

Getting More by Knowing Less: Bayesian Incentive Compatible Mechanisms for Fair Division

Vasilis Gkatzelis¹, Alexandros Psomas², Xizhi Tan¹ and Paritosh Verma²

¹Department of Computer Science, Drexel University.

²Department of Computer Science, Purdue University.

{gkatz,xizhi}@drexel.edu, {apsomas,verma136}@cs.purdue.edu

Abstract

We study fair resource allocation with strategic agents. It is well-known that, across multiple fundamental problems in this domain, truthfulness and fairness are incompatible. For example, when allocating indivisible goods, no truthful and deterministic mechanism can guarantee envy-freeness up to one item (EF1), even for two agents with additive valuations. Or, in cake-cutting, no truthful and deterministic mechanism always outputs a proportional allocation, even for two agents with piecewise constant valuations. Our work stems from the observation that, in the context of fair division, truthfulness is used as a synonym for Dominant Strategy Incentive Compatibility (DSIC), requiring that an agent prefers reporting the truth, no matter what other agents report.

In this paper, we instead focus on Bayesian Incentive Compatible (BIC) mechanisms, requiring that agents are better off reporting the truth in expectation over other agents' reports. We prove that, when agents know a bit less about each other, a lot more is possible: BIC mechanisms can guarantee fairness notions that are unattainable by DSIC mechanisms in both the fundamental problems of allocation of indivisible goods and cake-cutting. We prove that this is the case even for an arbitrary number of agents, as long as the agents' priors about each others' types satisfy a neutrality condition. Notably, for the case of indivisible goods, we significantly strengthen the state-of-the-art negative result for efficient DSIC mechanisms, while also highlighting the limitations of BIC mechanisms, by showing that a very general class of welfare objectives is incompatible with Bayesian Incentive Compatibility. Combined these results give a near-complete picture of the power and limitations of BIC and DSIC mechanisms for the problem of allocating indivisible goods.

1 Introduction

A central goal within the fair division literature is to design procedures that distribute indivisible or divisible goods

among groups of agents. For a resource allocation procedure to reach fair and efficient outcomes, however, it needs to have access to the preferences of the participating agents. Whenever this information is private and each agent can strategically misreport it, the literature is riddled with negative results: eliciting the true preferences of the agents while simultaneously guaranteeing fairness and efficiency is usually impossible. One of the main underlying reasons is that monetary payments are commonly infeasible or undesired in fair division, and designing “truthful” mechanisms in the absence of such payments poses often insurmountable obstacles.

Another factor that contributes to the plethora of impossibility results is that “truthfulness” in this literature has been used as a synonym for “Dominant Strategy Incentive Compatibility” (DSIC). This is a very demanding notion of incentive compatibility which requires that every agent prefers truth-telling to misreporting, *no matter what the other agents' reports are*. This ensures that no agent will have an incentive to lie, even if they know *exactly* what every other agent's preferences are. However, in most real-world applications it is unreasonable to assume that the agents know that much about each other, so this notion may be unnecessarily stringent. In fact, if we assume that an agent would strategically misreport their preferences only if they have enough information to suggest that this would be beneficial then, in some sense, the less an agent knows the less likely it is that they would misreport. Rather than going to the other extreme and assume the agents have no information about each other, in this paper we consider the well-established truthfulness notion of “Bayesian Incentive Compatibility” (BIC). This instead assumes that each agent's preference is drawn from a publicly known prior distribution and the requirement is that telling the truth is the optimal strategy for each agent *in expectation over the other agents' reports*. Simply put, our goal in this paper is to understand what can and what cannot be achieved by BIC mechanisms in the context of fair division, and the extent to which they can outperform DSIC mechanisms.

1.1 Our Contributions

We explore the power and limitations of incentive compatible mechanisms when allocating either indivisible or divisible goods among n strategic agents with additive valuations, i.e., such that an agent's value for a set of goods is equal to the sum of her values for each good separately.

Allocating indivisible goods. It is well-known that the only deterministic DSIC mechanism for allocating indivisible goods among two agents in a Pareto efficient way is the patently unfair “serial dictatorship” mechanism [Klaus and Miyagawa, 2002; Schummer, 1996], where one of the agents is allocated *all* the goods that they have a positive value for, and the other agent receives the leftovers. Since Pareto efficiency may be a lot to ask for, one could imagine that a more forgiving notion of efficiency based on stochastic dominance (SD) [Bogomolnaia and Moulin, 2001] may allow us to overcome this obstacle. Our main negative result shows that this is not the case: even for the more permissive notion of SD^+ efficiency (defined in Section 3.1) instead of Pareto efficiency, we prove that the only deterministic DSIC mechanism remains the serial dictatorship, and this holds even if for every agent i , their value for each good j can take one of only *three* possible values, i.e. $v_{i,j} \in \{0, x, y\}$ (**Theorem 3**). In fact, we show that our stronger negative result is tight in more than one ways: if we were to further relax SD^+ efficiency to SD efficiency (also defined in Section 3.1), or to further restrict the number of possible agent values to two instead of three, this would permit more deterministic DSIC mechanisms beyond serial dictatorship.

Surprisingly, it is also known that, even if we were to completely drop any efficiency requirement in order to avoid the extreme unfairness of the serial dictatorship, our ability to provide non-trivial fairness guarantees would still be severely limited. For example, if we wanted the outcome to be *envy-free up to one good* (EF1), which is a well-studied relaxation of envy-freeness [Lipton *et al.*, 2004; Budish, 2011], this is known to be impossible to achieve using *any* deterministic DSIC mechanism, even for instances with just two agents [Amanatidis *et al.*, 2017a]. Our main positive result shows that we can simultaneously overcome all of the aforementioned negative results if instead of DSIC we focus on BIC mechanisms. In fact, we can achieve this using a variation of the very practical Round-Robin procedure (where agents take turns choosing a single good each time). We prove that our variation, RR^{pass} , is BIC for any distribution over agent preferences that is “neutral” [Moulin, 1980], and it returns allocations that always combine fairness (in the form of EF1) with efficiency (in the form of SD^+ efficiency) (**Theorem 4**). Apart from providing a separation between BIC and DSIC mechanisms by bypassing both of these negative results, this also exhibits a third way in which our result from Theorem 3 is tight. Our definition of a “neutral” prior is satisfied by the common assumption in the stochastic fair division literature, that $v_{i,j}$ s are drawn i.i.d. from an agent-specific distribution $\mathcal{D}_i$ [Bai and Gözl, 2021; Manurangsi and Suksompong, 2020; Manurangsi and Suksompong, 2021; Amanatidis *et al.*, 2017b], going beyond, it also allows for certain priors where valuations are neither independently nor identically distributed (see Section 3.2).

Encouraged by our positive result, which exhibits a practical BIC mechanism that combines fairness and SD^+ efficiency, one may wonder whether BIC mechanisms can even optimize well-motivated welfare functions. For example, a lot of recent work in fair division has focused on maximizing the Nash social welfare (the geometric mean of

agents’ utilities). Apart from guaranteeing Pareto efficiency, this outcome is also known to be EF1 for agents with additive valuations [Caragiannis *et al.*, 2019], and it is used in Splidit, a popular platform for allocating indivisible goods [Goldman and Procaccia, 2015; Shah, 2017]. Our findings along this direction are negative: we show that even for a very simple neutral distribution (uniform over normalized valuations), a mechanism that maximizes any welfare function from a large well-studied family that satisfies the Pigou-Dalton principle (which, e.g., includes the Nash social welfare, the egalitarian social welfare, the utilitarian social welfare, and the leximin criterion) is not BIC (**Theorem 5**). Finally, we prove that if we aim for the strong efficiency guarantee of fractional Pareto optimality (fPO), requiring that the outcome is not Pareto dominated even by randomized allocations, then BIC mechanisms cannot even satisfy *fullfillment*, which is a fairness notion that is much weaker than EF1 (**Theorem 6**). Informally, an allocation is fulfilling if, whenever agent i values at least n goods strictly positively, her utility for her bundle is strictly positive.

