

Egalitarianism of Random Assignment Mechanisms

Haris Aziz^{1,2}, Jiashu Chen², Aris Filos-Ratsikas³,
Simon Mackenzie^{1,2}, and Nicholas Mattei^{1,2}

¹ NICTA, Sydney, Australia
{haris.aziz,simon.mackenzie,nicholas.mattei}@nicta.com.au

² UNSW, Sydney, Australia
jiashu.chen@student.unsw.edu.au

³ University of Oxford, Oxford, UK and Aarhus University, Aarhus, Denmark
filosra@cs.au.dk

Abstract. We consider the egalitarian welfare aspects of random assignment mechanisms when agents have unrestricted cardinal utilities over the objects. We give bounds on how well different random assignment mechanisms approximate the optimal egalitarian value and investigate the effect that different well-known properties like ordinality, envy-freeness, and truthfulness have on the achievable egalitarian value. Finally, we conduct detailed experiments analyzing the trade-offs between efficiency with envy-freeness or truthfulness using two prominent random assignment mechanisms — random serial dictatorship and the probabilistic serial mechanism — for different classes of utility functions and distributions.

1 Introduction

We explore the tradeoffs between fairness and efficiency for randomized mechanisms for the *assignment problem*. Specifically, we consider settings where n agents express preferences (cardinal or ordinal) over a set of m indivisible objects. The objective is to assign the objects to agents in a fair and mutually beneficial manner (see e.g., [5, 6, 10, 25]). This general setting has a number of important and significant applications including the assignment of tasks to cores in cloud computing, kidneys to patients in organ exchanges, runways to airplanes in transportation, and students to seats in schools.

While it is sometimes assumed that the agents’ preferences over the objects can be expressed fully through *ordinal rankings*, most of the classical literature on the assignment problem [10, 25, 41] assumes the existence of an underlying utility structure, where each agent assigns real values or *cardinal valuations* to the different objects; these are von Neumann-Morgenstern utilities that specify the intensity of the agents’ preferences. Some classical papers (see e.g. [10]) focus on *ordinal mechanisms*, i.e. mechanisms that operate solely on the rankings consistent with the cardinal valuations, but *cardinal mechanisms*, where agents explicitly report their numerical values, have also been considered in literature [25, 40, 41].

A well-established criterion for fairness is the Rawlsian concept of maximizing the happiness of the least satisfied agent [37]. Following the spirit of this idea, we quantify

the fairness of an allocation in terms of its *egalitarian value*: the minimum ratio of the value of objects assigned to an agent to his total valuation for all the objects. The *optimal egalitarian value* for a valuation profile of all agents is the best egalitarian value achievable over all assignments. The optimal egalitarian value is equivalent to the maximum egalitarian welfare if each agent has a total utility of one for the set of all objects. The advantage of considering the optimal egalitarian value is that it does not change if agents scale their relative values for the objects.

The egalitarian value is not the only criterion for desirable allocation mechanisms. Allocation mechanisms may have other goals and requirements such as envy-freeness or truthfulness. Crucially, both these properties are incompatible with optimizing the egalitarian value except in very restricted domains [11]. Thus, it is natural to examine the *tradeoffs* between optimizing the egalitarian value and achieving other desirable properties. In some settings, such as kidney exchanges, the tradeoff between fairness and efficiency is of the utmost concern [17].

Evaluating these tradeoffs also motivates the study of how established mechanisms with other desiderata perform in terms of the egalitarian value. For a given mechanism J , we examine the approximation ratio $guar(J)$, which is the minimum ratio (among all valuation profiles) of the egalitarian value of an allocation returned by the mechanism to the optimal egalitarian value. Our work falls under the umbrella of *approximate mechanism design without money*, a framework set by Procaccia and Tennenholtz [35] for the study of how well mechanisms with certain properties approximate some objective function of the agents' inputs.

In this paper, we study randomized assignment mechanisms for which achieving ex ante fairness is easier compared to deterministic mechanisms. Thus, to evaluate the performance of the mechanisms, we compare their egalitarian value with the optimal egalitarian value achieved by any randomized allocation. Note that computing the allocation with the optimal egalitarian value is an NP-hard problem when we restrict ourselves to deterministic allocations [16]. On the other hand, when we consider randomized allocations, the optimal egalitarian value can be computed in polynomial time via a linear program.

We give extra consideration to two randomized assignment mechanisms — *random serial dictatorship (RSD)* and *probabilistic serial (PS)*, which are probably the best-known and most-studied mechanisms in the random assignment literature. In RSD,⁴ a permutation over the agents is selected uniformly at random and each agent in the permutation picks the most preferred m/n units of object that are not yet allocated [4, 10, 38]. In PS, each object is considered to have an infinitely divisible probability weight of one. To compute an allocation, agents simultaneously and with the same speed eat the probability weight of their most preferred object which has not been completely consumed. Once an object has been completely eaten by a subset of agents, each of these agents moves on to eat their next most preferred object that has not been completely eaten. The procedure terminates after all the objects have been eaten. The random allocation of an agent by PS is the amount of each object he has eaten (see e.g., [10, 27]). PS satisfies stochastic dominance (SD) envy-freeness (envy-freeness with

⁴ The original definition of RSD is for n agents and n objects; the definition here is a straightforward adaptation for $n < m$.

respect to all cardinal utilities consistent with the ordinal preferences). We also formalize a mechanism which we refer to as *Optimal Egalitarian and Envy-Free Mechanism (OEEF)*, which maximizes the egalitarian value of an allocation under the constraint that the allocation is envy-free. Allocations under this mechanism can be computed in polynomial time via linear programming.

1.1 Our contributions

We present novel theoretical and empirical results regarding fairness in randomized mechanisms. Our main theoretical contributions are as follows.

- For any SD envy-free mechanism J : $guar(J) = O(n^{-1})$.
- For any envy-free mechanism J : $guar(J) = \Omega(n^{-1})$ and $guar(J) = O(n^{-1/5})$.
- For any truthful-in-expectation mechanism J : $guar(J) = O(n^{-1/5})$.
- For any ordinal mechanism J : $guar(J) = O(n^{-1})$.

As a result of our general bounds, we also get asymptotically tight bounds of $\Theta(n^{-1})$ for RSD and PS. As a result of our general bounds for envy-free mechanisms, we obtain bounds for well-known envy-free mechanisms such as *competitive equilibrium with equal incomes (CEEI)* [40] and the *pseudo-market* mechanism [25]. Since a random assignment of indivisible objects can also be interpreted as a fractional assignment of divisible objects, *our results apply as well to fair allocation of divisible objects*.

