents and m indivisible items. Agent i 's value for item j is X_{ij} , and all values are drawn independently from a distribution $\mathcal{D}$. The agents arrive one at a time. When an agent

arrives, their values for the remaining items are revealed, and an online algorithm must immediately decide which of these items to allocate to them. An offline algorithm, by contrast, sees the entire valuation matrix before choosing an allocation. If agent i receives the items $A_i \subseteq [m]$, their realized utility is $U_i = \sum_{j \in A_i} X_{ij}$.

There are two natural ways to evaluate welfare in this model. The *ex-ante* (ea) objective applies the p -mean to the agents' expected utilities, giving $w_p(\mathbb{E}[U_1], \dots, \mathbb{E}[U_n])$. The *ex-post* (ep) objective instead takes the expected p -mean of the realized utilities, giving $\mathbb{E}[w_p(U_1, \dots, U_n)]$. For egalitarian welfare, these objectives are $\min_i \mathbb{E}[U_i]$ and $\mathbb{E}[\min_i U_i]$, respectively.

For $s \in \{\text{ea}, \text{ep}\}$, let $\text{ALG}_{n,m}^s(p; \mathcal{D})$ and $\text{OPT}_{n,m}^s(p; \mathcal{D})$ denote the optimal online and offline values under objective s , and let $\text{CR}_{n,m}^s(p)$ be the corresponding competitive ratio.² We first observe that the ex-ante problem is equivalent to the single-item problem, for every number of items. Proofs for this section can be found in Appendix C.

Proposition 4.1. *For every $n \geq 2$, $m \geq 1$ and distribution $\mathcal{D}$, the sets of expected-utility profiles achievable by online and offline rules in the m -item problem are, respectively, exactly m times the corresponding sets in the single-item problem. Thus, for all $p \in [-\infty, 1]$, $\text{ALG}_{n,m}^{\text{ea}}(p; \mathcal{D}) = m \text{ALG}_n(p; \mathcal{D})$ and $\text{OPT}_{n,m}^{\text{ea}}(p; \mathcal{D}) = m \text{OPT}_n(\mathcal{D})$. Consequently, $\text{CR}_{n,m}^{\text{ea}}(p) = \text{CR}_n(p)$.*

One direction follows by applying a single-item rule independently to every item. For the other, start with an arbitrary many-item rule and focus on a uniformly random item. By internally simulating the values of all other items, we obtain a single-item rule whose expected-utility profile is the average profile of the original rule.

The same equivalence does not hold for ex-post welfare when the number of items is small. In fact, when there is exactly one item per agent, the prophet may give every agent positive utility with high probability even though an online algorithm has almost no chance of doing so.

Proposition 4.2. *For every $p \leq 0$, the worst-case ex-post competitive ratio with $m = n$ converges to zero as n grows. More specifically, $\text{CR}_{n,n}^{\text{ep}}(p) \in O(\log n/n)$.*

The construction takes $\mathcal{D} = \text{Bernoulli}(2 \log(n)/n)$. The prophet can give every agent a value-one item whenever the corresponding random bipartite graph contains a perfect matching, which happens with probability approaching one [FK15]. By contrast, for an online allocation to have positive p -mean welfare when $p \leq 0$, every agent must receive a value-one item. The final agent must therefore value the one remaining item, an event of probability only $2 \log(n)/n$.

This impossibility relies on the numbers of agents and items growing together. If we instead fix the number of agents and the distribution, ex-post welfare approaches ex-ante welfare as the number of items grows.

Proposition 4.3. *For every $n \geq 2$, $p \in [-\infty, 1]$, and distribution $\mathcal{D}$,*

$$\lim_{m \rightarrow \infty} \frac{\text{ALG}_{n,m}^{\text{ep}}(p; \mathcal{D})}{\text{OPT}_{n,m}^{\text{ep}}(p; \mathcal{D})} = \frac{\text{ALG}_n(p; \mathcal{D})}{\text{OPT}_n(\mathcal{D})}.$$

The lower bound is obtained by applying an optimal single-item rule independently to every item. The law of large numbers then makes each agent's average realized utility approach their expected utility under that rule. In the other direction, concavity of the p -mean bounds ex-post welfare by ex-ante welfare, and Proposition 4.1 applies. A finite mean is enough for this pointwise convergence, although it does not provide a uniform rate over distributions.

²For ex-post competitive ratios with $p \leq 0$, we restrict to $m \geq n$, since both optima are zero when $m < n$.

Stronger quantitative guarantees follow when the distribution has lighter tails using concentration and moment bounds. For example, suppose $X \leq B \cdot \mathbb{E}[X]$ almost surely, then for $m \in \Omega\left(\frac{nB}{\delta^2} \log\left(\frac{n}{\delta}\right)\right)$, $\frac{\text{ALG}_{n,m}^{\text{ep}}(p; \mathcal{D})}{\text{OPT}_{n,m}^{\text{ep}}(p; \mathcal{D})} \geq \frac{\text{ALG}_n(p; \mathcal{D})}{\text{OPT}_n(\mathcal{D})} - \delta$. Under the weaker assumption of a finite mean μ and finite variance σ^2 , we can instead require $m \in \Omega\left(\frac{n^3}{\delta^2} \left(1 + \frac{\sigma^2}{\mu^2}\right)\right)$. Details can be found in Appendix C.4.

These results explain why we focus on the single-item problem: this problem captures ex-ante fair division exactly for every number of items and captures ex-post fair division in the many-item limit. At the same time, the failure when $m = n$ shows that this convergence is not uniform when n and m grow together.

5 Proof of Main Result

5.1 Proof Overview

We begin with the lower bound. Scaling all values by a positive constant scales both $\text{ALG}_n(-\infty; \mathcal{D})$ and $\text{OPT}_n(\mathcal{D})$ by that constant, and therefore does not change their ratio. Hence, it suffices to prove the lower bound for distributions satisfying $\mathbb{E}[X] = 1$; we make this assumption throughout the lower-bound argument.

The lower bound has two main ingredients. First, we transform each distribution into a convex function whose objective value is exactly the prophet-to-online ratio. Second, we solve the resulting variational optimization problem.

Let $\mathcal{K}_{\text{fin}}$ denote the class of nonnegative, nondecreasing, convex functions $G: [0, \infty) \rightarrow [0, \infty]$ such that $G(0) = 0$ and $G(z) = +\infty$ for all sufficiently large z . For $G \in \mathcal{K}_{\text{fin}}$, define

$$\mathcal{I}(G) := \int_0^\infty e^{-G(z)} dz, \quad \mathcal{C}(G) := \int_1^\infty \frac{dz}{zG(z)},$$

with the conventions $e^{-\infty} = 0$, $1/(+\infty) = 0$, and $1/0 = +\infty$.

For a distribution $\mathcal{D}$ with $\mu = \mathbb{E}[X]$, recall that $S(q) = \mathbb{E}[F^{-1}(U)\mathbf{1}U \geq q]$. For $0 \leq y \leq \mu$, define its generalized inverse by

$$S^{-1}(y) := \sup\{q \in [0, 1] : S(q) \geq y\}.$$

The following lemma gives the reduction from the prophet problem to the variational problem.

Lemma 5.1. *Fix a distribution $\mathcal{D}$ satisfying $\mathbb{E}[X] = 1$, and let $c = \text{ALG}_n(-\infty; \mathcal{D})$. Define $H: [0, \infty) \rightarrow [0, \infty]$ by*

$$H(z) := \begin{cases} -(n-1) \log(S^{-1}(cz)), & 0 \leq z \leq 1/c, \\ +\infty, & z > 1/c. \end{cases}$$

Then $H \in \mathcal{K}_{\text{fin}}$ and

$$\mathcal{I}(H) = \frac{\text{OPT}_n(\mathcal{D})}{\text{ALG}_n(-\infty; \mathcal{D})} \quad \text{and} \quad \mathcal{C}(H) \leq 1.$$

The analytic core of the proof is the following variational inequality.

Lemma 5.2. *Every $G \in \mathcal{K}_{\text{fin}}$ satisfying $\mathcal{C}(G) \leq 1$ also satisfies $\mathcal{I}(G) \leq \frac{1}{\Gamma}$.*

Together, these lemmas bound the prophet-to-online ratio by $1/\Gamma$, giving $\text{CR}_n(-\infty) \geq \Gamma$.

We next explain where the function H , its objective $\mathcal{I}(H)$, and the constraint $\mathcal{C}(H) \leq 1$ come from. We then describe the main idea behind the variational inequality and, finally, the distributions giving the matching upper bound.

Where the variational reduction comes from. Recall that the pair $(S(q), q)$ describes everything that matters about using threshold q at a single agent. Conditional on the item still being available, the threshold gives the current agent expected utility $S(q)$ and preserves the item for the next agent with probability q . Lowering q gives more value to the current agent, but leaves less chance for everyone who comes later. The generalized inverse has an equally useful interpretation. If the current agent must receive conditional expected utility y , then $S^{-1}(y)$ is the largest threshold that provides this utility, and hence the largest probability with which the item can be preserved for the next agent. Thus, $S^{-1}(y)$ describes how much reach probability can be retained while providing a specified amount of utility now.

Let $c = \text{ALG}_n(-\infty; \mathcal{D})$. The function S^{-1} is useful in two ways. First, a direct calculation expresses the prophet-to-online ratio as

$$\frac{\text{OPT}_n(\mathcal{D})}{c} = \int_0^{1/c} (S^{-1}(cz))^{n-1} dz.$$

Second, S^{-1} describes the constraints imposed by the equalizing rule from Proposition 3.5. Let $q_1, \dots, q_{n-1}$ be its thresholds, and let r_i be the probability that agent i is reached. Because every agent receives expected utility c , we have $r_i S(q_i) = c$, so $q_i = S^{-1}(c/r_i)$, and hence $r_{i+1} = r_i S^{-1}(c/r_i)$. Moreover, $r_1 = 1$ and $r_n = c$, since the final agent receives the item whenever reached and $\mathbb{E}[X] = 1$.

Our goal is therefore to bound the prophet-to-online ratio over all functions S^{-1} whose associated reach probabilities can move from $r_1 = 1$ to $r_n = c$ in $n - 1$ steps. Working with these expressions directly is awkward because both have a multiplicative structure: the objective contains an $(n - 1)$ st power, while the recurrence repeatedly multiplies the reach probability by the next threshold. Taking logarithms simplifies both.

In the objective, we can write $(S^{-1}(cz))^{n-1} = \exp((n - 1)(\log S^{-1}(cz)))$. For the recurrence, set $x_i = -\log r_i$. Thus, x_i records the logarithmic decrease in reach probability by the time agent i arrives. The sequence begins at $x_1 = 0$ and ends at $x_n = -\log c$, and the recurrence becomes

$$x_{i+1} - x_i = -\log q_i = -\log S^{-1}(ce^{x_i}).$$

It is useful to regard the points (i, x_i) as a trajectory mapping arrival positions to log reach probabilities. For our purposes, it is more convenient to consider the inverse trajectory, which records the arrival position at which the reach probability falls to a given level. While this trajectory is only defined on n discrete points, we can create a more continuous trajectory (more amenable to constraints) by interpolating between these points in a piecewise linear fashion.

Formally, let $T: [0, -\log c] \rightarrow [1, n]$ be the piecewise-linear function satisfying $T(x_i) = i$. On each interval (x_i, x_{i+1}) , the slope (and hence derivative) is

$$T'(x) = \frac{1}{x_{i+1} - x_i} = \frac{1}{-\log S^{-1}(ce^{x_i})}.$$

The fundamental theorem of calculus gives a clean invariant of the trajectory that $\int_0^{-\log c} T'(x) dx = T(-\log c) - T(0) = n - 1$. This essentially captures the recurrence constraint: the inverse trajectory moves from agent 1 to agent n on the interval $[0, -\log c]$, so its total increase is exactly $n - 1$.

The next key idea is that we can approximate $T'(x)$ from below by a continuous function. Consider $1/[-\log S^{-1}(ce^x)]$. Since S^{-1} is decreasing, for $x \in (x_i, x_{i+1})$,

$$\frac{1}{-\log S^{-1}(ce^x)} \leq \frac{1}{-\log S^{-1}(ce^{x_i})} = T'(x).$$

This bound is close to equality when consecutive reach probabilities are close, which helps explain why the resulting continuous relaxation can be asymptotically tight.