Allocating divisible goods. Moving beyond indivisible goods, we then consider mechanisms for the fair allocation of continuous, divisible, heterogeneous resources, using the classic cake-cutting model. The cake, which captures a divisible resource, is represented as the interval $[0, 1]$ and it must be distributed among n agents with different valuations on different parts of the interval. Typically, agents’ valuation functions are described by probability density functions. Recently, [Tao, 2022] showed that, even for $n = 2$ agents with piecewise constant valuation functions, there is no deterministic and DSIC cake-cutting mechanism that always outputs a proportional allocation, i.e., one where each agent receives at least a $1/n$ fraction of their total value. Our main result for this problem circumvents this impossibility for any number of agents by relaxing to BIC mechanisms (**Theorem 8**). Specifically, we propose a deterministic cake-cutting mechanism that is proportional and BIC for all neutral priors. Our mechanism works over a sequence of n rounds: in round i , agent i arrives and agents 1 through $i - 1$ are asked to cut the pieces allocated to them so far into i equal-sized and equal-valued crumbs. Then, agent i takes one crumb from each of the first $i - 1$ agents. The fact that this operation can always be performed crucially relies on the result of [Alon, 1987] concerning the existence of a perfect partition.¹ Apart from piecewise constant valuation functions, our mechanism also provides the same guarantees for piecewise linear valuations, which can be succinctly represented, and for which perfect partitions can be computed efficiently.

1.2 Related Work

Bayesian incentive compatibility in fair division has been considered in the random assignment problem, with n goods and n agents with *ordinal* preferences. [Dasgupta and Mishra, 2022] consider ordinal BIC mechanisms (OBIC), and show that the probabilistic serial mechanism, which is ordinally efficient and satisfies equal treatment of equals, is OBIC with

¹Informally, a partition is said to be perfect if the value of every piece is the same for every agent.

respect to the uniform prior. The notion of OBIC of [Dasgupta and Mishra, 2022] has been used in voting theory, dating back to [d’Aspremont and Peleg, 1988], to escape infamous dictatorship results [Gibbard, 1973; Satterthwaite, 1975]. For cardinal valuations, such as ours, [Fujinaka, 2008] studies the case of a single indivisible good that must be allocated to one of n agents whose (private) values for the good are drawn from a known product distribution. When payments are possible, [Fujinaka, 2008] designs an envy-free, budget-balanced and BIC mechanism.

Another line of work considers other relaxations of DSIC. For example, in the random assignment problem, [Mennle and Seuken, 2021] introduce the notion of partial strategyproofness, a relaxation of DSIC, which is satisfied if truthful reporting is a dominant strategy for agents who have sufficiently different valuations for different objects, and show that in the context of school choice, this notion gives a separation between the classic and the adaptive Boston mechanism. Starting with [Troyan and Morrill, 2020], a number of recent papers [Ortega and Segal-Halevi, 2022; Aziz and Lam, 2021; Psomas and Verma, 2022] relax the DSIC requirement and explore the design of not obviously manipulable (NOM) mechanisms. Under NOM, an agent reports their true type, unless lying is obviously better, where the definition of “obvious” for a strategy follows recent work in Economics [Li, 2017]. Closer to our interest here, [Psomas and Verma, 2022] study the allocation of indivisible goods among additive agents, and prove that one can simultaneously guarantee EF1, PO, and NOM. [Ortega and Segal-Halevi, 2022] show that, in the context of cake-cutting, NOM is compatible with proportionality: an adaptation of the moving-knife procedure satisfies both properties.

Other ways to escape the aforementioned impossibility results in truthful fair division is restricting agents’ valuations, e.g. by focusing on dichotomous [Halpern *et al.*, 2020; Babaioff *et al.*, 2021; Benabbou *et al.*, 2021; Barman and Verma, 2021] or Leontief valuations [Ghodsi *et al.*, 2011; Friedman *et al.*, 2014; Parkes *et al.*, 2015], or by using money-burning (that is, leave resources unallocated) as a substitute for payments, while trying to minimize the inefficiency that these payment substitutes introduce (e.g., [Hartline and Roughgarden, 2008; Cole *et al.*, 2013; Fotakis *et al.*, 2016; Friedman *et al.*, 2019; Abebe *et al.*, 2020]). Finally, a recent research thread suggests the study of mechanisms that produce fair allocations in their equilibria [Amanatidis *et al.*, 2023b; Amanatidis *et al.*, 2023a]. Interestingly, for agents with additive valuations over indivisible goods, it is known that allocations obtained in equilibria of the Round-Robin algorithm are EF1 with respect to the agents’ true valuation functions [Amanatidis *et al.*, 2023a].

2 Preliminaries

Allocating indivisible goods. We are given a set $\mathcal{M}$ of m indivisible items indexed by $[m] := \{1, 2, \dots, m\}$ to be allocated among a set $\mathcal{N}$ of n agents indexed by $[n]$. An allocation $\mathbf{x} \in \{0, 1\}^{n \times m}$ assigns items to agents such that $x_{i,j} = 1$ if agent i gets item j and 0 otherwise. We use $\mathbf{x}_i = (x_{i,1}, \dots, x_{i,m}) \in \{0, 1\}^m$ to denote agent i ’s allo-

cation and $X_i = \{j : x_{i,j} = 1\}$ to denote agent i ’s *bundle* in $\mathbf{x}$. In an allocation, each item $j \in \mathcal{M}$ must be allocated to exactly one agent, i.e., $\sum_{i=1}^n x_{i,j} = 1$. Each agent $i \in \mathcal{N}$ has a private type which is a valuation vector $\mathbf{v}_i \in \mathbb{R}_{\geq 0}^m$ where $v_{i,j}$ is agent i ’s value for item $j \in \mathcal{M}$. Collectively, the valuations of all the agents are represented by a *valuation profile* $\mathbf{v} = (\mathbf{v}_1, \dots, \mathbf{v}_n)$. The *utility* of agent i for a given allocation $\mathbf{x}$, denoted by $u_i(\mathbf{x})$, is *additive* and defined as $u_i(\mathbf{x}) := \sum_{j \in \mathcal{M}} v_{i,j} x_{i,j}$. We will often overload notations and use $u_i(X_k)$ to denote the utility of agent $i \in \mathcal{N}$ for a specific bundle $X_k \in \mathcal{M}$.

An allocation $\mathbf{x}$ *Pareto dominates* another allocation $\mathbf{y}$ if for all agents $i \in \mathcal{N}$, we have $u_i(\mathbf{x}) \geq u_i(\mathbf{y})$ and $u_j(\mathbf{x}) > u_j(\mathbf{y})$ for some agent $j \in \mathcal{N}$. An allocation $\mathbf{x}$ is *Pareto efficient* (or Pareto optimal) if there is no other allocation $\mathbf{x}'$ that Pareto dominates it.

An allocation $\mathbf{x}$ is called *envy-free* (EF) if for every pair of agents i and $j \in \mathcal{N}$, agent i values their allocation at least as much as the allocation of agent j , i.e., $u_i(\mathbf{x}_i) \geq u_i(\mathbf{x}_j)$. However, for the indivisible item settings, an envy-free allocation is not guaranteed to exist. Therefore, we consider a relaxed notion called *envy-free up to one good* (EF1). Formally, an allocation $\mathbf{x}$ is EF1 if for every pair of agents i and $j \in \mathcal{N}$ with $X_j \neq \emptyset$, agent i values their allocation at least as much as the allocation of agent j without agent i ’s favorite item, i.e., $u_i(X_i) \geq u_i(X_j \setminus \{g\})$ for some $g \in X_j$.

