

Strategy-proof Pareto-improvement*

Samson Alva

University of Texas at San Antonio

samson.alva@gmail.com

Vikram Manjunath

University of Ottawa

vikramma@gmail.com

January 29, 2019

Abstract

We consider a model where each agent has an outside option of privately known value. At a given allocation, we call the set of agents who do not exercise their outside options the “participants.” We show that one strategy-proof and individually rational mechanism weakly Pareto-improves another if and only if, at each preference profile, it weakly expands (in terms of set inclusion) the set of participants. Corollaries include: a sufficient condition for a mechanism to be on the Pareto-efficient frontier of strategy-proof mechanisms; uniqueness of strategy-proof Pareto-improvements under true preferences over certain normatively meaningful benchmark allocation rules; and a characterization of the pivotal mechanism.

Keywords: strategy-proofness, Pareto-improvement, Pareto-constrained participation-maximality, school choice, pivotal mechanism

JEL Codes: C78; D47; D71; D82

1 Introduction

We consider mechanisms that choose an allocation as a function of agents’ preferences. In our model, each allocation is associated with a fixed set of participants.¹ A

*We thank Inácio Bó, Alexander Brown, Rossella Calvi, Lars Ehlers, Ben Golub, Guillaume Haeringer, Patrick Harless, Eun Jeong Heo, Sean Horan, Fuhito Kojima, Scott Kominers, Silvana Krasteva, Timo Mennle, Debasis Mishra, Thayer Morrill, Mallesh Pai, William Phan, Santanu Roy, Erling Skancke, Tayfun Sönmez, Alexander Teytelboym, William Thomson, Guoqiang Tian, Ryan Tierney, Juuso Toikka, Utku Ünver, Rodrigo Velez, Dan Walton, Alexander Westkamp, Alexander Wolitzky, and Bumin Yenmez for helpful comments and discussions. We also thank seminar audiences at the University of Rochester, the University of Cologne, the University of Bonn, Southern Methodist University, the University of Windsor, Harvard/MIT, the University of Ottawa, North Carolina State University, Carleton University, Université de Montréal, and participants at numerous conferences for their feedback. We are also grateful for the helpful suggestions of the editor and two anonymous referees.

¹ Instances of the model include object allocation (with priorities) (Hylland and Zeckhauser, 1979; Abdulkadiroğlu and Sönmez, 2003), matching with contracts (Hatfield and Milgrom, 2005), excludable public goods (Jackson and Nicolò, 2004), and more.

non-participant at an allocation consumes his outside option. A mechanism is strategy-proof if, at each profile of preferences, no agent is able to beneficially misreport his preferences. It is individually rational if, at each profile of preferences, no participant prefers his outside option to the chosen allocation.

We show the equivalence of two binary relations over strategy-proof and individually rational mechanisms. The first is the Pareto-improvement relation: mechanism φ (weakly) Pareto-improves mechanism φ' if, at each profile of preferences, each agent finds the choice of φ at least as desirable as the choice of φ' . The other is the participation-expansion relation: φ (weakly) participation-expands φ' if, at each profile of preferences, each participant at the choice of φ' is a participant at the choice of φ . Therefore, the requirements of strategy-proofness and individual rationality contain enough information about preferences to ensure that a comparison that makes no reference to preferences (participation-expansion) is equivalent to one that does (Pareto-improvement).²

We make three assumptions on preferences. Firstly, we rule out externalities on non-participants. So, if an agent participates at neither allocation α nor at allocation β , then he is indifferent between α and β . Secondly, only to show that participation-expansion implies Pareto-improvement, we assume the range of values that an agent's outside option may take are on the order of what the mechanism may offer him as a participant. We call this second assumption "richness of the outside option." It ensures that the (upward) movement of an agent's outside option in his preference is essentially unrestricted. Since how an agent compares his outside option to the various alternatives as a participant is often private information, this is natural in many applications. Finally, only to show that Pareto-improvement implies participation-expansion, we assume that no agent is indifferent between his outside option and any allocation at which he participates. We call this final assumption "no indifference with the outside option."

We point out three useful corollaries of our main result:

1. Say that a pair of mechanisms, φ and φ' , are participation-equivalent if, for each profile of preferences, the allocation chosen by φ has the same participants as the allocation chosen by φ' . *Given no externalities on non-participants and richness of the outside option, if a pair of individually rational and strategy-proof mechanisms are participation-equivalent, then they are also welfare-equivalent.* A further corollary is a simple, yet novel, characterization of the pivotal mechanism for the problem of strategy-proof social choice with continuous transfers.

² To our knowledge, the only other result relating the Pareto-improvement relation to participation is for the probabilistic allocation of objects: strict Pareto-improvement implies strict participation-expansion (Erdil, 2014).

2. Say that an allocation is Pareto-constrained participation-maximal if every allocation that strictly expands the set of participants makes some agent worse off. Say that a mechanism is strategy-proofness-constrained Pareto-efficient if no strategy-proof mechanism strictly Pareto-improves it. *Given our three assumptions, a sufficient condition for a strategy-proof and individually rational mechanism to be strategy-proofness-constrained Pareto-efficient is that it always selects a Pareto-constrained participation-maximal allocation.* This implies the strategy-proofness-constrained Pareto-efficiency of the student-optimal stable mechanism for school choice problems and of the cumulative offer mechanism for many matching problems with complex constraints.³
3. Suppose a normative rule associates each profile of preferences with some Pareto-constrained participation-maximal and individually rational allocation. *Given our three assumptions, there is at most one (in welfare terms) strategy-proof mechanism that selects, for each profile of preferences, an allocation that each agent finds at least as desirable as that prescribed by the benchmark under true preferences.*⁴ This result tells us that the student-optimal stable mechanism is the only strategy-proof mechanism satisfying various normative requirements that are weaker than stability.⁵

The rest of the paper is organized as follows. We introduce our model in Section 2. We define properties of allocations and mechanisms in Section 3. We state and prove our results in Section 4. We discuss applications in Section 5. We defer detailed discussion of related literature to Sections 4 and 5.

2 The Model

Let N be a finite set of **agents**. Let $\mathcal{F}$ be the nonempty set of **allocations**. Given $\alpha \in \mathcal{F}$, let $N(\alpha) \subseteq N$ be the **participants** at α . For instance, in the assignment of indivisible objects to agents, participants are those agents who receive an object. For each $i \in N$, let $\mathcal{F}_i \equiv \{\alpha \in \mathcal{F} : i \in N(\alpha)\}$ be the set of allocations that i participates at. If $\alpha \in \mathcal{F}$ is chosen and i is not a participant at α , then i consumes his **outside option**. We denote consumption of the outside option by $\emptyset$. For each $\alpha \in \mathcal{F}$, we denote by $\alpha(i)$ either α if $\alpha \in \mathcal{F}_i$ or $\emptyset$ otherwise. Note that the set of participants, $N(\alpha)$, at each α is a part of the specification

³ See Hirata and Kasuya (2017) for another sufficient condition.

⁴ However, this does not hold if the benchmark rule identifies allocations that are not individually rational or Pareto-constrained participation-maximal, as would be the case when the benchmark rule is a constant function, identifying a single benchmark allocation independent of preferences.

⁵ These include legality (Ehlers and Morrill, 2018) and partial fairness (Dur et al., forthcoming).

of the model. That is, the definition of α *presumes* the participation of $N(\alpha)$. In the next section, we define *individual rationality* to account for agents' preferences in regards to participation.

For each $i \in N$, his preference is a complete, reflexive, and transitive binary relation on $\mathcal{F}_i \cup \{\emptyset\}$. We denote it by R_i . Since i 's preference is over $\mathcal{F}_i \cup \{\emptyset\}$ rather than $\mathcal{F}$, we have assumed that i is indifferent between any pair of allocations that he does not participate at. Consequently, his welfare from such allocations is fully determined by his outside option. In effect, we assume **no externalities on non-participants**.

For each pair $\alpha, \beta \in \mathcal{F}$, we write $\alpha(i) R_i \beta(i)$ to mean that i finds $\alpha(i)$ to be at least as good as $\beta(i)$. We use P_i to denote strict preference and I_i to denote indifference, the asymmetric and symmetric components of R_i respectively. Let $\mathcal{R}_i$ be a set of preference relations for i . A *preference domain* is $\mathcal{R} \equiv \times_{i \in N} \mathcal{R}_i$.

Our analysis is for fixed N , $\mathcal{F}$, and $\mathcal{R}$. Thus, an economy is entirely described by $R \in \mathcal{R}$. A (direct) **mechanism**, $\varphi : \mathcal{R} \rightarrow \mathcal{F}$, associates each economy with an allocation. For each $R \in \mathcal{R}$ and each $i \in N$, instead of $\varphi(R)(i)$, we write $\varphi_i(R)$.

3 Properties of Allocations and Mechanisms

Individual rationality An allocation is individually rational if each agent finds it at least as desirable as not participating. That is, for each $R \in \mathcal{R}$ and each $\alpha \in \mathcal{F}$, α is individually rational at R if, for each $i \in N$, $\alpha(i) R_i \emptyset$. Since each $i \notin N(\alpha)$ consumes $\emptyset$, it is equivalent to say that for each $i \in N(\alpha)$, $\alpha(i) R_i \emptyset$. A mechanism, φ , is *individually rational* if, for each $R \in \mathcal{R}$, $\varphi(R)$ is individually rational at R . Individual rationality ensures that no agent has an incentive to exercise his outside option when the allocation chosen by the mechanism relies on his presence.

Participation-expansion An allocation participation-expands another if participation at the latter entails participation at the former. That is, for each pair $\alpha, \beta \in \mathcal{F}$, α participation-expands β if $N(\alpha) \supseteq N(\beta)$, and *strictly* so if $N(\alpha) \supsetneq N(\beta)$. If they have the same participants, then they are **participation-equivalent**. That is, α is participation-equivalent to β if $N(\alpha) = N(\beta)$. Given a pair of mechanisms φ and φ' , φ *participation-expands* φ' if, for each $R \in \mathcal{R}$, $\varphi(R)$ participation-expands $\varphi'(R)$. They are *participation-equivalent* if, for each $R \in \mathcal{R}$, $\varphi(R)$ and $\varphi'(R)$ are participation-equivalent.

