

Centre interuniversitaire de recherche
en économie quantitative

CIREQ

Cahier 15-2012

*Strategy-Proofness Makes the Difference :
Deferred-Acceptance with Responsive Priorities*

Lars EHLERS and
Bettina KLAUS



Le **Centre interuniversitaire de recherche en économie quantitative (CIREQ)** regroupe des chercheurs dans les domaines de l'économétrie, la théorie de la décision, la macroéconomie et les marchés financiers, la microéconomie appliquée et l'économie expérimentale ainsi que l'économie de l'environnement et des ressources naturelles. Ils proviennent principalement des universités de Montréal, McGill et Concordia. Le CIREQ offre un milieu dynamique de recherche en économie quantitative grâce au grand nombre d'activités qu'il organise (séminaires, ateliers, colloques) et de collaborateurs qu'il reçoit chaque année.

*The **Center for Interuniversity Research in Quantitative Economics (CIREQ)** regroups researchers in the fields of econometrics, decision theory, macroeconomics and financial markets, applied microeconomics and experimental economics, and environmental and natural resources economics. They come mainly from the Université de Montréal, McGill University and Concordia University. CIREQ offers a dynamic environment of research in quantitative economics thanks to the large number of activities that it organizes (seminars, workshops, conferences) and to the visitors it receives every year.*

Cahier 15-2012

Strategy-Proofness Makes the Difference : Deferred-Acceptance with Responsive Priorities

Lars EHLERS and Bettina KLAUS

CIREQ, Université de Montréal
C.P. 6128, succursale Centre-ville
Montréal (Québec) H3C 3J7
Canada

téléphone : (514) 343-6557
télécopieur : (514) 343-7221
cireq@umontreal.ca
<http://www.cireqmontreal.com>



Ce cahier a également été publié par le Département de sciences économiques de l'Université de Montréal sous le numéro (2012-12).

This working paper was also published by the Department of Economics of the University of Montreal under number (2012-12).

Dépôt légal - Bibliothèque nationale du Canada, 2012, ISSN 0821-4441

Dépôt légal - Bibliothèque et Archives nationales du Québec, 2012

ISBN-13 : 978-2-89382-637-0

Strategy-Proofness makes the Difference: Deferred-Acceptance with Responsive Priorities*

Lars Ehlers[†]

Bettina Klaus[‡]

September 2012

Abstract

In college admissions and student placements at public schools, the admission decision can be thought of as assigning indivisible objects with capacity constraints to a set of students such that each student receives at most one object and monetary compensations are not allowed. In these important market design problems, the agent-proposing deferred-acceptance (DA-)mechanism with responsive strict priorities performs well and economists have successfully implemented DA-mechanisms or slight variants thereof. We show that almost all real-life mechanisms used in such environments—including the large classes of priority mechanisms and linear programming mechanisms—satisfy a set of simple and intuitive properties. Once we add *strategy-proofness* to these properties, DA-mechanisms are the *only* ones surviving. In market design problems that are based on weak priorities (like school choice), generally multiple tie-breaking (MTB) procedures are used and then a mechanism is implemented with the obtained strict priorities. By adding stability with respect to the weak priorities, we establish the first normative foundation for MTB-DA-mechanisms that are used in NYC.

JEL Classification: D63, D70

Keywords: deferred-acceptance mechanism, indivisible objects allocation, multiple tie-breaking, school choice, strategy-proofness.

1 Introduction

In centralized entry-level medical markets students search for their first positions at hospitals. Each year first the finishing students submit a ranking of the available hospital positions and the

*Lars Ehlers acknowledges the hospitality of Toulouse School of Economics (INRA-LERNA), where an early version of this paper was written, and financial support from the SSHRC (Canada) and the FQRSC (Québec). Bettina Klaus acknowledges the hospitality of Harvard Business School, where part of this paper was written, and financial support from the Netherlands Organisation for Scientific Research (NWO) under grant VIDI-452-06-013. This article is a substantially revised version of Ehlers and Klaus (2009).

[†]*Corresponding author:* Département de Sciences Économiques and CIREQ, Université de Montréal, Montréal, Québec H3C 3J7, Canada; e-mail: lars.ehlers@umontreal.ca.

[‡]Faculty of Business and Economics, University of Lausanne, Internef 538, CH-1015, Lausanne, Switzerland; e-mail: bettina.klaus@unil.ch.

hospitals submit a ranking of the finishing students (and hospitals' preferences over sets of students are supposed to be "responsive" with respect to these ranking). Then, the centralized clearinghouse matches via a mechanism students and hospitals on the basis of the submitted lists. A matching is stable if (i) every student prefers his match to being unmatched, (ii) every hospital prefers each matched student to having the position unfilled, and (iii) there is no student-hospital pair that blocks the matching in the sense that the student prefers the hospital to his match and the hospital prefers the student to one of its matched students or having the position unfilled.

The American markets (Roth, 1984a) and the British markets in Edinburgh and Cardiff based the matching on a stable mechanism called hospital-proposing deferred-acceptance (DA-)mechanism (Gale and Shapley, 1962). Although the DA-mechanism is the predominant mechanism in the two-sided matching market literature (Roth, 2008), British markets have introduced and operated unstable mechanisms (Roth, 1991). The markets in Edinburgh (initially), Newcastle, and Birmingham have adopted priority mechanisms in 1966 and 1967, and the markets in Cambridge and the London Hospital Medical College have adopted linear programming mechanisms in the 1970's. The literature on two-sided matching markets has argued that in such markets stable mechanisms are more successful than others because unstable mechanisms may produce matchings which are blocked and unsustainable (e.g., Roth, 1991). This is especially due to the fact that both sides of the market have to be considered as active agents who, after the announced outcome of the centralized clearinghouse, can resign from positions (students) or fire students (hospitals) in order to engage in new matches.

In all these markets preferences are private information and agents need to reveal them. Ideally we would like to base the outcome on the true preferences and construct an incentive compatible mechanism. Unfortunately there does not exist any stable mechanism that, for all students and for all hospitals, makes truth-telling a weakly dominant strategy (Roth, 1982). However, the student-proposing DA-mechanism is strategy-proof for students (Dubins and Freedman, 1981; Roth, 1982), but there does not exist any stable mechanism that is strategy-proof for hospitals with several positions (Roth, 1984b). This is one of the main reasons why in the mid 1990's the American markets changed the mechanism from the hospital-proposing DA-mechanism to the student-proposing DA-mechanism.

Over the past couple of years, the two-sided matching problem has been adapted to address the assignment of students to public schools. In a so-called school choice problem, schools correspond to hospitals and students seek a slot at a school (Abdulkadiroğlu and Sönmez, 2003). The difference to a two-sided matching market problem is that a school's "preference relation" represents priorities of students to be admitted at that school. These priorities over individual agents are fixed (via exogenous criteria or objective tests) and cannot be changed. Furthermore, schools are not strategic agents. Now if priorities are strict (and priorities over sets of students are responsive with respect to the priorities over individual agents), then the DA-mechanism can be applied in a straightforward

manner. The same is true for priority mechanisms and linear programming mechanisms. Indeed, the famous Boston mechanism for school choice is equivalent to the Edinburgh 1967 priority mechanism (Ergin and Sönmez, 2006) and priority mechanisms are nowadays still in use in school choice.¹

In school choice, schools are not strategic agents and the motivation for stability needs to be slightly reconsidered (a school simply cannot deny enrollment to a student who has been assigned to it by the central authorities in order to rematch with a preferred student). The intuitive reinterpretation of stability for school choice problems is that stability property (i) still represents students' individual rationality, stability property (ii) corresponds to the non-wastefulness of school seats, and stability property (iii) captures the fairness property of no justified envy² (Balinski and Sönmez, 1999). Next, several contributions underline the policy importance of strategy-proofness for school choice mechanisms. The argument there is not simply that truthful revelation of preferences is desirable, but that it is important to level the playing field for students from different backgrounds. Abdulkadiroğlu, Pathak, Roth, and Sönmez (2006) and Pathak and Sönmez (2008) show that under the Boston mechanism sophisticated agents were much better off than unsophisticated agents: naïvely truth-telling students (or parents) tend to be the worst off students under the Boston mechanism, which is *not* strategy-proof. This is due to the fact that with a manipulable mechanisms it may be difficult to figure out the best (non-truthful) strategy whereas with a strategy-proof mechanisms each agent's simple strategy is to report preferences truthfully. In this paper we formulate some simple and intuitive properties that are satisfied by many mechanisms (using some arbitrary strict priorities) including priority mechanisms, linear programming mechanisms and responsive DA-mechanisms. However, we show that once strategy-proofness is added, only the student-proposing DA-mechanisms (with responsive priorities) survive among all mechanisms satisfying our set of simple and intuitive properties (Theorem 1). This provides a rationale for using the student-proposing responsive DA-mechanism if truthful revelation of preferences is important. Furthermore, the feature of responsiveness-to-individual-priorities makes this mechanism easily applicable in practice (Roth, 2008).³

For some school choice problems priorities may be coarse, e.g., in Boston priorities are divided into four classes: (i) siblings and walk zone, (ii) siblings, (iii) walk zone, and (iv) the rest. In extreme cases, no priorities are given and all students have equal rights for receiving any object. When priorities are weak, in order to apply a (student-proposing) DA-mechanism, one needs to resolve the ties in some way. Abdulkadiroğlu, Pathak, and Roth (2009) propose to use a fixed

¹In Boston this mechanism has now been changed to the student-proposing DA-mechanism, but many cities worldwide still use the so-called Boston (or direct acceptance) mechanism.

²A matching eliminates justified envy if whenever a student prefers another student's match to his own, he has a lower priority than the other student at that school.

³The search for "good" mechanisms is the subject of many recent contributions, but most of them deal with the house allocation model where exactly one object of each type is available (e.g., Ehlers, 2002; Ehlers and Klaus, 2006, 2007; Kesten, 2009; Pápai, 2000). In most papers that study the allocation of indivisible objects with capacity constraints, externally prescribed priorities are also specified; the corresponding class of problems is usually referred to as "school choice problems" or "student placement problems" (see Sönmez and Ünver, 2011, for a recent survey).

multiple tie-breaking procedure (MTB) (which is independent of students' preferences) and then apply the DA-mechanism. Erdil and Ergin (2008) propose to use different tie-breaking procedures for different preference profiles. Both these contributions regard stability (with respect to the weak priorities) as an important requirement. We show that in Theorem 1, we can replace individual rationality and weak non-wastefulness with stability with respect to weak priorities, and by doing so we establish the first normative foundation of responsive MTB-DA-mechanisms as used in NYC (Theorem 2).

The paper is organized as follows. In Section 2 we introduce a model of allocating indivisible objects with capacity constraints and exogenously given priorities. In Section 3 we introduce tie-breaking procedures for exogenously given priorities, priority mechanisms, linear programming mechanisms, and responsive DA-mechanisms. In Section 4 we state simple and intuitive properties for mechanisms. In Section 5 we show that all described real-life mechanisms satisfy our simple and intuitive properties. We then show that only (agent-proposing) responsive DA-mechanisms satisfy these simple and intuitive properties *and* strategy-proofness. For exogenously given priorities, in Section 6, we then obtain a characterization of responsive MTB-DA mechanisms when imposing stability. The Appendix contains all our proofs and in addition a second characterization of responsive DA-mechanisms with weak consistency.⁴

2 Allocation with Priorities and Variable Resources

Let $N = \{1, \dots, n\}$ denote a finite set of agents with $n \geq 2$. Let O denote a set of potential (real) object types or types for short. We assume that O contains at least two elements and that O is finite.⁵ Not receiving any real object is called “receiving the null object.” Let $\emptyset$ represent the *null object*.