Allocating divisible goods. We study the cake-cutting problem where the cake is represented by the interval $[0, 1]$ and is to be allocated among the set $\mathcal{N}$ of n agents. An allocation $\mathbf{x} = (X_1, \dots, X_n)$ is a collection of mutually disjoint subsets of $[0, 1]$ where $X_i \subseteq [0, 1]$ denotes the subset of the cake allocated to agent $i \in \mathcal{N}$. Each agent i ’s private type is a density (valuation) function $f_i : [0, 1] \rightarrow \mathbb{R}_{\geq 0}$ over the cake, such that $\int_0^1 f_i(x) dx = 1$ for all $i \in \mathcal{N}$. We use $\mathbf{f} = (f_1, \dots, f_n)$ to denote the collection of density functions of all agents. Given a subset $S \subseteq [0, 1]$, agent i ’s utility on S is defined as $u_i(S) = \int_S f_i(x) dx$.

We primarily focus on two classes of valuation functions, namely *piecewise-constant* and *piecewise-linear* valuations. For both families, the interval $[0, 1]$ can be partitioned into finite intervals. For the former class, the value of the function in each interval is a constant; for the latter class it is linear.

An allocation $\mathbf{x} = (X_1, \dots, X_n)$ is *proportional* if each agent receives her average share of the entire cake. Formally, for each agent $i \in \mathcal{N}$ we have $u_i(X_i) \geq \frac{1}{n} \cdot u_i([0, 1]) = \frac{1}{n}$.

Mechanisms and incentive compatibility. For both indivisible and divisible goods, each agent i has a private type t_i : the valuations $\mathbf{v}_i$ for the first and the functions f_i for the second. The agents, being strategic, may choose to misreport it if doing so increases their utility. We use $\mathbf{b} = (\mathbf{b}_1, \dots, \mathbf{b}_n)$ to represent the reported type $\mathbf{b}_i$ of all agent $i \in \mathcal{N}$. A mechanism is represented by an allocation function $\mathbf{x}(\cdot)$, which takes as input reports $\mathbf{b} = (\mathbf{b}_1, \dots, \mathbf{b}_n)$ and outputs a feasible allocation $\mathbf{x}(\mathbf{b})$. For notational simplicity, we often refer to a mechanism directly using its allocation function $\mathbf{x}(\cdot)$.

We say a mechanism is *incentive compatible* if it offers robustness guarantees against strategic behaviors of agents. In this paper, we consider two notions of incentive compatibil-

ity. A mechanism is *dominant strategy incentive compatible (DSIC)* if truthful reporting is a dominant strategy for every agent. Formally, for every agent $i \in \mathcal{N}$, every possible report $\mathbf{b}_i$, and all possible reports of all other agents $\mathbf{b}_{-i}$, we have

$$u_i(\mathbf{x}(\mathbf{t}_i, \mathbf{b}_{-i})) \geq u_i(\mathbf{x}(\mathbf{b}_i, \mathbf{b}_{-i})). \quad (\text{DSIC})$$

We also study the incentives in the Bayesian setting where agents have some prior knowledge of each other's type. We assume that the type $\mathbf{t}_i$ of each agent $i \in \mathcal{N}$ is drawn from some known prior distribution $\mathcal{D}_i$. A mechanism is *Bayesian incentive compatible (BIC)* for priors $\times_{i=1}^n \mathcal{D}_i$ if reporting truthfully is a Bayesian Nash equilibrium, i.e., no agent can increase her *expected* utility by unilaterally misreporting, the expectation being taken over the types of other agents. We use $\mathcal{D}_{-i}$ to denote the prior $\times_{j \neq i} \mathcal{D}_j$. Formally, for each agent $i \in \mathcal{N}$, and every possible report $\mathbf{b}_i$, we have

$$\mathbb{E}_{\mathbf{t}_{-i} \sim \mathcal{D}_{-i}} [u_i(\mathbf{x}(\mathbf{t}_i, \mathbf{t}_{-i}))] \geq \mathbb{E}_{\mathbf{t}_{-i} \sim \mathcal{D}_{-i}} [u_i(\mathbf{x}(\mathbf{b}_i, \mathbf{t}_{-i}))]. \quad (\text{BIC})$$

3 Allocating Indivisible Goods

Previous work on the allocation of indivisible goods has uncovered that DSIC mechanisms cannot be both fair and efficient. A notable example shows that requiring Pareto efficiency from a deterministic DSIC mechanism limits the available options to the patently unfair class of serial dictatorships.

Theorem 1 ([Klaus and Miyagawa, 2002]). *For $n = 2$ agents, a DSIC deterministic mechanism is Pareto efficient if and only if it is a serial dictatorship.*

Our first result, Theorem 3 in Section 3.1, strengthens Theorem 1, showing that even if we weaken the efficiency requirement from Pareto efficiency to SD^+ efficiency, a natural relaxation defined below, the only deterministic DSIC mechanisms that can guarantee SD^+ efficiency are serial dictatorships. Another notable result shows that no deterministic DSIC mechanism always outputs EF1 allocations.

Theorem 2 ([Amanatidis et al., 2017a]). *There is no deterministic DSIC mechanism that outputs EF1 allocations, even for instances involving just $n = 2$ agents.*

Theorem 4 in Section 3.2 shows that if we relax the notion of truthfulness from DSIC to BIC, then neither one of these two negative results holds anymore, i.e., there exists a non-dictatorial deterministic BIC mechanism that guarantees EF1. In Section 3.3 we also uncover the limitations of BIC mechanisms by proving that a wide family social choice functions are incompatible with BIC.

3.1 Strengthening the Characterization of Serial Dictatorships

For a given agent i let $\mathbf{p}_i$ be a ranking of items $j \in \mathcal{M}$ in decreasing order of i 's value for them, where $\mathbf{p}_i(k) = j$ if item j is agent i 's k^{th} favorite item.² For two bundles $\mathbf{x}_i, \mathbf{y}_i \in$

²For items with the same value, we break ties lexicographically, i.e., if $v_{i,j} = v_{i,k}$ and $j < k$, then item j is ranked before k in $\mathbf{p}_i$.

$\{0, 1\}^m$, $\mathbf{x}_i$ *weakly stochastically dominates* $\mathbf{y}_i$ for agent i , or $\mathbf{x}_i \succeq_i \mathbf{y}_i$, iff for all $\ell \in [m]$ we have

$$\sum_{k=1}^{\ell} x_{i, \mathbf{p}_i(k)} \geq \sum_{k=1}^{\ell} y_{i, \mathbf{p}_i(k)}. \quad (1)$$

Additionally, we say that $\mathbf{x}_i$ (*strictly*) *stochastically dominates* $\mathbf{y}_i$ for agent i , or $\mathbf{x}_i \succ_i \mathbf{y}_i$, iff $\mathbf{x}_i \succeq_i \mathbf{y}_i$ and $\mathbf{x}_i \neq \mathbf{y}_i$. If $\mathbf{x}_i \succ_i \mathbf{y}_i$, then there exists an $\ell \in [m]$ for which Inequality (1) is strict.