Pareto-improvement One allocation Pareto-improves another if each agent finds the first at least as desirable as the second. That is, for each $R \in \mathcal{R}$ and each pair $\alpha, \beta \in \mathcal{F}$,

α Pareto-improves β at R if, for each $i \in N$, $\alpha(i) R_i \beta(i)$.⁶ If α Pareto-improves β at R and there is $i \in N$ such that $\alpha(i) P_i \beta(i)$, then α **strictly Pareto-improves** β at R . If $\alpha \in \mathcal{F}$ is such that no allocation strictly Pareto-improves it at R , then α is **Pareto-efficient** at R . For each pair of mechanisms, φ and φ' , φ Pareto-improves φ' if, for each $R \in \mathcal{R}$, $\varphi(R)$ Pareto-improves $\varphi'(R)$ at R . If φ Pareto-improves φ' and for some $R \in \mathcal{R}$, $\varphi(R)$ strictly Pareto-improves $\varphi'(R)$ at R , then φ *strictly Pareto-improves* φ' . If, for each $R \in \mathcal{R}$ and each $i \in N$, $\varphi_i(R) I_i \varphi'_i(R)$, then φ and φ' are **welfare-equivalent**. If, for each $R \in \mathcal{R}$, $\varphi(R)$ is Pareto-efficient at R , then φ is *Pareto-efficient*.

Pareto-constrained participation-maximality An allocation is Pareto-constrained participation-maximal at a given $R \in \mathcal{R}$ if there is no other allocation that strictly expands the set of participants without harming anyone. That is, for each $R \in \mathcal{R}$ and each $\alpha \in \mathcal{F}$, α is Pareto-constrained participation-maximal at R if there is no $\beta \in \mathcal{F}$ such that (1) $N(\alpha) \subsetneq N(\beta)$ and (2) there is no $i \in N$, such that $\alpha(i) P_i \beta(i)$. Equivalently, an allocation is Pareto-constrained participation-maximal if there is no other allocation that simultaneously Pareto-improves it and strictly participation-expands it.⁷ By definition, each allocation that is individually rational and Pareto-efficient is Pareto-constrained participation-maximal, but the converse is not true.⁸ A mechanism, φ , is *Pareto-constrained participation-maximal* if, for each $R \in \mathcal{R}$, $\varphi(R)$ is Pareto-constrained participation-maximal at R .

IR-PCPM-Pareto-connectedness The Pareto-improvement relation is reflexive and transitive but not complete. Two allocations are **Pareto-comparable** if one Pareto-improves the other. Two individually rational and Pareto-constrained participation-maximal allocations are IR-PCPM-Pareto-connected if there is a sequence of individually rational and Pareto-constrained participation-maximal allocations starting at one and ending at the other such that successive allocations are Pareto-comparable. That is, two individually rational and Pareto-constrained participation-maximal allocations, $\alpha, \beta \in \mathcal{F}$, are IR-PCPM-Pareto-connected if there is a sequence $(\alpha_k)_{k=0}^K$, with $\alpha_0 \equiv \alpha$ and $\alpha_K \equiv \beta$, such that for each $k \in \{1, \dots, K\}$, α_k is individually rational, Pareto-constrained participation-maximal, and Pareto-comparable to α_{k-1} . Each pair of individually ratio-

⁶ In this case, some authors say that α *weakly* Pareto-improves β . However, since this is the main form of Pareto-improvement that we consider, we drop the qualifier.

⁷ We have chosen this terminology since an allocation satisfies this property when its set of participants is maximal subject to the constraint of Pareto-improvement.

⁸ In the special case of the object allocation model with strict preferences (see Section 5.1), where non-wastefulness is defined, non-wastefulness is stronger than Pareto-constrained participation-maximality.

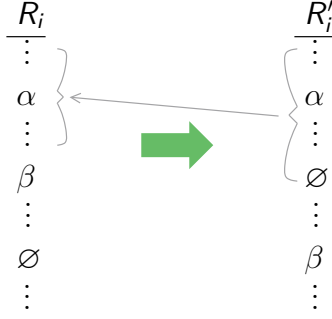


Figure 1: Suppose there is a preference relation, R_i , where i prefers α to β and finds β to be at least desirable as his outside option. Richness of the outside option says that there is another preference relation for i , R'_i , where α is better than his outside option, which is in turn better than β . Furthermore, any allocation that i finds at least as desirable as his outside option under R'_i , he prefers over β under R_i .

nal and Pareto-constrained participation-maximal mechanisms, φ and φ' , are *IR-PCPM-Pareto-connected* if, for each $R \in \mathcal{R}$, $\varphi(R)$ and $\varphi'(R)$ are IR-PCPM-Pareto-connected at R .

Strategy-proofness A mechanism, φ , is strategy-proof if no agent can benefit by misreporting his preferences, no matter what other agents do. That is, for each $R \in \mathcal{R}$, each $i \in N$, and each $R'_i \in \mathcal{R}_i$, $\varphi_i(R) R_i \varphi_i(R'_i, R_{-i})$.

Strategy-proofness-constrained Pareto-efficiency A strategy-proof mechanism is strategy-proofness-constrained Pareto-efficient if there is no strategy-proof mechanism that *strictly* Pareto-improves it.

4 Strategy-proof Pareto-improvement

We present here our main result and several corollaries, which we show under varying combinations of two further assumptions on preferences.

The first is **richness of the outside option** (Figure 1): for each $i \in N$, each $R_i \in \mathcal{R}_i$, and each pair $\alpha, \beta \in \mathcal{F}_i$ such that $\alpha(i) P_i \beta(i) R_i \emptyset$, there is $R'_i \in \mathcal{R}_i$ such that (1) $\alpha(i) P'_i \emptyset P'_i \beta(i)$, and (2) for each $\gamma \in \mathcal{F}_i$, if $\gamma(i) R'_i \emptyset$ then $\gamma(i) P_i \beta(i)$. The second assumption is **no indifference with the outside option**: for each $i \in N$ and each $R_i \in \mathcal{R}_i$, there is no $\alpha \in \mathcal{F}_i$ such that $\alpha(i) I_i \emptyset$. [Sönmez \(1999\)](#) makes two similar assumptions, where an agent's endowment plays the role that $\emptyset$ plays here. Also see [Erdil and Ergin \(2017\)](#) for an instance of no indifference with the outside option in a matching setting with weak preferences.

Remark 1. In many discrete private goods applications, it is typical to assume an agent's preference domain consists of all strict preference relations over own outcomes. This, however, is *far* more than what is implied by our assumptions. No indifference with the outside option only states that ties with $\emptyset$ are broken, but other indifferences may exist. More importantly, richness of the outside option is a great deal weaker than the assumption that all orderings are available. To see this, consider the fact that R_i and R'_i in the definition of richness need not preserve, entirely, the relative orderings of alternatives other than α , β , and $\emptyset$. This makes it even weaker than the analogous assumption in [Sönmez \(1999\)](#).

Our main result relates the Pareto-improvement and participation-expansion relations over strategy-proof and individually rational mechanisms. Under each of the above assumptions, we show that one of these relations is a refinement of the other. Thus, under both assumptions together, they are equivalent.

Theorem 1. *Consider the Pareto-improvement and the participation-expansion relations restricted to the set of strategy-proof and individually rational mechanisms.*

- (A) *If the preference domain satisfies no indifference with the outside option, then Pareto-improvement implies participation-expansion.*
- (B) *If the preference domain satisfies richness of the outside option, then participation-expansion implies Pareto-improvement.*

Proof.

- (A) Let φ and φ' be a pair of individually rational mechanisms such that φ' Pareto-improves φ . Then, for each $R \in \mathcal{R}$, and each $i \in N$, $\varphi'_i(R) R_i \varphi_i(R)$. By individual rationality of φ and no indifference with the outside option, if $i \in N(\varphi(R))$, then $\varphi_i(R) P_i \emptyset$, so $\varphi'_i(R) P_i \emptyset$. Thus, $i \in N(\varphi'(R))$. Since this holds at each R for each $i \in N(\varphi(R))$, φ' participation-expands φ .
- (B) Let φ and φ' be a pair of strategy-proof and individually rational mechanisms such that φ' participation-expands φ . If φ' does not Pareto-improve φ , then there are $i \in N$ and $R \in \mathcal{R}$ such that $\varphi_i(R) P_i \varphi'_i(R)$. Let $\alpha \equiv \varphi(R)$ and $\beta \equiv \varphi'(R)$. Since φ' participation-expands φ , $N(\beta) \supseteq N(\alpha)$. Since both φ and φ' are individually rational, $\alpha(i) R_i \emptyset$ and $\beta(i) R_i \emptyset$. Since $\alpha(i) P_i \beta(i) R_i \emptyset$, we deduce that $i \in N(\alpha)$. Thus, since $N(\beta) \supseteq N(\alpha)$, $i \in N(\beta)$.

Since $\alpha(i) P_i \beta(i) R_i \emptyset$, by richness of the outside option, there is $R'_i \in \mathcal{R}_i$ such that (1) $\alpha(i) P'_i \emptyset P'_i \beta(i)$, and (2) for each $\gamma \in \mathcal{F}_i$, if $\gamma(i) R'_i \emptyset$ then $\gamma(i) P_i \beta(i)$. Let $\gamma \equiv \varphi(R'_i, R_{-i})$ and $\gamma' \equiv \varphi(R'_i, R_{-i})$.

Since φ is strategy-proof, $\gamma(i) R'_i \alpha(i)$. Otherwise, i would have an incentive to misreport R_i if his true preferences are R'_i . Thus, by definition of R'_i , $\gamma(i) P'_i \emptyset$. So $i \in N(\gamma)$. Again, since φ' participation-expands φ , $i \in N(\gamma')$. Since φ' is individually rational, $\gamma'(i) R'_i \emptyset$. By definition of R'_i , $\gamma'(i) P_i \beta(i)$. This contradicts the strategy-proofness of φ' since i has an incentive to misreport his preference as R'_i if his true preferences are R_i . From this contradiction, we conclude that φ' does Pareto-improve φ . □

For individually rational mechanisms, Theorem 1(A) follows from the assumption of no indifference with the outside option and the definition of individual rationality. In fact, it does not even require the mechanisms being compared to be strategy-proof. The more novel and compelling aspect of Theorem 1 is Part (B). It says that for a pair of mechanisms that are both strategy-proof and individually rational, participation-expansion implies Pareto-improvement even if agents may be indifferent between participating and not participating, given richness of the outside option.