Integrating gives

$$\int_0^{-\log c} \frac{dx}{-\log S^{-1}(ce^x)} \leq \int_0^{-\log c} T'(x) dx = n - 1.$$

After the change of variables $z = e^x$, this becomes

$$\int_1^{1/c} \frac{dz}{-z \log S^{-1}(cz)} \leq n - 1.$$

At this point, the same expression $\log S^{-1}(cz)$ appears in both the prophet's value and the constraint. In the objective, it is multiplied by $n - 1$ in the exponent; in the constraint, the right-hand side is $n - 1$. This motivates defining

$$H(z) := \begin{cases} -(n - 1) \log S^{-1}(cz), & 0 \leq z \leq 1/c, \\ +\infty, & z > 1/c. \end{cases}$$

The prophet-to-online ratio and the constraint now take the simple forms

$$\mathcal{I}(H) := \int_0^\infty e^{-H(z)} dz = \frac{\text{OPT}_n(\mathcal{D})}{c}, \quad \mathcal{C}(H) := \int_1^\infty \frac{dz}{zH(z)} \leq 1.$$

To reemphasize, in the formal proof, we bypass the interpolated trajectory: we define H directly and prove $\mathcal{C}(H) \leq 1$ by summing over the intervals $[1/r_i, 1/r_{i+1}]$. The interpolated trajectory viewpoint is included only to motivate where this particular integral constraint arises.

The variational inequality. The proof of Lemma 5.2 primarily shows that any maximizer must take the form

$$G^*(z) = \begin{cases} 0, & 0 \leq z \leq J, \\ \sigma(z - J), & J < z < M, \\ +\infty, & z \geq M, \end{cases}$$

for some $0 \leq J < 1 < M$ and $\sigma > 0$. The key observation is that any such convex function can be represented as a mixture of *hinges*, functions of the form $h_a(z) = (z - a)_+$. We examine the first-order effect of adding a hinge and then rescaling to preserve the constraint. At a maximizer, each such perturbation has nonpositive first-order effect. However, the combined effect of the hinges already present in G^* is zero: adding them in their existing proportions simply rescales G^* , which the normalization undoes. Thus, every hinge in the support of the representing measure must have zero first-order effect. Analyzing these conditions shows that the measure is supported only at J , giving the form above.

The expression for $\rho(a, b)$ arises from evaluating $\mathcal{I}(G^*)$. Writing $a = \sigma J$ and $b = \sigma(M - J)$ gives $\mathcal{I}(G^*) = 1/\rho(a, b)$. Taking the infimum over $a, b > 0$ therefore gives $\mathcal{I}(G^*) \leq 1/\Gamma$ and establishes the desired lower bound.

The matching upper bound. To construct a hard distribution, for arbitrary $a, b > 0$, we consider a distribution that takes value $1/\varepsilon$ with probability ε , value n/a with probability b/n , and value 0 otherwise. The rare, very large value and the more common intermediate value produce the two factors appearing in $\rho(a, b)$.

A direct analysis of the equalizing threshold recurrence bounds the online value, while the prophet's value can be calculated explicitly. After first taking $\varepsilon \downarrow 0$ and then letting $n \rightarrow \infty$, the resulting competitive ratio is at most

$$\left[\left(1 - \frac{b}{a+b} e^{-a} \right) \left(1 + \frac{1 - e^{-b}}{a} \right) \right]^{-1} = \rho(a, b).$$

Since $a, b > 0$ are arbitrary, this gives $\limsup_{n \rightarrow \infty} \text{CR}_n(-\infty) \leq \Gamma$. Together with the variational lower bound, it follows that $\lim_{n \rightarrow \infty} \text{CR}_n(-\infty) = \Gamma$.

It is worth noting that the functions H corresponding to these distributions do indeed approach the form identified above. Thus, the variational problem is not merely a relaxation: its optimal value is exactly $1/\Gamma$, the inverse of the asymptotic competitive ratio.

5.2 Proof of Proposition 3.5

Fix a distribution $\mathcal{D}$ with mean μ and $n \geq 2$.

Starting with $i = n - 1$ and proceeding backward, choose $q_i \in (0, 1)$ such that

$$\frac{S(q_i)}{q_i} = \mu \prod_{j=i+1}^{n-1} q_j,$$

where the empty product is 1, and let π^* be the resulting quantile-threshold rule.

We first verify that this recursion is well-defined. Suppose that $q_{i+1}, \dots, q_{n-1}$ have already been chosen, and let $t_i = \mu \prod_{j=i+1}^{n-1} q_j \in (0, \mu]$. The function $q \mapsto S(q)/q$ is continuous on $(0, 1]$, tends to $+\infty$ as $q \downarrow 0$, and equals 0 at $q = 1$. The intermediate value theorem therefore gives some $q_i \in (0, 1)$ such that $S(q_i)/q_i = t_i$.

Moreover, every chosen threshold satisfies $q_i > F(0)$. This is immediate if $F(0) = 0$. Otherwise, for every $q \leq F(0)$, we have $S(q) = \mu$, and hence $S(q)/q = \mu/q \geq \mu/F(0) > \mu \geq t_i$. Thus no solution to the defining equation can lie in $[0, F(0)]$.

Let $r_i = \prod_{j=1}^{i-1} q_j$ be the probability that the item reaches agent i . For every $i < n$, the defining equation gives $S(q_i) = \mu \prod_{j=i}^{n-1} q_j$, and therefore

$$v_i(\pi^*, \mathcal{D}) = r_i S(q_i) = \left(\prod_{j=1}^{i-1} q_j \right) \left(\prod_{j=i}^{n-1} q_j \right) \cdot \mu = \mu r_n.$$

The final agent receives the item whenever reached, so $v_n(\pi^*, \mathcal{D}) = r_n \mathbb{E}[X] = \mu r_n$ as well. Thus π^* gives every agent expected utility μr_n .

Since $q_i > F(0)$ for every $i < n$, Proposition 3.4 implies that $\mathbf{v}(\pi^*, \mathcal{D}) = (\mu r_n, \dots, \mu r_n)$ is Pareto optimal. If $\text{ALG}_n(-\infty; \mathcal{D}) > \mu r_n$, then some online rule would give every agent expected utility strictly greater than μr_n , contradicting this Pareto optimality. Hence $\text{ALG}_n(-\infty; \mathcal{D}) = \mu r_n$, and π^* is optimal. $\square$

5.3 Proof of Lemma 5.1

Fix a distribution $\mathcal{D}$ satisfying $\mathbb{E}[X] = 1$ and $n \geq 2$. Let $q_1, \dots, q_{n-1}$ be the thresholds of the optimal equalizing rule from Proposition 3.5, and let r_i be its reach probabilities. Thus, $r_1 = 1$, $r_{i+1} = r_i q_i$, and $c = \text{ALG}_n(-\infty; \mathcal{D}) = r_n$. Because every $q_i \in (F(0), 1)$, we have $c \in (0, 1)$.

Recall that $S(q) = \int_q^1 F^{-1}(u) du$. Thus, S is continuous and nonincreasing on $[0, 1]$, constant and equal to 1 on $[0, F(0)]$, and strictly decreasing from 1 to 0 on $[F(0), 1]$. Consequently, S^{-1} is continuous and strictly decreasing on $[0, 1]$, with $S^{-1}(0) = 1$ and $S^{-1}(1) = F(0)$.

We first verify that $H \in \mathcal{K}_{\text{fin}}$. By construction, H is nonnegative, nondecreasing, satisfies $H(0) = 0$, and equals $+\infty$ for every $z > 1/c$. Moreover, S is concave because F^{-1} is nondecreasing. Its inverse S^{-1} is therefore concave and decreasing. Since $\log$ is increasing and concave, $\log S^{-1}$ is concave wherever S^{-1} is positive. It follows that $z \mapsto -(n-1) \log S^{-1}(cz)$ is convex on $[0, 1/c]$. Extending this function by $+\infty$ beyond $1/c$ preserves convexity, so $H \in \mathcal{K}_{\text{fin}}$.

We next compute $\mathcal{I}(H)$. Let $U_{-1}^{\max} = \max\{U_2, \dots, U_n\}$. Allocating the item to the agent with the largest quantile implements the prophet's allocation and treats the agents symmetrically. Conditioning on $U_{-1}^{\max}$ gives $\text{OPT}_n(\mathcal{D}) = \mathbb{E}[S(U_{-1}^{\max})]$. Since $0 \leq S(U_{-1}^{\max}) \leq 1$, the tail-integral formula yields

$$\begin{aligned} \frac{\text{OPT}_n(\mathcal{D})}{c} &= \int_0^{1/c} \Pr[S(U_{-1}^{\max}) \geq cz] dz \\ &= \int_0^{1/c} \Pr[U_{-1}^{\max} \leq S^{-1}(cz)] dz \\ &= \int_0^{1/c} (S^{-1}(cz))^{n-1} dz \\ &= \int_0^\infty e^{-H(z)} dz = \mathcal{I}(H). \end{aligned}$$

It remains to prove $\mathcal{C}(H) \leq 1$. Because the equalizing rule gives every agent expected utility c , we have $r_i S(q_i) = c$ for every $i < n$. Since $q_i > F(0)$ and S is strictly decreasing on $[F(0), 1]$, it follows that $q_i = S^{-1}(c/r_i)$ and hence $H(1/r_i) = -(n-1) \log q_i$. Moreover, $r_{i+1} = r_i q_i$, so $\log(r_i/r_{i+1}) = -\log q_i = H(1/r_i)/(n-1)$.

Because H is nondecreasing, for every $i = 1, \dots, n-1$,

$$\int_{1/r_i}^{1/r_{i+1}} \frac{dz}{zH(z)} \leq \frac{1}{H(1/r_i)} \int_{1/r_i}^{1/r_{i+1}} \frac{dz}{z} = \frac{\log(r_i/r_{i+1})}{H(1/r_i)} = \frac{1}{n-1}.$$

These intervals are consecutive and cover $[1, 1/c]$, since $r_1 = 1$ and $r_n = c$. Summing over i , and using $H(z) = +\infty$ for $z > 1/c$, gives $\mathcal{C}(H) \leq 1$. This completes the proof of Lemma 5.1. $\square$

5.4 Proof of Lemma 5.2

Fix an arbitrary $G \in \mathcal{K}_{\text{fin}}$ satisfying $\mathcal{C}(G) \leq 1$. Choose $R > 1$ sufficiently large, so that $G(z) = +\infty$ for every $z \geq R$. Let $\mathcal{F}_R$ be the class of all nonnegative, nondecreasing, convex functions $X: [0, \infty) \rightarrow [0, \infty]$ such that: (1) $X(0) = 0$, (2) $X(z) = +\infty$ for every $z \geq R$, and $\mathcal{C}(X) \leq 1$. In particular, $G \in \mathcal{F}_R$. Define

$$\Lambda_R := \sup_{X \in \mathcal{F}_R} \mathcal{I}(X).$$

The following lemma ensures that we may work with an actual maximizer and records the basic properties needed for later arguments. Its proof is deferred to Section D.1.

Lemma 5.3. *For every $R > 1$, the supremum defining Λ_R is attained. Moreover, $\Lambda_R > 1$, and every maximizer $G^* \in \mathcal{F}_R$ satisfies $\mathcal{C}(G^*) = 1$. Defining*

$$M := \sup\{z \geq 0 : G^*(z) < +\infty\}, \quad J := \sup\{z \geq 0 : G^*(z) = 0\},$$

then

$$0 \leq J < 1 < M \leq R.$$

Fix a maximizer $G^* \in \mathcal{F}_R$, and let J and M be as in Lemma 5.3. We may assume that $G^*(M) = +\infty$, since setting its value at M to $+\infty$ preserves convexity and changes neither $\mathcal{I}(G^*)$ nor $\mathcal{C}(G^*)$.

Hinges and hinge perturbations. For $a \in [0, M)$, define the *hinge* $h_a(z) := (z - a)_+ = \max\{z - a, 0\}$. Hinges are the elementary building blocks of convex functions: their coefficients record where and by how much the slope increases. To characterize G^* , we want to determine which hinges appear in its hinge representation. We begin by computing the first-order effect of adding a small hinge at each location a (and then renormalize to preserve the constraint $\mathcal{C} = 1$). The resulting first-order conditions will restrict which locations can have positive weight in the hinge representation of G^* .