To define our refinement of stochastic dominance, for each agent $i \in \mathcal{N}$ we let $\mathcal{M}_i^+$ denote the set of items that i has a positive value for, i.e., $\mathcal{M}_i^+ := \{j \in \mathcal{M} : v_{i,j} > 0\}$. Let $m_i = |\mathcal{M}_i^+|$. For two bundles $\mathbf{x}_i, \mathbf{y}_i \in \{0, 1\}^m$, we say that $\mathbf{x}_i$ *weakly stochastically and positively dominates* $\mathbf{y}_i$, denoted as $\mathbf{x}_i \succeq_i^+ \mathbf{y}_i$, iff Inequality (1) is satisfied for all $\ell \in \mathcal{M}_i^+$. Thus, partial order $\succeq_i^+$ is essentially the partial order $\succeq$ defined over the set of items $\mathcal{M}_i^+$. Similarly, $\mathbf{x}_i \succ_i^+ \mathbf{y}_i$ iff $\mathbf{x}_i \succeq_i^+ \mathbf{y}_i$ and $\mathbf{x}_i \neq \mathbf{y}_i$. Akin to the previous definition, if $\mathbf{x}_i \succ_i^+ \mathbf{y}_i$, then there exists an $\ell \in \mathcal{M}_i^+$ for which Inequality (1) is strict. Note that, if two bundles satisfy $\mathbf{x}_i \succeq_i \mathbf{y}_i$, then by definition, $\mathbf{x}_i \succeq_i^+ \mathbf{y}_i$ as well. The following lemma relates the partial orders defined above with utilities; its proof along with other missing proofs from this subsection appear in the full version [Gkatzelis et al., 2023].

Lemma 1. *For any agent $i \in \mathcal{N}$ and any two bundles $\mathbf{a}, \mathbf{b} \in \{0, 1\}^m$, the following statements hold: (1) If $\mathbf{a} \succeq_i^+ \mathbf{b}$, then $u_i(\mathbf{a}) \geq u_i(\mathbf{b})$, and (2) If $\mathbf{a} \succ_i^+ \mathbf{b}$, then $u_i(\mathbf{a}) > u_i(\mathbf{b})$.*

We say that an allocation $\mathbf{y}$ *stochastically dominates* $\mathbf{x}$, denoted as $\mathbf{y} \succ \mathbf{x}$ iff for all agents $i \in \mathcal{N}$ we have $\mathbf{y}_i \succeq_i \mathbf{x}_i$ and for at least one agent $j \in \mathcal{N}$ we have $\mathbf{y}_j \succ_j \mathbf{x}_j$. Similarly, we say that an allocation $\mathbf{y}$ *stochastically and positively dominates* $\mathbf{x}$, denoted as $\mathbf{y} \succ^+ \mathbf{x}$ iff for all agents $i \in \mathcal{N}$ we have $\mathbf{y}_i \succeq_i^+ \mathbf{x}_i$ and for at least one agent $j \in \mathcal{N}$ we have $\mathbf{y}_j \succ_j^+ \mathbf{x}_j$.

Definition 1. *An allocation $\mathbf{x}$ is SD efficient iff there is no other allocation $\mathbf{y}$ such that $\mathbf{y} \succ \mathbf{x}$.*

Definition 2. *An allocation $\mathbf{x}$ is SD^+ efficient iff there is no other allocation $\mathbf{y}$ such that $\mathbf{y} \succ^+ \mathbf{x}$.*

In the following lemma we show that SD^+ efficiency is implied by Pareto efficiency and SD^+ efficiency implies SD efficiency.

Lemma 2. *The following two implications hold for any number of items and agents with additive preferences: (1) If an allocation $\mathbf{x}$ is Pareto efficient then $\mathbf{x}$ is SD^+ efficient as well; (2) If an allocation $\mathbf{x}$ is SD^+ efficient then $\mathbf{x}$ is SD efficient as well. Moreover, there exist allocations that are SD^+ efficient and not Pareto efficient, and allocations that are SD efficient and not SD^+ efficient.*

Now we state the main result of this subsection.

Theorem 3. (Characterization of serial dictatorship) *For $n = 2$ agents, a deterministic DSIC mechanism is SD^+ efficient if and only if it is a serial dictatorship. This holds even when agents have ternary valuations, i.e., for all agents i and items j , we have $v_{i,j} \in \{0, x, y\}$ for some $y > x > 0$.*

Tightness of the above characterization Theorem 3 is tight in two orthogonal ways. Specifically, if we relax SD^+ efficiency to the weaker notion of SD efficiency, then serial dictatorship no longer remains the only DSIC mechanism — this is formally stated in Proposition 1. Note that although SD efficiency is widely used in the context of matchings (where all agents receive the same number of items), the notion of SD^+ efficiency is more appropriate when each agent may receive a different number of items. Specifically, an allocation can only be stochastically dominated by another allocation if all agents receive at least the same number of items; this is not true for stochastic and positive domination.

Proposition 1. *For $n = 2$ agents, there exists a deterministic mechanism other than the serial dictatorship that is DSIC and outputs SD efficient allocations.*

Second, the proof of Theorem 3 holds even for ternary preferences i.e., $v_{i,j} \in \{0, x, y\}$ for all $i \in \mathcal{N}$ and $j \in \mathcal{M}$. However, Theorem 3 does not hold if we consider the slightly more restricted class of bivalued preferences, i.e., $v_{i,j} \in \{x, y\}$ for fixed $x, y \in \mathbb{R}_{\geq 0}$; this essentially follows from known results.

Proposition 2. *For bivalued additive preferences, there exist DSIC, non-dictatorial mechanisms that output Pareto efficient allocations.*

3.2 Round Robin Is BIC

We now show that a variation of Round-Robin, which we call $\text{ROUND-ROBIN}^{\text{pass}}$ (RR^{pass}), is BIC for neutral priors. Moreover, RR^{pass} always outputs an SD^+ efficient and EF1 allocation,³ bypassing the negative results of Theorems 2 and 3. Missing proofs can be found in [Gkatzelis *et al.*, 2023].

The Bayesian Setting. We assume that the for each agent $i \in \mathcal{N}$, there is a neutral prior distribution $\mathcal{D}_i$ supported over $\mathbb{R}_{\geq 0}^m$ such that the valuation vector $\mathbf{v}_i = (v_{i,1}, \dots, v_{i,m}) \sim \mathcal{D}_i$. We say that $\mathcal{D}_i$ is neutral if for every permutation $\sigma : [m] \rightarrow [m]$, we have $\Pr_{\mathbf{v}_i \sim \mathcal{D}_i}[v_{i,\sigma(1)} > \dots > v_{i,\sigma(m)}] = \frac{1}{m!}$, and for each $j \in \mathcal{M}$, the marginal distribution of $\mathcal{D}_i$ wrt $v_{i,j}$ does not have point masses. Intuitively, neutral priors imply that all items are equivalent, in expectation. A natural example of a neutral prior is $v_{i,j} \sim \hat{\mathcal{D}}_i$ where $\hat{\mathcal{D}}_i$ is supported on $\mathbb{R}_{\geq 0}$ and does not have a point mass. This exact prior has been extensively studied in previous works including [Bai and Gözl, 2021; Manurangsi and Suksompong, 2020; Manurangsi and Suksompong, 2021; Amanatidis *et al.*, 2017b]; a stochastic fair division model based on neutrality is also studied in [Benadè *et al.*, 2023]. Note that, neutral priors are permissive enough to allow for correlated item values, for e.g., if $\mathbf{v}_i \sim \Delta^{m-1}$ where Δ^{m-1} denotes a uniform draw from the $m - 1$ -simplex,⁴ then the item values are identically distributed but not independent. Additionally, the item values for an agent don't even have to be identically drawn: for two items, the prior $v_{i,1} \sim \mathcal{U}[0, 1]$ and $v_{i,2} \sim \mathcal{U}[\frac{1}{2}, \frac{3}{4}]$ is neutral.⁵

³We note that items which have zero value for all the agents won't be allocated by RR^{pass} . Indeed, this does not affect the fairness and efficiency guarantees of RR^{pass} .