The constructions that provide the upper bounds for the $guar$ values can be considered as extreme examples that may not be common in real-life scenarios. In order to better understand how the mechanisms may perform in practice, we consider the approximation ratio achieved by RSD and PS. We also examine the effect of imposing the envy-freeness constraint. We generate ordinal profiles via a Mallow’s model for different levels of dispersion ϕ from a common reference ranking of objects, assigning cardinal utilities via the Borda and exponential scoring functions. Sweeping ϕ from 0, where all agents have the same preference, to 1.0, where all preference orders are equally likely (the Impartial Culture), allows us to make statements regarding situations where agent preferences are more or less correlated. We make the following observations.

- There is a negligible difference between the minimum and average achievable approximation ratios for PS and RSD under Borda utilities. While PS preforms slightly better than RSD when agents have more extreme (exponential) utilities, both mechanisms preform strictly worse when agents’ valuations are more similar, as they are under Borda utilities.
- When we require envy-freeness (as in OEEF) with exponential utilities, as ϕ increases towards 1.0 (i.e. Impartial Culture) the achievable approximation ratio first decreases slightly and then increases. Hence, as agents value more disparate objects highly, satisfying envy-freeness does not impose as stiff a penalty on the achievable approximation ratio.
- In our experiments, the requirement of envy-freeness as a constraint in itself (as in the OEEF mechanism) does not have a large impact on the OEV. However, since PS returns an SD envy-free (envy-free for all cardinal utilities consistent with the

ordinal preferences) allocation, its achievable approximation ratio is strictly less than OEEF.

1.2 Related work

The assignment problem has been in the center of attention in recent years in both computer science and economics [2, 13, 22, 24, 21, 23, 33]. Often, in the classical assignment literature, agents are assumed to have an underlying cardinal utility preference structure, even if they are not asked to report it explicitly. On the other hand, there are many examples of well-known cardinal mechanisms, such the pseudo-market (PM) mechanism of Hylland and Zeckhauser [25] and the competitive equilibrium with equal incomes (CEEI) mechanism [40]. Both mechanisms return allocations that are envy-free in expectation. The two prominent ordinal mechanisms in the literature are the probabilistic serial mechanism (PS) [10, 14, 39] and random serial dictator (RSD), a folklore mechanism that pre-existed the formulation of the assignment problem in [25]. Later, Che and Kojima [14] proposed a variant of PS called *multi-unit eating probabilistic serial (MPS)* that was formalised and axiomatically studied by Aziz [2].

The egalitarian welfare has received considerable interest within the computer science literature, especially for allocation of discrete objects in a deterministic manner. The problem is also referred to as the *Santa Claus problem* in which the goal is to compute an assignment which maximizes the utility of the agent that gets the least utility (see e.g., [1, 7, 18, 32]). For deterministic settings, Demko and Hill [16] proved that the problem is NP-hard. On the other hand, for randomized/fractional allocations, the problem can be solved via a linear program.⁵ Recently, another fairness constraint that has been considered is the maxmin fair share [12, 36]. The notion coincides with proportionality in the context of randomized/fractional allocations and hence is weaker than OEV.

Another popular objective is the maximization of the *utilitarian welfare*, i.e. the sum of agents' valuations for an assignment. Filos-Ratsikas et al. [19] proved that RSD guarantees $\Omega(n^{-1/2})$ of the total utilitarian welfare if the utilities are normalized to sum up to one for each agent, which is asymptotically optimal among all randomized truthful mechanisms. In a recent paper, Christodoulou et al. [15] proved similar results for the price of anarchy with respect to the utilitarian welfare of random assignment mechanisms, including RSD and PS. In this paper, we consider the effect on approximations of the egalitarian value from strategic aspects (truthful mechanisms), limited information (ordinal mechanisms), or additional fairness requirements (envy-free mechanisms). The egalitarian value does not require the agents' utilities to be normalized and does not require agents' utilities to be added.

Bhalgat et al. [8] determined the approximation ratio of RSD and PS when the objective is the maximization of a different notion, the ordinal social welfare, which is related to the “popularity” of an assignment [5, 26]. Caragiannis et al. [13] examined the issue of how much efficiency loss fairness requirements like envy-freeness incur but crucially, their objective is maximization of utilitarian welfare.

⁵ Even a lexicographic refinement of the OEV maximizing allocations (in which the value of the worst off agent, then the second worst off agent, and so on, are maximized) can be computed in polynomial time via a series of linear programs [20].

2 Preliminaries

An assignment problem is a triple (N, O, v) such that $N = \{1, \dots, n\}$ is the set of agents, $O = \{o_1, \dots, o_m\}$ is the set of indivisible objects, and $v = (v_1, \dots, v_n)$ is the valuation profile which specifies for each agent $i \in N$ utility or valuation function v_i where v_{ij} or $v_i(o_j)$ denotes the value of agent i for object o_j . We will denote by V^n the set of all possible valuation profiles.

A fractional or random allocation p is a $(n \times m)$ matrix $[p(i)(j)]$ such that $p(i)(j) \in [0, 1]$ for all $i \in N$, and $o_j \in O$. We denote by $\mathcal{P}$ the set of all feasible allocations. The term $p(i)(j)$ which we will also write as $p(i)(o_j)$ represents the probability of object o_j being allocated to agent i . Each row $p(i) = (p(i)(1), \dots, p(i)(m))$ represents the allocation of agent i . The set of columns correspond to the objects $o_1, \dots, o_m$. We will denote by $\hat{p}$ the m dimensional vector where the j -th entry is $\sum_{i \in N} p(i)(j)$ is the total probability that object j will be allocated to some agent.⁶ The utility of agent i from allocation p is $u_i(p(i)) = \sum_{j \in O} (p(i)(j))v_{ij}$. An allocation p is *proportional* if for all $i \in N$, $u_i(p(i)) \geq \frac{1}{n}u_i(O)$. An allocation p is *envy-free* if for all $i, j \in N$, $u_i(p(i)) \geq u_i(p(j))$. An allocation is *SD envy-free* if it is envy-free with respect to all cardinal utilities consistent with the ordinal preferences.⁷

We will consider randomized mechanisms that return a random allocation for each instance of an assignment problem. Note the connection between random assignments for indivisible objects and fractional assignments of divisible objects; a random assignment can be viewed as a fractional assignment when agents have additive utilities over the objects. In that sense, we can use well-known mechanisms for fractional assignments, like the CEEI mechanism, as randomized mechanisms for our setting.