Remark 2. The assumption that the preference domain satisfies no indifference with the outside option is critical for Theorem 1(A). Similarly, if the preference domain violates richness of the outside option, Theorem 1(B) does not hold. For counterexamples, where the corresponding assumptions fail, see Appendix A.

Remark 3. The proof of Theorem 1 is at the level of a single agent. Consequently, if strategic preference reporting is restricted to only a subset S of agents, a version of the theorem for S holds: if the preferences of agents in S satisfy richness of the outside option and no indifference with the outside option, then among mechanisms that are strategy-proof and individually rational for members of S , the versions of the Pareto-improvement and participation-inclusion relations restricted to S coincide.

For the probabilistic allocation of indivisible goods when each agent has unit demand and strict preferences, Erdil (2014) has shown that strict Pareto-improvement in the stochastic dominance sense implies a strict participation-expansion. However, we are heretofore unaware of other work showing an *equivalence* between the Pareto-improvement and participation-expansion relations.

Theorem 1 provides an understanding of the Pareto-improvement relation over the set of strategy-proof and individually rational mechanisms. The implications of this result,

some of which we consider in the remainder of this section, have many applications that we describe in Section 5.

By selecting an allocation for each profile of preferences, a mechanism selects a list of participants for each profile of preferences as well—those agents who participate at the allocation that it selects. The following corollary of Theorem 1(B) says that, under richness of the outside option, pinning down who participates at each profile pins down the entire mechanism in welfare terms. It does not require $\mathcal{R}$ to satisfy no indifference with the outside option, and so is relevant for economies with continuous preferences, as in the case where continuous transfers are possible.

Corollary 1. (Participation-equivalence) *Let the preference domain satisfy richness of the outside option. If a pair of strategy-proof and individually rational mechanisms are participation-equivalent, then they are welfare-equivalent.*

Under the assumption of no indifference with the outside option, the set of individually rational and Pareto-constrained participation-maximal allocations has a particularly nice structure. The IR-PCPM-Pareto-connectedness relation over this set is reflexive, symmetric, and transitive, so it is an equivalence relation. Therefore, it partitions this set into components that are IR-PCPM-Pareto-connected. The following lemma says that every allocation in the same component involves the same participants. Since the statement of the lemma is specific to a fixed profile of preferences, it does not rely on richness of the outside option.⁹

Lemma 1 (Structure Lemma). *Let the preference domain satisfy no indifference with the outside option. For each profile of preferences, if a pair of individually rational and Pareto-constrained participation-maximal allocations are IR-PCPM-Pareto-connected, then they are also participation-equivalent.*

Proof. Let $R \in \mathcal{R}$ and $\alpha \in \mathcal{F}$ such that α is individually rational and Pareto-constrained participation-maximal at R . We first show for each $\beta \in \mathcal{F}$, if β Pareto-improves α at R , then $N(\alpha) = N(\beta)$ and β is individually rational and Pareto-constrained participation-maximal at R as well. Then we show for each $\beta \in \mathcal{F}$, if β is individually rational, Pareto-constrained participation-maximal, and IR-PCPM-Pareto-connected to α , then α and β are participation-equivalent.

For each $i \in N(\alpha)$, by no indifference with the outside option, $\alpha(i) P_i \emptyset$. Since $\beta(i) R_i \alpha(i)$, $\beta(i) P_i \emptyset$, and so $i \in N(\beta)$, and β is individually rational at R . Thus, $N(\alpha) \subseteq N(\beta)$.

⁹ In Alva and Manjunath (2019), we draw stronger conclusions by adding structure to the model. Indeed, the Structure Lemma is closely related to the Rural Hospitals Theorem in two-sided matching (Roth, 1986).

Since β Pareto-improves α , Pareto-constrained participation-maximality of α implies it is not the case that $N(\beta) \supsetneq N(\alpha)$. Then together with $N(\alpha) \subseteq N(\beta)$, we have $N(\beta) = N(\alpha)$.

Let $\gamma \in \mathcal{F}$. If γ Pareto-improves β , then it Pareto-improves α . Since α is Pareto-constrained participation-maximal, it is not the case that $N(\gamma) \supsetneq N(\alpha)$. However, $N(\alpha) = N(\beta)$, so it is not the case that $N(\gamma) \supsetneq N(\beta)$. Since this holds for any γ that Pareto-improves β , β is Pareto-constrained participation-maximal at R .

Now, suppose β is individually rational, Pareto-constrained participation-maximal, and α and β are IR-PCPM-Pareto-connected. Then, there is a sequence of individually rational and Pareto-constrained participation-maximal allocations, $(\alpha^k)_{k=0}^K$, with $\alpha^0 = \alpha$ and $\alpha^K = \beta$, such that for each $k \in \{1, \dots, K\}$, either α^k Pareto-improves α^{k-1} or vice versa. In either case, by the argument above, $N(\alpha^k) = N(\alpha^{k-1})$, and so $N(\alpha) = N(\beta)$. $\square$

A consequence of Corollary 1 and the Structure Lemma is that no two welfare-distinct strategy-proof mechanisms select from the same component of the partition of the individually rational and Pareto-constrained participation-maximal allocations that we have described above at every preference profile.

Proposition 1. *Let the preference domain satisfy richness of the outside option and no indifference with the outside option. If a pair of strategy-proof, individually rational, and Pareto-constrained participation-maximal mechanisms are IR-PCPM-Pareto-connected, then they are welfare-equivalent.*

Proof. Let φ and φ' be a pair of strategy-proof, individually rational, and Pareto-constrained participation-maximal mechanisms. If they are IR-PCPM-Pareto-connected, then by the Structure Lemma, they are participation-equivalent. Thus, by the Corollary 1, they are welfare-equivalent. $\square$

Previous work on describing the Pareto frontier of strategy-proof mechanisms includes Anno and Kurino (2016) on object allocation with multi-unit demand and no transfers, Ohseto (2006) and Sprumont (2013) on object allocation with unit demand and transfers, and Anno and Sasaki (2013) on divisible goods allocation. Proposition 1 contributes by providing a useful sufficient condition for a mechanism to be on this frontier.

Corollary 2. *Let the preference domain satisfy richness of the outside option and no indifference with the outside option. A strategy-proof and individually rational mechanism is strategy-proofness-constrained Pareto-efficient if it is Pareto-constrained participation-maximal.¹⁰*

¹⁰ However, Pareto-constrained participation-maximality is not a necessary condition. See the Online Appendix for a counterexample.

A natural context where one might consider IR-PCPM-Pareto-connected mechanisms is when there is a normative benchmark rule that identifies a particular individually rational and Pareto-constrained participation-maximal allocation for every profile of preferences. If a designer is tasked with choosing a strategy-proof mechanism that Pareto-improves such a benchmark rule under agents' true preferences, Proposition 1 says that his problem has at most one solution.

Corollary 3. *Let the preference domain satisfy richness of the outside option and no indifference with the outside option. Consider an individually rational and Pareto-constrained participation-maximal benchmark rule. In welfare terms, there is at most one strategy-proof mechanism that Pareto-improves it under true preferences.*

The limitation of the above corollary is that it requires the benchmark to be individually rational and Pareto-constrained participation-maximal. Without these qualifications, the conclusion does not hold. Suppose, as is the case in many applications, it is possible to assign every agent his outside option. Then every individually rational mechanism Pareto-improves the benchmark that always assigns the outside option to every agent. Of course, there may be many strategy-proof and individually rational mechanisms. Yet, this does not contradict Corollary 3 since this benchmark is not Pareto-constrained participation-maximal.

5 Applications

We consider several applications of our results: the object allocation model, augmented with choice correspondences (Section 5.1); school choice (Section 5.2); excludable public goods (Section 5.3); and transferable utility (Section 5.4).

5.1 Object Allocation and Matching

We start by explaining how the object allocation model fits into the general model that we defined in Section 2.

The object allocation model Let O be a finite and nonempty set of objects, T be a nonempty set of terms under which an agent may be assigned an object, and $X \subseteq N \times O \times T$ be a nonempty set of possible triples. The triple $(i, o, t) \in N \times O \times T$ represents “ i consumes o under the terms t .” These are *contracts* in [Hatfield and Milgrom \(2005\)](#). For each $x \in X$, let $N(x)$ be the agent associated with x . For each $Y \subseteq X$, let $N(Y)$ be the set of agents

associated with triples in Y . For each $i \in N$, let $\mathbf{Y}(i)$ be the triples in Y associated with i . For each $o \in O$, let $\mathbf{Y}(o)$ be the triples in Y associated with o . Each object may only be allocated in certain ways. These constraints define, for each $o \in O$, the **feasible sets** for o , which is a collection of subsets of $X(o)$. We denote it by $\mathcal{F}_o$. In an allocation, each agent has one triple from $X(i)$ or consumes his outside option, $\emptyset$. One only cares about one's own triple, so i has a preference relation over $X(i) \cup \{\emptyset\}$. We also assume that this preference is *strict*—that is, a linear order over $X(i) \cup \{\emptyset\}$ —as is typical in the much of this literature.¹¹

This object allocation model can be embedded into our general model as follows. An *allocation* is a subset μ of X such that no two triples name the same agent and each object's assignment is a feasible set for it. If $\mu(i)$ is empty for agent i , he consumes his outside option, $\emptyset$. That is, $\mathcal{F}$ is a subset of 2^X such that for each $\mu \in \mathcal{F}$, each $i \in N$, and each $o \in O$, $|\mu(i)| \leq 1$ and $\mu(o) \in \mathcal{F}_o$. Thus, the participants at μ , $N(\mu)$, are the agents associated with some triple in μ . Each agent's preference relation in the object allocation model defines a preference relation over $\mathcal{F}$.