Each agent $i \in N$ is equipped with a *preference relation* R_i over all types $O \cup \{\emptyset\}$. The preference relation R_i is strict, i.e., R_i is a linear order over $O \cup \{\emptyset\}$. Given $x, y \in O \cup \{\emptyset\}$, $x P_i y$ means that agent i strictly prefers x to y (and $x \neq y$) and $x R_i y$ means that agent i weakly prefers x to y (and $x P_i y$ or $x = y$). Let $\mathcal{R}$ denote the *set of all preferences* over $O \cup \{\emptyset\}$, and $\mathcal{R}^N$ the *set of all (preference) profiles* $R = (R_i)_{i \in N}$ such that for all $i \in N$, $R_i \in \mathcal{R}$.

Given $R \in \mathcal{R}^N$ and $M \subseteq N$, let R_M denote the profile $(R_i)_{i \in M}$. It is the restriction of R to the set of agents M . We also use the notation $R_{-M} = R_{N \setminus M}$ and $R_{-i} = R_{N \setminus \{i\}}$.

Given $O' \subseteq O \cup \{\emptyset\}$, let $R_i|_{O'}$ denote the restriction of R_i to O' and $R|_{O'} = (R_i|_{O'})_{i \in N}$. Given $i \in N$ and $R_i \in \mathcal{R}$, type $x \in O$ is *acceptable under* R_i if $x P_i \emptyset$. Let $A(R_i) = \{x \in O : x P_i \emptyset\}$ denote the *set of acceptable types under* R_i .

⁴Our Supplementary Appendix also discusses Kojima and Manea (2010) in relation to our results and provides the independence of the properties used in our characterizations.

⁵Our results remain unchanged when O is infinite. For expositional convenience, finiteness of O is assumed.

For each type $x \in O$, at most $\bar{q}_x \in \mathbb{N}$ copies are available in any economy with $1 \leq \bar{q}_x \leq |N|$.⁶ Let q_x , $0 \leq q_x \leq \bar{q}_x$, denote the number of available objects or the *capacity* of type x . Let $q = (q_x)_{x \in O}$ denote a *capacity vector* and $\mathcal{Q} = \times_{x \in O} \{0, 1, \dots, \bar{q}_x\}$ denote the *set of all capacity vectors*. The null object is always available without scarcity and therefore we set $q_\emptyset = \infty$.

Given type x , let $\hat{\succsim}_x$ denote a (*weak*) *priority ordering* of type x on N . Here $i \hat{\succ}_x j$ means that agent i has higher priority than agent j for obtaining an object of type x and $i \hat{\sim}_x j$ means that agents i and j have the same priority for obtaining an object of type x . For each type x , the priority ordering $\hat{\succsim}_x$ is exogenously given. The priority ordering $\hat{\succsim}_x$ is strict if for any distinct $i, j \in N$, we have either $i \hat{\succ}_x j$ or $j \hat{\succ}_x i$. Let $\hat{\succsim} = (\hat{\succsim}_x)_{x \in O}$ denote a *priority structure* (e.g., determined via test/exam scores or other objective suitability criteria). Priority structure $\hat{\succsim}$ is strict if for all $x \in O$, priority ordering $\hat{\succsim}_x$ is strict. For a strict priority structure $\hat{\succsim}$ we sometimes write $\hat{\succ}$ instead of $\hat{\succsim}$. In the sequel, we denote an exogenously given priority structure by $\hat{\succsim}$ and an arbitrary strict priority structure by $\succ$. In some environments no priorities are given (no tests/exams are conducted). Then all agents have equal priority for any type. We denote the “no or equal priorities” priority structure by $\hat{\sim}$, i.e., for all x and for all $i, j \in N$, $i \hat{\sim}_x j$.

An (*allocation*) *problem (with capacity constraints and priorities)* consists of a preference profile $R \in \mathcal{R}^N$, a capacity vector $q \in \mathcal{Q}$, and a priority structure $\hat{\succsim}$. Since $\hat{\succsim}$ is exogenously given and remains fixed throughout, we often denote a *problem* by (R, q) and the *set of all problems* by $\mathcal{R}^N \times \mathcal{Q}$. Given a capacity vector q , let $O_+(q) = \{x \in O : q_x > 0\}$ denote the *set of available real types* under q . The *set of available types* is $O_+(q) \cup \{\emptyset\}$ and includes the null object, which is available at any problem.

Each agent i is to be allocated exactly one object in $O \cup \{\emptyset\}$ taking capacity constraints into account. An *allocation for (capacity vector) q* is a list $a = (a_i)_{i \in N}$ such that for all $i \in N$, $a_i \in O \cup \{\emptyset\}$, and any real type $x \in O$ is not assigned more than q_x times, i.e., for all $x \in O$, $|\{i \in N : a_i = x\}| \leq q_x$. Note that $\emptyset$, the null object, can be assigned to any number of agents and that not all real objects have to be assigned. Let $\mathcal{A}_q$ denote the *set of all allocations for q* and $\mathcal{A} = \bigcup_{q \in \mathcal{Q}} \mathcal{A}_q$ *set of all allocations*. Given a problem (R, q) and $a \in \mathcal{A}_q$, allocation a is *individually rational* if for all $i \in N$, $a_i R_i \emptyset$.

An *mechanism* is a function $\varphi : \mathcal{R}^N \times \mathcal{Q} \rightarrow \mathcal{A}$ such that for all problems $(R, q) \in \mathcal{R}^N \times \mathcal{Q}$, $\varphi(R, q) \in \mathcal{A}_q$. Given $i \in N$, we call $\varphi_i(R, q)$ the *allotment* of agent i at $\varphi(R, q)$.

3 Multiple Tie-Breaking and Real Life Mechanisms

For the application of the real-life mechanisms we introduce in this section, we first have to resolve or break ties. In other words, these mechanisms require the transformation of a weak priority

⁶By introducing “global upper bounds” via $\bar{q}_x$ ($x \in O$) we can, for instance, specify the *house allocation model* where at most one object of each type is available, i.e., for all $x \in O$, $\bar{q}_x = 1$.

structure $\widehat{\succsim}$ into a strict priority structure $\succ$. In applications this is typically done by breaking the ties for any type. A strict priority structure $\succ$ is a multiple tie-breaking (MTB) of the (weak) priority structure $\widehat{\succsim}$ if for any type x and any $i, j \in N$: if $i \widehat{\succsim}_x j$, then $i \succ_x j$. Note that then any tie in $\widehat{\succsim}$ is broken and $\succ$ respects all strict priorities of $\widehat{\succsim}$. Any MTB of $\widehat{\succsim}$ is strict and we often use the terminology “for any MTB $\succ$ of $\widehat{\succsim}$.”

3.1 Priority Mechanisms

Below we describe the large class of priority mechanisms. Any priority mechanism, for each preference profile R and any MTB $\succ$ of $\widehat{\succsim}$, matches sequentially all student-type pairs that have highest “priority” among all remaining pairs. First, the priority mechanism checks whether there are (1,1)-matches (i.e., an agent-type pair who ranks each other first). If there are any (1,1)-matches, then they are realized. Second, the priority mechanism checks whether there are matches that have the second highest priority (i.e., (1,2)- or (2,1)-matches) and matches any such pair. And so on. Of course, a match is only realized if the agent prefers his matched object to the null object. Because priority mechanisms differ in how they exactly prioritize matches, we have just sketched a class of mechanisms (note the indeterminacy in the above description when mentioning (1,2)- or (2,1)-matches as having the second highest priority). In particular, the Edinburgh 1967 mechanism ordered matches lexicographically according to the types’ priorities, that is, (1,1)-matches have first priority, followed by (2,1), (3,1), and so on. Only when all types’ first choices are exhausted, other matches ((1,2), (2,2), (3,2), ...) are considered. We will call any mechanism that orders matches lexicographically in types’ priorities or in students’ preferences, a *lexicographic priority mechanism*. The Boston mechanism for public school choice is a lexicographic priority mechanism (Ergin and Sönmez, 2006).

Given an agent’s preference relation R_i , we say that agent i ’s most preferred real type has rank 1, his second most preferred real type has rank 2, etc.. For each $x \in O$, we denote by $\text{rank}(x, R_i)$ the rank of x in R_i . Note that we do not rank the null object. Similarly we define $\text{rank}(i, \succ_x)$. Given $(R, \succ)$, we call (i, x) an (a, b) -match if $\text{rank}(x, R_i) = a$ and $\text{rank}(i, \succ_x) = b$.

A *priority function* is a one-to-one function $g : \mathbb{N} \times \mathbb{N} \rightarrow \mathbb{N}$ such that

- (i) for all $(a, b), (a', b) \in \mathbb{N} \times \mathbb{N}$, if $a < a'$, then $g(a, b) < g(a', b)$, and
- (ii) for all $(a, b), (a, b') \in \mathbb{N} \times \mathbb{N}$, if $b < b'$, then $g(a, b) < g(a, b')$.

The function g associates with each (a, b) -match its priority $g(a, b)$. Hence, in our terminology an (a, b) -match has higher priority than an (a', b') -match if $g(a, b) < g(a', b')$.

For each problem (R, q) , we will only consider the ranks of available real types under q . Thus, the priority function will take as inputs the rank information based on the reduced preference profile $R|_{O_+(q) \cup \{\emptyset\}}$. Then, for each problem (R, q) and any MTB $\succ$ of $\widehat{\succsim}$, the *MTB-priority mechanism*

based on g and $\succ$ operates on the restricted problem $(R|_{O_+(q) \cup \{\emptyset\}}, q)$ and matches sequentially all agent-type pairs that have the highest priority according to g subject to individual rationality and the quotas. Let $M^{(g, \succ)}(R, q)$ denote the resulting allocation and let $M^{(g, \succ)}$ denote the priority mechanism based on g and $\succ$.

3.2 Linear Programming Mechanisms

A linear programming mechanism assigns to each possible (a, b) -match a positive weight. These weights are strictly decreasing in both components. A *weighting function* is a function $h : (\mathbb{N} \times \mathbb{N}) \cup \{(0, 0)\} \rightarrow \mathbb{R}_+$ such that

- (i) $h(0, 0) = 0$,
- (ii) for all $(a, b), (a', b) \in \mathbb{N} \times \mathbb{N}$, if $a < a'$, then $h(a, b) > h(a', b)$, and
- (iii) for all $(a, b), (a, b') \in \mathbb{N} \times \mathbb{N}$, if $b < b'$, then $h(a, b) > h(a, b')$.

Recall that $\text{rank}(x, R_i)$ denotes the rank of real type $x \in O$ in R_i and $\text{rank}(i, \succ_x)$ the rank of $i \in N$ in $\succ_x$. For all $i \in N$, we now also define $\text{rank}(\emptyset, R_i) = 0$ and $\text{rank}(i, \succ_\emptyset) = 0$. Given a profile $R \in \mathcal{R}$ and an individually rational allocation a , the score of a is the sum of the weights of the matched agent-type pairs, i.e.,

$$s^{(h, \succ)}(a, R) = \sum_{i \in N} h(\text{rank}(a_i, R_i), \text{rank}(i, \succ_{a_i})).$$

Given a weighting function h , the linear programming mechanism $M^{(h, \succ)}$ chooses for each problem (R, q) an individually rational allocation for q with maximal score among all individually rational allocations for q . Let $\arg \max s^{(h, \succ)}(\cdot, R)$ denote the set of all individually rational allocations for q that maximize the weighted score $s^{(h, \succ)}$ among all individually rational allocations for q . We will assume that if there are several allocations maximizing the score under a submitted profile, then the linear programming mechanism breaks ties uniformly, i.e., each of these maximizers is chosen with equal probability. In such a case a linear programming mechanism chooses a lottery over allocations. It is easy to see that this does not create any significant problem and our analysis can be extended to mechanisms choosing a lottery for each profile.