We say that a mechanism is proportional if it always returns a proportional allocation. Similarly, a mechanism is envy-free if it always returns an envy-free allocation. A mechanism M is *truthful-in-expectation*, if for any agent $i \in N$, any valuation profile $v = (v_i, v_{-i})$ and any misreport v'_i of agent i it holds that $u_i(M(v_i, v_{-i})) \geq u_i(M(v'_i, v_{-i}))$, i.e. no agent has any incentive to misreport her true valuation.

Two valuations v_i and v'_i are *ordinally equivalent* if they induce the same ranking over objects, formally $v_i(o_j) \geq v_i(o_k)$ iff $v'_i(o_j) \geq v'_i(o_k)$. A profile v is ordinally equivalent to profile v' if for each $i \in N$, v_i and v'_i are ordinally equivalent. A mechanism J is *ordinal* if for any two preference profiles v and v' that are ordinally equivalent, $J(v) = J(v')$, i.e., the allocations are the same for any pair of ordinally equivalent profiles. We now define the main efficiency measures that we will examine in the paper.

- The *egalitarian value (EV)* of an allocation p with respect to valuation profile v is $EV(p, v) = \inf \{ \frac{u_i(p(i))}{u_i(O)} : i \in N \}$.
- For a given valuation profile, the *optimal egalitarian value (OEV)* is the maximum possible egalitarian value that can be achieved $OEV(v) = \sup_{\lambda} \{ \exists p \in \mathcal{P} : EV(p, v) = \lambda \}$.
- For a given valuation profile v , an allocation p achieves *approximation ratio* $\frac{EV(p, v)}{OEV(v)}$.

⁶ We assume that objects can also be left unallocated with some probability.

⁷ SD envy-freeness also applies to cardinal mechanisms e.g., one that maximize total welfare subject to SD envy-freeness constraints.

- For a given mechanism J and valuation profile v , we will say that J achieves *approximation ratio of J for valuation v* as $aar(J, v) = \frac{EV(J(v), v)}{OE(v)}$.
- An allocation rule J *guarantees an approximation ratio of $guar(J)$* where $guar(J)$ is defined as $guar(J) = \inf_{v \in V^n} \{aar(J, v)\}$.

The guaranteed approximation ratio $guar(J)$ is the worst-case guarantee over all instances of the problem that we will be looking to maximize in our theoretical results.

3 Theoretical Results

We first note that for a deterministic mechanism J , $guar(J) = 0$; in the worst case, if all the agents only value the same object then all agents get zero utility except the agent who gets the valued object. From now on, we will focus on randomized mechanisms. We start with the following lemma about proportional mechanisms.

Lemma 1. *For any mechanism J that is proportional, $guar(J) \geq n^{-1}$.*

Proof. If the allocation p is proportional then for each $i \in N$, $u_i(p(i)) \geq n^{-1}(u_i(O))$. Since $EV(p, v) \geq \inf_{i \in N} (u_i(p(i)) / (u_i(O)))$ and $(u_i(p(i)) / (u_i(O))) \geq n^{-1}$ for all $i \in N$, we get that $EV(x) \geq n^{-1}$. Since $OE(v) \leq 1$, $\frac{EV(J(p, v))}{OE(v)} \geq n^{-1}$. $\square$

Since both PS and RSD are proportional (see e.g., [3, 10]), we obtain the following guarantee on their approximation ratio.

Corollary 1. *$guar(PS) \geq n^{-1}$ and $guar(RSD) \geq n^{-1}$.*

A mechanism satisfies the *favourite share* property if whenever all the agents have the same most preferred object then each agent is assigned to it with probability n^{-1} . We obtain the following theorem.

Theorem 1. *For any mechanism J that satisfies the favourite share property, $guar(J) = O(n^{-1})$.*

Proof. Consider the following valuation profile with $n = m$, where ϵ is an arbitrarily small positive value.

$$v_1(o_j) = \begin{cases} 1, & \text{if } j = 1 \\ 0, & \text{otherwise.} \end{cases}$$

For $i \in \{2, \dots, n\}$,

$$v_i(o_j) = \begin{cases} 0.5 + \epsilon, & \text{if } j = 1 \\ 0.5 - \epsilon, & \text{if } j = i \\ 0, & \text{otherwise.} \end{cases}$$

Note that in the assignment which achieves $OE(v)$, agent 1 is assigned object o_1 and the rest of the agents get utility $0.5 - \epsilon$ and hence $OE(v) = 0.5 - \epsilon$. On the other hand, J gives $1/n$ of o_1 to each of the agents so that agent 1 gets utility $1/n$. Since ϵ can be arbitrarily small, it follows that $guar(J) = O(n^{-1})$. $\square$

We remark here that Theorem 1 holds even if agents have strict preferences; the utilities can be perturbed slightly to reflect strict preferences. Theorem 1 gives us the following.

Corollary 2. $\text{guar}(RSD) = O(n^{-1})$ and $\text{guar}(PS) = O(n^{-1})$.

Theorem 2. *For any mechanism J that satisfies SD envy-freeness, $\text{guar}(J) = O(n^{-1})$.*

Proof. SD envy-freeness implies the favourite share property. If an allocation does not satisfy the favourite share property, then the agent who gets less than $1/n$ of his most preferred object will be envious of another agent if he has extremely high utility for the object. $\square$

In the following, we will prove an upper bound on the approximation ratio of the OEV for two classes of mechanisms: envy-free mechanisms and truthful mechanisms. First we prove the following lemma that states that when looking for upper bounds on the approximation ratio, it suffices to only consider anonymous mechanisms. Similar lemmas have been proven before in literature [19, 22].

Lemma 2. *Let J be a mechanism with approximation ratio ρ . Then, there exists another mechanism J' which is anonymous and has approximation ratio at least ρ . Furthermore, if J is truthful or truthful-in-expectation, then J' is truthful-in-expectation and if J is envy-free, J' is envy-free.*

Proof. Let J' be the mechanism that on input valuation profile v first applies a uniformly random permutation to the set of agents and then runs mechanism J on v . Obviously, J' is anonymous. Additionally, since v can be an input to J and the approximation ratio is calculated over all possible instances, the ratio of J cannot be better than the ratio of J' . Finally, since the permutation is independent of valuations, if J is truthful or truthful-in-expectation, J' is truthful-in-expectation. $\square$

Now we state the following theorem, bounding the approximation ratio of any truthful-in-expectation mechanism. RSD and the uniform mechanism (that gives assignment probability of $1/n$ of each object to each agent) are strategyproof and ordinal mechanisms that both achieve a $\Theta(n^{-1})$ approximation of the OEV. We prove that for any truthful-in-expectation mechanism J , it holds that $\text{guar}(J) = O(n^{-1/5})$.