For $\varepsilon \geq 0$, let $G_{a,\varepsilon} := G^* + \varepsilon h_a$. This function is nonnegative, nondecreasing, and convex, and it has the same finite cutoff M as G^* . Moreover, $G_{a,\varepsilon} \geq G^*$, so $0 < \mathcal{C}(G_{a,\varepsilon}) \leq \mathcal{C}(G^*) = 1$. Define the normalized perturbation

$$\tilde{G}_{a,\varepsilon} := \mathcal{C}(G_{a,\varepsilon}) G_{a,\varepsilon}.$$

Then $\tilde{G}_{a,\varepsilon} \in \mathcal{F}_R$ and $\mathcal{C}(\tilde{G}_{a,\varepsilon}) = \frac{\mathcal{C}(G_{a,\varepsilon})}{\mathcal{C}(G_{a,\varepsilon})} = 1$. In particular, since $\mathcal{C}(G^*) = 1$, $\tilde{G}_{a,0} = G^*$.

We now compute the first-order effect of this perturbation. Since G^* maximizes $\mathcal{I}$ over $\mathcal{F}_R$ and $\tilde{G}_{a,\varepsilon} \in \mathcal{F}_R$ for every $\varepsilon \geq 0$, $\mathcal{I}(\tilde{G}_{a,\varepsilon}) \leq \mathcal{I}(G^*)$. Consequently, if the right derivative exists, then

$$\Psi(a) := \left. \frac{d}{d\varepsilon} \mathcal{I}(\tilde{G}_{a,\varepsilon}) \right|_{\varepsilon=0+} \leq 0.$$

The following lemma computes this derivative; its proof in Appendix D.2 gives the formal details.

Lemma 5.4. *For $a \in [0, M)$, the right derivative $\Psi(a)$ exists. Define*

$$\begin{aligned} A(a) &:= - \left. \frac{d}{d\varepsilon} \mathcal{C}(G^* + \varepsilon h_a) \right|_{\varepsilon=0+} = \int_1^M \frac{h_a(z)}{z G^*(z)^2} dz, \\ B(a) &:= - \left. \frac{d}{d\varepsilon} \mathcal{I}(G^* + \varepsilon h_a) \right|_{\varepsilon=0+} = \int_0^M h_a(z) e^{-G^*(z)} dz, \\ \lambda &:= - \left. \frac{d}{d\varepsilon} \mathcal{I}((1 + \varepsilon)G^*) \right|_{\varepsilon=0} = \int_0^M G^*(z) e^{-G^*(z)} dz. \end{aligned}$$

Then $\Psi(a) = \lambda A(a) - B(a)$. All three quantities are finite and strictly positive.

We will also use the following derivatives of Ψ , derived in Appendix D.3.

Lemma 5.5. *The function Ψ is continuously differentiable on $[0, M)$, with*

$$\Psi'(a) = \begin{cases} -\lambda \int_1^M \frac{dz}{z G^*(z)^2} + \int_a^M e^{-G^*(z)} dz, & 0 \leq a \leq 1, \\ \int_a^M \left(e^{-G^*(z)} - \frac{\lambda}{z G^*(z)^2} \right) dz, & 1 < a < M. \end{cases}$$

Moreover, Ψ is twice continuously differentiable away from $a = 1$, with

$$\Psi''(a) = \begin{cases} -e^{-G^*(a)}, & 0 < a < 1, \\ \frac{\lambda}{a G^*(a)^2} - e^{-G^*(a)}, & 1 < a < M. \end{cases}$$

The hinge representation and its support. We next show that G^* is linear on (J, M) . We use the standard hinge representation of a convex function. Because G^* is finite and convex on $[0, M)$, we have $G^*(J) = 0$. There is a nonnegative, locally finite measure μ^* on $[J, M)$ such that

$$G^*(z) = \int_{[J, M)} h_a(z) \mu^*(da) \quad (0 \leq z < M). \quad (1)$$

Here the atom of μ^* at J equals the initial right slope $G_+'(J)$, while its mass after J records subsequent increases in the slope. The following lemma, proved in Appendix D.4, shows that the aggregate first-order effect of the hinges appearing in G^* is zero. Intuitively, if these hinges had a strictly negative aggregate effect, slightly subtracting them from G^* would preserve convexity and yield a first-order improvement. Furthermore, perturbing G^* in the aggregate direction of all its constituent hinges simply rescales G^* , an effect that is undone by the normalization. Its aggregate first-order effect must therefore be zero.

Lemma 5.6. *For the measure μ^* in Equation (1), $\int_{[J, M)} \Psi(a) \mu^*(da) = 0$.*

Since $\Psi(a) \leq 0$ for every a , Lemma 5.6 implies that $\mu^*(\{a \in [J, M) : \Psi(a) < 0\}) = 0$. Recall that the support of μ^* is the set

$$\text{supp}(\mu^*) := \{a \in [J, M) : \mu^*((a - \delta, a + \delta) \cap [J, M)) > 0 \text{ for every } \delta > 0\}.$$

In other words, a point belongs to the support if every neighborhood of that point receives positive μ^* -measure. We claim that Ψ vanishes everywhere on this support. Indeed, suppose that $a_0 \in \text{supp}(\mu^*)$ but $\Psi(a_0) < 0$. By continuity, there are $\delta, \eta > 0$ such that $\Psi(a) \leq -\eta$ whenever $a \in (a_0 - \delta, a_0 + \delta) \cap [J, M)$. This neighborhood has positive μ^* -measure because a_0 lies in the support. It would therefore make a strictly negative contribution to the integral of Ψ . Since $\Psi \leq 0$ everywhere, this would imply

$$\int_{[J, M)} \Psi(a) \mu^*(da) < 0,$$

contradicting the equality above. Hence

$$\Psi(a) = 0 \quad \text{for every } a \in \text{supp}(\mu^*). \quad (2)$$

Moreover, $G^*(z) > 0$ for every $z > J$. Thus Equation (1) implies $\mu^*([J, z)) > 0$ for every $z > J$. Therefore, $J \in \text{supp}(\mu^*)$, and hence $\Psi(J) = 0$.

We have shown that $J \in \text{supp}(\mu^*)$ and that Ψ vanishes on $\text{supp}(\mu^*)$. To prove that G^* consists of a single hinge, it therefore remains to show that μ^* has no other support points. For $(J, 1]$, this is relatively straightforward: By Lemma 5.5, $\Psi''(a) = -e^{-G^*(a)} < 0$ for $J < a < 1$. Thus Ψ is strictly concave on $[J, 1]$. Since $\Psi \leq 0$ and $\Psi(J) = 0$, strict concavity gives $\Psi(a) < 0$ for every $a \in (J, 1]$: otherwise, Ψ would be positive between J and a . Together with Equation (2), this implies $\text{supp}(\mu^*) \cap (J, 1] = \emptyset$. The range $(1, M)$ is more involved, and we defer the details to Appendix D.5.

Lemma 5.7. $\text{supp}(\mu^*) \cap (J, M) = \emptyset$.

The above arguments combined show that $\text{supp}(\mu^*) = \{J\}$. Thus $\mu^* = \sigma \delta_J$ for some $\sigma > 0$, and,

$$G^*(z) = \begin{cases} 0, & 0 \leq z \leq J, \\ \sigma(z - J), & J < z < M, \\ +\infty, & z \geq M. \end{cases} \quad (3)$$

Evaluating the single-hinge maximizer. It remains to evaluate this function. First suppose that $J > 0$, and define $a := \sigma J$ and $b := \sigma(M - J)$. Then $a, b > 0$. Since $\mathcal{C}(G^\star) = 1$,

$$1 = \frac{1}{\sigma} \int_1^M \frac{dz}{z(z - J)} = \frac{1}{\sigma J} \log \left(\frac{M - J}{M(1 - J)} \right) = \frac{1}{a} \log \left(\frac{b}{(a + b)(1 - J)} \right).$$

Therefore, $J = 1 - \frac{b}{a+b}e^{-a}$. Since $1/\sigma = J/a$, we get

$$\begin{aligned} \mathcal{I}(G^\star) &= J + \frac{1 - e^{-b}}{\sigma} \\ &= \left(1 - \frac{b}{a+b}e^{-a}\right) \left(1 + \frac{1 - e^{-b}}{a}\right) \\ &= \frac{1}{\rho(a, b)} \leq \frac{1}{\Gamma}. \end{aligned}$$

If $J = 0$, then $\mathcal{C}(G^\star) = 1$ gives $\sigma = 1 - \frac{1}{M}$. Setting $b = \sigma M = M - 1 > 0$, we get $\mathcal{I}(G^\star) = \frac{1 - e^{-b}}{\sigma} = \frac{b+1}{b}(1 - e^{-b})$. Here, the parameter $a = \sigma J$ from the preceding case equals zero, which lies outside the domain $a > 0$ in the definition of ρ . However, this case is obtained by taking $a \downarrow 0$: $\mathcal{I}(G^\star) = \lim_{a \downarrow 0} \frac{1}{\rho(a, b)}$; and hence again $\mathcal{I}(G^\star) \leq \frac{1}{\Gamma}$.

Finally, the original function G belongs to $\mathcal{F}_R$, so $\mathcal{I}(G) \leq \mathcal{I}(G^\star) \leq \frac{1}{\Gamma}$. This proves Lemma 5.2. $\square$

5.5 Proof of Theorem 3.1

5.5.1 Lower bound: $\text{CR}_n(-\infty) \geq \Gamma$

Fix any distribution $\mathcal{D}$ with finite, positive mean, and $n \geq 2$. Scaling all values by a positive constant scales both $\text{ALG}_n(-\infty; \mathcal{D})$ and $\text{OPT}_n(\mathcal{D})$ by that constant, and therefore does not change their ratio. We may thus assume that $\mathbb{E}[X] = 1$.

By Lemma 5.1, there is a function $H \in \mathcal{K}_{\text{fin}}$ satisfying

$$\frac{\text{OPT}_n(\mathcal{D})}{\text{ALG}_n(-\infty; \mathcal{D})} = \mathcal{I}(H) \quad \text{and} \quad \mathcal{C}(H) \leq 1.$$

Applying Lemma 5.2 gives

$$\frac{\text{OPT}_n(\mathcal{D})}{\text{ALG}_n(-\infty; \mathcal{D})} = \mathcal{I}(H) \leq \frac{1}{\Gamma}.$$

Equivalently,

$$\frac{\text{ALG}_n(-\infty; \mathcal{D})}{\text{OPT}_n(\mathcal{D})} \geq \Gamma.$$

Since this holds for every distribution $\mathcal{D}$ with finite, positive mean, taking the infimum over $\mathcal{D}$ proves $\text{CR}_n(-\infty) \geq \Gamma$.

5.5.2 Upper bound: $\limsup_{n \rightarrow \infty} \text{CR}_n(-\infty) \leq \Gamma$

Next, we prove our upper bound on $\text{CR}_n(-\infty)$.

Fix arbitrary $a, b > 0$. We construct a sequence of distributions whose competitive ratios converge to at most $\rho(a, b)$ (as defined in the beginning of Section 3). Fix $n > b$, and pick a sufficiently small $\varepsilon > 0$, so that $\varepsilon < 1 - \frac{b}{n}$ and $\frac{1}{\varepsilon} > \frac{n}{a}$. Let $\mathcal{D}_{n, \varepsilon}$ take value $1/\varepsilon$ with probability ε , value n/a with probability b/n , and value 0 otherwise. Its mean is $\mathbb{E}[X] = \frac{1}{\varepsilon}\varepsilon + \frac{n}{a}\frac{b}{n} = 1 + \frac{b}{a}$. The constraint on ε

ensures that this is well-defined, and that $1/\varepsilon$ truly is the highest value. Abbreviating $S_{\mathcal{D}_{n,\varepsilon}}$ to S , we have

$$S(q) = \begin{cases} 1 + \frac{b}{a}, & 0 \leq q \leq 1 - \varepsilon - b/n, \\ 1 + \frac{n}{a}(1 - \varepsilon - q), & 1 - \varepsilon - b/n < q \leq 1 - \varepsilon, \\ \frac{1 - q}{\varepsilon}, & 1 - \varepsilon < q \leq 1. \end{cases}$$

We obtain the desired bound by first letting $\varepsilon \downarrow 0$ and then letting $n \rightarrow \infty$. The argument below computes the online and prophet values directly, although the construction can also be viewed as realizing the maximizers from the optimization problem.³

Bounding ALG_n . Let $c = \text{ALG}_n(-\infty; \mathcal{D}_{n,\varepsilon})$. Let $q_1, \dots, q_{n-1}$ be the thresholds of the rule in Proposition 3.5 and let r_i be the probability that the item is still available when agent i arrives. Thus $r_1 = 1$, $r_{i+1} = r_i q_i$, and $r_i S(q_i) = c$ for every $i < n$. The last agent's utility is $c = r_n \mathbb{E}[X] = r_n (1 + \frac{b}{a})$, so $r_n = \frac{ac}{a+b}$.