Let M denote a *random* mechanism choosing for each problem (R, q) a lottery over $\mathcal{A}_q$. Given an allocation a for q , let $\Pr\{M(R, q) = a\}$ denote the probability that M assigns to allocation a . Similarly, given $i \in N$ and $y \in O \cup \{\emptyset\}$, let $\Pr\{M_i(R, q) = y\}$ denote the probability that i is assigned to y under the lottery $M(R, q)$.

For each problem (R, q) , we will only consider the ranks of available real types under q . Thus, the weighting function will take as inputs the rank information based on the reduced preference

profile $R|_{O_+(q) \cup \{\emptyset\}}$. Then, for each problem (R, q) and any MTB $\succ$ of $\widehat{\succeq}$, the *MTB-linear programming mechanism based on h and $\succ$* operates on the restricted problem $(R|_{O_+(q) \cup \{\emptyset\}}, q)$ and for all individually rational allocations a for q such that

$$s^{(h, \succ)}(a, R|_{O_+(q) \cup \{\emptyset\}}) = \max_{a' \in \mathcal{A}_q \text{ is individually rational}} s^{(h, \succ)}(a', R|_{O_+(q) \cup \{\emptyset\}}), \text{ we have}$$

$$\Pr\{M^{(h, \succ)}(R, q) = a\} = \frac{1}{|\arg \max s^{(h, \succ)}(\cdot, R|_{O_+(q) \cup \{\emptyset\}})|} \text{ (where } |S| \text{ denotes the cardinality of set } S).$$

3.3 Stability and Agent-Proposing Responsive DA-Mechanisms

Given the priority structure $\widehat{\succeq}$ and a problem (R, q) , we can interpret $(R, \widehat{\succeq}, q)$ as a college admissions problem (Gale and Shapley, 1962; Roth and Sotomayor, 1990) where the set of agents N corresponds to the set of students, the set of real types O corresponds to the set of colleges, the capacity vector q describes colleges' quota, preferences R correspond to students' preferences over colleges and not receiving any college, and the priority structure $\widehat{\succeq}$ is taken to represent colleges' responsive preferences over students. Stability is an important requirement for many real-life matching markets and it turns out to be essential in our context of allocating indivisible objects to agents.

Stability under $\widehat{\succeq}$: Given $(R, q) \in \mathcal{R}^N \times \mathcal{Q}$, an allocation $a \in \mathcal{A}_q$ is *stable under $\widehat{\succeq}$* if there exists no agent-type pair $(i, x) \in N \times O \cup \{\emptyset\}$ such that $x P_i a_i$ and (s1) $|\{j \in N : a_j = x\}| < q_x$ or (s2) there exists $k \in N$ such that $a_k = x$ and $i \widehat{\succ}_x k$.

Note that (s1) implies *individual rationality* because $q_\emptyset = \infty$. Furthermore, *mechanism φ is stable* if there exists a priority structure $\succeq$ such that for each problem $(R, q) \in \mathcal{R}^N \times \mathcal{Q}$, $\varphi(R, q)$ is stable under $\succeq$. When $\succeq$ is strict, we suppose that colleges' priorities over sets of students are responsive with respect to $\succeq$ (we discuss responsiveness in detail below in Remark 1).

For any college admissions problem with strict and responsive priorities $(R, \succ, q)$, we denote by $DA^\succ(R, q)$ the *agent-optimal stable allocation* that is obtained by using Gale and Shapley's (1962) *agent-proposing deferred-acceptance algorithm*, DA-algorithm for short:⁷ let $(R, q, \succ)$ be given. Then,

- at the first step of the DA-algorithm, every agent applies to his favorite acceptable type. For each real type x , the q_x applicants who have the highest priority for x (all applicants if there are fewer than q_x) are placed on the waiting list of x , and all others are rejected.
- At the r -th step of the DA-algorithm, those applicants who were rejected at step $r - 1$ apply to their next best acceptable type. For each real type x , the q_x applicants among the new

⁷Analogously we may define the type-proposing deferred-acceptance algorithm where each type applies at each step to as many agents as there are objects of this type available.

applicants and those on the waiting list who have the highest priority for x are placed on the updated waiting list of x , and all others are rejected.

The DA-algorithm terminates when every agent is on a waiting list or has applied to all acceptable types. Once the algorithm ends, objects are assigned to the agents on the respective type waiting lists and the null object is assigned to the agents who are not on any waiting list. The resulting allocation is the agent-optimal stable allocation for the problem $(R, q, \succ)$, denoted by $DA^\succ(R, q)$.

Responsive DA-Mechanisms: A mechanism φ is a *responsive DA-mechanism* if there exists a strict priority structure $\succ$ such that for each $(R, q) \in \mathcal{R}^N \times \mathcal{Q}$, $\varphi(R, q) = DA^\succ(R, q)$.

In real life we often take a MTB $\hat{\succ}$ of $\hat{\preceq}$ before applying the DA-mechanism $DA^\succ$. We will call such a mechanism a *responsive MTB-DA-mechanism*.

Note that the above condition of stability under $\succ$ takes only blocking by individual agents and by agent-type pairs into account. This is sometimes referred to as “pairwise” stability. However, one may also consider group stability where blocking is allowed by arbitrary groups of agents and types. For college admissions problems with responsive preferences, pairwise stability and group stability coincide. In such environments it suffices to know the priority orderings over individual agents and the implementation of responsive DA-mechanisms is much easier than it would be for more general college preferences (e.g., substitutable preferences). From now on, we will assume that priorities are *responsive*, i.e., we assume that college preferences are responsive in the associated college admissions problem.⁸

Remark 1. Responsiveness of Priorities

Let $x \in O$, $0 < q_x \leq |N|$, and $\succ_x$ be a priority ordering. Let $2_{q_x}^N$ denote the set of all subsets of N that do not exceed the capacity q_x , i.e., $2_{q_x}^N \equiv \{S \subseteq N : |S| \leq q_x\}$. Let P_x denote a *priority relation* on $2_{q_x}^N$, i.e., P_x strictly orders all sets in $2_{q_x}^N$. Then, P_x is responsive to $\succ_x$ if the following two conditions hold: (r1) for all $S \in 2_{q_x}^N$ such that $|S| < q_x$ and all $i \in N \setminus S$, $S \cup \{i\} P_x S$ and (r2) for all $S \in 2_{q_x}^N$ such that $|S| < q_x$ and all $i, j \in N \setminus S$, $(S \cup \{i\}) P_x (S \cup \{j\})$ if and only if $i \succ_x j$. When formulating (r1) we implicitly assume that each object finds all agents acceptable. Let $\mathcal{P}(\succ_x)$ denote the set of all priority relations that are responsive to $\succ_x$.

Now, given $(R, q) \in \mathcal{R}^N \times \mathcal{Q}$ and a priority relation profile $P_O \in \times_{x \in O} \mathcal{P}(\succ_x)$, an allocation $a \in \mathcal{A}_q$ is *group stable under P_O* if there exists no coalition (consisting possibly of both agents and types) that blocks allocation a .⁹

⁸Correspondingly, priorities would be *substitutable* if college preferences are substitutable in the associated college admissions problem.

⁹Formally, given (R, q) and P_O , a coalition $S \subseteq N \cup O$ blocks $a \in \mathcal{A}_q$ if there exists an allocation $b \in \mathcal{A}_q$ such that (g1) for all $i \in S \cap N$, $b_i \in S \cap O$, (g2) for all $i \in S \cap N$, $b_i R_i a_i$, (g3) for all $x \in S \cap O$, $i \in \{j \in N : b_j = x\}$ implies $i \in S \cup \{j \in N : a_j = x\}$, (g4) for all $x \in S \cap O$, $\{j \in N : b_j = x\} R_x \{j \in N : a_j = x\}$, and (g5) for at least one member of S , (g2) or (g4) holds with strict preference.

It is known that given $(R, q) \in \mathcal{R}^N \times \mathcal{Q}$ and a responsive priority relation profile $P_O \in \times_{x \in O} \mathcal{P}(\succ_x)$, an allocation $a \in \mathcal{A}_q$ is group stable under P_O if and only if $a \in \mathcal{A}_q$ is stable under $\succ$. In other words, group stability is identical with stability for responsive priorities and the set of group stable matchings is invariant with respect to the responsive preference extensions of $\succ$. This implies that for the implementation of any responsive DA-mechanism, we only need to know the priority orderings over individual agents and not the complete priority relations over sets of agents. This makes the application of responsive DA-mechanisms very easy and convenient in real-life matching markets. $\diamond$

4 Simple and Intuitive Properties

A natural requirement for a mechanism is that the chosen allocation depends only on preferences over the set of available types. Given a capacity vector q , a type x is *unavailable* if $q_x = 0$.

Unavailable Type Invariance: For all $(R, q) \in \mathcal{R}^N \times \mathcal{Q}$ and all $R' \in \mathcal{R}^N$ such that $R|_{O_+(q) \cup \{\emptyset\}} = R'|_{O_+(q) \cup \{\emptyset\}}$, $\varphi(R, q) = \varphi(R', q)$.

By *individual rationality* each agent should weakly prefer his allotment to the null object (which may represent an outside option such as off-campus housing in the context of university housing allocation, or private schools or home schooling in the context of student placement in public schools).

Individual Rationality: For all $(R, q) \in \mathcal{R}^N \times \mathcal{Q}$ and all $i \in N$, $\varphi_i(R, q) R_i \emptyset$.

Next, we introduce two properties that require a mechanism to not waste any resources. First, *non-wastefulness* (Balinski and Sönmez, 1999) requires that no agent prefers an available object that is not assigned to his allotment. Note that *non-wastefulness* corresponds to the first part (s1) of stability (for any arbitrary priority structure).

Non-Wastefulness: For all $(R, q) \in \mathcal{R}^N \times \mathcal{Q}$, all $x \in O_+(q)$, and all $i \in N$, if $x P_i \varphi_i(R, q)$, then $|\{j \in N : \varphi_j(R, q) = x\}| = q_x$.

Non-wastefulness is *not* a weak efficiency requirement in environments where priorities can be interpreted as preferences because types (or colleges or schools) may be worse off after the departure of some agents. Or in other words, each type may prefer to have assigned as many objects of this type as possible (not caring of the identities of the students receiving this type).¹⁰ Next, we weaken *non-wastefulness* by requiring that no agent receives the null object while he prefers an available object that is not assigned.

Weak Non-Wastefulness: For all $(R, q) \in \mathcal{R}^N \times \mathcal{Q}$, all $x \in O_+(q)$, and all $i \in N$, if $x P_i \varphi_i(R, q)$ and $\varphi_i(R, q) = \emptyset$, then $|\{j \in N : \varphi_j(R, q) = x\}| = q_x$.

¹⁰This may be due to avoiding the closure of a school or a college and if more positions are filled, then this may increase the prestige or visibility of the institution.

Weak non-wastefulness is a limited efficiency requirement. For example, suppose that a central agency registers all agents who did not receive anything (or are unemployed) and all those agents report all real types (or jobs) which are acceptable to them. Then it should not be the case that some agent who receives nothing prefers one of the available real objects (or available jobs) to the null object.

Many mechanisms that are used in real life ignore agents' preferences below their allotments (e.g., all mechanisms mentioned in Section 3). That is, an allocation does not change if an agent changes his reported preferences below his allotment. We formulate a weaker version of this invariance property by restricting agents' changes below their allotments to truncations.

A truncation strategy is a preference relation that ranks the real types in the same way as the corresponding original preference relation and each real type which is acceptable under the truncation strategy is also acceptable under the original preference relation. Formally, given $i \in N$ and $R_i \in \mathcal{R}$, a strategy $\bar{R}_i \in \mathcal{R}$ is a *truncation (strategy) of R_i* if (t1) $\bar{R}_i|_O = R_i|_O$ and (t2) $A(\bar{R}_i) \subseteq A(R_i)$. Loosely speaking, a truncation strategy of R_i is obtained by moving the null object "up."