Theorem 3. *For any truthful-in-expectation mechanism J , $\text{guar}(J) = O(n^{-1/5})$.*

Proof. Let J be a truthful-in-expectation mechanism; by Lemma 2, we can assume without loss of generality that J is anonymous. Consider the following valuation profile v (summarized in Figure 1) with $n = n_1 + n_1^2 + n_1^{5/2}$ agents and $n_1^2 + 1$ objects, where ϵ will be defined later:

- For every agent $i \in A = \{1, \dots, n_1\}$, it holds that $v_i(1) = 1$ and $v_i(j) = 0$ for every object $j \neq 1$.
- For every agent $i \in B = \{n_1 + 1, n_1^2\}$, it holds that $v_i(1) = 1 - \epsilon$, $v_i(i - n_1 + 1) = \epsilon$ and $v_i(j) = 0$ for all objects $j \in O \setminus \{1, i - n_1 + 1\}$.
- For every $\ell = 1, \dots, n_1^2$ and agent $i \in C_\ell = \{n_1^2 + (\ell - 1)\sqrt{n_1} + 1, n_1^2 + \ell\sqrt{n_1}\}$, it holds that $v_i(\ell + 1) = 1$ and $v_i(j) = 0$ for all objects $j \neq \ell$.

$n_1^2 + 1$ Objects	1	2	3	4	...	$n_1^2 + 1$	
n_1 agents, set A	1	0	0	0	...	0	
	1	0	0	0	...	0	
	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\ddots$	$\vdots$	
	1	0	0	0	...	0	
n_1^2 agents, set B	$1 - \epsilon$	ϵ	0	0	...	0	
	$1 - \epsilon$	0	ϵ	0	...	0	
	$1 - \epsilon$	0	0	ϵ	...	0	
	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\ddots$	$\vdots$	
	$1 - \epsilon$	0	0	0	...	ϵ	
$n_1^{5/2}$ agents, set C	0	1	0	0	...	0	
	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\ddots$	$\vdots$	$\sqrt{n_1}$ agents, set C_1
	0	1	0	0	...	0	
	0	0	1	0	...	0	
	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\ddots$	$\vdots$	$\sqrt{n_1}$ agents, set C_2
	0	0	1	0	...	0	
	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\ddots$	$\vdots$	$\vdots$
	0	0	0	0	...	1	
	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\ddots$	$\vdots$	$\sqrt{n_1}$ agents, set $C_{n_1^2}$
	0	0	0	0	...	1	

Fig. 1. Valuation profile v .

In other words, the instance consists of n_1 agents that have value 1 for the first object (set A) and 0 for everything else, n_1^2 agents that value object 1 at $1 - \epsilon$ and another object at value ϵ (set B) and $n_1^{5/2}$ agents that value a single object at 1 (set $C = \cup_l C_l$), such that $\sqrt{n_1}$ agents value the object that some agent in the set B has value ϵ for.

Now if we let $\epsilon = 1/(n_1 - \sqrt{n_1})$, then the egalitarian value of the optimal allocation is at least $1/n_1$; an allocation with such a value is the following:

- Every agent $i \in A$ is allocated $1/n_1$ of object 1.
- Every agent $i \in B$ is allocated $(n_1 - \sqrt{n_1})/n_1$ of the object they value at ϵ .
- Every agent $i \in C_l$ is allocated $1/n_1$ of object $l + 1$.

Next consider a family of valuation profiles $\mathcal{V}$, consisting of profiles where all agents have the same valuations as in v , except one agent from B that has value 1 for the object that she had value ϵ in v and 0 for all other objects. Formally, for $\ell = 1, \dots, n_1^2$, we define a profile $v^\ell \in \mathcal{V}$ as follows:

- For every agent $i \neq n_1 + \ell$, it holds that $v_i^\ell(j) = v_i(j)$ for all objects $j \in M$.
- For agent $n_1 + \ell$, it holds that $v_{n_1+\ell}^\ell(\ell + 1) = 1$ and $v_{n_1+\ell}^\ell(j) = 0$, for all objects $j \neq \ell + 1$.

Consider now any ℓ and the corresponding valuation profile v^ℓ . Since J is anonymous and agents in $C_\ell \cup \{n_1 + \ell\}$ have identical valuations and since $|C_\ell| = \sqrt{n_1}$, the probability that agent $n_1 + \ell$ is allocated object $\ell + 1$ is at most $1/(\sqrt{n_1} + 1)$ and her utility is hence at most $1/(\sqrt{n_1} + 1)$. Now consider valuation profile v and consider the probability $p(n_1 + \ell)(\ell + 1)$ that agent $n_1 + \ell$ is allocated object $\ell + 1$. By truthfulness, and since $v_{n_1+\ell}$ could be a misreport from $v_{n_1+\ell}^\ell$, it must hold that $p(n_1 + \ell)(\ell + 1) \leq$

$1/(\sqrt{n_1}+1) < 1/\sqrt{n}$. This implies that the contribution to the expected utility of agent $n_1 + \ell$ from object $\ell + 1$ is at most $\epsilon/\sqrt{n_1}$, which is at most $1/(n_1\sqrt{n_1} - n_1)$.

Now consider the probability $p(n_1 + \ell)(1)$ that agent $n_1 + \ell$ is allocated object 1. From the arguments above, if $p(n_1 + \ell)(1) < 1/(n_1\sqrt{n_1} - n_1)$, then the expected utility of agent $n_1 + \ell$ is at most $2/(n_1\sqrt{n_1} - n_1)$ and the ratio is $O(1/\sqrt{n_1})$. Since $n = n_1 + n_1^2 + n_1^{5/2}$, that would mean that the theorem is proven. Hence, for J to achieve a better ratio than $O(n^{-1/5})$, it has to be the case that for every $\ell = 1, \dots, n_1^2$, it holds that $p(n_1 + \ell)(1) \geq 1/(n_1\sqrt{n_1} - n_1)$. This is not possible however, since then $\sum_{\ell=1}^{n_1^2} p(n_1 + \ell)(1) \geq n_1/(\sqrt{n_1} - 1) > 1$. This completes the proof. $\square$

Note that for utilitarian welfare maximization, Filos-Ratsikas et al. [19] proved that an ordinal mechanism, RSD achieves the best approximation ratio among all truthful mechanisms. We conjecture that this is the case for the maximization of the egalitarian value as well, i.e. for any truthful mechanism J , $\text{guar}(J) = O(n^{-1})$.

We now turn our attention to envy-free mechanisms. For this class, we will prove an $O(n^{-1/5})$ upper bound as well; the proof actually uses the same valuation profile as the proof of Theorem 3.