We claim that

$$c \leq \frac{1}{1 - \frac{b}{a+b} \left(\frac{n}{n+a} \right)^{n-1}}.$$

The right hand side is greater than 1, so the claim is immediate if $c \leq 1$. It remains to consider $c > 1$.

Since $r_i \leq 1$, we have $S(q_i) = \frac{c}{r_i} > 1$ for every $i < n$. Moreover, the optimal thresholds satisfy $q_i > F(0) = 1 - \varepsilon - \frac{b}{n}$. Since $S(q) \leq 1$ for $q \geq 1 - \varepsilon$, the inequality $S(q_i) > 1$ implies $q_i < 1 - \varepsilon$. Together with $q_i > 1 - \varepsilon - b/n$, this shows that $1 - \varepsilon - \frac{b}{n} < q_i < 1 - \varepsilon$. Substituting $S(q_i) = 1 + \frac{n}{a}(1 - \varepsilon - q_i)$ into $r_i S(q_i) = c$ gives $r_i(1 - \varepsilon - q_i) = \frac{a}{n}(c - r_i)$. It follows that

$$\begin{aligned} c - r_{i+1} &= c - r_i q_i \\ &= \left(1 + \frac{a}{n}\right) (c - r_i) + \varepsilon r_i \\ &\geq \left(1 + \frac{a}{n}\right) (c - r_i). \end{aligned}$$

Since $r_1 = 1$, by induction we have $c - r_n \geq \left(1 + \frac{a}{n}\right)^{n-1} (c - 1)$. On the other hand, $c - r_n = c - \frac{ac}{a+b} = c \frac{b}{a+b}$. Therefore,

$$\left(1 + \frac{a}{n}\right)^{n-1} (c - 1) \leq c \frac{b}{a+b}.$$

Rearranging gives

$$c \leq \frac{1}{1 - \frac{b}{a+b} \left(1 + \frac{a}{n}\right)^{-(n-1)}} = \frac{1}{1 - \frac{b}{a+b} \left(\frac{n}{n+a}\right)^{n-1}}.$$

³To see this, write $\mu = 1 + b/a$. After taking $\varepsilon \downarrow 0$, the recurrence below becomes an equality and gives $c_n = \left(1 - \frac{b}{a+b} (1 + a/n)^{-(n-1)}\right)^{-1} \rightarrow c_\infty$, where $c_\infty = \left(1 - \frac{b}{a+b} e^{-a}\right)^{-1}$. Moreover, $S^{-1}(y) = 1$ for $y \leq 1$ and $S^{-1}(y) = 1 - \frac{a}{n}(y - 1)$ for $1 < y \leq \mu$. Consequently, after normalizing the distribution by μ , the associated functions H converge almost everywhere to the function that is zero on $[0, J]$, has slope $\sigma = ac_\infty$ on (J, M) , and is infinite for $z > M$, where $J = 1/c_\infty$ and $M = \mu/c_\infty$. In particular, $\sigma J = a$ and $\sigma(M - J) = b$. This exactly gives the objective value $1/\rho(a, b)$.

Bounding OPT_n . The maximum equals $1/\varepsilon$ if at least one draw takes the high value. This event has probability $1 - (1 - \varepsilon)^n$. If no draw takes the high value but at least one draw takes the middle value, the maximum equals n/a . This event has probability $(1 - \varepsilon)^n - (1 - \varepsilon - \frac{b}{n})^n$. Therefore

$$\text{OPT}_n(\mathcal{D}_{n,\varepsilon}) = \frac{1}{n} \left[\frac{1 - (1 - \varepsilon)^n}{\varepsilon} + \frac{n}{a} \left((1 - \varepsilon)^n - \left(1 - \varepsilon - \frac{b}{n}\right)^n \right) \right].$$

Taking $\varepsilon \rightarrow 0$ gives

$$\lim_{\varepsilon \rightarrow 0^+} \text{OPT}_n(\mathcal{D}_{n,\varepsilon}) = 1 + \frac{1}{a} \left[1 - \left(1 - \frac{b}{n}\right)^n \right].$$

Putting it all together. Since $\text{CR}_n(-\infty)$ is the infimum over distributions, combining the online and offline bounds gives

$$\text{CR}_n(-\infty) \leq \frac{1}{\left(1 - \frac{b}{a+b} \left(\frac{n}{n+a}\right)^{n-1}\right) \left(1 + \frac{1}{a} \left[1 - \left(1 - \frac{b}{n}\right)^n\right]\right)}.$$

Finally, $\left(\frac{n}{n+a}\right)^{n-1} = \left(1 + \frac{a}{n}\right)^{-(n-1)} \rightarrow e^{-a}$, and $\left(1 - \frac{b}{n}\right)^n \rightarrow e^{-b}$. It follows that

$$\limsup_{n \rightarrow \infty} \text{CR}_n(-\infty) \leq \frac{1}{\left(1 - \frac{b}{a+b} e^{-a}\right) \left(1 + \frac{1 - e^{-b}}{a}\right)} = \rho(a, b).$$

Because $a, b > 0$ were arbitrary,

$$\limsup_{n \rightarrow \infty} \text{CR}_n(-\infty) \leq \inf_{a, b > 0} \rho(a, b) = \Gamma.$$

Combining this upper bound with $\text{CR}_n(-\infty) \geq \Gamma$ for every n completes the proof of Theorem 3.1. $\square$

AI Disclosure We worked on this problem for several months and obtained weaker bounds before using GPT-5.6 Sol. Its initial attempts did not yield the tight bound, but over extended interactions it proposed recasting the remaining problem as an optimization problem, and suggested a variational inequality. In subsequent exchanges, it assisted in developing the statement and proof of the current version of this inequality (Lemma 5.2). We then verified the proof and developed the exposition, with particular emphasis on readability and intuition. The authors take full responsibility for the correctness of the results and all other content in the paper.

References

- [AK22] Makis Arsenis and Robert Kleinberg. Individual fairness in prophet inequalities. In *Proceedings of the 23rd ACM Conference on Economics and Computation*, 2022. Extended abstract; full version available as arXiv:2205.10302.
- [ALMT25] Georgios Amanatidis, Alexandros Lolos, Evangelos Markakis, and Victor Turmel. Online fair division for personalized 2-value instances. In *International Symposium on Algorithmic Game Theory*, pages 209–227. Springer, 2025.

- [BBKS20] Siddharth Barman, Umang Bhaskar, Anand Krishna, and Ranjani G. Sundaram. Tight approximation algorithms for p -mean welfare under subadditive valuations. In *28th Annual European Symposium on Algorithms (ESA 2020)*, volume 173 of *Leibniz International Proceedings in Informatics*, pages 11:1–11:17, 2020.
- [BFT23] Siddhartha Banerjee, Giannis Fikioris, and Eva Tardos. Robust pseudo-markets for reusable public resources. In *Proceedings of the 24th ACM Conference on Economics and Computation*, EC '23, page 241, New York, NY, USA, 2023. Association for Computing Machinery.
- [BGJ22] Siddhartha Banerjee, Vasilis Gkatzelis, Artur Gorokh, and Billy Jin. Online nash social welfare maximization with predictions. In *Proceedings of the 2022 Annual ACM-SIAM Symposium on Discrete Algorithms (SODA)*, pages 1–19. SIAM, 2022.
- [BHP25] Gerdus Benadè, Daniel Halpern, and Alexandros Psomas. Dynamic fair division with partial information. *Operations Research*, 73(4):1876–1896, 2025.
- [BHS26] Sid Banerjee, Chamsi Hssaine, and Sean R Sinclair. Online fair allocation of perishable resources. *Operations Research*, 2026.
- [BKM22] Siddharth Barman, Arindam Khan, and Arnab Maiti. Universal and tight online algorithms for generalized-mean welfare. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 36, pages 4793–4800, 2022.
- [BKP⁺24] Gerdus Benade, Aleksandr M Kazachkov, Ariel D Procaccia, Alexandros Psomas, and David Zeng. Fair and efficient online allocations. *Operations Research*, 72(4):1438–1452, 2024.
- [BX23] Santiago R. Balseiro and Shangzhou Xia. Uniformly bounded regret in dynamic fair allocation. In *Proceedings of the 3rd ACM Conference on Equity and Access in Algorithms, Mechanisms, and Optimization*, EAAMO '23, New York, NY, USA, 2023. Association for Computing Machinery.
- [Car00] N. L. Carothers. *Real analysis*. Cambridge University Press, 1 edition, 2000.
- [CFH⁺17] José Correa, Patricio Foncea, Ruben Hoeksma, Tim Oosterwijk, and Tjark Vredeveld. Posted price mechanisms for a random stream of customers. In *Proceedings of the 2017 ACM Conference on Economics and Computation*, pages 169–186, 2017.
- [FAT24] Giannis Fikioris, Rachit Agarwal, and Éva Tardos. Incentives in dominant resource fair allocation under dynamic demands. In *International Symposium on Algorithmic Game Theory*, pages 108–125. Springer, 2024.
- [FBT25] Giannis Fikioris, Siddhartha Banerjee, and Eva Tardos. Beyond worst-case online allocation via dynamic max-min fairness. In *Proceedings of the 26th ACM Conference on Economics and Computation*, EC '25, page 92, New York, NY, USA, 2025. Association for Computing Machinery.
- [FK15] Alan Frieze and Michał Karoński. *Introduction to random graphs*. Cambridge University Press, 2015.

- [GPK21] Yuan Gao, Alex Peysakhovich, and Christian Kroer. Online market equilibrium with application to fair division. *Advances in Neural Information Processing Systems*, 34:27305–27318, 2021.
- [GPT21] Vasilis Gkatzelis, Alexandros Psomas, and Xizhi Tan. Fair and efficient online allocations with normalized valuations. In *Proceedings of the AAAI conference on artificial intelligence*, volume 35, pages 5440–5447, 2021.
- [HK82] Theodore P. Hill and Robert P. Kertz. Comparisons of stop rule and supremum expectations of I.I.D. random variables. *The Annals of Probability*, 10(2):336–345, 1982.
- [HK90] Theodore P. Hill and Doug P. Kennedy. Optimal stopping problems with generalized objective functions. *Journal of Applied Probability*, 27(4):828–838, 1990.
- [HLSW23] Zhiyi Huang, Minming Li, Xinkai Shu, and Tianze Wei. Online nash welfare maximization without predictions. In *International Conference on Web and Internet Economics*, pages 402–419. Springer, 2023.
- [HLSW25] Zhiyi Huang, Chui Shan Lee, Xinkai Shu, and Zhaozi Wang. The long arm of nashian allocation in online p -mean welfare maximization. *arXiv preprint arXiv:2504.13430*, 2025.
- [HPPZ19] Jiafan He, Ariel D. Procaccia, Alexandros Psomas, and David Zeng. Achieving a fairer future by changing the past. In *Proceedings of the 28th International Joint Conference on Artificial Intelligence, IJCAI’19*, page 343–349. AAAI Press, 2019.
- [HPVX25] Daniel Halpern, Alexandros Psomas, Paritosh Verma, and Daniel Xie. Online envy minimization and multicolor discrepancy: Equivalences and separations. In *Proceedings of the 26th ACM Conference on Economics and Computation, EC ’25*, page 188, New York, NY, USA, 2025. Association for Computing Machinery.
- [IMMP20] Sungjin Im, Benjamin Moseley, Kamesh Munagala, and Kirk Pruhs. Dynamic weighted fairness with minimal disruptions. *Proceedings of the ACM on Measurement and Analysis of Computing Systems*, 4(1):1–18, 2020.
- [Ker86] Robert P. Kertz. Stop rule and supremum expectations of i.i.d. random variables: A complete comparison by conjugate duality. *Journal of Multivariate Analysis*, 19(1):88–112, 1986.
- [KMNP26] Pooja Kulkarni, Ruta Mehta, Vishnu V. Narayan, and Tomasz Ponitka. Online fair division with subsidy: When do envy-free allocations exist, and at what cost? In *Proceedings of the 25th International Conference on Autonomous Agents and Multiagent Systems, AAMAS ’26*, page 348–356, Richland, SC, 2026. International Foundation for Autonomous Agents and Multiagent Systems.
- [KMS25] Pooja Kulkarni, Ruta Mehta, and Parnian Shahkar. Online fair division: Towards ex-post constant mms guarantees. In *Proceedings of the 26th ACM Conference on Economics and Computation, EC ’25*, page 638, New York, NY, USA, 2025. Association for Computing Machinery.
- [KPS14] Ian Kash, Ariel D Procaccia, and Nisarg Shah. No agent left behind: Dynamic fair division of multiple resources. *Journal of Artificial Intelligence Research*, 51:579–603, 2014.