Since preferences are strict, let $\mathcal{P} \equiv \times_{i \in N} \mathcal{P}_i$, where $\mathcal{P}_i \subseteq \mathcal{R}_i$ is the set of preferences of i over $\mathcal{F}_i$ that correspond to the strict preferences over $X(i) \cup \{\emptyset\}$. For each $R_i \in \mathcal{P}_i$, I_i is trivial and P_i completely identifies R_i . So we refer to $P_i \in \mathcal{P}_i$. Notice that $\mathcal{P}$ necessarily satisfies no indifference with the outside option. The richness assumption is much weaker than requiring $\mathcal{P}_i$ to contain all strict preferences, which is a standard assumption in such contexts.

When T is a singleton, each $x \in X$ is fully identified by the associated agent and object. In such cases, for each $i \in N$, each triple in $X(i)$ is identified by an element of O , while for each $o \in O$, each triple in $X(o)$ is identified by an element of N . Also, each object's feasible set is identified by a collection of subsets of N , while each agent's preference relation is identified by an ordering of $O \cup \{\emptyset\}$.

As an example, consider the *classical* object allocation model, where each $o \in O$ is an object with capacity $q_o \in \mathbb{Z}_+$. In this model, there is only one term under which an agent can be assigned to an object. For each $o \in O$, $\mathcal{F}_o$ consists of all subsets of $X(o)$ containing no more than q_o elements. In general, we say that $\mathcal{F}_o$ is **capacity-based** if there is $q_o \in \mathbb{R}_+$ such that for each $Y \subseteq X(o)$, $Y \in \mathcal{F}_o$ if and only if $|Y| \leq q_o$.

This model accommodates economies with cross-object constraints: $\mathcal{F}$ can be any subset of 2^X so long as each allocation contains at most one triple associated with each agent. In the absence of such cross-object constraints, the only constraints are that the triples chosen for each object are feasible. That is, for each $\mu \subseteq X$, $\mu \in \mathcal{F}$ if and only if (1) for each

¹¹ See Bogomolnaia et al. (2005) and Erdil and Ergin (2017) for difficulties that arise with indifferences.

$i \in N$, $|\mu(i)| \leq 1$, and (2) for each $o \in O$, $\mu(o) \in \mathcal{F}_o$. In such cases, we say that $\mathcal{F}$ is **Cartesian**. If, in addition to $\mathcal{F}$ being Cartesian, for each $o \in O$, $\mathcal{F}_o$ is capacity-based, then we say that $\mathcal{F}$ is **capacity-based**. Since we do not require $\mathcal{F}$ to be Cartesian, our results are applicable even with distributional constraints (Kamada and Kojima, 2015; Goto et al., 2017).

Non-wastefulness For the classical object allocation model, where feasibility is capacity-based, a natural requirement is that an agent ought not to prefer an object that has remaining capacity to his assignment. If he were to, we could allow him to consume this available resource at no expense to the other agents.

We define non-wastefulness for our general object allocation model as follows: Given $P \in \mathcal{P}$, $\mu \in \mathcal{F}$ is *wasteful at P* if there are $o \in O$, $i \in N$, and $\nu \in \mathcal{F}$, such that (1) $|\nu(o)| > |\mu(o)|$, so that ν allocates o to more agents than μ does, (2) $\nu(i) P_i \mu(i)$, so that i prefers his assignment at ν to that at μ , and (3) for each $j \in N \setminus \{i\}$, $\nu(j) R_i \mu(j)$, so that no agent is worse off at ν compared to μ . If it is not wasteful at P , then μ is *non-wasteful at P* . A mechanism, φ , is *non-wasteful* if, for each $P \in \mathcal{P}$, $\varphi(P)$ is non-wasteful at P .

Remark 4. *If an allocation is non-wasteful, then it is Pareto-constrained participation-maximal. The converse is not true, even for the classical object allocation model.*¹²

For the classical object allocation model, Balinski and Sönmez (1999) define non-wastefulness as follows: $\mu \in \mathcal{F}$ is non-wasteful at $P \in \mathcal{P}$ if there is no $o \in O$ such that $|\mu(o)| < q_o$ and $i \in N$ such that $o P_i \mu(i)$. For this narrower setting, our definition of non-wastefulness is equivalent.^{13,14}

Choice In many applications, there is more information available about each object than just the feasible sets. These might be priorities over agents as in school choice, objectives of the army in cadet-branching, and so on. We model the extra information about how feasible sets are prioritized by associating each $o \in O$ with a **choice correspondence**, $C_o : 2^{X(o)} \rightrightarrows 2^{X(o)}$, such that (1) for each $Y \subseteq X(o)$, $C_o(Y) \subseteq 2^Y$, and (2) the range of C_o , $\bigcup_{Y \subseteq X(o)} C_o(Y)$, is $\mathcal{F}_o$. Condition (1) says that from any set, C_o picks only subsets of it. Condition (2) says that the feasible sets are exactly those that are chosen from *some* set. To satisfy Condition (2), it would suffice, for instance, to select each feasible set from itself. An alternative approach is to start with C_o as the primitive and define $\mathcal{F}_o$ to be its range. Let $C \equiv (C_o)_{o \in O}$.

¹² See the Online Appendix for a proof.

¹³ See the Online Appendix for a proof.

¹⁴ For capacity-based $\mathcal{F}$ and singleton T , Ehlers and Klaus (2014) define an even weaker property that they call *weak non-wastefulness*. However, even in that specific setting, a result like the Structure Lemma does not hold for such a weak version of non-wastefulness.

We associate each object with a choice correspondence rather than a *choice function*, since applications like school choice with weak priorities (Erdil and Ergin, 2008; Abdulkadiroğlu et al., 2009) are better modeled with choice correspondences.

Stability An allocation is stable if no set of agents prefers to drop their assignments in favor of being assigned to a new object under some terms that the object would “choose.” That is, for each $\mu \in \mathcal{F}$ and $P \in \mathcal{P}$, μ is stable at P if it is individually rational¹⁵ at P and there are no $o \in O$ and $Y \subseteq X(o) \setminus \mu(o)$ such that (1) for each $i \in N$, $|Y(i)| \leq 1$, (2) for each $y \in Y$, $y P_{N(y)} \mu(N(y))$, (3) $\mu(o) \notin C_o(\mu(o) \cup Y)$, (4) there is $Z \in C_o(\mu(o) \cup Y)$ such that $Y \subseteq Z$, and (5) $(\mu \setminus (\mu(o) \cup \mu(N(Y)))) \cup Z \in \mathcal{F}$. Condition (1) says that Y contains at most one triple per agent. Condition (2) says that every agent associated with a triple in Y finds it preferable to his triple in μ . These are familiar conditions from the definition of stability for choice functions. Since we are concerned with choice correspondences, the next part of the definition needs to be broken into two parts. The first, Condition (3), says that $\mu(o)$ is not among what is chosen by o when Y is available. The second, Condition (4), says that there is some chosen set, Z , that contains Y . That is, Condition (3) and Condition (4) together say that Y is contained in some Z that is *revealed* by C_o to have a higher priority than $\mu(o)$. The standard definition of stability typically does not include Condition (3) since it is implied by Condition (4) when choice correspondences are single-valued. Condition (5) is only relevant if $\mathcal{F}$ is not Cartesian. It requires that replacing the assignments of o and $N(Y)$ at μ with Z is feasible. A mechanism, φ , is *stable* if, for each $P \in \mathcal{P}$, $\varphi(P)$ is stable at P . This definition is equivalent to the standard one if, for each $o \in O$, C_o is single-valued.

Given a profile of preferences P , if a stable allocation μ Pareto-improves every other stable allocation at P , then it is the **agent-optimal stable** allocation at P . If a stable allocation μ is Pareto-improved by every other stable allocation at P , then it is the **agent-pessimal stable** allocation at P . If C is such that an agent-optimal stable allocation exists for each profile of preferences, then we denote by φ^{AOS} the mechanism that selects this allocation. Similarly, let φ^{APS} select the agent-pessimal stable allocation at each profile of preferences.

Stability is relevant if the choice correspondences represent more than feasibility constraints: they may represent the rights of agents with regards to the objects or particular design goals of the policy maker. The constraints imposed by this information may keep a normative benchmark rule below the Pareto frontier. Stability is a natural requirement

¹⁵ Individual rationality accounts for agents’ preferences while feasibility, along with the requirement that, for each $o \in O$, $\mathcal{F}_o$ be the range of C_o , accounts for objects’ choice correspondences.

for such a benchmark that the mechanism designer may need to Pareto-improve under true preferences. The designer's choice of mechanism might then be constrained to those that select, at every preference profile, an allocation that Pareto-improves *some* stable allocation. Such a mechanism is **stable-dominating**.¹⁶ Since we do not insist on strict Pareto-improvement, every stable mechanism is stable-dominating.

Results Since Corollary 3 only applies to Pareto-constrained participation-maximal benchmarks, we need stability to imply Pareto-constrained participation-maximality to invoke it here. In many applications, like school choice, where priority rankings define the choice correspondences, it is easy to see that stability implies non-wastefulness, and so, by Remark 4, Pareto-constrained participation-maximality. However, this may not be the case without any restrictions on choice correspondences, even if they are single-valued.¹⁷

We place two restrictions on choice correspondences to address this. The first says that the choices from each set should be at least as large as each choice from each of its subsets. That is, C is **size monotonic** if, for each $o \in O$, each $Y \subseteq X(o)$, each finite $Y' \subseteq Y$, each $Z \in C_o(Y)$, and each $Z' \in C_o(Y')$, $|Z| \geq |Z'|$.^{18,19} The second restriction is a mild consistency requirement. It says that if a set is among those chosen from a larger set, it ought to be among what is chosen from itself. That is, C is **idempotent** if, for each $o \in O$, and each $Y \in \text{range}(C_o)$, $Y \in C_o(Y)$.²⁰ Unlike most of the literature on matching with contracts, we do not assume that choice correspondences satisfy a condition like *irrelevance of rejected contracts* (IRC) (Aygün and Sönmez, 2013).²¹

The assumptions that preferences are strict and that choice correspondences are size monotonic and idempotent ensure that every stable allocation is non-wasteful. The non-triviality of the proof, in Appendix B, is because we have not assumed IRC.