If an agent truncates his preference relation in a way such that his allotment remains acceptable under the truncated preference relation, then *truncation invariance* requires that the allocation is the same under both profiles. The property is quite natural on its own in the sense that the chosen allocations do not depend on where any agent, who receives a real object, ranks the null object below his allotment.

Truncation Invariance: For all $(R, q) \in \mathcal{R}^N \times \mathcal{Q}$, all $i \in N$, and all $\bar{R}_i \in \mathcal{R}_i$, if $\bar{R}_i$ is a truncation of R_i and $\varphi_i(R, q)$ is acceptable under $\bar{R}_i$ (i.e., $\varphi_i(R, q) \in A(\bar{R}_i)$) and $\bar{R} = (\bar{R}_i, R_{-i})$, then $\varphi(\bar{R}, q) = \varphi(R, q)$.

Our last property applies to problems when two agents compete for the same object in a maximal conflict situation, i.e., they have the same preference relation with only one acceptable real type $x \in O$ and for the two problems under consideration the preference profiles are identical. *Two-agent consistent conflict resolution* then requires that if in these problems one of them receives the object, the conflict is resolved consistently in that it has to be the same agent in both problems who "wins the conflict" and receives the object.

We denote a preference relation with only one acceptable type $x \in O$ by R^x , i.e., $A(R^x) = \{x\}$. We denote the set of all preference relations that have $x \in O$ as the unique acceptable type by $\mathcal{R}^x$.

Two-Agent Consistent Conflict Resolution: For all $R \in \mathcal{R}^N$, all $q, q' \in \mathcal{Q}$, and all $R^x \in \mathcal{R}^x$, if $R_i = R_j = R^x$ and $\{\varphi_i(R, q), \varphi_j(R, q)\} = \{\varphi_i(R, q'), \varphi_j(R, q')\} = \{x, \emptyset\}$, then for $k \in \{i, j\}$, $\varphi_k(R, q) = \varphi_k(R, q')$.

One can interpret this property as a weak tie-breaking property because it only imposes that the tie between two agents in very specific maximal conflict situations is broken consistently in favor

of the same agent. In particular, when considering two different preference profiles $R, R' \in \mathcal{R}^N$ where agents i and j have a maximal conflict for type x (i.e., $R_i = R_j = R'_i = R'_j \in \mathcal{R}^x$ and $\{\varphi_i(R, q), \varphi_j(R, q)\} = \{\varphi_i(R', q'), \varphi_j(R', q')\} = \{x, \emptyset\}$), *two-agent consistent conflict resolution* does not imply that the same agent (out of i and j) receives object x because the assignment of x may depend on the other agents' preferences (i.e., for $k \in \{i, j\}$, $\varphi_k(R, q) \neq \varphi_k(R', q')$ is possible).

5 Strategy-Proofness makes the Difference

The properties we have introduced up to now are simple and intuitive. Indeed, our first result confirms this since all the real-life mechanisms we have described so far satisfy them.

Proposition 1. *Let $\succ$ be a strict priority structure.*

- (i) *For any priority function g , the priority mechanism based on $(g, \succ)$ satisfies unavailable type invariance, individual rationality, weak non-wastefulness, two-agent consistent conflict resolution, and truncation invariance.*
- (ii) *For any weighting function h , the linear programming mechanism based on $(h, \succ)$ satisfies unavailable type invariance, individual rationality, weak non-wastefulness, two-agent consistent conflict resolution, and truncation invariance.*
- (iii) *The responsive DA-mechanism based on $\succ$ satisfies unavailable type invariance, individual rationality, weak non-wastefulness, two-agent consistent conflict resolution, and truncation invariance.*

Note that all the properties introduced in Section 4 and used in Proposition 1 ignore the priority structure $\hat{\succ}$. We prove Proposition 1 in Appendix A.

Next we characterize the class of all responsive DA-mechanisms (and not only MTB-DA-mechanisms). For our characterization of responsive DA-mechanisms we need one more property.

The well-known non-manipulability property *strategy-proofness* requires that no agent can ever benefit from misrepresenting his preferences.

Strategy-Proofness: For all $(R, q) \in \mathcal{R}^N \times \mathcal{Q}$, all $i \in N$, and all $\bar{R}_i \in \mathcal{R}$, $\varphi_i(R, q) R_i \varphi_i((\bar{R}_i, R_{-i}), q)$.

It turns out that the class of mechanisms satisfying our simple and intuitive properties *and strategy-proofness* is very small. With the following characterization, we establish the first normative foundation for the class of responsive DA-mechanisms.

Theorem 1. *Responsive DA-mechanisms are the only mechanisms satisfying unavailable type invariance, individual rationality, weak non-wastefulness, two-agent consistent conflict resolution, truncation invariance, and strategy-proofness.*

It is important to note that priority mechanisms and linear programming mechanisms satisfy all properties of Theorem 1 except for *strategy-proofness*. This demonstrates how weak all properties in Theorem 1 are, except for *strategy-proofness*. It is the incentive compatibility property that favors responsive DA-mechanisms over all other mechanisms. We prove Theorem 1 in Appendix B.¹¹ Appendix C contains a second characterization (Theorem 4) of responsive DA-mechanisms in a variable agents environment based on *individual rationality*, *weak non-wastefulness*, *weak consistency*, and *strategy-proofness*. Here *weak consistency* requires that agents receiving the null object should still receive the null object if some agents leave the economy with their allotments.

Remark 2. (*Weak*) *Non-Wastefulness*

Note that if a mechanism is *non-wasteful*, then (s1) in the definition of stability can never occur and therefore, *non-wastefulness* is a partial stability property. However, *weak non-wastefulness*, even together with *individual rationality*, does not imply (s1) and therefore, it is *not* a partial stability property. Any mechanism satisfying the properties of Theorem 1 must be *stable* (under a strict priority structure). The combination of the properties in Theorem 1 includes *weak non-wastefulness* as a partial efficiency requirement, but not as a partial stability property. However, in the above characterization (and in Theorem 4, Appendix C), *weak non-wastefulness* may be replaced by *non-wastefulness* (without shortening the proofs or altering the independence of the axioms). Note that similarly to responsive DA-mechanisms, priority mechanisms and linear programming mechanisms satisfy *weak non-wastefulness*, but in contrast to them, they *do not satisfy non-wastefulness*. $\diamond$

6 Stability makes a Difference as well

We have seen that among all the mechanisms discussed in this article, responsive DA-mechanisms are the only *strategy-proof* mechanisms. However, they are in fact also the only ones that respect the *stability* with respect to priorities. Another strategy-proof mechanism that has been discussed in school choice redesigns is the so-called top-trading cycles mechanism (Shapley and Scarf, 1974). While in most school choice redesigns (e.g., Boston and New York City) stable DA-mechanisms are implemented, only San Francisco so far has announced plans to adopt an unstable top-trading cycles based mechanism. In many applications, stability seems to be an extremely important criterium as well and it differentiates DA-mechanisms from top-trading mechanisms.

As already mentioned in our introduction, many real life applications of market design, e.g., school choice problems, have weak priority structures as inputs (see Abdulkadiroğlu, Pathak, and Roth, 2009; Erdil and Ergin, 2008). Given a weak priority structure $\widehat{\succeq}$, stability under $\widehat{\succeq}$ immediately implies *non-wastefulness* and *individual rationality*. Thus, we obtain the following first normative foundation for responsive MTB-DA-mechanisms.

¹¹We discuss the independence of properties in Theorem 1 in the Supplementary Appendix.

Theorem 2. *Let $\hat{\succsim}$ be a (weak) priority structure. Responsive MTB-DA-mechanisms are the only mechanisms satisfying stability under $\hat{\succsim}$, unavailable type invariance, two-agent consistent conflict resolution, truncation invariance, and strategy-proofness.*

A special case of a weak priority structure is the one where all agents have equal priorities for all objects (or there are no priorities given), which we denoted by $\hat{\sim}$. Then any problem reduces to an allocation problem with no priorities. Furthermore, it is easy to see that stability under $\hat{\sim}$ is equivalent to non-wastefulness and individual rationality (and the properties in Theorem 2 are independent because the properties in Theorem 1 are independent).

When applying responsive DA-mechanisms to problems with weak priorities, ties have to be broken (in the special case of no priorities, all ties have to be broken). Essentially, whenever ties are present in school choice, we face trade-offs between efficiency, stability, and strategy proofness. Abdulkadiroğlu, Pathak, and Roth (2009) show that simulations with field data from the New York City high school match and the theory favor single tie-breaking rules, i.e., indifferences are broken the same way at every school. Formally, a MTB $\succ$ of $\hat{\succsim}$ is a *single tie-breaking (STB)* if for all $x, y \in O$ and all $i, j \in N$, if $i \hat{\sim}_x j$ and $i \hat{\sim}_y j$, then either $[i \succ_x j \text{ and } i \succ_y j]$ or $[j \succ_x i \text{ and } j \succ_y i]$.

Now, it is interesting to ask what exactly characterizes the subclass of responsive STB-DA-mechanisms. Clearly, *two-agent consistent conflict resolution* in Theorem 2 adds an element of decisiveness that, together with all other properties, helps to pins down the priority structure for the underlying DA-mechanism. In order to characterize responsive STB-DA-mechanisms, the priority ordering will have to be the same for all objects and therefore we will strengthen the *two-agent consistent conflict resolution* property by extending the decisiveness requirement to two (possibly distinct) objects.

Strong two-agent consistent conflict resolution applies to two problems when two agents compete for two (possibly distinct) objects in maximal conflict situations, i.e., in each problem they have the same preference relation such that in one problem they find only $x \in O$ acceptable and in another problem they find only $y \in O$ acceptable: it then requires that if in these problems one agent receives the acceptable object and the other receives the null object, the conflict is resolved consistently in that it has to be the same agent in both problems who “wins the conflict” and receives the acceptable object.

Strong Two-Agent Consistent Conflict Resolution: For all $R, R' \in \mathcal{R}^N$, all $q, q' \in \mathcal{Q}$, all $R^x \in \mathcal{R}^x$, and all $R^y \in \mathcal{R}^y$, if $R_i = R_j = R^x$, $i \hat{\sim}_x j$ and $\{\varphi_i(R, q), \varphi_j(R, q)\} = \{x, \emptyset\}$, and $R'_i = R'_j = R^y$, $i \hat{\sim}_y j$ and $\{\varphi_i(R', q'), \varphi_j(R', q')\} = \{y, \emptyset\}$, then for $k \in \{i, j\}$, $\varphi_k(R, q) = \emptyset \Leftrightarrow \varphi_k(R', q') = \emptyset$.

Now a straightforward consequence of Theorem 2 is the following normative foundation for responsive STB-DA-mechanisms.

Theorem 3. *Let $\hat{\succeq}$ be a (weak) priority structure. Responsive STB-DA-mechanisms are the only mechanisms satisfying stability under $\hat{\succeq}$, unavailable type invariance, strong two-agent consistent conflict resolution, truncation invariance, and strategy-proofness.*

7 Conclusion

Our model is that of assigning types (school seats) to a set of agents (students). Agents have strict preferences over types and whenever we consider agents' strict priorities for a certain type, we capture them by an ordering of all agents. Of course, since sets of agents are assigned to an type, in general priorities may depend on the whole set of agents and not only on the ordering of individual agents. However, if priorities over sets of agents are responsive with respect to the strict priority ordering over individual agents, then in determining stable allocations we only need to know the priority orderings over individual agents. It is exactly this responsiveness-to-individual-priorities feature that makes the agent-proposing deferred-acceptance mechanism easily applicable in practice (Roth, 2008).