- [KS77] Ulrich Krengel and Louis Sucheston. Semiamarts and finite values. *Bulletin of the American Mathematical Society*, 83(4):745–747, 1977.
- [KS22] Yasushi Kawase and Hanna Sumita. Online max-min fair allocation. In *International Symposium on Algorithmic Game Theory*, pages 526–543. Springer, 2022.
- [LGK22] Luofeng Liao, Yuan Gao, and Christian Kroer. Nonstationary dual averaging and online fair allocation. *Advances in Neural Information Processing Systems*, 35:37159–37172, 2022.
- [LLL18] Bo Li, Wenyang Li, and Yingkai Li. Dynamic fair division problem with general valuations. In *Proceedings of the 27th International Joint Conference on Artificial Intelligence, IJCAI’18*, page 375–381. AAAI Press, 2018.
- [Mou03] Hervé Moulin. *Fair Division and Collective Welfare*. MIT Press, 2003.
- [MPV24] Marios Mertzaniadis, Alexandros Psomas, and Paritosh Verma. Automating food drop: The power of two choices for dynamic and fair food allocation. In *Proceedings of the 25th ACM Conference on Economics and Computation, EC ’24*, page 243, New York, NY, USA, 2024. Association for Computing Machinery.
- [PSZ24] Ariel D Procaccia, Benjamin Schiffer, and Shirley Zhang. Honor among bandits: No-regret learning for online fair division. *Advances in Neural Information Processing Systems*, 37:13183–13227, 2024.
- [SC84] Ester Samuel-Cahn. Comparison of threshold stop rules and maximum for independent nonnegative random variables. *The Annals of Probability*, 12(4):1213–1216, 1984.
- [SHPK22] Max Springer, MohammadTaghi Hajiaghayi, Debmalya Panigrahi, and Mohammad Khani. Online algorithms for the santa claus problem. *Advances in Neural Information Processing Systems*, 35:30732–30743, 2022.
- [SJB23] Sean R Sinclair, Gauri Jain, Siddhartha Banerjee, and Christina Lee Yu. Sequential fair allocation: Achieving the optimal envy-efficiency trade-off curve. *Operations Research*, 71(5):1689–1705, 2023.
- [SZ25] Benjamin Schiffer and Shirley Zhang. Improved regret bounds for online fair division with bandit learning. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 39, pages 14079–14086, 2025.
- [VFA⁺23] Midhul Vuppapapati, Giannis Fikioris, Rachit Agarwal, Asaf Cidon, Anurag Khandelwal, and Éva Tardos. Karma: Resource allocation for dynamic demands. In *17th USENIX Symposium on Operating Systems Design and Implementation (OSDI 23)*, pages 645–662, 2023.
- [VPF22] Shai Vardi, Alexandros Psomas, and Eric Friedman. Dynamic fair resource division. *Mathematics of Operations Research*, 47(2):945–968, 2022.
- [YKK26] Zongjun Yang, Rachitesh Kumar, and Christian Kroer. Online generalized-mean welfare maximization: Achieving near-optimal regret from samples. *arXiv preprint arXiv:2602.10469*, 2026.

- [YLK24] Zongjun Yang, Luofeng Liao, and Christian Kroer. Greedy-based online fair allocation with adversarial input: enabling best-of-many-worlds guarantees. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 38, pages 9960–9968, 2024.

A Proof of Proposition 3.3

First consider $n = 2$, and normalize so that $\mathbb{E}[X] = 1$. If the first agent uses quantile threshold q , then the two agents receive expected utilities $S(q)$ and q . The optimal egalitarian threshold therefore satisfies $S(q) = q$, and hence $\text{ALG}_2(-\infty; \mathcal{D}) = q$ and $S^{-1}(q) = q$. Moreover, concavity of S , together with $S(0) = 1$ and $S(1) = 0$, implies $S(q) \geq 1 - q$, so $q \geq 1/2$.

The prophet’s expected utility per agent is $\text{OPT}_2(\mathcal{D}) = \int_0^1 S(t) dt = \int_0^1 S^{-1}(s) ds$, where the second equality follows by integrating the region under S in the opposite order. Let $-a$ be a supergradient of the concave function S^{-1} at q . Then $S^{-1}(s) \leq \min\{1, q - a(s - q)\}$. The endpoint conditions $S^{-1}(0) = 1$ and $S^{-1}(1) \geq 0$ further imply $(1 - q)/q \leq a \leq q/(1 - q)$. Therefore,

$$\begin{aligned} \text{OPT}_2(\mathcal{D}) &\leq \int_0^1 \min\{1, q - a(s - q)\} ds \\ &= 1 - \frac{(1 - q)^2}{2} \left(a + 2 + \frac{1}{a} \right) \leq 1 - 2(1 - q)^2, \end{aligned}$$

where the last inequality follows from $a + 1/a \geq 2$. Therefore,

$$\frac{\text{ALG}_2(-\infty; \mathcal{D})}{\text{OPT}_2(\mathcal{D})} \geq \frac{q}{1 - 2(1 - q)^2} \geq \frac{2 + \sqrt{2}}{4},$$

where the last inequality follows by minimizing over $q \in [1/2, 1]$. Since $\mathcal{D}$ was arbitrary, $\text{CR}_2(-\infty) \geq (2 + \sqrt{2})/4$. Hill and Kertz show that the classical two-agent ratio satisfies $\text{CR}_2(1) = (2 + \sqrt{2})/4$ [HK82]. Monotonicity in p now gives the claimed equality for every $p \in [-\infty, 1]$.

For $n = 3$, apply the hard-distribution construction (Section 5.5.2) with $a = 1$ and $b = 3/2$. Equivalently, for sufficiently small $\varepsilon > 0$, consider the distribution that places probabilities ε , $1/2$, and $1/2 - \varepsilon$ on the values $1/\varepsilon$, 3 , and 0 , respectively. The bounds from that construction give

$$\text{CR}_3(-\infty) \leq \frac{1}{\left(1 - \frac{3}{5} \left(\frac{3}{4}\right)^2\right) \left(1 + \left[1 - \left(\frac{1}{2}\right)^3\right]\right)} = \frac{128}{159} \approx 0.80503.$$

By comparison, the sharp classical three-agent ratio is $\text{CR}_3(1) \approx 1/1.22138 \approx 0.81875$ [HK82]. Thus $\text{CR}_3(-\infty) < \text{CR}_3(1)$, completing the proof. $\square$

B Proof of Proposition 3.4

Let R_i be the event that the item is still available when agent i arrives, and let $r_i = \Pr(R_i)$ be the probability that this event occurs. For an online rule π , for each $i < n$ with $r_i > 0$, let $\alpha_i = \Pr(\pi = i \mid R_i)$ and set $q_i = 1 - \alpha_i$. If $r_i = 0$, then agent i and all later agents have zero utility, and the remaining thresholds may be chosen arbitrarily.

Since R_i depends only on previous values and the rule’s internal randomness, it is independent of U_i . Thus, conditional on R_i , the quantile U_i remains uniform on $[0, 1]$. Averaging over the histories that lead to R_i , the rule accepts agent i on a total quantile mass $\alpha_i = 1 - q_i$. Among all such

acceptance rules, accepting the upper interval $[q_i, 1]$ maximizes the expected utility of i . Therefore, $v_i(\pi, \mathcal{D}) \leq r_i S(q_i)$.

The item reaches agent $i + 1$ precisely when it reaches agent i and π does not allocate to agent i . Therefore, $r_{i+1} = r_i(1 - \alpha_i) = r_i q_i$. A threshold rule π' with thresholds $q_1, \dots, q_{n-1}$ also rejects agent i , conditional on reaching them, with probability q_i . Since both π and π' reach agent 1 with probability 1, the preceding recurrence shows inductively that π and π' have the same reach probability r_i for every agent. Therefore, $v_i(\pi', \mathcal{D}) = r_i S(q_i) \geq v_i(\pi, \mathcal{D})$ for all $i < n$. Furthermore, since π' always allocates to the final agent when reached, $v_n(\pi', \mathcal{D}) = r_n \mathbb{E}[X] \geq v_n(\pi, \mathcal{D})$. This proves the first claim.

For the second claim, fix a threshold rule π with thresholds $q_i \geq F(0)$. Suppose that some online rule Pareto dominates π . By the first claim, we may assume that the dominating rule is also a threshold rule, with thresholds $\tilde{q}_1, \dots, \tilde{q}_{n-1}$. We may also assume that $\tilde{q}_i \geq F(0)$: replacing a threshold below $F(0)$ by $F(0)$ only rejects zero-valued realizations and cannot decrease any agent's utility. Let $r_i = \prod_{j < i} q_j$, and $\tilde{r}_i = \prod_{j < i} \tilde{q}_j$ be the two rules' reach probabilities.

We will show by induction that $\tilde{r}_i \leq r_i$ for every $i \in [n]$. The statement is trivial for $i = 1$, since $\tilde{r}_1 = r_1 = 1$. Suppose that $\tilde{r}_i \leq r_i$ (and $i < n$). If $r_i = 0$, then $\tilde{r}_i = 0$, and hence $\tilde{r}_{i+1} = r_{i+1} = 0$. Now suppose that $r_i > 0$. If $q_i < 1$, then $S(q_i) > 0$, so Pareto domination implies $\tilde{r}_i > 0$ and $S(\tilde{q}_i) \geq \frac{r_i}{\tilde{r}_i} S(q_i) \geq S(q_i)$. Since S is strictly decreasing on $[F(0), 1]$, it follows that $\tilde{q}_i \leq q_i$. If $q_i = 1$, then $\tilde{q}_i \leq q_i$ holds trivially. In either case, $\tilde{r}_{i+1} = \tilde{r}_i \tilde{q}_i \leq r_i q_i = r_{i+1}$, completing the induction.

We now show that Pareto domination forces the two threshold rules to have identical utility profiles. First, suppose that $q_i > 0$ for every $i \in [n - 1]$. The preceding induction also gives $\tilde{q}_i \leq q_i$ for every $i \in [n - 1]$. Thus, whenever $\tilde{r}_i < r_i$, we have $\tilde{r}_{i+1} \leq \tilde{r}_i q_i < r_i q_i = r_{i+1}$. Any strict inequality therefore persists until agent n , giving $v_n(\tilde{\pi}, \mathcal{D}) = \tilde{r}_n \mathbb{E}[X] < r_n \mathbb{E}[X] = v_n(\pi, \mathcal{D})$, contrary to Pareto domination. Hence $\tilde{r}_i = r_i$ for all $i \in [n]$. It follows (by the definition of $\tilde{r}_i$ and r_i) that $\tilde{q}_i = q_i$ for every $i \in [n - 1]$, and therefore the two threshold rules have the same utility profile.

Finally, suppose that $q_k = 0$, and let k be the first such index. Then $F(0) = 0$ and $S(q_k) = S(0) = \mathbb{E}[X]$. Pareto domination requires that $\tilde{r}_k S(\tilde{q}_k) \geq r_k \mathbb{E}[X]$. Since $\tilde{r}_k \leq r_k$ and $S(\tilde{q}_k) \leq \mathbb{E}[X]$, equality must hold throughout. Hence $\tilde{r}_k = r_k$ and $S(\tilde{q}_k) = \mathbb{E}[X]$. Since $F(0) = 0$ and S is strictly decreasing on $[0, 1]$, this implies $\tilde{q}_k = 0 = q_k$. Moreover, if any earlier inequality $\tilde{r}_i \leq r_i$ were strict, it would persist until agent k , since $q_i > 0$ for every $i < k$. This would contradict $\tilde{r}_k = r_k$. Thus the two rules agree through agent k , and both have zero reach probability afterward. Their utility profiles again coincide. Therefore no online rule can strictly Pareto dominate π . $\square$

C Proofs for the Fair Division Results

C.1 Proof of Proposition 4.1

The key idea of this proof is showing that for every allocation rule for the optimization problem, there is a rule for the fair division problem that achieves exactly m times the expected value for each agent, and for every allocation rule for the fair division problem, there is a rule for the optimization problem that achieves exactly $\frac{1}{m}$ times the expected value for each agent. Further, for any online allocation rule for either problem, the corresponding rule for the other problem is also online.