Theorem 4. *For any mechanism J that satisfies envy-freeness, $\text{guar}(J) = O(n^{-1/5})$.*

Proof. Consider the valuation profile v used in the proof of Theorem 3 and again consider the probability $p(n_1 + \ell)(\ell + 1)$ that agent $n_1 + \ell$ is allocated object $\ell + 1$. Recall the definition of sets A, B and C from the proof of Theorem 3. By envy-freeness, it holds that $p(n_1 + \ell)(\ell + 1) \leq 1/(\sqrt{n} + 1) \leq 1/\sqrt{n}$ otherwise some agent $j \in C_\ell$ (who only values object $\ell + 1$) would be envious of agent $n_1 + \ell$.

The rest of the steps are the same as in the proof of Theorem 3. Again, consider the probability $p(n_1 + \ell)(1)$ that agent $n_1 + \ell$ is allocated object 1. Since $p(n_1 + \ell)(1) < 1/(n_1\sqrt{n_1} - n_1)$, if $p(n_1 + \ell)(1) < 1/(n_1\sqrt{n_1} - n_1)$ then for the same reasons mentioned in the last paragraph of the proof of Theorem 3, we are done. Hence, we can assume that for every $\ell = 1, \dots, n_1^2$, it holds that $p(n_1 + \ell)(1) \geq 1/(n_1\sqrt{n_1} - n_1)$. This is not possible however, since then $\sum_{\ell=1}^{n_1^2} p(n_1 + \ell)(1) \geq n_1/(\sqrt{n_1} - 1) > 1$. $\square$

From Theorem 4, we obtain the following corollary.

Corollary 3. $\text{guar}(CEEI) = O(n^{-1/5})$ and $\text{guar}(PM) = O(n^{-1/5})$.

It would be interesting to provide a better bound for Theorem 4 or show it is optimal, i.e. come up with an envy-free mechanism that actually achieves the ratio. Finally, we consider the *OEV* guarantees of ordinal mechanisms.

Theorem 5. *For any mechanism J that is ordinal, $\text{guar}(J) = O(n^{-1})$.*

Proof. Consider the setting with n agents and $n + 1$ objects $\{o^*, o_0, \dots, o_{n-1}\}$. The preferences are as follows: each agent values o^* the most. Agent 1 has preference order $o^*, o_0, \dots, o_{n-2}, o_{n-1}$. The preference of each agent $i \in N \setminus \{1\}$ over the objects $O \setminus \{o^*\}$ are obtained as follows: take agent $i - 1$ preference order over $O \setminus \{o^*\}$ and move the most preferred object of $i - 1$ among $O \setminus \{o^*\}$ to the end of the preference order for agent i .

$$\begin{aligned}
1 &: o^*, o_0, \dots, o_{n-2}, o_{n-1} \\
2 &: o^*, o_1, \dots, o_{n-1}, o_0 \\
&\vdots \\
i &: o^*, o_{i-1}, \dots, o_{n-i+1}, o_{i-2}
\end{aligned}$$

By Lemma 2, we can assume without loss of generality that J is anonymous. Furthermore, since J is ordinal, due to the preference profile, the mechanism cannot differentiate among the agents even though they may have different valuations over the objects. Assume that there is some agent that is allocated at most $1/n$ of the universally most preferred object o^* . In this case case, consider the scenario where this agent has utility almost 1 for o^* and the other agents i have utility $0.5 + \epsilon$ for o^* and utility $0.5 - \epsilon$ for o_{i-1} where ϵ is an arbitrarily small positive value.. In this case, the egalitarian value achieved is $1/n$ whereas the OEV is almost 0.5. Hence $\text{guar}(J) = O(1/n)$. $\square$

Since MPS is an ordinal mechanism, it follows that $\text{guar}(MPS) = O(n^{-1})$.

4 Experimental Results

The results in Section 3 give us worst case bounds on the *guaranteed approximation ratios* ($\text{guar}(J)$) for a number of prominent randomized mechanisms including RSD and PS. Hence, in this section we present experimental results which provide a perspective on what may happen “in practice.” Since PS can be considered as the most efficient SD envy-free mechanism (in view of various characterizations [9, 39]), the results for PS can also be viewed as the effect of enforcing SD envy-freeness. In order to test the quality of RSD and PS we need to generate both preferences and cardinal utilities for the agents. There are a number of generative statistical cultures that are commonly used to generate ordinal preferences over objects and the choice of model can have significant impact on the outcome of an experimental study (see e.g. [34]).

Since our focus is fairness, and fairness is often hard to achieve when agents have similar valuations, we employ the *Mallows model* [29] and use the generator from WWW.PREFLIB.ORG [31] in our study. Mallows models are often used in machine learning and preference handling as they allow us to easily control the correlation between the preferences of the agents; a common phenomena in preference data [31, 30, 28]. A Mallows model has two parameters: (1) a **Reference Order** (σ), the preference order at the center of the distribution, and (2) a **Dispersion Parameter** (ϕ), the variance in the distribution which controls the level of similarity of the agent preference orders. When $\phi = 0$ all agents have the same ordinal preference; when $\phi = 1$ then the ordinal preferences are drawn uniformly at random from the space of all preference orders.

Formally, the probability of observing an ordering r is inversely proportional to the Kendall Tau distance between σ and r . This probability is weighted by ϕ , which allows us to control the shape of the distribution. For a given ordinal preference, we superimpose cardinal utilities for the agents using two well-established scoring functions: (1) **Borda Utilities**, each agent has valuation of $m - i$ for his i -th preferred object, and (2) **Exponential Utilities**, each agent has valuation of 2^{m-i} for his i -th preferred object.

In our experiments we generate 10,000 valuation profiles (instances) for each combination of parameters with the number of agents $n \in \{2, \dots, 9\}$, number of objects $m \in \{2, \dots, 9\}$, and dispersion parameter $\phi \in \{0.0, 0.1, \dots, 1.0\}$. We draw σ i.i.d. for each instance.

4.1 Experiments: The Performance of RSD and PS

For each instance v generated, and each mechanism J , we examined the achieved approximation ratio, $aar(J, v) = \frac{EV(J(v), v)}{OE(v)}$, of the RSD and PS mechanisms. Among all such values computed, we examined the minimum and average ratio achieved for a given set of parameters. The results of our experiments for Borda Utilities are shown in Figure 2 while the results for exponential utilities are shown in Figure 3. All our figures are aggregations over all values of m for particular combinations of n and ϕ . Note that since $aar(J, v)$ is normalized over the total utility we can aggregate these terms as it is invariant to this scaling. This allows us to draw more general conclusions as we range over the number of agents and objects. Empirically we found that increasing the number of agents has a greater impact on the approximation performance of the mechanisms compared to an increase in the number of objects, hence the decision to aggregate the graphs in the manner chosen. This empirical result is in line with our theoretical results showing that the worst case approximation ratio is a function of n . The results for both mechanisms in Figure 2 strictly dominate the results in Figure 3. Hence, we observe that the achieved approximation ratio is better for Borda utilities than for exponential utilities.