⁴Recall that, an $m - 1$ -simplex $\Delta^{m-1} = \{(x_1, x_2, \dots, x_m) : \sum_i x_i = 1 \text{ and } x_i \geq 0 \text{ for all } i\}$

⁵ $\mathcal{U}[a, b]$ represents a uniform distribution over the interval $[a, b]$.

Our Mechanism. We begin by defining the RR^{pass} mechanism: intuitively, we execute the standard Round-Robin procedure but allow agents to *pass* their turn if all the remaining items are zero-valued. Formally, given the report $\mathbf{b}_i$ for each agent $i \in \mathcal{N}$, the mechanism first computes a tuple $(\mathbf{p}_i, m_i)$, where $\mathbf{p}_i$ is the strict preference order of items for agent i (ties are broken lexicographically), and $m_i = |\mathcal{M}_i^+|$, where $\mathcal{M}_i^+ = \{j : b_{i,j} > 0\}$ is the set of items agent i values positively. Then the mechanism constructs the output allocation over rounds. In each round, agents arrive in the fixed order $1, 2, \dots, n$, and each agent $i \in \mathcal{N}$ is allocated the item with the smallest (according to $\mathbf{p}_i$) ranking among the set of remaining items; if in agent i 's turn all items in $\mathcal{M}_i^+$ have already been allocated, then RR^{pass} skips i 's turn. This process is repeated until either all the items are allocated or all remaining items are undesirable for all the agents. In the latter case, the remaining items are allocated deterministically and arbitrarily among the agents, resulting in a complete allocation. RR^{pass} is almost ordinal, i.e., to execute it we need the (strict) preference order $\mathbf{p}_i$, but also the number of zero-valued items m_i for each $i \in \mathcal{N}$ (noting that one can construct $\mathcal{M}_i^+$ by looking at the bottom m_i elements of $\mathbf{p}_i$). Thus, in the subsequent discussion, we often interchange $\mathbf{v}_i$ with $(\mathbf{p}_i, m_i)$. Our goal is to prove the following theorem.

Theorem 4. *The $\text{ROUND-ROBIN}^{\text{pass}}$ mechanism is Bayesian incentive compatible for neutral priors, and always outputs SD^+ efficient and EF1 allocations.*

First, we prove that RR^{pass} is BIC for neutral priors. The high-level idea of our proof is similar to the arguments of [Dasgupta and Mishra, 2022] for showing that the probabilistic serial mechanism is BIC in a setting with complete ordinal preferences. Unlike the ordinal setting, our cardinal setting allows agents to have undesirable, zero-valued items. As hinted by Theorem 3 and Proposition 1, allowing for zero values makes it harder to establish incentive guarantees. To mitigate this, our mechanism RR^{pass} (contrary to Round-Robin) never allocates an item to an agent that doesn't value it, unless the item is valued at zero by everyone. This special treatment of zero values is required to prove key structural properties about the *interim allocation* (defined below), which are used throughout our proofs.

For each agent $i \in \mathcal{N}$ and item $j \in \mathcal{M}$ we define the **interim allocation** $q_{i,j}(\mathbf{b}_i) \in [0, 1]$ of a mechanism with allocation function $\mathbf{x}(\cdot)$ to be the probability that agent i is allocated item j when she reports $\mathbf{b}_i$. Formally, $q_{i,j}(\mathbf{b}_i) = \mathbb{P}_{\mathbf{v}_i \sim \mathcal{D}_i}[x_{i,j}(\mathbf{b}_i, \mathbf{v}_i)]$.

We show that, for each agent $i \in \mathcal{N}$, the interim allocation of the RR^{pass} mechanism, given neutral priors, is very structured. For each agent $i \in \mathcal{N}$, there is a (fixed) **positional interim allocation** $\mathbf{q}_i^{\text{pos}} = (q_{i,1}^{\text{pos}}, q_{i,2}^{\text{pos}}, \dots, q_{i,m}^{\text{pos}})$ such that, if agent i reports valuation $\mathbf{b}_i$ to RR^{pass} , then her interim allocation $q_{i,j}(\mathbf{b}_i)$ is exactly $q_{i,k}^{\text{pos}}$ if $b_{i,j} > 0$ and $\mathbf{p}_i(k) = j$; otherwise, $q_{i,j}(\mathbf{b}_i) = 0$ if $b_{i,j} = 0$ (Lemma 3). Intuitively, the interim allocation of a positively valued item depends on only its position in the preference order. Furthermore, the positional interim allocations are monotone: $q_{i,1}^{\text{pos}} \geq \dots \geq q_{i,m}^{\text{pos}}$ (Lemma 4).

Lemma 3. For each agent $i \in \mathcal{N}$, there exists a positional interim allocation $\mathbf{q}_i^{\text{pos}} = (q_{i,1}^{\text{pos}}, \dots, q_{i,m}^{\text{pos}})$ such that if agent i reports any valuation $\mathbf{b}_i$, then her interim allocation $q_{i,j}(\mathbf{b}_i) = q_{i,k}^{\text{pos}}$ if (i) $b_{i,j} > 0$, and (ii) $\mathbf{p}_i(k) = j$; otherwise, $q_{i,j}(\mathbf{b}_i) = 0$ if $b_{i,j} = 0$.

Lemma 4. Positional interim allocations are monotone, i.e., $q_{i,1}^{\text{pos}} \geq q_{i,2}^{\text{pos}} \geq \dots \geq q_{i,m}^{\text{pos}}$ for all $i \in \mathcal{N}$.

Using these two lemmas we show that RR^{pass} is BIC for neutral priors. The high-level intuition is as follows. If agent i reports valuation $\mathbf{b}_i$ having a corresponding preference order $\mathbf{p}_i = (j_1, j_2, \dots, j_m)$, then Lemma 3 implies that the interim allocations $q_{i,j_1}(\mathbf{b}_i), q_{i,j_2}(\mathbf{b}_i), \dots, q_{i,j_m}(\mathbf{b}_i)$ will be (a prefix of) the positional interim allocation $\mathbf{q}_i^{\text{pos}}$, which, as per Lemma 4 is monotone, $q_{i,1}^{\text{pos}} \geq \dots \geq q_{i,m}^{\text{pos}}$. That is, the interim allocations are non-increasing in the reported preference order and are always a permutation of $\mathbf{q}_i^{\text{pos}}$. Note that, the expected utility of agent i as per her true preference $\mathbf{v}_i$ is $\sum_{k=1}^m q_{i,j_k}(\mathbf{b}_i) \cdot v_{i,j_k}$, which essentially is a vector dot product of the true valuation and the interim allocation. This dot product is maximized when the true valuations and the interim allocations are *aligned* (i.e., are in the same order). Given that the interim allocation is always a permutation of $\mathbf{q}_i^{\text{pos}}$, this alignment happens when agent i 's reported preference order matches her true preference order, i.e., when agent i reports $\mathbf{v}_i$, establishing that RR^{pass} is BIC. This intuition is formalized in the following lemma.

Lemma 5. The $\text{ROUND-ROBIN}^{\text{pass}}$ mechanism is Bayesian incentive compatible for neutral priors.

The following lemma proves that RR^{pass} also satisfies SD^+ efficiency and EF1.

Lemma 6. The $\text{ROUND-ROBIN}^{\text{pass}}$ mechanism outputs SD^+ efficient and EF1 allocations.

Combined, Lemmas 5 and 6 imply Theorem 4.

3.3 Welfare Functions and BIC

In this section, we prove that a large family of “welfare maximizing” mechanisms is not Bayesian incentive compatible. All proofs of this section appear in [Gkatzelis et al., 2023].