Kojima and Manea (2010) provide two characterizations of deferred acceptance mechanisms with so-called acceptant substitutable priorities (a larger class of mechanisms than the class of responsive DA-mechanisms which is based on priorities that are determined by a choice function that reflects substitutability in priorities over sets of agents; see also Hatfield and Milgrom, 2005). For this class of DA-mechanisms, priority orderings over individual agents do not suffice—the priorities over sets of agents must be known. Furthermore, their characterizations use two new monotonicity properties which are violated by priority mechanisms and linear programming mechanisms. Indeed, it can be seen that those classes of mechanisms violate all their properties in their first characterization and two of the three properties in their second characterization. While the class of DA-mechanisms with acceptant substitutable priorities is sometimes employed in real-life markets, many applications involving DA-mechanisms are based on responsive priorities or slight generalizations thereof. Thus, providing a rationale for this class of mechanisms is important. We characterized responsive DA-mechanisms using very simple and intuitive properties plus strategy-proofness.

References

- Abdulkadiroğlu, A., Pathak, P., and A.E. Roth (2009): Strategy-Proofness versus Efficiency in Matching with Indifferences: Redesigning the NYC High School Match, *American Economic Review* 99, 1954–1978.
- Abdulkadiroğlu, A., Pathak, P., Roth, A.E., and T. Sönmez (2006): Changing the Boston School Choice Mechanism: Strategy-Proofness as Equal Access, Working Paper, Harvard University.

- Abdulkadiroğlu, A., T. Sönmez (2003): School Choice: a Mechanism Design Approach, *American Economic Review* 93, 729–747.
- Balinski, M., and T. Sönmez (1999): A Tale of Two Mechanisms: Student Placement, *Journal of Economic Theory* 84, 73–94.
- Dubins, L.E., and D.A. Freedman (1981): Machiavelli and the Gale-Shapley Algorithm, *American Mathematical Monthly* 88, 485–494.
- Ehlers, L. (2002): Coalitional Strategy-Proof House Allocation, *Journal of Economic Theory* 105, 298–317.
- Ehlers, L. (2008): Truncation Strategies in Matching Markets, *Mathematics of Operations Research* 33, 327–335.
- Ehlers, L., and A. Erdil (2010): Efficient Assignment Respecting Priorities, *Journal of Economic Theory* 145, 1269–1282.
- Ehlers, L., and B. Klaus (2006): Efficient Priority Rules, *Games and Economic Behavior* 55, 372–384.
- Ehlers, L., and B. Klaus (2007): Consistent House Allocation, *Economic Theory* 30, 561–574.
- Ehlers, L., and B. Klaus (2009): Allocation via Deferred Acceptance, CIREC Cahier 17-2009.
- Erdil, A., and H. Ergin (2008): What’s the Matter with Tie-Breaking? Improving Efficiency in School Choice, *American Economic Review* 98, 669–689.
- Ergin, H.İ. (2002): Efficient Resource Allocation on the Basis of Priorities, *Econometrica* 70, 2489–2497.
- Ergin, H.İ., and T. Sönmez (2006): Games of School Choice under the Boston Mechanism, *Journal of Public Economics* 90, 215–237.
- Gale, D., and L. Shapley (1962): College Admissions and the Stability of Marriage, *American Mathematical Monthly* 69, 9–15.
- Hatfield, J.W., and P. Milgrom (2005): Matching with Contracts, *American Economic Review* 95, 913–935.
- Kesten, O. (2009): Coalitional Strategy-Proofness and Resource Monotonicity for House Allocation Problems, *International Journal of Game Theory* 38, 17–21.
- Kojima, F., and M. Manea (2010): Axioms for Deferred Acceptance, *Econometrica* 78, 633–653.
- Pápai, S. (2000): Strategyproof Assignment by Hierarchical Exchange, *Econometrica* 68, 1403–1433.
- Pathak, P., and T. Sönmez (2008): Leveling the Playing Field: Sincere and Sophisticated Players in the Boston Mechanism, *American Economic Review* 98, 1636–52.
- Roth, A.E. (1982): The Economics of Matching: Stability and Incentives, *Mathematics of Operations Research* 7, 617–628.
- Roth, A.E. (1984a): The Evolution of the Labor Market for Medical Interns and Residents: A Case Study in Game Theory, *Journal of Political Economy* 92, 991–1016.

- Roth, A.E. (1984b): Misrepresentation and Stability in the Marriage Problem, *Journal of Economic Theory* 34, 383–387.
- Roth, A.E. (1991): A Natural Experiment in the Organization of Entry Level Labor Markets: Regional Markets for New Physicians and Surgeons in the U.K., *American Economic Review* 81, 415–440.
- Roth, A.E. (2008): Deferred Acceptance Algorithm: History, Theory, Practice, and Open Questions, *International Journal of Game Theory* 36, 537–569.
- Roth, A.E., and M.A.O. Sotomayor (1990): Two-Sided Matching: A Study in Game-Theoretic Modeling and Analysis. Econometric Society Monograph Series. New York: Cambridge University Press.
- Shapley, L., and H. Scarf (1974): On Cores and Indivisibility, *Journal of Mathematical Economics* 1, 23–37.
- Sönmez, T., and M.U. Ünver (2011): Matching, Allocation, and Exchange of Discrete Resources, Jess Benhabib, Alberto Bisin, and Matthew Jackson (eds.) *Handbook of Social Economics*, Elsevier.
- Thomson, W. (2009): Consistent Allocation Rules, book manuscript.

Appendix

A Proof of Proposition 1

Let $\succ$ be a strict priority structure.

In showing (i), let g be a priority function and consider the priority mechanism $M^{(g, \succ)}$ based on g and $\succ$. By definition, $M^{(g, \succ)}$ satisfies *unavailable type invariance* and *individual rationality*.

For *weak non-wastefulness*, suppose that for (R, q) , $x \in O_+(q)$ and $i \in N$ we have $x \in P_i M_i^{(g, \succ)}(R, q)$ and $M_i^{(g, \succ)}(R, q) = \emptyset$. If $|\{j \in N : M_j^{(g, \succ)}(R, q) = x\}| < q_x$, then (i, x) is an individually rational (a, b) -match and the priority mechanism $M^{(g, \succ)}$ could not have stopped without matching i and x .

For *two-agent consistent conflict resolution*, let (R, q) and (R, q') be such that $R_i = R_j = R^x$ and $\{M_i^{(g, \succ)}(R, q), M_j^{(g, \succ)}(R, q)\} = \{M_i^{(g, \succ)}(R, q'), M_j^{(g, \succ)}(R, q')\} = \{x, \emptyset\}$. Since $\text{rank}(x, R_i) = 1 = \text{rank}(x, R_j)$, (i, x) is an $(1, b)$ -match and (j, x) is an $(1, b')$ -match where $b = \text{rank}(i, \succ_x)$ and $b' = \text{rank}(j, \succ_x)$. Because $\succ_x$ is strict, $b \neq b'$ and we obtain from $\{M_i^{(g, \succ)}(R, q), M_j^{(g, \succ)}(R, q)\} = \{M_i^{(g, \succ)}(R, q'), M_j^{(g, \succ)}(R, q')\} = \{x, \emptyset\}$ both $M_i^{(g, \succ)}(R, q) = M_i^{(g, \succ)}(R, q')$ and $M_j^{(g, \succ)}(R, q) = M_j^{(g, \succ)}(R, q')$, the desired conclusion.

For *truncation invariance*, let (R, q) , $i \in N$, and $\bar{R}_i$ be such that $M_i^{(g, \succ)}(R, q) \in A(\bar{R}_i)$ and $\bar{R}_i$ is a truncation of R_i . Note that this implies $M_i^{(g, \succ)}(R, q) \neq \emptyset$. Hence, under (R, q) and

$(\bar{R}_i, R_{-i}, q)$ the priority mechanism $M^{(g, \succ)}$ matches identically all agent-type pairs and we obtain $M^{(g, \succ)}(R, q) = M^{(g, \succ)}(\bar{R}_i, R_{-i}, q)$.

In showing (ii), note that all properties can be easily formulated for random mechanisms (for instance by putting positive probability only on *individually rational* allocations). Let h be a weighting function and consider the linear programming mechanism $M^{(h, \succ)}$ based on h and $\succ$. By definition, $M^{(h, \succ)}$ satisfies *unavailable type invariance* and *individual rationality*.

For *weak non-wastefulness*, suppose that for (R, q) , $x \in O_+(q)$ and $i \in N$ we have $x P_i M_i^{(h, \succ)}(R, q)$ and $M_i^{(h, \succ)}(R, q) = \emptyset$. If $|\{j \in N : M_j^{(h, \succ)}(R, q) = x\}| < q_x$, then (i, x) is some individually rational (a, b) -match with $h(a, b) > 0$. Now $M^{(h, \succ)}(q, R)$ could not have maximized the score among all individually rational allocations.

For *two-agent consistent conflict resolution*, let (R, q) and (R, q') be such that $R_i = R_j = R^x$ and $\{M_i^{(h, \succ)}(R, q), M_j^{(h, \succ)}(R, q)\} = \{M_i^{(h, \succ)}(R, q'), M_j^{(h, \succ)}(R, q')\} = \{x, \emptyset\}$. Since $\text{rank}(x, R_i) = 1 = \text{rank}(x, R_j)$, (i, x) is an $(1, b)$ -match and (j, x) is an $(1, b')$ -match where $b = \text{rank}(i, \succ_x)$ and $b' = \text{rank}(j, \succ_x)$. Because $\succ_x$ is strict, $b \neq b'$ and $h(1, b) \neq h(1, b')$. Thus, we obtain from $\{M_i^{(h, \succ)}(R, q), M_j^{(h, \succ)}(R, q)\} = \{M_i^{(h, \succ)}(R, q'), M_j^{(h, \succ)}(R, q')\} = \{x, \emptyset\}$ both $M_i^{(h, \succ)}(R, q) = M_i^{(h, \succ)}(R, q')$ and $M_j^{(h, \succ)}(R, q) = M_j^{(h, \succ)}(R, q')$, the desired conclusion.

For *truncation invariance*, let (R, q) , $i \in N$, and $\bar{R}_i$ be such that $M_i^{(h, \succ)}(R, q) \in A(\bar{R}_i)$ and $\bar{R}_i$ is a truncation of R_i . Note that this implies $M_i^{(h, \succ)}(R, q) \neq \emptyset$. Hence, under (R, q) and $(\bar{R}_i, R_{-i}, q)$ the allocation $M^{(h, \succ)}(R, q)$ maximizes the score among all individually rational allocations and we obtain $M^{(h, \succ)}(R, q) = M^{(h, \succ)}(\bar{R}_i, R_{-i}, q)$ (we refer to Ehlers (2008) for random mechanisms satisfying *truncation invariance*).

In showing (iii), consider the responsive DA-mechanism based on $\succ$. It is easy to see that $DA^\succ$ satisfies *unavailable type invariance*, *individual rationality*, *weak non-wastefulness*, *two-agent consistent conflict resolution*, and *truncation invariance*.

B Proof of Theorem 1

It is easy to verify that responsive DA-mechanisms satisfy the properties of Theorem 1. Conversely, let φ be a mechanism satisfying the properties of Theorem 1. First, we “calibrate/construct the priority structure using maximal conflict preference profiles.”

Let $x \in O$ and let $R^x \in \mathcal{R}^x$ (i.e., $A(R^x) = \{x\}$). Let $R^\emptyset \in \mathcal{R}$ be such that $A(R^\emptyset) = \emptyset$.

For any $S \subseteq N$, let $R_S^x = (R_i^x)_{i \in S}$ such that for all $i \in S$, $R_i^x = R^x$, and similarly $R_S^\emptyset = (R_i^\emptyset)_{i \in S}$ such that for all $i \in S$, $R_i^\emptyset = R^\emptyset$.

Let 1_x denote the capacity vector q such that $q_x = 1$ and for all $z \in O \setminus \{x\}$, $q_z = 0$. Similarly, for $y \in O \setminus \{x\}$ let 1_{xy} denote the capacity vector q such that $q_x = 1$, $q_y = 1$, and for all $z \in O \setminus \{x, y\}$, $q_z = 0$.