When $\phi = 0.0$ (not shown in our graphs), the achieved approximation ratio is 1 for both PS and RSD. Both of these mechanisms return the uniform allocation, assigning probability of $1/n$ for each object to each agent, when all the agents have identical preferences. In general, sweeping the value of ϕ from completely correlated to completely uncorrelated preferences has little impact on the overall achievable approximation ratio, though for both models the achievable approximation ratio did strictly decrease as we increased ϕ . Comparing Figures 2 and 3, the impact of changing ϕ was strictly greater for the exponential utility model than it was for the Borda utility model, highlighting again that, as the difference between the valuations of the objects grows large, it becomes harder to achieve fair allocations.

Interestingly, in Figure 2 there appears to be almost no difference between the minimum and average ratios for PS and RSD under Borda utilities. Furthermore, these ratios appear to be very high compared to our theoretical results. Finally, PS consistently performs slightly better than RSD for the minimum and average ratios under exponential utilities and on par with RSD for Borda utilities. This provides more empirical support to the argument that PS is superior to RSD in terms of fairness.

4.2 Experiments: The Effect of Envy-freeness

In order to evaluate the effect that envy-freeness has on the allocations we turn to the OEEF mechanism. To understand the worst case effects of adding envy-freeness as a hard constraint has on small instances we exhaustively tested the parameter space with agents $n \in \{2, \dots, 6\}$, number of objects $m \in \{3, \dots, 4\}$ under Borda and exponential utilities. In this entire parameter space, the worse case achievable approximation

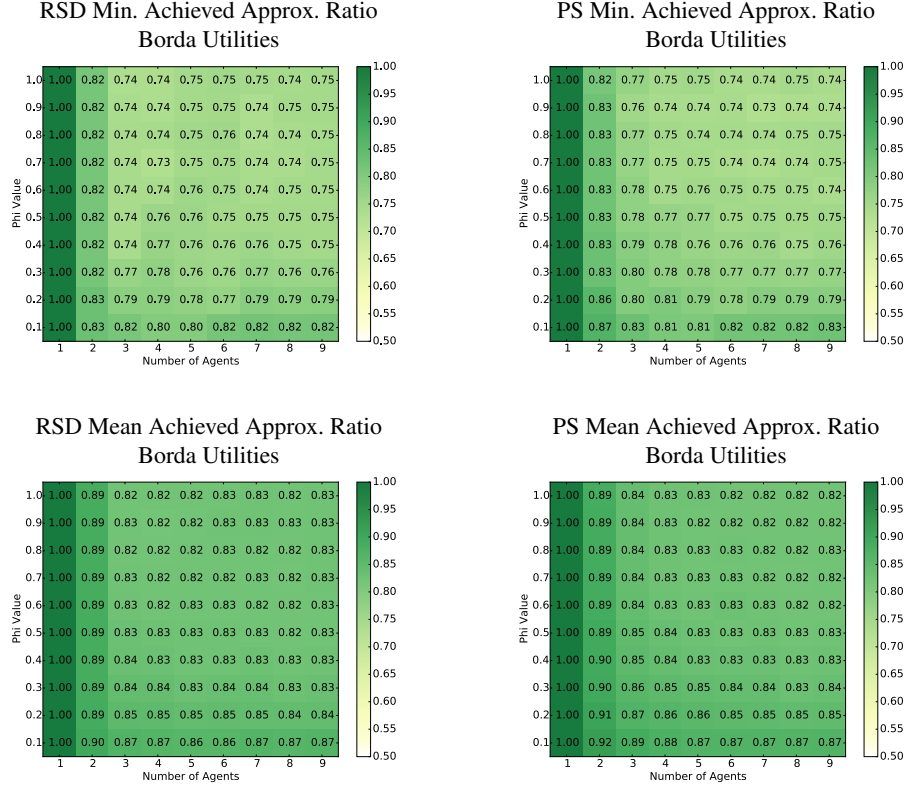


Fig. 2. Minimum (top) and average (bottom) achieved approximation ratio for the RSD (left) and PS (right) mechanisms with Borda utilities. Observe that both mechanisms preform similarly and significantly better than the derived $guar(J)$. Both mechanisms are relatively invariant to the level of dispersion in the underlying valuation profiles. For each $n = \{1, \dots, 9\}$ the graphs are aggregated over the complete range of objects. For example, the cell $(n = 4, \phi = 0.2)$ is the minimum (resp. average) achieved approximation ratio over all instances with $m \in \{1, \dots, 9\}$.

ratio was 0.87, significantly higher than the theoretical worst case. This shows that for smaller instances and some standard utility models, the requirement of envy-freeness does not have a significant negative impact on the achievable approximation ratio. To get an understanding of the performance of OEEF in a larger parameter space we repeated the experiments from the previous section here, evaluating the performance of OEEF across a large parameter space with number of agents $n \in \{2, \dots, 9\}$, number of objects $m \in \{2, \dots, 9\}$, and dispersion parameter $\phi \in \{0.0, 0.1, \dots, 1.0\}$. The results of these tests, again aggregated by m and ϕ , are shown in Figure 4.

When agents have exponential utilities, the achieved approximation ratio, much like in the last section, is strictly worse. Additionally, when we have exponential utilities, as ϕ increases, the approximation ratio for the envy-free mechanisms first decreases slightly and then increases for higher number of agents. Since $\phi = 1.0$ means that agents preferences are drawn uniformly at random, it is more likely that each agent has