A welfare function $f : \mathbb{R}_{\geq 0}^n \rightarrow \mathbb{R}$ is a function that, given the utilities of agents $u_1(\mathbf{x}), \dots, u_n(\mathbf{x})$ in allocation $\mathbf{x}$, outputs $f(u_1(\mathbf{x}), \dots, u_n(\mathbf{x}))$, that represents the collective welfare of all the agents in allocation $\mathbf{x}$. An allocation $\mathbf{x}$ is said to maximize f if for all other feasible allocations $\mathbf{y}$ we have $f(u_1(\mathbf{x}), \dots, u_n(\mathbf{x})) \geq f(u_1(\mathbf{y}), \dots, u_n(\mathbf{y}))$.

Definition 3 (Anonymous). A welfare function f is called anonymous iff $f(x_1, x_2, \dots, x_n) = f(x_{\pi_1}, x_{\pi_2}, \dots, x_{\pi_n})$ for any permutation $\pi : [n] \rightarrow [n]$.

Definition 4 (Strictly monotone). A welfare function f is strictly monotone iff $f(x_1, x_2, \dots, x_n) > f(y_1, y_2, \dots, y_n)$ whenever we have $x_i \geq y_i$ for all $i \in [n]$ and $x_j > y_j$ for some $j \in [n]$.

Definition 5 (The Pigou-Dalton principle). The Pigou-Dalton Principle (PDP) states that a welfare function should weakly prefer a profile that reduces the inequality between two agents, assuming that the utilities of all other agents stay

unchanged. Formally, a welfare function f satisfies PDP iff for any two vectors $(x_1, \dots, x_n)$ and $(y_1, \dots, y_n)$ such that there are two distinct $i, j \in [n]$ satisfying (i) $x_i + x_j = y_i + y_j$, (ii) $x_i > x_j$, (iii) $x_i > y_i > x_j$ (equivalently $x_i > y_j > x_j$), and (iv) $x_k = y_k$ for all $k \in [n] \setminus \{i, j\}$, we have $f(x_1, x_2, \dots, x_n) \leq f(y_1, y_2, \dots, y_n)$.

Welfare functions satisfying anonymity, strict monotonicity, and the PDP constitute a broad set of functions that include, e.g., the class of all the p -mean welfare functions⁶ for $p \leq 1$ [Moulin, 2004]. The class of p -mean welfare functions, in turn, includes Nash social welfare, egalitarian social welfare, utilitarian social welfare, and the leximin criterion. In Theorem 5, we show that, for any deterministic (or randomized) mechanism that outputs allocations that (ex-post) maximize such a welfare function, f , is not BIC. We prove this fact for a very simple neutral prior distribution according to which each agent i 's values are drawn uniformly from the $m-1$ -simplex, i.e., $\mathbf{v}_i = (v_1, v_2, \dots, v_m) \sim \Delta^{m-1}$. Notably, Theorem 4 also holds for this neutral prior supported over normalized valuations.⁷

Theorem 5. There is no deterministic (or randomized) mechanism that is Bayesian incentive compatible for the uniform prior and which always outputs allocations that ex-post maximize welfare function f , where f is any anonymous, strictly monotone function satisfying the Pigou-Dalton principle.

A slightly modified version of the above proof can be used to show the non-existence of deterministic (or randomized) mechanisms that are Bayesian incentive compatible and output allocations that are ex-post fPO and EF1. In fact, we can rule out BIC mechanisms that output fPO allocations that satisfy a fairness criterion we call *fulfillment*, that is much weaker than EF1. Specifically, an allocation $\mathbf{x}$ is fulfilling iff the following condition is satisfied for all agents $i \in \mathcal{N}$: if i has at least n items that she values positively, then agent i gets a bundle of positive utility. It is easy to see that an EF1 allocation must be fulfilling. Formally, this result is stated as Theorem 6.

Theorem 6. There is no deterministic (or randomized) Bayesian incentive compatible mechanism that outputs fPO allocations that are fulfilling.

4 Allocating Divisible Goods: Cake Cutting

Recently, [Tao, 2022] proved a strong incompatibility between fairness and incentives in the cake-cutting problem: no deterministic DSIC cake-cutting mechanism can always output proportional allocations (see Theorem 7), for the case of $n = 2$ agents and piecewise constant density functions.

⁶For a given $p \in \mathbb{R}_{\geq 0}$, the p -mean welfare function is defined as $f_p(x_1, x_2, \dots, x_n) := (\frac{1}{n} \sum_{i=1}^n x_i^p)^{1/p}$. Notably, f_p captures utilitarian welfare if $p = 1$, Nash welfare if $p \rightarrow 0$, and egalitarian welfare if $p \rightarrow -\infty$.

⁷It is necessary to consider normalized valuations when studying welfare objectives that aren't scale-free, for e.g. social welfare. In particular, if the valuations aren't normalized, then agents could (mis)report large values for all the items, thereby forcing the mechanism to allocate all the items to them.

Theorem 7 ([Tao, 2022]). *There is no deterministic, DSIC, and proportional cake-cutting mechanism for piecewise-constant valuations, even for $n = 2$ agents.*

The Bayesian Setting. The *private* piecewise-constant valuation function f_i^* of each agent i is drawn from a *neutral* prior distribution $\mathcal{D}_i$. We say a prior distribution $\mathcal{D}$ is neutral, iff for all integers $k > 1$, and disjoint pieces of cake $X_1, X_2, \dots, X_k$ satisfying $|X_1| = |X_2| = \dots = |X_k|$, we have $\Pr_{f^* \sim \mathcal{D}} [j = \arg \max_{i=1}^k f^*(X_i)] = 1/k$. Intuitively, neutral priors represent a distribution of density functions such that the probability of an agent preferring an interval over another same-length interval is the same, i.e., same-length intervals are *equivalent* in expectation. Towards bypassing Theorem 7, we provide a deterministic mechanism (Mechanism 1), which we refer to as INCREMENTALACCOMMODATION (IA), and we show that IA is Bayesian incentive compatible for neutral priors and it always outputs proportional allocations (Theorem 8). Moreover, IA can be executed in polynomial time for piecewise constant density functions.

Theorem 8. *The INCREMENTALACCOMMODATION mechanism is Bayesian incentive compatible for neutral product distributions, and always outputs proportional allocations.*

At a high level, given reported preferences $f_1, f_2, \dots, f_n$, IA operates as follows. Initially, the entire cake $[0, 1]$ is allocated to agent 1. Other agents arrive one-by-one following the order $2, 3, \dots, n$. Before the arrival of agent i , the entire cake is always allocated among the agents $\{1, 2, \dots, i-1\}$; this allocation is denoted by $X_1, X_2, \dots, X_{i-1}$. When agent i arrives, each agent $j \in \{1, 2, \dots, i-1\}$ cuts their piece of cake, X_j , into i pieces, $C_j^1, C_j^2, \dots, C_j^i$ using the SPLITEQUAL operation (Subroutine ??) described below.