Fix n, m , and $\mathcal{D}$. Given a single-item allocation rule π , apply it independently to every item. Each item contributes $v_i(\pi, \mathcal{D})$ to agent i 's expected utility, so linearity of expectation gives total expected utility $mv_i(\pi, \mathcal{D})$. This construction preserves online feasibility.

Conversely, fix a many-item allocation rule, and let A_i denote the random bundle it allocates to agent i . Construct a single-item rule π as follows. Choose an item K uniformly from $[m]$, independently of all values and internal randomness. Use the real single-item values $X_1, \dots, X_n$ as

the values for item K , and independently sample the values for all other items from $\mathcal{D}$. Simulate the many-item rule on this valuation matrix, and allocate the real item to whichever agent receives item K , leaving it unallocated if the simulated rule does so.

Conditional on any choice of K , the simulated valuation matrix has the same distribution as the original many-item instance. Therefore,

$$v_i(\pi, \mathcal{D}) = \frac{1}{m} \sum_{j=1}^m \mathbb{E}[X_{ij} \mathbf{1}_{j \in A_i}] = \frac{\mathbb{E}[U_i]}{m}.$$

If the many-item rule is online, supply the simulated values as the corresponding agents arrive. Its decision at agent i then uses only the real values observed so far and internally generated values, so the resulting single-item rule is also online. The construction applies to offline rules as well.

Thus, the online and offline expected-utility sets in the many-item problem are exactly m times their single-item counterparts. Since w_p is positively homogeneous (scaling all entries by λ scales the output by λ), it follows that $\text{ALG}_{n,m}^{\text{ea}}(p; \mathcal{D}) = m \text{ALG}_n(p; \mathcal{D})$ and $\text{OPT}_{n,m}^{\text{ea}}(p; \mathcal{D}) = m \text{OPT}_n(\mathcal{D})$. Taking ratios and then the infimum over $\mathcal{D}$ gives $\text{CR}_{n,m}^{\text{ea}}(p) = \text{CR}_n(p)$. $\square$

C.2 Proof of Proposition 4.2

Fix $p \leq 0$ and $n \geq 2$, and consider the distribution $\mathcal{D} = \text{Bernoulli}(\lambda)$, where $\lambda = 2 \log(n)/n$. If any agent has zero utility, the p -mean is zero. Thus, positive welfare requires every agent to receive at least one value-one item. Since $m = n$, this means that every agent receives exactly one item, which they value at one. Consequently, the p -mean is either zero or one, and its expectation equals the probability that every agent receives a value-one item.

Prophet. The prophet can give every agent a value-one item if there is a permutation π of the items such that $X_{i,\pi(i)} = 1$ for every agent i . Define a bipartite graph G whose two parts represent the agents and items, with an edge between agent i and item j if and only if $X_{ij} = 1$. Finding such an allocation is equivalent to finding a perfect matching in G .

Since the valuations are independent Bernoulli random variables, G is a random bipartite graph $G(n, n, \lambda)$. For $\lambda = 2 \log(n)/n$, this graph contains a perfect matching with probability $1 - o(1)$ [FK15]. Therefore, $\text{OPT}_{n,n}^{\text{ep}}(p; \mathcal{D}) = 1 - o(1)$.

Online algorithm. Fix any online allocation rule and condition on the history before the final agent arrives. Positive welfare is possible only if each of the first $n - 1$ agents has received exactly one value-one item and exactly one item remains. On any such history, the remaining item's value to the final agent is an independent Bernoulli random variable with parameter λ . Thus, the conditional probability of positive welfare is at most λ . On every other history, it is zero.

It follows that every online rule has expected welfare at most λ , so $\text{ALG}_{n,n}^{\text{ep}}(p; \mathcal{D}) \leq \lambda$. Therefore,

$$\text{CR}_{n,n}^{\text{ep}}(p) \leq \frac{\text{ALG}_{n,n}^{\text{ep}}(p; \mathcal{D})}{\text{OPT}_{n,n}^{\text{ep}}(p; \mathcal{D})} \leq \frac{2 \log(n)}{n(1 - o(1))} = O\left(\frac{\log n}{n}\right).$$

In particular, $\text{CR}_{n,n}^{\text{ep}}(p) \rightarrow 0$ as $n \rightarrow \infty$ for every $p \leq 0$. $\square$

C.3 Proof of Proposition 4.3

Fix a distribution $\mathcal{D}$, $n \geq 2$ and $p \in [-\infty, 1]$. Fix a single-item allocation rule π , and apply it independently to every item, using independent internal randomness for each copy. Write agent i 's

total utility as $U_i = \sum_{j=1}^m Y_{ij}$, where Y_{ij} is their utility from item j . For each agent i , the variables Y_{ij} are i.i.d. across items, with expectation $v_i(\pi, \mathcal{D})$ and finite mean since $0 \leq Y_{ij} \leq X_{ij}$. By the strong law of large numbers, $U_i/m \rightarrow v_i(\pi, \mathcal{D})$ almost surely. Since there are finitely many agents and the p -mean is continuous on nonnegative utility profiles, this gives $w_p(\mathbf{U}/m) \rightarrow w_p(\mathbf{v}(\pi, \mathcal{D}))$ almost surely.

We next show that the expectations also converge. Since the p -mean is nonnegative, Fatou's lemma gives

$$w_p(\mathbf{v}(\pi, \mathcal{D})) \leq \liminf_{m \rightarrow \infty} \mathbb{E}[w_p(\mathbf{U}/m)].$$

On the other hand, concavity of the p -mean and Jensen's inequality give $\mathbb{E}[w_p(\mathbf{U}/m)] \leq w_p(\mathbb{E}[\mathbf{U}/m]) = w_p(\mathbf{v}(\pi, \mathcal{D}))$ for every m . Combining these bounds and using positive homogeneity, we obtain

$$\lim_{m \rightarrow \infty} \frac{\mathbb{E}[w_p(\mathbf{U})]}{m} = w_p(\mathbf{v}(\pi, \mathcal{D})).$$

This provides a lower bound on the optimal ex-post welfare. Indeed, for every online single-item rule π , independent repetition is a feasible online many-item rule, so

$$\liminf_{m \rightarrow \infty} \frac{\text{ALG}_{n,m}^{\text{ep}}(p; \mathcal{D})}{m} \geq w_p(\mathbf{v}(\pi, \mathcal{D})).$$

Taking the supremum over online single-item rules gives a lower bound of $\text{ALG}_n(p; \mathcal{D})$. The same argument for offline rules gives a lower bound of $\text{OPT}_n(\mathcal{D})$.

For the upper bounds, Jensen's inequality shows that the ex-post welfare of any rule is at most its ex-ante welfare. Applying Proposition 4.1, we therefore have

$$\text{ALG}_{n,m}^{\text{ep}}(p; \mathcal{D}) \leq m \text{ALG}_n(p; \mathcal{D}), \quad \text{OPT}_{n,m}^{\text{ep}}(p; \mathcal{D}) \leq m \text{OPT}_n(\mathcal{D}).$$

Together, the upper and lower bounds show that

$$\lim_{m \rightarrow \infty} \frac{\text{ALG}_{n,m}^{\text{ep}}(p; \mathcal{D})}{m} = \text{ALG}_n(p; \mathcal{D}), \quad \lim_{m \rightarrow \infty} \frac{\text{OPT}_{n,m}^{\text{ep}}(p; \mathcal{D})}{m} = \text{OPT}_n(\mathcal{D}).$$

Since $\text{OPT}_n(\mathcal{D}) > 0$, taking the ratio gives

$$\lim_{m \rightarrow \infty} \frac{\text{ALG}_{n,m}^{\text{ep}}(p; \mathcal{D})}{\text{OPT}_{n,m}^{\text{ep}}(p; \mathcal{D})} = \frac{\text{ALG}_n(p; \mathcal{D})}{\text{OPT}_n(\mathcal{D})}.$$

This completes the proof. □

C.4 Finite Convergence Bounds for Restricted Distributions

We now give quantitative bounds on the number of items needed for p -mean welfare to approach the single-item guarantee.

Proposition C.1. *Fix $n \geq 2$, $p \in [-\infty, 1]$, $\delta \in (0, 1)$, and a distribution $\mathcal{D}$ with mean μ . Each of the following conditions is sufficient to guarantee*

$$\frac{\text{ALG}_{n,m}^{\text{ep}}(p; \mathcal{D})}{\text{OPT}_{n,m}^{\text{ep}}(p; \mathcal{D})} \geq \frac{\text{ALG}_n(p; \mathcal{D})}{\text{OPT}_n(\mathcal{D})} - \delta.$$

1. $\mathcal{D}$ has finite variance σ^2 , and

$$m \geq \frac{9n^3}{\delta^2} \left(1 + \frac{\sigma^2}{\mu^2} \right).$$

2. $X \leq B\mu$ almost surely, and

$$m \geq \frac{24nB}{\delta^2} \log \left(\frac{2n}{\delta} \right).$$

Proof. For $p = 1$, the ex-ante and ex-post objectives coincide, so the result follows from Proposition 4.1. Henceforth, assume $p < 1$.

Fix an optimal single-item rule for the p -mean objective, and let $\mathbf{v} = (v_1, \dots, v_n)$ be its expected-utility profile. Such a rule exists by Proposition 3.4, since threshold profiles depend continuously on thresholds in the compact set $[0, 1]^{n-1}$. For $p = -\infty$, choose the equalizing rule from Proposition 3.5.

Write $a = w_p(\mathbf{v}) = \text{ALG}_n(p; \mathcal{D})$ and $V = \text{OPT}_n(\mathcal{D})$. Applying the rule independently to each item, including independent internal randomness, gives total utilities $U_i = \sum_{j=1}^m Y_{ij}$, where the Y_{ij} are independent across items, satisfy $0 \leq Y_{ij} \leq X_{ij}$, and have expectation v_i . Let $Z_i = U_i/m$.

By Jensen's inequality and Proposition 4.1, $\text{OPT}_{n,m}^{\text{ep}}(p; \mathcal{D}) \leq mV$. Consequently,

$$\frac{\text{ALG}_{n,m}^{\text{ep}}(p; \mathcal{D})}{\text{OPT}_{n,m}^{\text{ep}}(p; \mathcal{D})} \geq \frac{\mathbb{E}[w_p(\mathbf{Z})]}{V}.$$

Since $a/V \leq 1$, it suffices to prove $\mathbb{E}[w_p(\mathbf{Z})] \geq (1 - \delta)a$.

Properties of an optimal single-item rule. Allocating to a uniformly random agent gives every agent expected utility μ/n , while every agent's expected utility is at most μ . Thus, $\mu/n \leq a \leq \mu$.

For finite $p < 1$, every v_i is positive. For $p \leq 0$, this follows from $a > 0$. For $0 < p < 1$, maximizing w_p is equivalent to maximizing $\sum_i v_i^p$. If some coordinate is zero, mixing toward $(\mu/n, \dots, \mu/n)$ with probability ε increases the contribution of the zero coordinates by $\Omega(\varepsilon^p)$, while changing the contribution of the positive coordinates by at most $O(\varepsilon)$. Since $\varepsilon^p/\varepsilon \rightarrow \infty$ as $\varepsilon \downarrow 0$, this contradicts optimality.

Define

$$g_i = \frac{\partial w_p}{\partial v_i}(\mathbf{v}) = \frac{1}{n} a^{1-p} v_i^{p-1}.$$

This formula also holds for $p = 0$, and $\sum_i g_i v_i = a$. Allocating the item to agent i regardless of its value gives the feasible profile $\mu \mathbf{e}_i$. By optimality of $\mathbf{v}$, mixing toward this profile cannot have a positive directional derivative. Therefore, $\mu g_i \leq a$ and hence $v_i^{1-p} \geq \frac{\mu}{na^p}$.

Let $p_0 = 1/(1 + \log n)$. If $p \leq 0$, the preceding bound and $a \geq \mu/n$ give $v_i \geq \mu/n$. If $0 < p \leq p_0$, using $a \leq \mu$ instead gives

$$v_i \geq \mu n^{-1/(1-p)} \geq \frac{\mu}{en} \geq \frac{\mu}{3n}.$$

For $p = -\infty$, the equalizing rule also satisfies $v_i = a \geq \mu/n$. Thus, whenever $p \leq p_0$, every agent has expected utility at least $\mu/(3n)$.