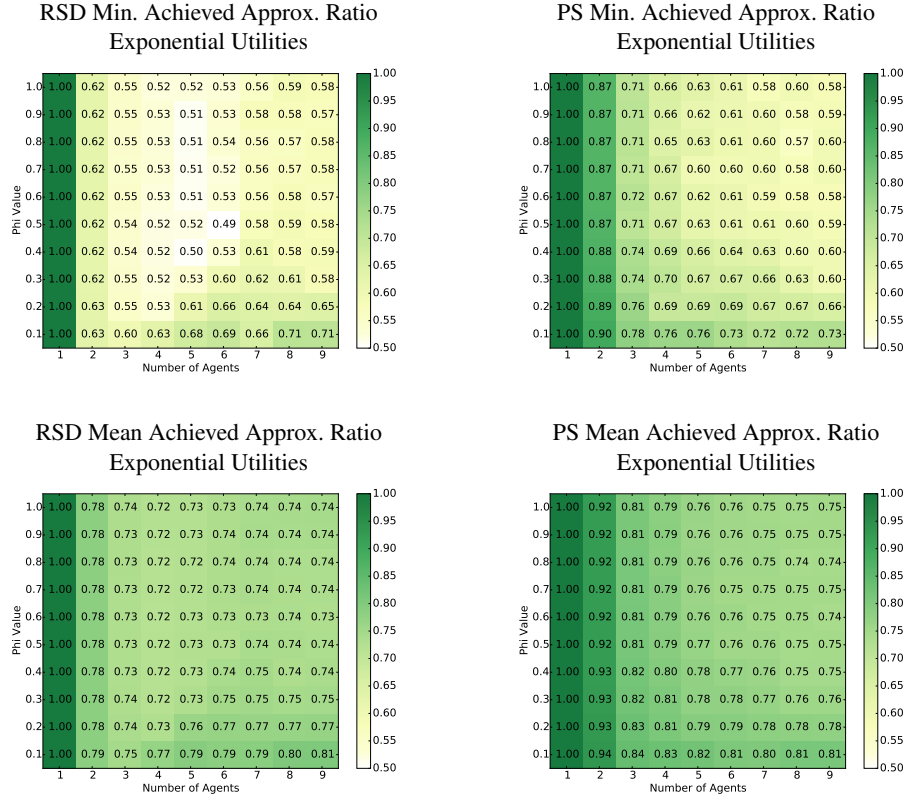


Fig. 3. Minimum (top) and average (bottom) achieved approximation ratio for the for the RSD (left) and PS (right) mechanism with exponential utilities. These results are strictly dominated by those in Figure 2; implying that as difference between the valuations of the objects grows the achieved approximation ratio decreases. For each $n = \{1, \dots, 9\}$ these graphs are aggregated over the complete range of objects. For example, the cell $(n = 4, \phi = 0.2)$ is the minimum (resp. average) approximation ratio achieved over all instances with $m \in \{1, \dots, 9\}$.

high valuation for different objects. Hence, as the preferences move from concentrated to dispersed, there seems to be an interesting transition from high to low and back to high in terms of the achievable approximation ratio. As in the previous subsection, we observe that when all agents have the same preferences, the uniform allocation is both envy-free and has maximal achieved approximation ratio. Hence, when $\phi = 0$, the ratio is 1.0 (not shown in Figure 4). We note that RSD preforms much more poorly across the board compared to OEEF. The results in Figure 4 strictly dominate the results for PS in Figures 2 and 3. Hence, SD envy-freeness that is satisfied by PS has a significant impact on the achieved approximation ratio.

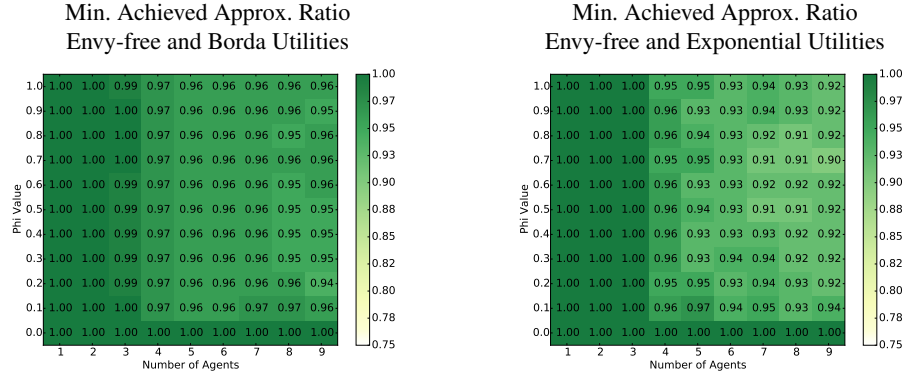


Fig. 4. Minimum achieved approximation ratio when we enforce envy-freeness as a hard constraint. Contrasting these results with those from Figures 2 and 3 for the PS mechanism demonstrates that the stronger notion of envy-freeness that is satisfied by PS has a significant impact on the achieved approximation ratio while the more common notion of envy-freeness does not. For each $n = \{1, \dots, 9\}$ these graphs are aggregated over the complete range of objects. For example, the cell $(n = 4, \phi = 0.2)$ is the minimum (resp. average) approximation ratio achieved over all instances with $m \in \{1, \dots, 9\}$.

5 Conclusion

We present theoretical and experimental results concerning how well different randomized mechanisms approximate the optimal egalitarian value. It has been well-known that egalitarianism can be incompatible with envy-free or truthfulness. In this paper, we quantified how much egalitarianism is affected by such properties. In a recent paper, Christodoulou et al. [15] proved results for the utilitarian welfare of the Nash equilibria of assignment mechanisms. It will be interesting to adopt a similar approach with respect to the egalitarian value and study the price of anarchy of randomized mechanisms with respect to that objective. We assume *additive* cardinal utilities in this paper, it would be interesting to consider other assumptions. To conclude, we mention an open problem: what is the *best* OEV approximation guaranteed by truthful mechanisms?

Acknowledgments

NICTA is funded by the Australian Government through the Department of Communications and the Australian Research Council through the ICT Centre of Excellence Program. Aris Filos-Ratsikas was supported by the Sino-Danish Center for the Theory of Interactive Computation, funded by the Danish National Research Foundation and the National Science Foundation of China (under the grant 61061130540), and by the Center for research in the Foundations of Electronic Markets (CFEM), supported by the Danish Strategic Research Council.

References

1. A. Asadpour and A. Saberi. An approximation algorithm for max-min fair allocation of indivisible goods. *SIAM Journal on Computing*, 39(7):2970–2989, 2010.