SplitEqual operation. Given a piece of cake X , a function f , and a number k , the SPLITEQUAL operation (see [Gkatzelis et al., 2023]) partitions X into pieces $C^1, C^2, \dots, C^k$ having equal sizes and equal values, i.e., $|C^1| = |C^2| = \dots = |C^k|$ and $f(C^1) = f(C^2) = \dots = f(C^k)$. Note that this operation can always be performed given “any” density function f and $k > 1$ since the existence of such a split is directly implied by the existence of a *perfect partition*.⁸ Specifically, if we consider k density functions where the first function, $y_1(X) = |X|$ while the other $k-1$ density functions are f , then a perfect partition $C^1, C^2, \dots, C^k$ of X — which always exists — will satisfy the two required conditions of equal value (since $f(C^r) = f(C^s)$ for all $r, s \in [k]$) and equal size (since $|C^r| = y_1(C^r) = y_1(C^s) = |C^s|$ for all $r, s \in [k]$). The SPLITEQUAL operation can be efficiently performed for piecewise constant valuations by simply dividing each subinterval within X having a constant-value w.r.t. f into k pieces of equal size; the equal value condition will follow from the fact the subinterval being divided has a constant-value. Furthermore, it is known that a perfect partition can be computed efficiently for the significantly more general class of piecewise linear valuations [Chen et al., 2013], allowing for efficient execution of IA for this class of valuations as well.

⁸[Alon, 1987] proved that given arbitrary density functions $y_1, y_2, \dots, y_n$, there always exists a partition $X_1, X_2, \dots, X_n$ of $[0, 1]$ such that $y_i(X_j) = 1/n$ for all $i, j \in [n]$.

After performing the SPLITEQUAL operation, each agent $j \in \{1, \dots, i-1\}$ offers agent i one piece amongst $C_j^1, C_j^2, \dots, C_j^i$, and agent i picks the one that maximizes her value, as per the reported density function f_i . The resulting allocation, after the arrival of the n^{th} agent, is then returned as the final allocation. In Lemma 8, we show that IA is BIC for neutral priors. In Lemma 9, we show that IA always outputs proportional allocations. The proof of Lemma 9 resembles a proof of [Fink, 1964] wherein proportionality is established for a similar recursive cake-cutting algorithm. All proofs of this section appear in the full version [Gkatzelis et al., 2023].

Mechanism 1: INCREMENTALACCOMMODATION

```

1 Input:  $(f_1, \dots, f_n)$ , the reported valuations of the
   agents
2 set  $X_1 \leftarrow [0, 1]$ , and  $X_i \leftarrow \emptyset$  for all  $i \in \{2, \dots, n\}$ 
3 for  $i \leftarrow 2, 3, \dots, n$  do
4   for  $j \leftarrow 1, 2, \dots, i-1$  do
5     set  $(C_j^1, C_j^2, \dots, C_j^i) \leftarrow$ 
       SPLITEQUAL( $X_j, f_j, i$ ); // defined
       in [Gkatzelis et al., 2023]
6     let  $k^* \leftarrow \arg \max_k f_i(C_j^k)$ 
7     set  $X_i \leftarrow X_i \cup C_j^{k^*}$ 
8     set  $X_j \leftarrow X_j \setminus C_j^{k^*}$ 
9 return  $(X_1, X_2, \dots, X_n)$ 
    
```

Lemma 7. *Let X_j^t, X_j be the pieces of cake that agent j is allocated at the beginning and the end of iteration t , respectively, for some $1 \leq j < t \leq n$.⁹ Then, for any valuation function $\hat{f}$, the expected value of $\hat{f}(X_j)$ is equal to $\frac{t-1}{t} \cdot \hat{f}(X_j^t)$, where the expectation is over the randomness of the valuation function of agent t , f_t^* , drawn from a neutral prior $\mathcal{D}$. Formally, $\mathbb{E}_{f_t^* \sim \mathcal{D}} [\hat{f}(X_j)] = \frac{t-1}{t} \cdot \hat{f}(X_j^t)$.*

Now using Lemma 7, we will show that IA is Bayesian incentive compatible for neutral priors.

Lemma 8. *The INCREMENTALACCOMMODATION is Bayesian incentive compatible with respect to product distributions that are neutral.*

Proof Sketch. At a high level, we show that if agent i has piece X_i^i after her arrival, then, eventually, in the final allocation, her utility, as per her true valuation f_i^* , will be proportional to $f_i^*(X_i^i)$. Specifically, let X_i^n be the piece that agent i has at the end of the mechanism’s execution. Through repeated applications of Lemma 7 we prove that $f_i^*(X_i^n) = \frac{i}{n} \cdot f_i^*(X_i^i)$ for any reported valuation function f_i . Therefore, it suffices to prove that truthfully reporting f_i^* maximizes $f_i^*(X_i^i)$.

The following lemma proves that IA outputs proportional allocations.

Lemma 9. *The INCREMENTALACCOMMODATION mechanism outputs proportional allocations.*

Combining Lemmas 8 and 9 implies Theorem 8.

⁹Note that X_j is a random variable which depends on the valuation function of agent t , f_t^* , drawn from $\mathcal{D}$.

Acknowledgements

Alexandros Psomas and Paritosh Verma are supported in part by the NSF CAREER award CCF-2144208, a Google Research Scholar Award, and a Google AI for Social Good award. Vasilis Gkatzelis and Xizhi Tan are partially supported by the NSF CAREER award CCF-2047907.

References

- [Abebe *et al.*, 2020] Rediet Abebe, Richard Cole, Vasilis Gkatzelis, and Jason D Hartline. A truthful cardinal mechanism for one-sided matching. In *Proceedings of the Fourteenth Annual ACM-SIAM Symposium on Discrete Algorithms*, pages 2096–2113. SIAM, 2020.
- [Alon, 1987] Noga Alon. Splitting necklaces. *Advances in Mathematics*, 63(3):247–253, 1987.
- [Amanatidis *et al.*, 2017a] Georgios Amanatidis, Georgios Birmpas, George Christodoulou, and Evangelos Markakis. Truthful allocation mechanisms without payments: Characterization and implications on fairness. In *Proceedings of the 2017 ACM Conference on Economics and Computation*, pages 545–562, 2017.
- [Amanatidis *et al.*, 2017b] Georgios Amanatidis, Evangelos Markakis, Afshin Nikzad, and Amin Saberi. Approximation algorithms for computing maximin share allocations. *ACM Transactions on Algorithms (TALG)*, 13(4):1–28, 2017.
- [Amanatidis *et al.*, 2023a] Georgios Amanatidis, Georgios Birmpas, Federico Fusco, Philip Lazos, Stefano Leonardi, and Rebecca Reiffenhäuser. Allocating indivisible goods to strategic agents: Pure nash equilibria and fairness. *Mathematics of Operations Research*, 2023.
- [Amanatidis *et al.*, 2023b] Georgios Amanatidis, Georgios Birmpas, Philip Lazos, Stefano Leonardi, and Rebecca Reiffenhäuser. Round-robin beyond additive agents: Existence and fairness of approximate equilibria. *arXiv preprint arXiv:2301.13652*, 2023.
- [Aziz and Lam, 2021] Haris Aziz and Alexander Lam. Obvious manipulability of voting rules. In *Algorithmic Decision Theory: 7th International Conference, ADT 2021, Toulouse, France, November 3–5, 2021, Proceedings 7*, pages 179–193. Springer, 2021.
- [Babaioff *et al.*, 2021] Moshe Babaioff, Tomer Ezra, and Uriel Feige. Fair and truthful mechanisms for dichotomous valuations. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 35, pages 5119–5126, 2021.
- [Bai and Gözl, 2021] Yushi Bai and Paul Gözl. Envy-free and pareto-optimal allocations for agents with asymmetric random valuations. *arXiv preprint arXiv:2109.08971*, 2021.
- [Barman and Verma, 2021] Siddharth Barman and Paritosh Verma. Truthful and fair mechanisms for matroid-rank valuations. *arXiv preprint arXiv:2109.05810*, 2021.
- [Benabbou *et al.*, 2021] Nawal Benabbou, Mithun Chakraborty, Ayumi Igarashi, and Yair Zick. Finding fair and efficient allocations for matroid rank valuations. *ACM Transactions on Economics and Computation*, 9(4):1–41, 2021.
- [Benadè *et al.*, 2023] Gerdus Benadè, Daniel Halpern, Alexandros Psomas, and Paritosh Verma. On the existence of envy-free allocations beyond additive valuations. *arXiv preprint arXiv:2307.09648*, 2023.
- [Bogomolnaia and Moulin, 2001] Anna Bogomolnaia and Hervé Moulin. A new solution to the random assignment problem. *Journal of Economic Theory*, 100(2):295–328, 2001.
- [Budish, 2011] Eric Budish. The combinatorial assignment problem: approximate competitive equilibrium from equal incomes. *Journal of Political Economy*, 119(6):1061–1103, 2011.
- [Caragiannis *et al.*, 2019] Ioannis Caragiannis, David Kurokawa, Hervé Moulin, Ariel D Procaccia, Nisarg Shah, and Junxing Wang. The unreasonable fairness of maximum nash welfare. *ACM Transactions on Economics and Computation (TEAC)*, 7(3):1–32, 2019.
- [Chen *et al.*, 2013] Yiling Chen, John K Lai, David C Parkes, and Ariel D Procaccia. Truth, justice, and cake cutting. *Games and Economic Behavior*, 77(1):284–297, 2013.
- [Cole *et al.*, 2013] Richard Cole, Vasilis Gkatzelis, and Gagan Goel. Mechanism design for fair division: allocating divisible items without payments. In *Proceedings of the fourteenth ACM conference on Electronic commerce*, pages 251–268, 2013.
- [Dasgupta and Mishra, 2022] Sulagna Dasgupta and Debasis Mishra. Ordinal Bayesian incentive compatibility in random assignment model. *Review of Economic Design*, pages 1–14, 2022.
- [d’Aspremont and Peleg, 1988] Claude d’Aspremont and Bezalel Peleg. Ordinal Bayesian incentive compatible representations of committees. *Social Choice and Welfare*, 5(4):261–279, 1988.
- [Fink, 1964] AM Fink. A note on the fair division problem. *Mathematics Magazine*, 37(5):341–342, 1964.
- [Fotakis *et al.*, 2016] Dimitris Fotakis, Dimitris Tsipras, Christos Tzamos, and Emmanouil Zampetakis. Efficient money burning in general domains. *Theory of Computing Systems*, 59(4):619–640, 2016.
- [Friedman *et al.*, 2014] Eric Friedman, Ali Ghodsi, and Christos-Alexandros Psomas. Strategyproof allocation of discrete jobs on multiple machines. In *Proceedings of the fifteenth ACM conference on Economics and computation*, pages 529–546, 2014.
- [Friedman *et al.*, 2019] Eric J Friedman, Vasilis Gkatzelis, Christos-Alexandros Psomas, and Scott Shenker. Fair and efficient memory sharing: Confronting free riders. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 33, pages 1965–1972, 2019.