Lemma 2. *Let C be size monotonic and idempotent. For each profile of preferences, if an allocation is stable, then it is non-wasteful.*

Lemma 2 allows us to link our results to the literature on stable mechanisms. In particular, we have the following corollary of Theorem 1, Lemma 2, and Remark 4.

¹⁶ We define this property for allocations as well: $\mu \in \mathcal{F}$ is *stable-dominating* if there is $\nu \in \mathcal{F}$ such that ν is stable and μ Pareto-improves ν .

¹⁷ See the Online Appendix for an example.

¹⁸ This is an extension to correspondences of a condition defined for choice functions (Alkan, 2002; Alkan and Gale, 2003; Fleiner, 2003; Hatfield and Milgrom, 2005).

¹⁹ For finite Y , setting $Y' = Y$, size monotonicity implies that for each pair $Z, Z' \in C_o(Y)$, $|Z| = |Z'|$.

²⁰ This rules out, for instance, C_o such that $C_o(\{x, y, z\}) = \{\{x, y\}\}$ but $C_o(\{x, y\}) = \{\{x\}\}$.

²¹ Aygün and Sönmez (2013) define IRC for choice functions. Alva (2018) extends the definition of IRC to choice correspondences and shows its equivalence to the weak axiom of revealed preference.

Corollary 4. *Let C be size monotonic and idempotent. For each stable-dominating benchmark rule, there is at most one strategy-proof mechanism that Pareto-improves it under true preferences.*

Discussion An immediate implication of Corollary 4 is that every strategy-proof and stable-dominating mechanism is strategy-proofness-constrained Pareto-efficient. Hirata and Kasuya (2017) define a version of non-wastefulness weaker than ours, but is logically independent from Pareto-constrained participation-maximality. They show that for matching with contracts models with single-valued choice functions satisfying IRC, every non-wasteful and strategy-proof mechanism is strategy-proofness-constrained Pareto-efficient. They also show that for such choice functions, there is at most one stable and strategy-proof mechanism. Under our choice assumptions, there may be more than one such mechanism.²²

Kamada and Kojima (2015), in a model with distributional constraints, define a strategy-proof mechanism that Pareto-improves on a deferred acceptance mechanism. This does not contradict Corollary 4, which takes $\mathcal{F}$ as fixed. By making flexible the constraints that define $\mathcal{F}$, they obtain a strategy-proof Pareto-improvement on a benchmark that is no longer Pareto-constrained participation-maximal under the redefined $\mathcal{F}$.²³ In a similar setting with more general constraints, Goto et al. (2017) show that an adaptation of the deferred acceptance mechanism is strategy-proofness-constrained Pareto-efficient.

A simple yet policy-relevant measure of a mechanism’s performance is the *expected match rate*: the expected number, with respect to a prior distribution over preference profiles, of agents to whom the mechanism assigns an object. Theorem 1 implies that comparisons of strategy-proof mechanisms by Pareto-improvement translate to comparisons of expected match rates. An interesting application concerns the proposal of Kamada and Kojima (2015) to address doctor shortages in rural areas (Roth, 1986). Their strategy-proof mechanism assigns more doctors to rural areas if only caps on the number of doctors assigned to non-rural regions are tightened. However, this tightening leads to a Pareto-worse outcome from the doctors’ point of view. Thus, Theorem 1 provides a cautionary message: an expected increase of one doctor matched to a rural region leads to an expected decrease of more than one doctor matched to other regions if the prior distribution over preference profiles has full support.

²² See the Online Appendix for an example.

²³ In particular, the benchmark rule in their study is the deferred acceptance mechanism where objects have target capacities that satisfy distributional constraints. Taking these targets as given, this benchmark has no strategy-proof Pareto-improvement. However, by allowing a mechanism the flexibility to exceed these targets while satisfying the distributional constraints, which changes the set of feasible allocations $\mathcal{F}$, Kamada and Kojima (2015) obtain a strategy-proof Pareto-improvement.

Many of the mechanisms that solve market design problems are based on the cumulative offer algorithm (Hatfield and Milgrom, 2005). Under size-monotonic and idempotent choice functions,²⁴ these mechanisms are strategy-proofness-constrained Pareto-efficient, by Corollary 4. However, these choice conditions are not necessary. To apply Corollaries 2 or 3, one only needs to ensure non-wastefulness or Pareto-constrained participation-maximality. Typically, it is easy to show that these mechanisms are non-wasteful even without size-monotonicity.²⁵ Thus, they are also strategy-proofness-constrained Pareto-efficient.

5.2 School Choice

We consider here the more specialized *school choice* model, which is the classical object allocation model augmented with weak priorities. In this model, T is a singleton and $\mathcal{F}$ is capacity-based. Additionally, each $o \in O$ is associated with a **priority** over N denoted by $\succsim_o$, which is a complete, transitive, and reflexive binary relation. Let $\succsim \equiv (\succsim_o)_{o \in O}$.

Given priorities, $\succsim$, for each $o \in O$, we define $C_o^{\succsim}$ as follows. For each $Y \subseteq N$,

$$C_o^{\succsim}(Y) \equiv \begin{cases} \{Y\} & \text{if } |Y| \leq q_o, \\ \{Z \subseteq Y : |Z| = q_o \text{ and for each } i \in Z \text{ and each } j \in Y \setminus Z, i \succsim_o j\} & \text{otherwise.} \end{cases}$$

That is, for each subset of agents, if it contains no more than q_o elements, then the entire set is the only one that is chosen. If not, then all subsets that contain exactly q_o elements are chosen, except for ones that include agents with *strictly* lower priority than an excluded agent. This $C_o^{\succsim}$ is size monotonic and idempotent.

Suppose that an agent prefers a particular object o to the one that he is assigned. Under the interpretation of priorities as consumption “rights,” if o is assigned to someone else who has *strictly* lower priority, then i has the right to protest this allocation. For each $\mu \in \mathcal{F}$, μ **respects priorities** if no agent can protest on such grounds. That is, there is no pair $i, j \in N$ and $o \in O$ such that $o P_j \mu(j)$, $\mu(i) = o$, and $j \succ_o i$.

Stability and dominating stable allocations as fairness Interpreting respect for the priorities as a fairness constraint (Balinski and Sönmez, 1999; Abdulkadiroğlu and Sönmez, 2003), we are interested in mechanisms that are individually rational, non-wasteful, and fair. Respect for priorities alongside individual rationality and non-

²⁴ See, for instance, the problems of cadet-branch matching (Sönmez and Switzer, 2013; Sönmez, 2013) and controlled school choice (Hafalir et al., 2013; Ehlers et al., 2014).

²⁵ These include problems with slot-specific priorities (Kominers and Sönmez, 2016) or hidden substitutes (Hatfield and Kominers, 2016).

wastefulness is equivalent to stability with respect the choice correspondence for each object defined from these priorities.

Remark 5. *For each profile of preferences, an allocation is stable with respect to $C^\succsim$ if and only if it is individually rational, non-wasteful, and respects $\succsim$.*²⁶

Since the combination of individual rationality, non-wastefulness, and respect for priorities is equivalent to stability, we speak of an allocation being stable rather than it respecting priorities and being individually rational and non-wasteful.

When priorities are strict, φ^{AOS} is well defined (Gale and Shapley, 1962). However, when priorities contain ties, there may not exist a single stable allocation that Pareto-improves every other stable allocation. Since φ^{AOS} may not be well defined, a common approach to handling weak priorities is to use tie breakers to form strict priorities. Let $\tau \equiv (\tau_o)_{o \in O}$ be a profile of linear orders over N , one for each object. For each such τ , let $\succsim^\tau$ be the priorities **tie broken by τ** . That is, for each distinct pair $i, j \in N$, $i \succ_o^\tau j$ if either (1) $i \succ_o j$ or (2) $i \sim_o j$ and $i \tau_o j$. Let $\mathcal{T}$ be the set of all profiles of tie breakers. Since the agent-optimal stable mechanism for strict priorities is well defined, given $\succsim$ and $\tau \in \mathcal{T}$, we define the agent-optimal stable mechanism for the priorities tie broken by τ as $\varphi^{AOS\tau}$. These are the mechanisms studied by Abdulkadiroğlu et al. (2009).

For strict priorities, unless they satisfy a restrictive acyclicity condition, φ^{AOS} is not Pareto-efficient (Ergin, 2002). So when priorities are weak, arbitrarily breaking ties could cause Pareto-inefficiency. In fact, for any $\tau \in \mathcal{T}$, $\varphi^{AOS\tau}$ may select an allocation that is Pareto-improvable by another stable allocation (Erdil and Ergin, 2008). Furthermore, some priorities permit a stable, (group) strategy-proof, and Pareto-efficient mechanism, even while $\varphi^{AOS\tau}$ is Pareto-inefficient for every $\tau \in \mathcal{T}$ (Ehlers and Erdil, 2010).

Since $\varphi^{AOS\tau}$ may not be Pareto-efficient, Abdulkadiroğlu et al. (2009) consider the following question: for each $\tau \in \mathcal{T}$, is it possible to find a strategy-proof mechanism that Pareto-improves $\varphi^{AOS\tau}$? They show that the answer is negative.²⁷ Suppose, however, that the answer were positive. The question posed by Abdulkadiroğlu et al. (2009) is important because if such a mechanism existed, though not stable, it would satisfy a natural notion of fairness: it would *Pareto-improve a stable mechanism*, namely $\varphi^{AOS\tau}$. As we have explained above, there is nothing special about $\varphi^{AOS\tau}$, other than strategy-proofness, when priorities are weak. In fact, they may even select allocations that are Pareto-improved by other stable allocations.

In Alva and Manjunath (2019), we suggest the choice of a mechanism can be justified on the grounds that it Pareto-improves *some* stable allocation at each profile of prefer-

²⁶ See the Online Appendix for a proof.