Consider the problem $(R_N^x, 1_x)$. By *weak non-wastefulness*, for some $j \in N$, $\varphi_j(R_N^x, 1_x) = x$, say $j = 1$. Then, for all $i \in N \setminus \{1\}$, we set $1 \succ_x i$.

Next consider the problem $((R_1^\emptyset, R_{-1}^x), 1_x)$. By *weak non-wastefulness* and *individual rationality*, for some $j \in N \setminus \{1\}$, $\varphi_j((R_1^\emptyset, R_{-1}^x), 1_x) = x$, say $j = 2$. Then, for all $i \in N \setminus \{1, 2\}$, we set $2 \succ_x i$.

By induction, we obtain $\succ_x$ for any type x and thus a priority structure $\succ = (\succ_x)_{x \in O}$.

Lemma 1. *For all $R \in \mathcal{R}^N$ and all $x \in O$, if for some $j \in N$, $\varphi_j(R, 1_x) = x$, then for all $i \in N \setminus \{j\}$, $x \in A(R_i)$ implies $j \succ_x i$.*

Proof. Let $R \in \mathcal{R}^N$ and $x \in O$. Without loss of generality, suppose $1 \succ_x 2 \succ_x \dots \succ_x n$. Let $S = \{i \in N : x \in A(R_i)\}$ and let $j = \min S$. We prove Lemma 1 by showing that $\varphi_j(R, 1_x) = x$. In the sequel, when using *two-agent consistent conflict resolution* we often also implicitly apply *weak non-wastefulness* and *individual rationality*.

Note that for all $i \in N \setminus S$, $\emptyset P_i x$. We partition the set $N \setminus S$ into the “lower” set $L = \{1, \dots, j-1\}$ (possibly $L = \emptyset$) and the “upper” set $U = N \setminus (L \cup S)$ (possibly $U = \emptyset$). Note that by *unavailable type invariance*, $\varphi(R, 1_x) = \varphi((R_L^\emptyset, R_S^x, R_U^\emptyset), 1_x)$.

By the construction of $\succ_x$, $\varphi_j((R_L^\emptyset, R_{S \cup U}^x), 1_x) = x$. Hence, if $U = \emptyset$, then $\varphi_j(R, 1_x) = x$ and for all $i \in N$, $x \in A(R_i)$ implies $j \succ_x i$.

Step 1: Let $k \in U$. We prove that $\varphi_j((R_L^\emptyset, R_{S \cup (U \setminus \{k\})}^x, R_k^\emptyset), 1_x) = \varphi_j((R_L^\emptyset, R_{S \cup U}^x), 1_x) = x$.

Let $y \in O \setminus \{x\}$. By *two-agent consistent conflict resolution*, $\varphi_j((R_L^\emptyset, R_{S \cup U}^x), 1_{xy}) = x$. Let $R'_k \in \mathcal{R}$ be such that $R'_k : y \emptyset x \dots$. By *strategy-proofness*, $\varphi_k((R_L^\emptyset, R_{S \cup (U \setminus \{k\})}^x, R'_k), 1_{xy}) \neq x$. If $S \cup (U \setminus \{k\}) = \{j\}$, then by *weak non-wastefulness* and *individual rationality*, $\varphi_j((R_L^\emptyset, R_{S \cup (U \setminus \{k\})}^x, R'_k), 1_{xy}) = x$. Otherwise, by *two-agent consistent conflict resolution*, $\varphi_j((R_L^\emptyset, R_{S \cup (U \setminus \{k\})}^x, R'_k), 1_{xy}) = x$. By *weak non-wastefulness* and *individual rationality*, $\varphi_k((R_L^\emptyset, R_{S \cup (U \setminus \{k\})}^x, R'_k), 1_{xy}) = y$.

Let $R''_k \in \mathcal{R}$ be such that $R''_k : y \emptyset x \dots$ and $R''_k|_O = R'_k|_O$. By *strategy-proofness*, $\varphi_k((R_L^\emptyset, R_{S \cup (U \setminus \{k\})}^x, R''_k), 1_{xy}) = y$. By *weak non-wastefulness* and *individual rationality*, for some $l \in S \cup (U \setminus \{k\})$, $\varphi_l((R_L^\emptyset, R_{S \cup (U \setminus \{k\})}^x, R''_k), 1_{xy}) = x$.

Now R''_k is a truncation of R'_k and both $y \in A(R''_k)$ and $\varphi_k((R_L^\emptyset, R_{S \cup (U \setminus \{k\})}^x, R'_k), 1_{xy}) = y$. Thus, $\varphi((R_L^\emptyset, R_{S \cup (U \setminus \{k\})}^x, R''_k), 1_{xy}) = \varphi((R_L^\emptyset, R_{S \cup (U \setminus \{k\})}^x, R'_k), 1_{xy})$. Hence, $j = l$ and $\varphi_j((R_L^\emptyset, R_{S \cup (U \setminus \{k\})}^x, R''_k), 1_{xy}) = x$.

By *individual rationality*, $\varphi_k((R_L^\emptyset, R_{S \cup (U \setminus \{k\})}^x, R''_k), 1_x) \neq x$. If $S \cup (U \setminus \{k\}) = \{j\}$, then by *weak non-wastefulness* and *individual rationality*, $\varphi_j((R_L^\emptyset, R_{S \cup (U \setminus \{k\})}^x, R''_k), 1_x) = x$. Otherwise, by *two-agent consistent conflict resolution*, $\varphi_j((R_L^\emptyset, R_{S \cup (U \setminus \{k\})}^x, R''_k), 1_x) = x$. Thus, by *unavailable type invariance*, $\varphi_j((R_L^\emptyset, R_{S \cup (U \setminus \{k\})}^x, R_k^\emptyset), 1_x) = \varphi_j((R_L^\emptyset, R_{S \cup (U \setminus \{k\})}^x, R''_k), 1_x) = x$.

Steps 2, ..., l: Let $U = \{k_1, \dots, k_l\}$. Then using the same arguments as above, it follows that $x = \varphi_j((R_L^\emptyset, R_{S \cup U}^x), 1_x) = \varphi_j((R_L^\emptyset, R_{S \cup (U \setminus \{k_1\})}^x, R_{k_1}^\emptyset), 1_x) = \varphi_j((R_L^\emptyset, R_{S \cup (U \setminus \{k_1, k_2\})}^x, R_{\{k_1, k_2\}}^\emptyset), 1_x) =$

$\dots = \varphi_j((R_L^\emptyset, R_{S \cup \{k_l\}}^x, R_{U \setminus \{k_l\}}^\emptyset), 1_x) = \varphi_j((R_L^\emptyset, R_S^x, R_U^\emptyset), 1_x) = \varphi_j(R, 1_x)$. Hence, we obtain the desired result that $\varphi_j(R, 1_x) = x$. $\square$

For any number $q_x \in \{0, 1, \dots, \bar{q}_x\}$, let $q_x \circ 1_x = (q_x, 0_{O \setminus \{x\}})$ denote the capacity vector with exactly q_x copies of object x . The following lemma is the extension of Lemma 1 to the general capacity case.

Lemma 2. *For all $R \in \mathcal{R}^N$ and all $q_x \in \{1, \dots, \bar{q}_x\}$, if $T = \{j \in N : \varphi_j(R, q_x \circ 1_x) = x\}$, then for all $j \in T$ and all $i \in N \setminus T$, $x \in A(R_i)$ implies $j \succ_x i$.*

Proof. Let $R \in \mathcal{R}^N$, $q_x \in \{1, \dots, \bar{q}_x\}$, $T = \{j \in N : \varphi_j(R, q_x \circ 1_x) = x\}$, and $S = \{i \in N : x \in A(R_i)\}$. By individual rationality, $T \subseteq S$. Let $j \in T$ and $i \in S \setminus T$. We prove Lemma 2 by showing that $j \succ_x i$. In the sequel, when using *two-agent consistent conflict resolution* we often also implicitly apply *weak non-wastefulness* and *individual rationality*.

If $q_x = 1$, then by Lemma 1, $j \succ_x i$.

Next, suppose that $q_x = 2$. Note that by *unavailable type invariance*, $\varphi(R, q_x \circ 1_x) = \varphi((R_S^x, R_{-S}^\emptyset), q_x \circ 1_x)$. Let $k \in S$ be such that $\varphi_k((R_S^x, R_{-S}^\emptyset), 1_x) = x$. By Lemma 1, $k \succ_x i$. By *two-agent consistent conflict resolution*, $\varphi_k((R_S^x, R_{-S}^\emptyset), q_x \circ 1_x) = x$. Hence, $k \in T$ and $k \succ_x i$.

By *weak non-wastefulness* and *individual rationality*, there exists $h \in T \setminus \{k\}$ such that $\varphi_h((R_S^x, R_{-S}^\emptyset), q_x \circ 1_x) = x$. Let $y \in O \setminus \{x\}$ and $R'_k \in \mathcal{R}$ be such that $R'_k : y \ x \ \emptyset \dots$. By *unavailable type invariance*, $\varphi((R_{S \setminus \{k\}}^x, R_{-S}^\emptyset, R'_k), q_x \circ 1_x) = \varphi((R_S^x, R_{-S}^\emptyset), q_x \circ 1_x)$. Hence, $\varphi_h((R_{S \setminus \{k\}}^x, R_{-S}^\emptyset, R'_k), q_x \circ 1_x) = x$.

Next, we replace one object of type x by one object of type y . By *weak non-wastefulness* and *individual rationality*, $\varphi_k((R_{S \setminus \{k\}}^x, R_{-S}^\emptyset, R'_k), 1_{xy}) = y$. By *two-agent consistent conflict resolution*, $\varphi_h((R_{S \setminus \{k\}}^x, R_{-S}^\emptyset, R'_k), 1_{xy}) = x$. By *truncation invariance*, $\varphi_h((R_{S \setminus \{k\}}^x, R_{-S}^\emptyset, R'_k), 1_{xy}) = x$. By *two-agent consistent conflict resolution*, $\varphi_h((R_{S \setminus \{k\}}^x, R_{-S}^\emptyset, R'_k), 1_x) = x$. Hence, by Lemma 1, $h \succ_x i$.

Now by induction on q_x the conclusion of Lemma 2 follows. $\square$

Lemma 3. *For all $(R, q) \in \mathcal{R}^N \times \mathcal{Q}$, $\varphi(R, q)$ is stable under $\succ$.*

Proof. Let $(R, q) \in \mathcal{R}^N \times \mathcal{Q}$ and assume that $\varphi(R, q)$ is not stable under $\succ$. Then, there exists an agent-object pair $(i, x) \in N \times O \cup \{\emptyset\}$ such that $x \bar{P}_i \varphi_i(R, q)$ and (s1) $|\{j \in N : \varphi_j(R, q) = x\}| < q_x$ or (s2) there exists $k \in N$ such that $\varphi_k(R, q) = x$ and $i \succ_x k$. By *individual rationality*, $x \neq \emptyset$.

Let $\bar{R} \in \mathcal{R}^N$ be such that (a) for all $j \in N$ such that $\varphi_j(R, q) \neq \emptyset$, $\bar{R}_j$ is a truncation of R_j such that there exists no $y \in O \setminus \{\varphi_j(R, q)\}$ with $\varphi_j(R, q) \bar{R}_j y \bar{R}_j \emptyset$ and (b) for all $j \in N$ such that $\varphi_j(R, q) = \emptyset$, $\bar{R}_j = R_j$. (By *individual rationality*, $\bar{R}_j$ in (a) is well-defined as truncation of R_j .) By *truncation invariance*, $\varphi(\bar{R}, q) = \varphi(R, q)$ and $(i, x) \in N \times O$ is such that $x \bar{P}_i \varphi_i(\bar{R}, q)$ and (s1') $|\{j \in N : \varphi_j(\bar{R}, q) = x\}| < q_x$ or (s2') there exists $k \in N$ such that $\varphi_k(\bar{R}, q) = x$ and $i \succ_x k$.