For $p_0 < p < 1$, we will instead use the following bound, valid for every nonnegative profile $\mathbf{z}$:

$$a - w_p(\mathbf{z}) \leq \frac{1}{p} \sum_i g_i (v_i - z_i)_+.$$

To verify it, let $d_i = (v_i - z_i)_+$ and $t = \sum_i g_i d_i / a \in [0, 1]$. Since $(1 - s)^p \geq 1 - s$ for $s \in [0, 1]$,

$$w_p(\mathbf{z})^p \geq \frac{1}{n} \sum_i (v_i - d_i)^p \geq a^p (1 - t).$$

Taking the p th root and using $(1 - t)^{1/p} \geq 1 - t/p$ proves the bound.

Finite variance. Suppose that $\mathcal{D}$ has finite variance σ^2 , and write $M_2 = \mathbb{E}[X^2] = \sigma^2 + \mu^2$. Since $0 \leq Y_{ij} \leq X_{ij}$ and the contributions are independent across items, $\text{Var}(Z_i) \leq M_2/m$.

First suppose that $p \leq p_0$. Define $T = \max_i (v_i - Z_i)_+/v_i$. Then $T \in [0, 1]$ and $\mathbf{Z} \geq (1 - T)\mathbf{v}$ coordinatewise, so $w_p(\mathbf{Z}) \geq (1 - T)a$. Moreover,

$$\mathbb{E}[T] \leq \sqrt{\sum_i \frac{\mathbb{E}[(v_i - Z_i)_+]^2}{v_i^2}} \leq \sqrt{\sum_i \frac{\text{Var}(Z_i)}{v_i^2}} \leq 3\sqrt{\frac{n^3 M_2}{m\mu^2}}.$$

Here the first inequality follows from bounding the maximum by the Euclidean norm and applying Jensen's inequality.

Now suppose that $p_0 < p < 1$. Using the bound from the preceding paragraph, $g_i \leq a/\mu$, and $\mathbb{E}[(v_i - Z_i)_+] \leq \sqrt{\text{Var}(Z_i)}$, we obtain

$$a - \mathbb{E}[w_p(\mathbf{Z})] \leq \frac{1}{p} \sum_i g_i \mathbb{E}[(v_i - Z_i)_+] \leq \frac{na}{p\mu} \sqrt{\frac{M_2}{m}} \leq 2a\sqrt{\frac{n^3 M_2}{m\mu^2}},$$

where the last inequality uses $1/p < 1 + \log n \leq 2\sqrt{n}$.

Thus, in either case,

$$\mathbb{E}[w_p(\mathbf{Z})] \geq a \left(1 - 3\sqrt{\frac{n^3}{m} \left(1 + \frac{\sigma^2}{\mu^2} \right)} \right).$$

The stated bound on m makes the error at most δ .

Bounded values. Now suppose that $X \leq B\mu$ almost surely. First consider $p \leq p_0$. Since $Y_{ij} \in [0, B\mu]$ and $v_i \geq \mu/(3n)$, a multiplicative Chernoff bound gives

$$\Pr[Z_i < (1 - \delta/2)v_i] \leq \exp\left(-\frac{\delta^2 m v_i}{8B\mu}\right) \leq \exp\left(-\frac{\delta^2 m}{24nB}\right).$$

Under the stated bound on m , a union bound shows that $\mathbf{Z} \geq (1 - \delta/2)\mathbf{v}$ coordinatewise with probability at least $1 - \delta/2$. Nonnegativity, monotonicity, and positive homogeneity therefore give

$$\mathbb{E}[w_p(\mathbf{Z})] \geq (1 - \delta/2)^2 a \geq (1 - \delta)a.$$

It remains to consider $p_0 < p < 1$. We use a moment bound to control the expected welfare directly. For any nonnegative random variable Z with mean $v > 0$ and finite second moment, Hölder's inequality gives

$$\mathbb{E}[Z^p] \geq \frac{v^{2-p}}{(\mathbb{E}[Z^2])^{1-p}} = v^p \left(1 + \frac{\text{Var}(Z)}{v^2} \right)^{-(1-p)} \geq v^p - (1-p)v^{p-2} \text{Var}(Z).$$

The last inequality follows from convexity of $t \mapsto (1+t)^{-(1-p)}$.

Since $Y_{ij}^2 \leq B\mu Y_{ij}$, $\text{Var}(Z_i) \leq B\mu v_i/m$. Applying the moment bound and using $\sum_i g_i \leq na/\mu$, we obtain

$$\begin{aligned} \frac{1}{n} \sum_i \mathbb{E}[Z_i^p] &\geq a^p - \frac{(1-p)B\mu}{nm} \sum_i v_i^{p-1} \\ &= a^p - \frac{(1-p)B\mu}{m} a^{p-1} \sum_i g_i \\ &\geq a^p \left(1 - \frac{(1-p)nB}{m} \right). \end{aligned}$$

The assumed bound on m implies $m \geq nB$. Since $x \mapsto x^{1/p}$ is convex, Jensen's inequality gives

$$\begin{aligned}\mathbb{E}[w_p(\mathbf{Z})] &\geq \left(\frac{1}{n} \sum_i \mathbb{E}[Z_i^p] \right)^{1/p} \\ &\geq a \left(1 - \frac{(1-p)nB}{pm} \right) \\ &\geq a \left(1 - \frac{nB \log n}{m} \right),\end{aligned}$$

where the final inequality uses $(1-p)/p < \log n$. The stated bound on m makes this error at most δ , completing the proof. $\square$

D Technical Details for the Proof of Lemma 5.2

D.1 Proof of Lemma 5.3

Observe that $\Lambda_R > 1$. Indeed, define

$$L_R(z) = \begin{cases} \left(1 - \frac{1}{R}\right)z, & 0 \leq z < R, \\ +\infty, & z \geq R. \end{cases}$$

Then $\mathcal{C}(L_R) = \frac{1}{1-1/R} \int_1^R \frac{dz}{z^2} = 1$, whereas

$$\mathcal{I}(L_R) = \int_0^R e^{-(1-1/R)z} dz = \frac{1 - e^{-(R-1)}}{1 - 1/R} > 1.$$

The final inequality follows from $e^{R-1} > R$.

Existence of a maximizer. We first show that the supremum defining Λ_R is attained. Let (G_i) be a sequence in $\mathcal{F}_R$ such that $\mathcal{I}(G_i) \rightarrow \Lambda_R$, and define $\varphi: [0, \infty] \rightarrow [0, 1]$, as $\varphi(x) = \frac{x}{1+x}$, with $\varphi(+\infty) = 1$. Let

$$G'_i = \varphi \circ G_i.$$

Each G'_i is nondecreasing and takes values in $[0, 1]$. We use the following form of Helly's selection theorem.

Lemma D.1 (Helly's selection theorem [Car00, Lemma 13.15]). *Let (f_i) be a sequence of nondecreasing functions $f_i: [0, T] \rightarrow [0, 1]$. Then there exists a nondecreasing function $f: [0, T] \rightarrow [0, 1]$ and a subsequence (f_{i_j}) such that $\lim_{j \rightarrow \infty} f_{i_j}(x) = f(x)$ for every $x \in [0, T]$.*

Applying Helly's selection theorem (Lemma D.1) to (G'_i) on $[0, R]$, there exist indices $i_1 < i_2 < \dots$ and a nondecreasing function $G': [0, R] \rightarrow [0, 1]$ such that $G'_{i_j}(z) \rightarrow G'(z)$ for every $z \in [0, R]$. Since $G'_i(z) = 1$ for every $z \geq R$, we extend G' by setting $G'(z) = 1$ for $z \geq R$.

Define

$$G^*(z) = \varphi^{-1}(G'(z)) = \frac{G'(z)}{1 - G'(z)},$$

where the right-hand side is interpreted as $+\infty$ when $G'(z) = 1$. Then $G_{i_j}(z) \rightarrow G^*(z)$ for every $z \geq 0$. Pointwise limits preserve nonnegativity, monotonicity, and convexity. Moreover, $G^*(0) = 0$,

and since $G_{i_j}(z) = +\infty$ for every $z \geq R$, we also have $G^*(z) = +\infty$ for every $z \geq R$. Fatou's lemma allows us to upper-bound the integral of the pointwise limit by the liminf of the integrals:

$$\mathcal{C}(G^*) = \int_1^\infty \lim_{j \rightarrow \infty} \frac{1}{zG_{i_j}(z)} dz \leq \liminf_{j \rightarrow \infty} \int_1^\infty \frac{dz}{zG_{i_j}(z)} \leq 1.$$

Thus $G^* \in \mathcal{F}_R$.

It remains to verify that the objective value is preserved in the limit. We have that $0 \leq e^{-G_{i_j}(z)} \leq \mathbf{1}\{z < R\}$, and the function on the right is integrable. Therefore, the dominated convergence theorem gives

$$\mathcal{I}(G^*) = \int_0^\infty e^{-G^*(z)} dz = \lim_{j \rightarrow \infty} \int_0^\infty e^{-G_{i_j}(z)} dz = \Lambda_R.$$

Hence G^* is a maximizer over $\mathcal{F}_R$, and since the original function G belongs to $\mathcal{F}_R$, $\mathcal{I}(G) \leq \mathcal{I}(G^*)$.

Basic properties of the maximizer. Define $M := \sup\{z \geq 0 : G^*(z) < +\infty\}$ and $J := \sup\{z \geq 0 : G^*(z) = 0\}$. Since $G^* \in \mathcal{F}_R$, we have $M \leq R$.

First, $M > 1$. Indeed, we showed above that $\mathcal{I}(G^*) = \Lambda_R > 1$, so if $M \leq 1$, then

$$\mathcal{I}(G^*) = \int_0^M e^{-G^*(z)} dz \leq M \leq 1,$$

which is a contradiction.

We next show that $J < 1$. Suppose instead that $J \geq 1$. Since G^* is nonnegative and nondecreasing, this implies $G^*(1) = 0$. Choose any $z_0 \in (1, M)$. If $G^*(z_0) = 0$, then $\mathcal{C}(G^*) = +\infty$. Otherwise, convexity gives, for $1 < z < z_0$, $G^*(z) \leq \frac{z-1}{z_0-1} G^*(z_0)$. Consequently,

$$\mathcal{C}(G^*) \geq \int_1^{z_0} \frac{dz}{zG^*(z)} \geq \frac{z_0-1}{G^*(z_0)} \int_1^{z_0} \frac{dz}{z(z-1)} = +\infty,$$

contradicting $\mathcal{C}(G^*) \leq 1$. Hence $0 \leq J < 1 < M \leq R$.

We also claim that $\mathcal{C}(G^*) = 1$. Since $J < 1 < M$, the function G^* is finite and strictly positive on $(1, M)$, so $\mathcal{C}(G^*) > 0$. Suppose that $c := \mathcal{C}(G^*) < 1$. Then $cG^* \in \mathcal{F}_R$ and $\mathcal{C}(cG^*) = \frac{\mathcal{C}(G^*)}{c} = 1$. Moreover, since $0 < c < 1$ and $G^*(z) > 0$ for $1 < z < M$, $e^{-cG^*(z)} > e^{-G^*(z)}$. It follows that $\mathcal{I}(cG^*) > \mathcal{I}(G^*)$, contradicting the maximality of G^* over $\mathcal{F}_R$. Therefore $\mathcal{C}(G^*) = 1$. $\square$

D.2 Proof of Lemma 5.4

Fix $a \in [0, M)$, abbreviate G^* to G , and write

$$h = h_a, \quad G_\varepsilon = G + \varepsilon h, \quad c(\varepsilon) = \mathcal{C}(G_\varepsilon).$$

Recall that $J < 1 < M$, G is finite on $[0, M)$, and $\mathcal{C}(G) = 1$. In particular, $G(1) > 0$, and monotonicity gives $G(z) \geq G(1)$ for $z \geq 1$.