²⁷ Kesten and Kurino (2016) shows the same result even without outside options.

ences. If an agent were to protest the violation of his priority at some object, we offer a move to the Pareto-improved stable allocation so this protest would be moot: the agent would not be better off at *this* stable allocation. That is, by requiring the chosen mechanism to just be *stable-dominating*, rather than stable, we may enlarge the options for a strategy-proof Pareto-improvement.

For strict priorities, since φ^{APS} is well defined (Gale and Shapley, 1962), Corollary 4 implies the following, which is stronger than the known result that φ^{AOS} is the only stable and strategy-proof mechanism (Alcalde and Barberà, 1994).

Corollary 5. *If $\succsim$ consists of strict priorities, then φ^{AOS} is the unique stable-dominating and strategy-proof mechanism.*

On the other hand, for weak priorities, there may be more than one stable-dominating and strategy-proof mechanism. Nevertheless, Proposition 1 yields the following corollary.

Corollary 6. *Each stable-dominating and strategy-proof mechanism (including, for each $\tau \in \mathcal{T}$, $\varphi^{AOS\tau}$) is strategy-proofness-constrained Pareto-efficient. Furthermore, no two stable-dominating and strategy-proof mechanisms are IR-PCPM-Pareto-connected.*

While, for each $\tau \in \mathcal{T}$, $\varphi^{AOS\tau}$ is strategy-proofness-constrained Pareto-efficient, $\varphi^{AOS\tau}$ s are not the only stable and strategy-proof mechanisms (Ehlers and Erdil, 2010). Corollary 6 extends the result of Abdulkadiroğlu et al. (2009) to all of these.

Beyond stability as fairness While we have focused on stability and stable-domination as notions of fairness, our results allow us to draw conclusions about other fairness concepts as well. Take for instance the *legal set* of allocations under strict priorities (Ehlers and Morrill, 2018), where only *harmful* and *redressable* priority violations are ruled out. The legal set includes the stable set and always has a Pareto-worst member that is non-wasteful. Consequently, Proposition 1 implies that φ^{AOS} is the unique strategy-proof selection from the legal set.

Corollary 7. *If $\succsim$ consists solely of strict priorities, then φ^{AOS} is the unique strategy-proof mechanism selecting from the legal set.*

5.3 Excludable Public Goods

In this section, we illustrate how to accommodate excludable public goods in our model. We focus on a simple setting for expositional clarity, since our purpose is to highlight how to apply our results to obtain novel insights. A thorough analysis of the efficient

frontier of strategy-proof mechanisms for a general excludable public goods model could use the full power of our results, but would take us too far afield.

The following model is similar to that of [Jackson and Nicolò \(2004\)](#) and [Cantala \(2004\)](#). Suppose a public facility is to be located on the interval $[0, 1]$ and a set of users chosen. The set of agents is partitioned into two sets: agents in N_L live at 0 and agents in N_R live at 1. Each $i \in N$ prefers to have the facility located as close to his own home as possible. Furthermore, he is unwilling to travel beyond a certain **threshold** $t_i \in (0, 1)$. That is, if i lives at 0, then he is unwilling to travel to the right of t_i and if he lives at 1, then he is unwilling to travel to the left of t_i . Since we know exactly how each agent ranks the locations, the only private information for i is t_i . So a preference profile is identified by the threshold profile t . We assume that each $i \in N$ prefers to travel to t_i and enjoy the facility over opting out. However, he prefers to opt out rather than travel beyond t_i . Thus, no indifference with the outside option is satisfied. Since $t_i \in (0, 1)$, there is always some positive distance he is willing to travel. Yet, richness of the outside option is satisfied since t_i may be arbitrarily close to i 's home (0 if $i \in N_L$ or 1 if $i \in N_R$).

An *allocation* consists of two parts: the location of a public facility in the interval $[0, 1]$ and the set of users. That is, $\mathcal{F} \equiv [0, 1] \times 2^N$. A mechanism maps threshold profiles to allocations. It is **dictatorial** if it ignores the thresholds and always locates the facility at 0 or always locates the facility at 1.

Even in this stark model, if we insist on strategy-proofness, individually rationality, and Pareto-efficiency, then the only two options are the dictatorial mechanisms.²⁸ Thus, the requirement of Pareto-efficiency alongside strategy-proofness and individual rationality precludes the possibility that any compromise is ever reached.

Are there attractive mechanisms that compromise if we give up Pareto-efficiency? Consider the following family of mechanisms that select a compromise location at some threshold profiles. A *one-sided unanimous compromise* mechanism is defined by a compromise $x \in [0, 1]$ and a side $J \in \{L, R\}$. If agents in J unanimously find x acceptable, then the mechanism selects the location x . Otherwise, it selects the home location of the other side, denoted K . The set of users is the set of agents willing to travel to the chosen location.

Each one-sided unanimous compromise mechanism is individually rational by definition. It is also strategy-proof: only threshold reports of J are used to determine the location, and it is clear that misreports can only hurt an agent in J , independently of others' reports. However, it is not Pareto-efficient: if every agent in K finds compromise x unacceptable yet every agent in J finds it acceptable, then x is chosen by the mechanism even though the home location of agents in J is a Pareto-improvement.

²⁸ See [the Online Appendix](#) for a proof.

When the chosen location is x , every agent in N_J participates. An agent in N_K who does not participate would do so only if the location were moved some distance towards his home location. But this would make every member of N_J worse off. On the other hand, if the chosen location is the home location of K , then any other location would make members of N_K worse off. Therefore, every one-sided compromise mechanism is Pareto-constrained participation-maximal. So, by Corollary 3, it is on the Pareto-frontier of strategy-proof mechanisms.

Corollary 8. *Each one-sided unanimous compromise mechanism is strategy-proofness-constrained Pareto-efficient.*

5.4 Transferable Utility

Here we consider the problem of making a social decision and assigning payments to agents when preferences are quasilinear in payments. Consequently, there is indifference with the outside option. Nonetheless, Corollary 1 applies if richness of the outside option is satisfied. Below we study its implications.

Suppose $\mathcal{D}$ is a set of **social decisions**. At each $\delta \in \mathcal{D}$, the agents participating in δ are $N(\delta)$. Let $\mathcal{T} \subseteq \mathbb{R}^N$ be the set of **possible payment profiles**. At each $\tau \in \mathcal{T}$, for each $i \in N$, τ_i is the payment that i makes. Like $\mathcal{F}_i$, let $\mathcal{D}_i$ be the set of decisions that i participates at.

To fit this into our general model, an allocation in $\mathcal{F}$ is a pair $(\delta, \tau) \in \mathcal{D} \times \mathcal{T}$. Furthermore, $N(\delta, \tau) = N(\delta)$. Since a non-participant consumes his outside option, we require that for each $i \notin N(\delta)$, $\tau_i = 0$. That is, no transfers are made to or from non-participants.

We assume that each agent's preferences are quasilinear in the payment that he makes. Thus, for each $i \in N$, each $R_i \in \mathcal{R}_i$ is identified by a **valuation**, $v_i \in \mathbb{R}^{\mathcal{D}_i \cup \{\emptyset\}}$. Let i 's **valuation space**, $\mathcal{V}_i$, be the set of all possible valuations for i . Thus, i 's preference relation is represented by $v_i(\delta) - \tau_i$, where $v_i(\delta)$ is the δ th coordinate of v_i if $i \in N(\delta)$ and the $\emptyset$ th coordinate otherwise. The set of all **valuation profiles** is $\mathcal{V} \equiv \times_{i \in N} \mathcal{V}_i$. For each $v \in \mathcal{V}$, $i \in N$, and $\delta \in \mathcal{D}$, denote by $u_i^v(\delta)$ the **net value** of δ to i , relative to being excluded. That is, $u_i^v(\delta) \equiv v_i(\delta) - v_i(\emptyset)$. Interpreting $v_i(\emptyset)$ as i 's opportunity cost of participating, $u_i^v(\delta)$ is i 's valuation of δ net of this cost. Of course, if $i \notin N(\delta)$, then $u_i^v(\delta) = 0$.

The richness of the outside option condition that we described in Section 2 was for preference relations over $\mathcal{F}_i \cup \{\emptyset\}$. It would be met if, for instance, $\mathcal{V}_i$ were such that for each $v_i \in \mathcal{V}_i$ there were $v'_i \in \mathcal{V}_i$, such that the valuations of each of the decisions were unchanged, but the valuation of the outside option were increased. That is, for each $\kappa > v_i(\emptyset)$, there is v'_i such that $v'_i(\emptyset) = \kappa$ and for each $\delta \in \mathcal{D}_i$, $v'_i(\delta) = v_i(\delta)$.²⁹

²⁹ We could state this with a bound on κ if there were a bound on valuations of the decisions in $\mathcal{D}_i$ and

A mechanism in this setting consists of two parts: a **decision rule**, $d : \mathcal{V} \rightarrow \mathcal{D}$, and a **payment rule**, $t : \mathcal{V} \rightarrow \mathcal{T}$. If (d, t) is strategy-proof, we say that t **implements** d . If there is a payment rule that implements d , then d is **implementable**. As defined, implementation says nothing about individual rationality. We define parallel concepts with this requirement added: t **IR-implements** d if (d, t) is not only strategy-proof but also individually rational. In this case, we say that d is **IR-implementable**. If there is a unique payment rule that IR-implements d , we say that d is **uniquely IR-implementable**.

An **efficient** decision rule maximizes the net value of the decision to its participants. That is, a decision rule d is efficient if, for each $v \in \mathcal{V}$,

$$d(v) \in \operatorname{argmax}_{\delta \in \mathcal{D}} \sum_{i \in N} u_i^v(\delta).$$

An efficient mechanism is one that has an efficient decision rule. In a context where the status quo is that each agent enjoys his outside option, an efficient decision maximizes the total surplus over the status quo.

With these concepts in hand, we obtain the following counterpart to Corollary 1:

Corollary 9. *Every IR-implementable decision rule is uniquely IR-implementable.*

In fact we can say more than Corollary 9. For any pair of decision rules d and d' that are participation-equivalent, if t IR-implements d and t' IR-implements d' , then (d, t) and (d', t') are welfare-equivalent.