Let $S = \{j \in N : x \bar{P}_j \varphi_j(\bar{R}, q)\}$. Note that $i \in S$. Let $T = \{j \in N : \varphi_j(\bar{R}, q) = x\}$. By *unavailable type invariance*, $\varphi(\bar{R}, q_x \circ 1_x) = \varphi((\bar{R}_{S \cup T}^x, \bar{R}_{N \setminus (S \cup T)}^\emptyset), q_x \circ 1_x)$.

Step 1: Assume that for $j \in S$, $\varphi_j(\bar{R}, q_x \circ 1_x) = x$. Then, by *strategy-proofness* and *individual rationality*, $\varphi_j((\bar{R}_{-j}, R_j^x), q_x \circ 1_x) = x$ and $\varphi_j((\bar{R}_{-j}, R_j^x), q) = \emptyset$. By *weak non-wastefulness*, $|\{k \in N : \varphi_k((\bar{R}_{-j}, R_j^x), q) = x\}| = q_x$. By *individual rationality* there exists $l \in N$ such that $x \in A(\bar{R}_l)$, and $\varphi_l((\bar{R}_{-j}, R_j^x), q) = x$ and $\varphi_l((\bar{R}_{-j}, R_j^x), q_x \circ 1_x) = \emptyset$. By Lemma 2, $j \succ_x l$. Next, by *strategy-proofness*, $\varphi_l((\bar{R}_{-j, l}, R_j^x, R_l^x), q) = x$. By *unavailable type invariance*, $\varphi((\bar{R}_{-j, l}, R_j^x, R_l^x), q_x \circ 1_x) = \varphi((\bar{R}_{-j}, R_j^x), q_x \circ 1_x)$; in particular, $\varphi_j((\bar{R}_{-j, l}, R_j^x, R_l^x), q_x \circ 1_x) = x$ and $\varphi_l((\bar{R}_{-j, l}, R_j^x, R_l^x), q_x \circ 1_x) = \emptyset$. Thus, $\varphi_j((\bar{R}_{-j, l}, R_j^x, R_l^x), q) = \emptyset$ would contradict *two-agent consistent conflict resolution* because then $\{\varphi_j((\bar{R}_{-j, l}, R_j^x, R_l^x), q), \varphi_l((\bar{R}_{-j, l}, R_j^x, R_l^x), q)\} = \{\varphi_j((\bar{R}_{-j, l}, R_j^x, R_l^x), q_x \circ 1_x), \varphi_l((\bar{R}_{-j, l}, R_j^x, R_l^x), q_x \circ 1_x)\} = \{x, \emptyset\}$. Hence, $\varphi_j((\bar{R}_{-j, l}, R_j^x, R_l^x), q) = x$. Thus, there exists an agent $j_2 \in N \setminus \{j, l\}$ such that $\varphi_{j_2}((\bar{R}_{-j, l}, R_j^x, R_l^x), q) = \emptyset$ and $\varphi_{j_2}((\bar{R}_{-j, l}, R_j^x, R_l^x), q_x \circ 1_x) = x$. However, $\bar{R}_{j_2} = R_{j_2}^x$ would contradict *two-agent consistent conflict resolution* because then $\{\varphi_{j_2}((\bar{R}_{-j, l}, R_j^x, R_l^x), q), \varphi_l((\bar{R}_{-j, l}, R_j^x, R_l^x), q)\} = \{\varphi_{j_2}((\bar{R}_{-j, l}, R_j^x, R_l^x), q_x \circ 1_x), \varphi_l((\bar{R}_{-j, l}, R_j^x, R_l^x), q_x \circ 1_x)\} = \{x, \emptyset\}$. Hence, $\bar{R}_{j_2} \neq R_{j_2}^x$.

Steps 2, ...: Step 2 replicates Step 1 with the starting preference profile $(\bar{R}_{-j, l}, R_j^x, R_l^x)$ and with agent j_2 in the role of agent j . Throughout the step, we strictly increase the number of agents with preferences R^x (at least by agent j_2). Furthermore, the step ends with the existence of another agent j_3 with $\bar{R}_{j_3} \neq R_{j_3}^x$ with whom we proceed to Step 3. Since the number of agents is finite, we obtain a contradiction in finitely many steps.

Final Step: With the previous steps we have established that for all $j \in S$, $\varphi_j(\bar{R}, q_x \circ 1_x) = \emptyset$. Note that by construction of $\bar{R}$, for all $j \in N \setminus (S \cup T)$, $\emptyset \bar{P}_j x$. Hence, by *individual rationality*, for all $j \in N \setminus (S \cup T)$, $\varphi_j(\bar{R}, q_x \circ 1_x) = \emptyset$. Note that by *individual rationality*, for all $j \in T$, $x \in A(\bar{R}_j)$. Hence, by *weak non-wastefulness*, for all $j \in T$, $\varphi_j(\bar{R}, q_x \circ 1_x) = x$ and $|T| = q_x$. By Lemma 2, we have that (*) for all $j \in T$ and all $l \in S$, $j \succ_x l$.

Recall that $T = \{j \in N : \varphi_j(\bar{R}, q) = x\}$. Hence, (s1') cannot be the case. Thus, by (s2') there exists $k \in N$ such that $\varphi_k(\bar{R}, q) = x$ and $i \succ_x k$. Recall that $i \in S$ and that $\varphi_k(\bar{R}, q) = x$ implies that $k \in T$ so that $i \succ_x k$ contradicts (*). $\square$

So far we have established that for any mechanism φ that satisfies the properties of Theorem 1, there exists a priority ordering $\succ$ such that for any $(R, q) \in \mathcal{R}^N \times \mathcal{Q}$, $\varphi(R, q)$ is stable under $\succ$. Hence, in the terminology of two-sided matching, the mechanism φ picks a stable allocation for the many-to-one two-sided market where types have responsive preferences over sets of agents who consume the objects based on the priority structure $\succ$ and agents have strict preferences over objects based on preferences R (see Roth and Sotomayor, 1990, Chapter 5). For these markets it is well-known that the responsive DA-mechanism is the only *strategy-proof* stable matching mechanism.

For completeness, we provide a proof which uses some standard results from many-to-one two-sided markets with responsive preferences.

Lemma 4. *For all $(R, q) \in \mathcal{R}^N \times \mathcal{Q}$, $\varphi(R, q) = DA^\succ(R, q)$.*

Proof. Let $(R, q) \in \mathcal{R}^N \times \mathcal{Q}$ and assume that $\varphi(R, q) \neq DA^\succ(R, q)$. By Lemma 3, $\varphi(R, q)$ is stable under $\succ$. Thus, since $DA^\succ(R, q)$ is the agent-optimal stable matching (Roth and Sotomayor, 1990, Corollary 5.9), for all $i \in N$, $DA_i^\succ(R, q) R_i \varphi_i(R, q)$. Since $\varphi(R, q) \neq DA^\succ(R, q)$, there exists $j \in N$ such that $DA_j^\succ(R, q) P_j \varphi_j(R, q)$. By individual rationality, $\varphi_j(R, q) R_j \emptyset$. Thus, $DA_j^\succ(R, q) \neq \emptyset$.

Let $\bar{R}_j$ be such that $\bar{R}_j|_O = R_j|_O$ and there is no $x \in O \setminus \{DA_j^\succ(R, q)\}$ such that $DA_j^\succ(R, q) \bar{R}_j x \bar{R}_j \emptyset$. Since $\bar{R}_j$ is a truncation of R_j and $DA^\succ$ satisfies *truncation invariance*, $DA^\succ((\bar{R}_j, R_{-j}), q) = DA^\succ(R, q)$. Thus, $DA_j^\succ((\bar{R}_j, R_{-j}), q) \neq \emptyset$.

By Lemma 3, $\varphi_j((\bar{R}_j, R_{-j}), q)$ is stable under $\succ$. By Roth and Sotomayor (1990, Corollary 5.9), $DA_j^\succ((\bar{R}_j, R_{-j}), q) R_j \varphi_j((\bar{R}_j, R_{-j}), q)$. By $DA_j^\succ((\bar{R}_j, R_{-j}), q) \neq \emptyset$ and the fact that the set of agents receiving the null object is identical for any two stable matchings (Roth and Sotomayor, 1990, Theorem 5.12), we have $\varphi_j((\bar{R}_j, R_{-j}), q) \neq \emptyset$. Now, by the definition of $\bar{R}_j$, $\varphi_j((\bar{R}_j, R_{-j}), q) = DA_j^\succ((\bar{R}_j, R_{-j}), q)$. Hence, $\varphi_j((\bar{R}_j, R_{-j}), q) = DA_j^\succ((\bar{R}_j, R_{-j}), q) = DA_j^\succ(R, q) P_j \varphi_j(R, q)$, which contradicts *strategy-proofness*. $\square$

C Variable Agents and Weak Consistency

Let N again denote the finite set of agents, but the set of agents who are present in a problem can vary. We define the set of all nonempty subsets of N by $\mathcal{N} \equiv \{M : M \subseteq N \text{ and } M \neq \emptyset\}$. An *(allocation) problem (with capacity constraints)* now consists of a set $N' \in \mathcal{N}$ of agents, a preference profile $R \in \mathcal{R}^{N'}$, and a capacity vector q . We denote the set of all problems by $\bigcup_{N' \in \mathcal{N}} \mathcal{R}^{N'} \times \mathcal{Q}$. We adjust our previous model, definitions, and properties by simply replacing the domain of problems $\mathcal{R}^N \times \mathcal{Q}$ by the variable population domain $\bigcup_{N' \in \mathcal{N}} \mathcal{R}^{N'} \times \mathcal{Q}$. Given $N' \in \mathcal{N}$ and $R \in \mathcal{R}^{N'}$, for any $M' \in \mathcal{N}$ such that $M' \subseteq N'$, let $R_{M'}$ denote the profile $(R_i)_{i \in M'}$.

A requirement for a mechanism that is very much in the spirit of *unavailable type invariance* is *unassigned type invariance*: the chosen allocation does not depend on the unconsumed or unassigned objects. Given a problem $(R, q) \in \bigcup_{N' \in \mathcal{N}} \mathcal{R}^{N'} \times \mathcal{Q}$, we define by $q(\varphi(R, q))$ the *capacity vector of assigned objects*: for all $x \in O$, $q_x(\varphi(R, q)) = |\{j \in N' : \varphi_j(R, q) = x\}|$.

Unassigned Type Invariance: For all $(R, q) \in \bigcup_{N' \in \mathcal{N}} \mathcal{R}^{N'} \times \mathcal{Q}$, $\varphi(R, q) = \varphi(R, q(\varphi(R, q)))$.

Consistency is one of the key properties in many frameworks with variable population scenarios. Thomson (2009) provides an extensive survey of consistency in various applications. Consistency requires that if some agents leave a problem with their allotments, then the mechanism should allocate the remaining objects among the agents who did not leave in the same way as in the original problem.

Consistency: For all $M', N' \in \mathcal{N}$ such that $M' \subseteq N'$, all $R \in \mathcal{R}^{N'}$, all $q \in \mathcal{Q}$, and all $i \in M'$, $\varphi_i(R, q) = \varphi_i(R_{M'}, \tilde{q})$ where $\tilde{q}_x = q_x - |\{j \in N' \setminus M' : \varphi_j(R, q) = x\}|$ for all $x \in O$.

It follows from Ergin (2002, Theorem 1)—for a fixed capacity vector q —that the only responsive DA-mechanisms satisfying *consistency* are the ones with an “acyclic” priority structure.¹²

The following property is a weak consistency property that all responsive DA-mechanisms satisfy. It requires that if some agents leave a problem with their allotments, then an agent who did not leave and who received the null object, still receives the null object. In other words, allocations only need to be consistent with respect to the agents who receive the null object.