We first justify the derivative of the constraint. Since $G_\varepsilon = +\infty$ on (M, ∞) , for every $\varepsilon > 0$,

$$\frac{c(\varepsilon) - c(0)}{\varepsilon} = - \int_1^M \frac{h(z)}{zG(z)(G(z) + \varepsilon h(z))} dz.$$

The integrand converges pointwise as $\varepsilon \downarrow 0$ to $-h(z)/(zG(z)^2)$. Moreover,

$$0 \leq \frac{h(z)}{zG(z)(G(z) + \varepsilon h(z))} \leq \frac{h(z)}{zG(z)^2} \leq \frac{M}{G(1)} \frac{1}{zG(z)}.$$

The final function is integrable because $\mathcal{C}(G) = 1$. Dominated convergence therefore gives

$$c'(0+) = - \int_1^M \frac{h(z)}{zG(z)^2} dz = -A(a).$$

We next differentiate the unnormalized objective. For $\varepsilon > 0$,

$$\frac{\mathcal{I}(G_\varepsilon) - \mathcal{I}(G)}{\varepsilon} = - \int_0^M e^{-G(z)} \frac{1 - e^{-\varepsilon h(z)}}{\varepsilon} dz.$$

Since $0 \leq 1 - e^{-t} \leq t$ for $t \geq 0$, the absolute value of the integrand is bounded by

$$h(z)e^{-G(z)} \leq M \mathbf{1}\{z < M\},$$

which is integrable. Dominated convergence thus yields

$$\left. \frac{d}{d\varepsilon} \mathcal{I}(G + \varepsilon h) \right|_{\varepsilon=0+} = - \int_0^M h(z) e^{-G(z)} dz = -B(a).$$

The same argument applies to the scaling perturbation. For $|\varepsilon| \leq 1/2$, the derivative of $e^{-(1+\varepsilon)G(z)}$ is bounded in absolute value by $G(z)e^{-G(z)/2}$, which is uniformly bounded for $G(z) \geq 0$. Hence differentiation under the integral sign is valid and gives

$$\left. \frac{d}{d\varepsilon} \mathcal{I}((1 + \varepsilon)G) \right|_{\varepsilon=0} = - \int_0^M G(z) e^{-G(z)} dz = -\lambda.$$

It remains to differentiate the normalized perturbation. By definition,

$$\tilde{G}_{a,\varepsilon} = c(\varepsilon)G_\varepsilon = c(\varepsilon)(G + \varepsilon h).$$

Because $c(0) = 1$ and $c'(0+) = -A(a)$, for every $z < M$,

$$\frac{\tilde{G}_{a,\varepsilon}(z) - G(z)}{\varepsilon} \longrightarrow h(z) - A(a)G(z).$$

We verify that this pointwise derivative may be passed through the objective integral. Since $c(\varepsilon) \rightarrow 1$, for all sufficiently small $\varepsilon > 0$ we have $1/2 \leq c(\varepsilon) \leq 2$, and $(c(\varepsilon) - 1)/\varepsilon$ is bounded. The mean value theorem therefore gives

$$\left| \frac{e^{-\tilde{G}_{a,\varepsilon}(z)} - e^{-G(z)}}{\varepsilon} \right| \leq (KG(z) + 2M)e^{-G(z)/2}$$

for some constant K independent of z and ε . The right-hand side is bounded on $[0, M)$, since both $xe^{-x/2}$ and $e^{-x/2}$ are bounded on $[0, \infty)$, and it vanishes outside the finite interval $[0, M)$. It is therefore an integrable dominating function. Dominated convergence now gives

$$\begin{aligned} \Psi(a) &= \left. \frac{d}{d\varepsilon} \mathcal{I}(\tilde{G}_{a,\varepsilon}) \right|_{\varepsilon=0+} \\ &= - \int_0^M (h(z) - A(a)G(z)) e^{-G(z)} dz \\ &= \lambda A(a) - B(a). \end{aligned}$$

Finally, all three quantities are finite and strictly positive. Indeed,

$$A(a) \leq \frac{M}{G(1)} \int_1^M \frac{dz}{zG(z)} = \frac{M}{G(1)} < \infty, \quad B(a) \leq M^2 < \infty,$$

and $\lambda < \infty$ because $xe^{-x} \leq 1/e$ for $x \geq 0$. Because $a < M$ and $M > 1$, the integrand defining $A(a)$ is strictly positive on $(\max\{1, a\}, M)$, while the integrand defining $B(a)$ is strictly positive on (a, M) . Finally, $G(z) > 0$ on (J, M) , so the integrand defining λ is strictly positive there. Hence $A(a), B(a), \lambda > 0$. $\square$

D.3 Proof of Lemma 5.5

Abbreviate G^* to G , and define

$$f(z) := \frac{1}{zG(z)^2} \quad (1 < z < M), \quad g(z) := e^{-G(z)} \quad (0 < z < M).$$

Both functions are integrable on their respective domains. Indeed, $0 \leq g \leq 1$, while $G(z) \geq G(1) > 0$ for $z \geq 1$, and hence

$$\int_1^M f(z) dz \leq \frac{1}{G(1)} \int_1^M \frac{dz}{zG(z)} = \frac{1}{G(1)} < \infty.$$

Moreover, because G is finite and convex on $[0, M)$, it is continuous on $(0, M)$. Thus g is continuous on $(0, M)$, and f is continuous on $(1, M)$.

We will repeatedly use the following elementary differentiation formula. If w is integrable and nonnegative on (ℓ, M) , then

$$K_w(a) := \int_a^M (z - a)w(z) dz$$

satisfies

$$K'_w(a) = - \int_a^M w(z) dz.$$

To see this without invoking an unproved Leibniz rule, Tonelli's theorem gives

$$K_w(a) = \int_a^M \int_a^z w(z) dt dz = \int_a^M \left(\int_t^M w(z) dz \right) dt.$$

The tail integral $t \mapsto \int_t^M w(z) dz$ is continuous, so the fundamental theorem of calculus gives the asserted derivative. If w is continuous at a , differentiating the tail integral once more gives

$$K''_w(a) = w(a).$$

Apply this formula first to

$$B(a) = \int_a^M (z - a)g(z) dz.$$

For every $a \in [0, M)$, where the derivative at $a = 0$ is understood as a right derivative,

$$B'(a) = - \int_a^M e^{-G(z)} dz.$$

This derivative is continuous in a . Since g is continuous on $(0, M)$, we also have

$$B''(a) = e^{-G(a)} \quad (0 < a < M).$$

We next differentiate A . If $0 \leq a \leq 1$, then $z \geq 1$ implies $h_a(z) = z - a$, so

$$A(a) = \int_1^M (z - a)f(z) dz$$

is affine in a , with

$$A'(a) = - \int_1^M f(z) dz = - \int_1^M \frac{dz}{zG(z)^2}.$$

In particular, $A''(a) = 0$ for $0 < a < 1$.

If $1 < a < M$, then

$$A(a) = \int_a^M (z - a)f(z) dz.$$

The preceding differentiation formula gives

$$A'(a) = - \int_a^M f(z) dz = - \int_a^M \frac{dz}{zG(z)^2},$$

and, since f is continuous on $(1, M)$,

$$A''(a) = f(a) = \frac{1}{aG(a)^2}.$$

The two formulas for A' agree at $a = 1$, because

$$\lim_{a \downarrow 1} \left(- \int_a^M f(z) dz \right) = - \int_1^M f(z) dz.$$

Thus A is continuously differentiable on $[0, M)$, with a one-sided derivative at 0, and is twice continuously differentiable separately on $(0, 1)$ and $(1, M)$.

By Lemma 5.4, $\Psi(a) = \lambda A(a) - B(a)$. Combining the preceding derivative formulas gives

$$\Psi'(a) = \begin{cases} -\lambda \int_1^M \frac{dz}{zG(z)^2} + \int_a^M e^{-G(z)} dz, & 0 \leq a \leq 1, \\ \int_a^M \left(e^{-G(z)} - \frac{\lambda}{zG(z)^2} \right) dz, & 1 < a < M. \end{cases}$$

The expressions agree at $a = 1$, and the continuity of tail integrals shows that Ψ' is continuous on $[0, M)$. Therefore Ψ is continuously differentiable there.

Finally, differentiating once more on either side of $a = 1$ yields

$$\Psi''(a) = \begin{cases} -e^{-G(a)}, & 0 < a < 1, \\ \frac{\lambda}{aG(a)^2} - e^{-G(a)}, & 1 < a < M. \end{cases}$$

Because G is continuous on $(0, M)$, these expressions are continuous on $(0, 1)$ and $(1, M)$, respectively. Hence Ψ is twice continuously differentiable away from $a = 1$, as claimed. $\square$

D.4 Proof of Lemma 5.6

We average the first-order effect $\Psi(a)$ over all the hinges appearing in the representation of G^* . The key point is that averaging these hinges against μ^* reproduces G^* itself.

We begin with A . Recalling its definition and applying Tonelli's theorem to the nonnegative integrand, we obtain

$$\begin{aligned}
\int_{[J,M)} A(a) \mu^*(da) &= \int_{[J,M)} \left(\int_1^M \frac{h_a(z)}{zG^*(z)^2} dz \right) \mu^*(da) \\
&= \int_1^M \frac{\int_{[J,M)} h_a(z) \mu^*(da)}{zG^*(z)^2} dz \\
&= \int_1^M \frac{G^*(z)}{zG^*(z)^2} dz \\
&= \int_1^M \frac{dz}{zG^*(z)} \\
&= \mathcal{C}(G^*) = 1.
\end{aligned}$$

The analogous calculation for B gives

$$\begin{aligned}
\int_{[J,M)} B(a) \mu^*(da) &= \int_{[J,M)} \left(\int_0^M h_a(z) e^{-G^*(z)} dz \right) \mu^*(da) \\
&= \int_0^M \left(\int_{[J,M)} h_a(z) \mu^*(da) \right) e^{-G^*(z)} dz \\
&= \int_0^M G^*(z) e^{-G^*(z)} dz \\
&= \lambda.
\end{aligned}$$

Again, Tonelli's theorem applies because every integrand being exchanged is nonnegative.

Since $\Psi(a) = \lambda A(a) - B(a)$, these two identities imply

$$\begin{aligned}
\int_{[J,M)} \Psi(a) \mu^*(da) &= \lambda \int_{[J,M)} A(a) \mu^*(da) - \int_{[J,M)} B(a) \mu^*(da) \\
&= \lambda \mathcal{C}(G^*) - \lambda \\
&= 0.
\end{aligned}$$

D.5 Proof of Lemma 5.7

Abbreviate G^* to G . It remains to exclude support in $(1, M)$. Define $m(z) = zG(z)^2 e^{-G(z)}$. By Lemmas 5.4 and 5.5, for $1 < a, z < M$,

$$\Psi''(z) = \frac{\lambda - m(z)}{zG(z)^2}, \quad \Psi(a) = \int_a^M (z - a) \Psi''(z) dz.$$

We first obtain a boundary condition on m . Since G is nondecreasing, the limit $L = \lim_{z \uparrow M} G(z)$ exists in $(0, \infty]$. Consequently, $m_M = \lim_{z \uparrow M} m(z)$ also exists, with $m_M = 0$ if $L = +\infty$. If $m_M < \lambda$, then $\Psi''(z) > 0$ for all z sufficiently close to M . The integral identity above would therefore give $\Psi(a) > 0$ for a sufficiently close to M , contradicting $\Psi \leq 0$. Hence $m_M \geq \lambda$, and in particular $L < \infty$.

The key observation is that m has at most one peak. On the interval where $G(z) \leq 2$, it is strictly increasing, since $x \mapsto x^2 e^{-x}$ is nondecreasing on $(0, 2]$. On the interval where $G(z) \geq 2$, we have

$$\log m(z) = \log z + 2 \log G(z) - G(z).$$

The function $x \mapsto 2 \log x - x$ is concave and nonincreasing on $[2, \infty)$, so its composition with the convex function G is concave. Adding $\log z$ makes $\log m$ strictly concave on this interval. Thus, if $m(a) \geq \lambda$ for some $a \in (1, M)$, then $m(z) > \lambda$ for every $a < z < M$: this follows from strict increase up to $G = 2$ and strict log-concavity thereafter, using $m_M \geq \lambda$ at the right endpoint.

Now suppose that $a \in \text{supp}(\mu^*) \cap (1, M)$. By Equation (2), $\Psi(a) = 0$. Since $\Psi \leq 0$, this is a local maximum, so $\Psi''(a) \leq 0$ and hence $m(a) \geq \lambda$. The preceding observation gives $m(z) > \lambda$, and therefore $\Psi''(z) < 0$, for every $a < z < M$. But the integral identity then gives $\Psi(a) < 0$, a contradiction. Thus $\text{supp}(\mu^*) \cap (1, M) = \emptyset$, completing the proof. $\square$