This leads to a characterization of **pivotal mechanisms**: efficient mechanisms with *pivotal* payment rules. A pivotal payment rule assigns to each agent a payment equal to the externality that his participation imposes on others: t is a pivotal payment rule if, for each $v \in \mathcal{V}$ and each $i \in N$,

$$t_i(v) \equiv \max_{\delta \in \mathcal{D} \setminus \mathcal{D}_i} \left\{ \sum_{j \neq i} u_j^v(\delta) \right\} - \sum_{j \neq i} u_j^v(d(v)).$$

A pivotal payment rule is a particular Groves scheme (Groves, 1973) and pivotal mechanisms are well understood to be strategy-proof (Vickrey, 1961; Clarke, 1971). Moreover, for each $v \in \mathcal{V}$ and each $i \in N$, if $i \notin N(d(v))$, then t prescribes a zero payment for i . Therefore, (d, t) is feasible. Finally, the efficiency of d ensures that (d, t) is also individually rational. Thus, from Corollary 9 we have:

on the payments that i may make in $\mathcal{T}$.

Corollary 10. *A mechanism (d, t) is efficient, strategy-proof, and individually rational if and only if it is a pivotal mechanism. That is, t IR-implements efficient d if and only if it is the pivotal payment rule.*

When we have neither the restriction that non-participants receive zero payments nor the requirement of individual rationality, Groves schemes are the only ones to implement efficient decision rules (Green and Laffont, 1977; Holmström, 1979). For private goods economies, pivotal rules are the only Groves schemes that are individually rational without making payments to non-participants (Chew and Serizawa, 2007). For pure public goods, the counterpart of individual rationality requires that agents have no incentive to free-ride. Substituting this property for individual rationality similarly characterizes pivotal rules (Moulin, 1986). Corollary 10 neither implies nor is implied by existing characterizations of the pivotal mechanism since our assumption that $\mathcal{V}$ satisfy richness of the outside option is logically independent from the domain assumptions made in existing results.

Appendices

A Necessity of Preference Domain Assumptions

The two assumptions, no indifference with the outside option and richness of the outside option, are required for different parts of our results. The following table summarizes where each plays a role.

	Theorem 1 (A) Pareto $\Rightarrow$ Participation	Theorem 1 (B) Pareto $\Leftarrow$ Participation
Richness	Not required	Required
No indifference	Required	Not required

We present below counterexamples demonstrating a failure of the result without the corresponding assumption.

Example 1. *Failure of Theorem 1 (A) without the no indifference with the outside option assumption.*

Let $N \equiv \{i_1, i_2\}$ and $\mathcal{F} \equiv \{\alpha, \beta\}$ such that $N(\alpha) = \{i_1\}$ and $N(\beta) = \{i_2\}$. Then, i_2 identifies α with $\emptyset$, and i_1 identifies β with $\emptyset$. Let $\mathcal{R}_{i_1}$ consist only of R_{i_1} and let $\mathcal{R}_{i_2}$ consist of R_{i_2}

and R'_{i_2} , defined below.

R_{i_1}	R_{i_2}	R'_{i_2}
α	β	$\beta, \emptyset$
$\emptyset$	$\emptyset$	

These preference domains trivially satisfy richness of the outside option.

Consider φ that always selects β . That is, $\varphi(R_{i_1}, R_{i_2}) = \varphi(R_{i_1}, R'_{i_2}) = \beta$. It is strategy-proof, individually rational, and Pareto-constrained participation-maximal. Let φ' be the mechanism that selects β if i_2 prefers it to $\emptyset$, but α when i_2 is indifferent. That is, $\varphi'(R_{i_1}, R_{i_2}) = \beta$ and $\varphi'(R_{i_1}, R'_{i_2}) = \alpha$. Not only is φ' strategy-proof and individually rational, but it also *strictly* Pareto-improves φ , yet $\varphi(R_{i_1}, R'_{i_2})$ and $\varphi'(R_{i_1}, R'_{i_2})$ cannot be compared by participation-expansion.

Thus, in contrast to Theorem 1 (A), we have that, without the no indifference with the outside option assumption, Pareto-improvement does not imply participation-expansion. Nevertheless, it is straightforward to see that, given individual rationality, Pareto-improvement implies participation-expansion or participation-incomparability. That is, Pareto-improvement refines participation-expansion.

The failure of Theorem 1 (B) without richness of the outside option is easy to see by considering the model of [Shapley and Scarf \(1974\)](#), where richness fails because not receiving an object is always ranked at the bottom. Given the endowment profile, the core mechanism strictly Pareto-improves the no-trade mechanism but both are strategy-proof, individually rational, and participation-equivalent.

We provide another simple example to demonstrate this failure.

Example 2. *Failure of Theorem 1 (B) without richness of the outside option.*

Let $N \equiv \{i_1, i_2\}$ and $\mathcal{F} \equiv \{\alpha, \beta, \gamma\}$ such that $N(\alpha) = \{i_1, i_2\}$, $N(\beta) = \{i_1\}$, and $N(\gamma) = \emptyset$. Then, i_2 identifies both β and γ with $\emptyset$ and i_1 identifies only γ with $\emptyset$. Let $\mathcal{R}_{i_1}$ consist of R_{i_1} , R'_{i_1} , and R''_{i_1} and $\mathcal{R}_{i_2}$ consist of R_{i_2} and R'_{i_2} , defined below.

R_{i_1}	R'_{i_1}	R''_{i_1}	R_{i_2}	R'_{i_2}
α	β	α	α	$\emptyset$
β	α	$\emptyset$	$\emptyset$	α
$\emptyset$	$\emptyset$	β		

Let φ^1 and φ^2 be such that for each $R \in \mathcal{R}$,

$$\varphi^1(R) = \begin{cases} \alpha & \text{if } \alpha P_{i_2} \emptyset \\ \beta & \text{if } \emptyset P_{i_2} \alpha \text{ and } \beta P_{i_1} \emptyset \\ \gamma & \text{otherwise} \end{cases}$$

and

$$\varphi^2(R) = \begin{cases} \beta & \text{if } \beta P_{i_1} \alpha \\ \alpha & \text{otherwise if } \alpha P_{i_1} \emptyset \text{ and } \alpha P_{i_2} \emptyset \\ \gamma & \text{otherwise.} \end{cases}$$

Both φ^1 and φ^2 are strategy-proof, individually rational, and *Pareto-efficient*. Nonetheless, φ^1 participation-expands φ^2 , so Theorem 1 (B) does not hold. Notice that both agents' preferences satisfy no indifference with the outside option, but i_1 's preferences fail the requirement of richness of the outside option: conditional on i_1 preferring β to α , he prefers α to $\emptyset$.

B Proof of Lemma 2

Before showing that stability implies non-wastefulness under the assumptions of size monotonicity and idempotence, we start with a definition and a lemma. For each $P \in \mathcal{P}$, each $\mu \in \mathcal{F}$, and each $o \in O$, let $\mathbf{Y}_o^\mu(P)$ be the triples in $X(o)$ that are associated with an agent who prefers it to what he is assigned at μ . That is, $Y_o^\mu(P) \equiv \{x \in X(o) : x P_{N(x)} \mu(N(x))\}$.

Lemma 3. *Let C be size monotonic and idempotent. For each $P \in \mathcal{P}$, each stable $\mu \in \mathcal{F}$, each $o \in O$, each finite $Y \subseteq Y_o^\mu(P)$, and each $Z \in C_o(\mu(o) \cup Y)$, $|Z| = |\mu(o)|$.*

Proof. We proceed by induction over subsets of $Y_o^\mu(P)$. Let $Y \subseteq Y_o^\mu(P)$.

For the base case, where $Y = \emptyset$, since $\mu(o) \in \text{range}(C_o)$, by idempotence of C , $\mu(o) \in C_o(\mu(o) \cup Y)$ and by size monotonicity of C , for each $Z \in C_o(\mu(o))$, $|Z| = |\mu(o)|$.

As an induction hypothesis, assume that for each $Y' \subsetneq Y$ and each $Z \in C_o(\mu(o) \cup Y')$, $|Z| = |\mu(o)|$. Equivalently, for each $T \subseteq \mu(o) \cup Y$ such that $Y \not\subseteq T$, for each $Z \in C_o(\mu(o) \cup T)$, $|Z| = |\mu(o)|$.

The induction step is to show that, for each $Z \in C_o(\mu(o) \cup Y)$, $|Z| = |\mu(o)|$. Let $Z \in C_o(\mu(o) \cup Y)$. By idempotence of C , $Z \in C_o(Z)$. Thus, since $Z \subseteq \mu(o) \cup Z$, by size monotonicity of C ,

$$\text{for each } Z' \in C_o(\mu(o) \cup Z), |Z| \leq |Z'|. \quad (1)$$

By stability of μ , either $Y \not\subseteq Z$ or $\mu(o) \in C_o(\mu(o) \cup Y)$. If $\mu(o) \in C_o(\mu(o) \cup Y)$, then by size monotonicity of C , for each $Z \in C_o(\mu(o) \cup Y)$, $|Z| = |\mu(o)|$. Instead, if $Y \not\subseteq Z$, the induction hypothesis implies

$$\text{for each } Z' \in C_o(\mu(o) \cup Z), |Z'| = |\mu(o)|. \quad (2)$$

By (1) and (2), $|Z| \leq |\mu(o)|$. Since C is idempotent and $\mu \in \mathcal{F}$, $\mu(o) \in C_o(\mu(o))$. Then, since C is size monotonic and $Z \in C_o(\mu(o) \cup Y)$, $|\mu(o)| \leq |Z|$. Thus, we conclude that $|Z| = |\mu(o)|$. $\square$