- [Fujinaka, 2008] Yuji Fujinaka. A Bayesian incentive compatible mechanism for fair division. 2008.
- [Ghodsi *et al.*, 2011] Ali Ghodsi, Matei Zaharia, Benjamin Hindman, Andy Konwinski, Scott Shenker, and Ion Stoica. Dominant resource fairness: Fair allocation of multiple resource types. In *Nsdi*, volume 11, pages 24–24, 2011.
- [Gibbard, 1973] Allan Gibbard. Manipulation of voting schemes: a general result. *Econometrica: journal of the Econometric Society*, pages 587–601, 1973.
- [Gkatzelis *et al.*, 2023] Vasilis Gkatzelis, Alexandros Psomas, Xizhi Tan, and Paritosh Verma. Getting more by knowing less: Bayesian incentive compatible mechanisms for fair division. *arXiv preprint arXiv:2306.02040*, 2023.
- [Goldman and Procaccia, 2015] Jonathan Goldman and Ariel D Procaccia. Spliddit: Unleashing fair division algorithms. *ACM SIGecom Exchanges*, 13(2):41–46, 2015.
- [Halpern *et al.*, 2020] Daniel Halpern, Ariel D Procaccia, Alexandros Psomas, and Nisarg Shah. Fair division with binary valuations: One rule to rule them all. In *International Conference on Web and Internet Economics*, pages 370–383. Springer, 2020.
- [Hartline and Roughgarden, 2008] Jason D Hartline and Tim Roughgarden. Optimal mechanism design and money burning. In *Proceedings of the fortieth annual ACM symposium on Theory of computing*, pages 75–84, 2008.
- [Klaus and Miyagawa, 2002] Bettina Klaus and Eiichi Miyagawa. Strategy-proofness, solidarity, and consistency for multiple assignment problems. *International Journal of Game Theory*, 30(3):421–435, 2002.
- [Li, 2017] Shengwu Li. Obviously strategy-proof mechanisms. *American Economic Review*, 107(11):3257–87, 2017.
- [Lipton *et al.*, 2004] R. J. Lipton, E. Markakis, E. Mossel, and A. Saberi. On approximately fair allocations of indivisible goods. In *Proceedings 5th ACM Conference on Electronic Commerce (EC-2004)*, pages 125–131, 2004.
- [Manurangsi and Suksompong, 2020] Pasin Manurangsi and Warut Suksompong. When do envy-free allocations exist? *SIAM Journal on Discrete Mathematics*, 34(3):1505–1521, 2020.
- [Manurangsi and Suksompong, 2021] Pasin Manurangsi and Warut Suksompong. Closing gaps in asymptotic fair division. *SIAM Journal on Discrete Mathematics*, 35(2):668–706, 2021.
- [Mennle and Seuken, 2021] Timo Mennle and Sven Seuken. Partial strategyproofness: Relaxing strategyproofness for the random assignment problem. *Journal of Economic Theory*, 191:105144, 2021.
- [Moulin, 1980] Hervé Moulin. Implementing efficient, anonymous and neutral social choice functions. *Journal of Mathematical Economics*, 7(3):249–269, 1980.
- [Moulin, 2004] Hervé Moulin. *Fair division and collective welfare*. MIT press, 2004.
- [Ortega and Segal-Halevi, 2022] Josué Ortega and Erel Segal-Halevi. Obvious manipulations in cake-cutting. *Social Choice and Welfare*, 59(4):969–988, 2022.
- [Parkes *et al.*, 2015] David C Parkes, Ariel D Procaccia, and Nisarg Shah. Beyond dominant resource fairness: Extensions, limitations, and indivisibilities. *ACM Transactions on Economics and Computation (TEAC)*, 3(1):1–22, 2015.
- [Psomas and Verma, 2022] Alexandros Psomas and Paritosh Verma. Fair and efficient allocations without obvious manipulations. In *Advances in Neural Information Processing Systems*, 2022.
- [Satterthwaite, 1975] Mark Allen Satterthwaite. Strategy-proofness and arrow’s conditions: Existence and correspondence theorems for voting procedures and social welfare functions. *Journal of economic theory*, 10(2):187–217, 1975.
- [Schummer, 1996] James Schummer. Strategy-proofness versus efficiency on restricted domains of exchange economies. *Social Choice and Welfare*, 14(1):47–56, 1996.
- [Shah, 2017] Nisarg Shah. Spliddit: two years of making the world fairer. *XRDS: Crossroads, The ACM Magazine for Students*, 24(1):24–28, 2017.
- [Tao, 2022] Biaoshuai Tao. On existence of truthful fair cake cutting mechanisms. In *Proceedings of the 23rd ACM Conference on Economics and Computation*, pages 404–434, 2022.
- [Troyan and Morrill, 2020] Peter Troyan and Thayer Morrill. Obvious manipulations. *Journal of Economic Theory*, 185:104970, 2020.