2. H. Aziz. Random assignment with multi-unit demands. Technical Report 1401.7700, arXiv.org, 2014.
3. H. Aziz and C. Ye. Cake cutting algorithms for piecewise constant and piecewise uniform valuations. In *Proc. of 10th WINE*, pages 1–14, 2014.
4. H. Aziz, F. Brandt, and M. Brill. The computational complexity of random serial dictatorship. *Economics Letters*, 121(3):341–345, 2013.
5. H. Aziz, F. Brandt, and P. Stursberg. On popular random assignments. In *Proc. of 6th International Symposium on Algorithmic Game Theory (SAGT)*, volume 8146 of *LNCS*, pages 183–194. Springer, 2013.
6. H. Aziz, S. Gaspers, S. Mackenzie, and T. Walsh. Fair assignment of indivisible objects under ordinal preferences. In *Proc. of 13th AAMAS Conference*, pages 1305–1312, 2014.
7. I. Bezáková and V. Dani. Allocating indivisible goods. *SIGecom Exchanges*, 5(3):11–18, 2005.
8. A. Bhargat, D. Chakrabarty, and S. Khanna. Social welfare in one-sided matching markets without money. In *Proceedings of APPROX-RANDOM*, pages 87–98, 2011.
9. A. Bogomolnaia and E. J. Heo. Probabilistic assignment of objects: Characterizing the serial rule. *Journal of Economic Theory*, 147:2072–2082, 2012.
10. A. Bogomolnaia and H. Moulin. A new solution to the random assignment problem. *Journal of Economic Theory*, 100(2):295–328, 2001.
11. A. Bogomolnaia and H. Moulin. Random matching under dichotomous preferences. *Econometrica*, 72(1):257–279, 2004.
12. S. Bouveret and M. Lemaître. Characterizing conflicts in fair division of indivisible goods using a scale of criteria. In *Proc. of 13th AAMAS Conference*, pages 1321–1328. IFAAMAS, 2014.
13. I. Caragiannis, C. Kaklamanis, P. Kanellopoulos, and M. Kyropoulou. The efficiency of fair division. *Theory of Computing Systems*, 50(4):589–610, 2012.
14. Y.-K. Che and F. Kojima. Asymptotic equivalence of probabilistic serial and random priority mechanisms. *Econometrica*, 78(5):1625–1672, 2010.
15. G. Christodoulou, A. Filos-Ratsikas, S. K. S. Frederiksen, P. W. Goldberg, J. Zhang, and J. Zhang. Welfare ratios in one-sided matching mechanisms. Technical report, arXiv.org, 2015.
16. S. Demko and T. P. Hill. Equitable distribution of indivisible objects. *Mathematical Social Sciences*, 16:145–158, 1988.
17. J. P. Dickerson, A. D. Procaccia, and T. Sandholm. Price of fairness in kidney exchange. In *Proceedings of the 13th International Joint Conference on Autonomous Agents and Multi-Agent Systems (AAMAS)*, pages 1013–1020, 2014.
18. D. Ferraioli, L. Gourves, and J. Monnot. On regular and approximately fair allocations of indivisible goods. In *Proc. of 13th AAMAS Conference*, pages 997–1004. IFAAMAS, 2014.
19. A. Filos-Ratsikas, S. K. S. Frederiksen, and J. Zhang. Social welfare in one-sided matchings: Random priority and beyond. In *Proc. of 7th International Symposium on Algorithmic Game Theory (SAGT)*, pages 1–12, 2014.
20. N. Garg, T. Kavitha, A. Kumar, K. Mehlhorn, and J. Mestre. Assigning papers to referees. *Algorithmica*, 58:119–136, 2010.
21. A. Ghodsi, M. Zaharia, B. Hindman, A. Konwinski, S. Shenker, and I. Stoica. Dominant resource fairness: Fair allocation of multiple resource types. In *Proceedings of the 8th USENIX Symposium on Networked Systems Design and Implementation*, 2011.
22. M. Guo and V. Conitzer. Strategy-proof allocation of multiple items without payments or priors. In *Proc. of 9th AAMAS Conference*, pages 881–888. IFAAMAS, 2010.
23. A. Gutman and N. Nisan. Fair allocation without trade. In *Proc. of 11th AAMAS Conference*, pages 719–728. IFAAMAS, 2012.

24. L. Han, C. Su, L. Tang, and H. Zhang. On strategy-proof allocation without payments or priors. In *Proc. of 7th WINE*, LNCS, pages 182–193, 2011.
25. A. Hylland and R. Zeckhauser. The efficient allocation of individuals to positions. *The Journal of Political Economy*, 87(2):293–314, 1979.
26. T. Kavitha, J. Mestre, and M. Nasre. Popular mixed matchings. *Theoretical Computer Science*, 412(24):2679–2690, 2011.
27. F. Kojima. Random assignment of multiple indivisible objects. *Mathematical Social Sciences*, 57(1):134–142, 2009.
28. T. Lu and C. Boutilier. Learning Mallows models with pairwise preferences. In *Proc. of 28th ICML*, pages 145–152, 2011.
29. C. L. Mallows. Non-null ranking models. *Biometrika*, 44(1/2):114–130, 1957.
30. N. Mattei. Empirical evaluation of voting rules with strictly ordered preference data. In *Proc. of 2nd ADT*, pages 165–177. 2011.
31. N. Mattei and T. Walsh. PrefLib: A library for preference data. In *Proc. of 3rd ADT*, volume 8176 of LNCS, pages 259–270. Springer, 2013.
32. T. T. Nguyen, M. Roos, and J. Rothe. A survey of approximability and inapproximability results for social welfare optimization in multiagent resource allocation. *Annals of Mathematics and Artificial Intelligence*, pages 1–26, 2013.
33. D. C. Parkes, A. D. Procaccia, and N. Shah. Beyond dominant resource fairness: Extensions, limitations, and indivisibilities. In *Proc. of 13th ACM-EC Conference*, pages 808–825, 2013.
34. A. Popova, M. Regenwetter, and N. Mattei. A behavioral perspective on social choice. *Annals of Mathematics and Artificial Intelligence*, 68(1–3):135–160, 2013.
35. A. D. Procaccia and M. Tennenholtz. Approximate mechanism design without money. *ACM Transactions on Economics and Computation*, 1(4), 2013.
36. A. D. Procaccia and J. Wang. Fair enough: Guaranteeing approximate maximin shares. In *Proc. of 15th ACM-EC Conference*, pages 675–692. ACM Press, 2014.
37. J. Rawls. *A Theory of Justice*. Harvard University Press, 1971.
38. L.-G. Svensson. Strategy-proof allocation of indivisible goods. *Social Choice and Welfare*, 16(4):557–567, 1999.
39. M. U. Ünver, O. Kesten, M. Kurino, T. Hashimoto, and D. Hirata. Two axiomatic approaches to the probabilistic serial mechanism. *Theoretical Economics*, 9:253–277, 2014.
40. V. V. Vazirani. Combinatorial algorithms for market equilibria. In N. Nisan, T. Roughgarden, É. Tardos, and V. Vazirani, editors, *Algorithmic Game Theory*, chapter 5, pages 103–134. Cambridge University Press, 2007.
41. L. Zhou. On a conjecture by Gale about one-sided matching problems. *Journal of Economic Theory*, 52(1):123–135, 1990.