Proof of Lemma 2. Suppose that μ is wasteful. Then there are $o \in O$ and $v \in \mathcal{F}$ such that $|\nu(o)| > |\mu(o)|$ and, for each $y \in \nu(o) \setminus \mu(o)$, $y P_{N(y)} \mu(N(y))$. Let $Y \equiv \nu(o) \setminus \mu(o)$. Since $Y \subseteq Y_o^\mu(P)$ and Y is finite, by Lemma 3, for each $Z' \in C_o(\mu(o) \cup Y)$, $|Z'| = |\mu(o)|$. However, $\nu(o) \subseteq \mu(o) \cup Y$. So by size monotonicity of C , for each $Z \in C_o(\nu(o))$ and each $Z' \in C_o(\mu(o) \cup Y)$, $|Z| \leq |Z'| = |\mu(o)|$. Since $\nu \in \mathcal{F}$, $\nu \in \mathcal{F}_o = \text{range}(C_o)$. So by idempotence of C , $\nu(o) \in C_o(\nu(o))$. Thus, $|\nu(o)| \leq |\mu(o)|$. This contradicts the definition of ν . $\square$

References

- Abdulkadiroğlu, Atila, Parag A Pathak, and Alvin E Roth (2009) “Strategy-Proofness versus Efficiency in Matching with Indifferences: Redesigning the NYC High School Match,” *American Economic Review*, Vol. 99, No. 5, pp. 1954–1978. [14], [18], [19]
- Abdulkadiroğlu, Atila and Tayfun Sönmez (2003) “School Choice: A Mechanism Design Approach,” *American Economic Review*, Vol. 93, No. 3, pp. 729–747. [1], [17]
- Alcalde, José and Salvador Barberà (1994) “Top Dominance and the Possibility of Strategy-Proof Stable Solutions to Matching Problems,” *Economic Theory*, Vol. 4, No. 3, pp. 417–35. [19]
- Alkan, Ahmet (2002) “A class of multipartner matching markets with a strong lattice structure,” *Economic Theory*, Vol. 19, No. 4, pp. 737–746. [15]
- Alkan, Ahmet and David Gale (2003) “Stable Schedule Matching Under Revealed Preference,” *Journal of Economic Theory*, Vol. 112, No. 2, pp. 289–306. [15]
- Alva, Samson (2018) “WARP and combinatorial choice,” *Journal of Economic Theory*, Vol. 173, pp. 320–333. [15]
- Alva, Samson and Vikram Manjunath (2019) “Stable-dominating Rules,” working paper, University of Texas at San Antonio. [9], [18]
- Anno, Hidekazu and Morimitsu Kurino (2016) “On the operation of multiple matching markets,” *Games and Economic Behavior*, Vol. 100, pp. 166–185. [10]

- Anno, Hidekazu and Hiroo Sasaki (2013) “Second-best efficiency of allocation rules: Strategy-proofness and single-peaked preferences with multiple commodities,” *Economic Theory*, Vol. 54, No. 3, pp. 693–716. [10]
- Aygün, Orhan and Tayfun Sönmez (2013) “Matching with Contracts: Comment,” *American Economic Review*, Vol. 103, No. 5, pp. 2050–51. [15]
- Balinski, Michel and Tayfun Sönmez (1999) “A Tale of Two Mechanisms: Student Placement,” *Journal of Economic Theory*, Vol. 84, No. 1, pp. 73–94. [13], [17]
- Bogomolnaia, Anna, Rajat Deb, and Lars Ehlers (2005) “Strategy-proof assignment on the full preference domain,” *Journal of Economic Theory*, Vol. 123, No. 2, pp. 161–186. [12]
- Cantala, David (2004) “Choosing the level of a public good when agents have an outside option,” *Social Choice and Welfare*, Vol. 22, No. 3, pp. 491–514. [20]
- Chew, Soo Hong and Shigehiro Serizawa (2007) “Characterizing the Vickrey Combinatorial Auction by Induction,” *Economic Theory*, Vol. 33, pp. 393–406. [23]
- Clarke, Edward H. (1971) “Multipart pricing of public goods,” *Public Choice*, Vol. 11, No. 1, pp. 17–33. [22]
- Dur, Umut, A. Arda Gitmez, and Özgür Yılmaz (forthcoming) “School Choice under Partial Fairness,” *Theoretical Economics*. [3]
- Ehlers, Lars and Aytok Erdil (2010) “Efficient assignment respecting priorities,” *Journal of Economic Theory*, Vol. 145, No. 3, pp. 1269–1282. [18], [19]
- Ehlers, Lars, Isa E. Hafalir, M. Bumin Yenmez, and Muhammed A. Yildirim (2014) “School choice with controlled choice constraints: Hard bounds versus soft bounds,” *Journal of Economic Theory*, Vol. 153, pp. 648 – 683. [17]
- Ehlers, Lars and Bettina Klaus (2014) “Strategy-Proofness Makes the Difference: Deferred-Acceptance with Responsive Priorities,” *Mathematics of Operations Research*, Vol. 39, No. 4, pp. 949–966. [13]
- Ehlers, Lars and Thayer Morrill (2018) “(Il)legal Assignments in School Choice,” working paper, North Carolina State University. [3], [19]
- Erdil, Aytok (2014) “Strategy-proof stochastic assignment,” *Journal of Economic Theory*, Vol. 151, pp. 146 – 162. [2], [8]
- Erdil, Aytok and Haluk I. Ergin (2008) “What’s the Matter with Tie-Breaking? Improving Efficiency in School Choice,” *American Economic Review*, Vol. 98, No. 3, pp. 669–89. [14], [18]
- (2017) “Two-Sided Matching with Indifferences,” *Journal of Economic Theory*, Vol. 171, pp. 268 – 292. [6], [12]

- Ergin, Haluk I. (2002) “Efficient Resource Allocation on the Basis of Priorities,” *Econometrica*, Vol. 70, No. 6, pp. 2489–2497. [18]
- Fleiner, Tamás (2003) “A fixed-point approach to stable matchings and some applications,” *Mathematics of Operations Research*, Vol. 28, pp. 103–126. [15]
- Gale, David and Lloyd Shapley (1962) “College Admissions and the Stability of Marriage,” *American Mathematical Monthly*, Vol. 69, pp. 9–15. [18], [19]
- Goto, Masahiro, Fuhito Kojima, Ryoji Kurata, Akihisa Tamura, and Makoto Yokoo (2017) “Designing Matching Mechanisms under General Distributional Constraints,” *American Economic Journal: Microeconomics*, Vol. 9, No. 2, pp. 226–262. [13], [16]
- Green, Jerry and Jean-Jacques Laffont (1977) “Characterization of Satisfactory Mechanisms for the Revelation of Preferences for Public Goods,” *Econometrica*, Vol. 45, pp. 417–438. [23]
- Groves, Theodore (1973) “Incentives in Teams,” *Econometrica*, Vol. 41, No. 4, pp. 617–631. [22]
- Hafalir, Isa E., M. Bumin Yenmez, and Muhammed A. Yildirim (2013) “Effective Affirmative Action in School Choice,” *Theoretical Economics*, Vol. 8, No. 2, pp. 325–363. [17]
- Hatfield, John William and Scott Duke Kominers (2016) “Hidden Substitutes,” working paper, University of Texas at Austin. [17]
- Hatfield, John William and Paul Milgrom (2005) “Matching with Contracts,” *American Economic Review*, Vol. 95, No. 4, pp. 913–935. [1], [11], [15], [17]
- Hirata, Daisuke and Yusuke Kasuya (2017) “On Stable and Strategy-Proof Rules in Matching Markets with Contracts,” *Journal of Economic Theory*, Vol. 168, pp. 27–43. [3], [16]
- Holmström, Bengt (1979) “Groves’ Scheme on Restricted Domains,” *Econometrica*, Vol. 47, No. 5, pp. 1137–1144. [23]
- Hylland, Aanund and Richard Zeckhauser (1979) “The Efficient Allocation of Individuals to Positions,” *Journal of Political Economy*, Vol. 87, No. 2, pp. 293–314. [1]
- Jackson, Matthew O. and Antonio Nicolò (2004) “The strategy-proof provision of public goods under congestion and crowding preferences,” *Journal of Economic Theory*, Vol. 115, No. 2, pp. 278–308. [1], [20]
- Kamada, Yuichiro and Fuhito Kojima (2015) “Efficient Matching under Distributional Constraints: Theory and Applications,” *American Economic Review*, Vol. 105, No. 1, pp. 67–99. [13], [16]
- Kesten, Onur and Morimitsu Kurino (2016) “Do outside options matter in matching? A new perspective on the trade-offs in school choice,” working paper, Carnegie Mellon University. [18]

- Kominers, Scott Duke and Tayfun Sönmez (2016) “Matching with Slot-Specific Priorities: Theory,” *Theoretical Economics*, Vol. 11, No. 2, pp. 683–710. [17]
- Moulin, Hervé (1986) “Characterizations of the Pivotal Mechanism,” *Journal of Public Economics*, Vol. 31, No. 1, pp. 53–78. [23]
- Ohseto, Shinji (2006) “Characterizations of strategy-proof and fair mechanisms for allocating indivisible goods,” *Economic Theory*, Vol. 29, No. 1, pp. 111–121. [10]
- Roth, Alvin E. (1986) “On the Allocation of Residents to Rural Hospitals: A General Property of Two-Sided Matching Markets,” *Econometrica*, Vol. 54, No. 2, pp. 425–427. [9], [16]
- Shapley, Lloyd and Herbert Scarf (1974) “On cores and indivisibility,” *Journal of Mathematical Economics*, Vol. 1, No. 1, pp. 23–37. [24]
- Sönmez, Tayfun (1999) “Strategy-Proofness and Essentially Single-Valued Cores,” *Econometrica*, Vol. 67, No. 3, pp. 677–690. [6], [7]
- Sönmez, Tayfun (2013) “Bidding for Army Career Specialties: Improving the ROTC Branching Mechanism,” *Journal of Political Economy*, Vol. 121, No. 1, pp. 186–219. [17]
- Sönmez, Tayfun and Tobias B. Switzer (2013) “Matching With (Branch-of-Choice) Contracts at the United States Military Academy,” *Econometrica*, Vol. 81, No. 2, pp. 451–488. [17]
- Sprumont, Yves (2013) “Constrained-optimal strategy-proof assignment: Beyond the Groves mechanisms,” *Journal of Economic Theory*, Vol. 148, No. 3, pp. 1102–1121. [10]
- Vickrey, William (1961) “Counterspeculation, Auctions, and Competitive Sealed Tenders,” *Journal of Finance*, Vol. 16, No. 1, pp. 8–37. [22]