Weak Consistency: For all $M', N' \in \mathcal{N}$ such that $M' \subseteq N'$, all $R \in \mathcal{R}^{N'}$, all $q \in \mathcal{Q}$, and all $i \in M'$, if $\varphi_i(R, q) = \emptyset$, then $\varphi_i(R_{M'}, \tilde{q}) = \emptyset$ where $\tilde{q}_x = q_x - |\{j \in N' \setminus M' : \varphi_j(R, q) = x\}|$ for all $x \in O$.

Theorem 4. *Responsive DA-mechanisms are the only mechanisms satisfying unassigned type invariance, individual rationality, weak non-wastefulness, weak consistency, and strategy-proofness.*

Below we adjust the definition of priority mechanisms and linear programming mechanisms to the variable population framework such that they satisfy all properties in Theorem 4 except for *strategy-proofness* (like Proposition 1 for Theorem 1).

For a strict priority structure $\succ$ and a priority function g , the priority mechanism based on $(g, \succ)$ matches for any problem $(R, q) \in \bigcup_{N' \in \mathcal{N}} \mathcal{R}^{N'} \times \mathcal{Q}$ all agent-type pairs that have the highest priority according to g subject to individual rationality and the quotas. Note that the ranking of unavailable objects is maintained in the agents’ preferences and the ranking of agents belonging to $N \setminus N'$ is maintained in $\succ$.

Similarly, for a strict priority structure $\succ$ and a weighting function h , the linear programming mechanism based on $(h, \succ)$ chooses for any problem $(R, q) \in \bigcup_{N' \in \mathcal{N}} \mathcal{R}^{N'} \times \mathcal{Q}$ an individually rational allocation with maximal score among all individually rational allocations. Note that here again the ranking of unavailable objects is maintained in the agents’ preferences and the ranking of agents belonging to $N \setminus N'$ is maintained in $\succ$.

Analogously to Theorem 2 we obtain a foundation of responsive MTB-DA-mechanisms using the previous characterization (Theorem 4).

Theorem 5. *Let $\hat{\succ}$ be a (weak) priority structure. Responsive MTB-DA-mechanisms are the only mechanisms satisfying stability under $\hat{\succ}$, unassigned type invariance, weak consistency, and strategy-proofness.*

¹²Ehlers and Erdil (2010) generalize this result from strict priorities to weak priorities.

Proof of Theorem 4

It is easy to verify that responsive DA-mechanisms satisfy the properties of Theorem 4. Conversely, let φ be a mechanism satisfying these properties. First, we “calibrate/construct the priority structure using maximal conflict preference profiles.”

Let $x \in O$ and let $R^x \in \mathcal{R}^x$ (i.e., $A(R^x) = \{x\}$). Let $R^\emptyset \in \mathcal{R}$ be such that $A(R^\emptyset) = \emptyset$.

For any $S \subseteq N$, let $R_S^x = (R_i^x)_{i \in S}$ such that for all $i \in S$, $R_i^x = R^x$, and similarly $R_S^\emptyset = (R_i^\emptyset)_{i \in S}$ such that for all $i \in S$, $R_i^\emptyset = R^\emptyset$.

Let 1_x denote the capacity vector q such that $q_x = 1$ and for all $z \in O \setminus \{x\}$, $q_z = 0$.

Consider the problem $(R_N^x, 1_x)$. By *weak non-wastefulness*, for some $j \in N$, $\varphi_j(R_N^x, 1_x) = x$, say $j = 1$. Then, for all $i \in N \setminus \{1\}$, we set $1 \succ_x i$.

Next consider the problem $(R_{N \setminus \{1\}}^x, 1_x)$. By *weak non-wastefulness* and *individual rationality*, for some $j \in N \setminus \{1\}$, $\varphi_j(R_{N \setminus \{1\}}^x, 1_x) = x$, say $j = 2$. Then, for all $i \in N \setminus \{1, 2\}$, we set $2 \succ_x i$.

By induction, we obtain $\succ_x$ for any type x and thus a priority structure $\succ = (\succ_x)_{x \in O}$.

Lemma 5. *For all $N' \in \mathcal{N}$, all $R' \in \mathcal{R}^{N'}$, and all $x \in O$, if for some $j \in N$, $\varphi_j(R', 1_x) = x$, then for all $i \in N' \setminus \{j\}$, $x \in A(R'_i)$ implies $j \succ_x i$.*

Proof. Let $N' \in \mathcal{N}$, $R' \in \mathcal{R}^{N'}$, and $x \in O$. Without loss of generality, suppose $1 \succ_x 2 \succ_x \dots \succ_x n$. Let $S = \{i \in N' : x \in A(R'_i)\}$ and let $j = \min S$. We prove Lemma 5 by showing that $\varphi_j(R', 1_x) = x$.

By *weak non-wastefulness* and *individual rationality*, for some $l \in S$, $\varphi_l(R', 1_x) = x$. Assume, for a contradiction, that $l \neq j$. Thus, $\varphi_j(R', 1_x) = \emptyset$. Hence, by *weak consistency*, $\varphi_j(R'_{\{j,l\}}, 1_x) = \emptyset$. By *weak non-wastefulness*, $\varphi_l(R'_{\{j,l\}}, 1_x) = x$. By *strategy-proofness*, $\varphi_l((R'_j, R_l^x), 1_x) = x$ and $\varphi_j((R_j^x, R_l^x), 1_x) = \emptyset$. By *weak non-wastefulness*, $\varphi_l((R_j^x, R_l^x), 1_x) = x$.

Let $L = \{1, \dots, j-1\}$ (possibly $L = \emptyset$). By the construction of $\succ_x$, $\varphi_j(R_{N \setminus L}^x, 1_x) = x$. Thus, $\varphi_l(R_{N \setminus L}^x, 1_x) = \emptyset$. Hence, by *weak consistency*, $\varphi_l(R_{\{j,l\}}^x, 1_x) = \emptyset$. By *weak non-wastefulness*, $\varphi_j(R_{\{j,l\}}^x, 1_x) = x$. Since $R_{\{j,l\}}^x = (R_j^x, R_l^x)$ and $l \neq j$, we have established a contradiction. Thus, $l = j$ and $\varphi_j(R', 1_x) = x$. $\square$

Let $\hat{\mathcal{R}} = \{R_i \in \mathcal{R} : |A(R_i)| \leq 1\}$.

Lemma 6.

- (a) *For all $N' \in \mathcal{N}$ and all $(R, q) \in \mathcal{R}^{N'} \times \mathcal{Q}$, if $|N'| = 2$, then $\varphi(R, q)$ is stable under $\succ$.*
- (b) *For all $N' \in \mathcal{N}$ and all $(R, q) \in \hat{\mathcal{R}}^{N'} \times \mathcal{Q}$, $\varphi(R, q)$ is stable under $\succ$.*

Proof. In order to show (a), suppose that $\varphi(R, q)$ is not stable under $\succ$. Then, there exists an agent-object pair $(i, x) \in N' \times O \cup \{\emptyset\}$ such that $x P_i \varphi_i(R, q)$ and (s1) $|\{j \in N' : \varphi_j(R, q) = x\}| < q_x$ or (s2) there exists $k \in N'$ such that $\varphi_k(R, q) = x$ and $i \succ_x k$. By *individual rationality*, $x \neq \emptyset$.

Without loss of generality, let i be the agent ranked highest according to $\succ_x$ to form such an agent-object blocking pair.

Let $N' = \{i, k\}$ and $R_i^x \in \mathcal{R}^x$ as in the calibration/construction step used to define $\succ_x$. By *strategy-proofness*, $\varphi_i((R_i^x, R_k), q) = \emptyset$. By *weak non-wastefulness* and *individual rationality*, $\varphi_k((R_i^x, R_k), q) = x$, $x P_k \emptyset$, and $q_x = 1$. Let $R_k^x = R_i^x$. By *strategy-proofness*, $\varphi_k((R_i^x, R_k^x), q) = x$. By *individual rationality*, $\varphi_i((R_i^x, R_k^x), q) = \emptyset$. Then, by Lemma 5 it follows that $k \succ_x i$. Hence, for $\varphi(R, q)$ we must have (s1) $|\{j \in N' : \varphi_j(R, q) = x\}| < q_x$.

Recall that we obtained $q_x = 1$ and $k \succ_x i$. Hence, (s1) implies that object x is not assigned at $\varphi(R, q)$. Because i is the agent ranked highest according to $\succ_x$ to form an agent-object blocking pair, $\varphi_k(R, q) P_k x$. Let $\varphi_k(R, q) = y \in O$. Recall that $\varphi_k((R_i^x, R_k), q) = x$. Then, by *strategy-proofness*, $\varphi_k(R_i^x, R_k^y, q) = \emptyset$. By *individual rationality*, $\varphi_i(R_i^x, R_k^y, q) \neq y$. But this is a contradiction to *weak non-wastefulness* for k and y .

In order to show **(b)**, suppose that $\varphi(R, q)$ is not stable under $\succ$. Then, there exists an agent-object pair $(i, x) \in N' \times O \cup \{\emptyset\}$ such that $x P_i \varphi_i(R, q)$ and (s1) $|\{j \in N' : \varphi_j(R, q) = x\}| < q_x$ or (s2) there exists $k \in N'$ such that $\varphi_k(R, q) = x$ and $i \succ_x k$. By *individual rationality*, $x \neq \emptyset$. Without loss of generality, let i be the agent ranked highest according to $\succ_x$ to form such an agent-object blocking pair.

Since $R_i \in \hat{\mathcal{R}}$, we have $A(R_i) = \{x\}$ and $\varphi_i(R, q) = \emptyset$. By *weak non-wastefulness*, (s1) cannot occur, i.e., we have (s2) there exists $k \in N'$ such that $\varphi_k(R, q) = x$ and $i \succ_x k$.

By *weak consistency*, $\varphi_i(R_{\{i, k\}}, \tilde{q}) = \emptyset$ (where $\tilde{q}_z = q_z - |\{j \in N' \setminus \{i, k\} : \varphi_j(R, q) = z\}|$ for all $z \in O$). By *weak non-wastefulness*, $\varphi_k(R_{\{i, k\}}, \tilde{q}) = x$. Hence, $\varphi(R_{\{i, k\}}, \tilde{q})$ is not stable under $\succ$, which contradicts (a). $\square$

Lemma 7. For all $N' \in \mathcal{N}$ and all $(R, q) \in \mathcal{R}^{N'} \times \mathcal{Q}$, $\varphi(R, q)$ is stable under $\succ$.

Proof. For any profile $R \in \mathcal{R}^{N'}$, let $\hat{N}(R) = \{i \in N' : R_i \notin \hat{\mathcal{R}}\}$. We prove that $\varphi(R, q)$ is stable under $\succ$ by induction on $|\hat{N}(R)|$.

Induction Basis: For $|\hat{N}(R)| = 0$, Lemma 6 (b) implies that $\varphi(R, q)$ is stable under $\succ$.

Induction Hypothesis: Assume that $\varphi(R, q)$ is stable under $\succ$ for any $R \in \mathcal{R}^{N'}$ such that $|\hat{N}(R)| \leq k$.

Induction Step: Let $R \in \mathcal{R}^{N'}$ be such that $|\hat{N}(R)| = k + 1$. Suppose that $\varphi(R, q)$ is not stable under $\succ$. Then, there exists an agent-object pair $(i, x) \in N' \times O \cup \{\emptyset\}$ such that $x P_i \varphi_i(R, q)$ and (s1) $|\{j \in N' : \varphi_j(R, q) = x\}| < q_x$ or (s2) there exists $l \in N'$ such that $\varphi_l(R, q) = x$ and $i \succ_x l$. By *individual rationality*, $x \neq \emptyset$. Without loss of generality, let i be the agent ranked highest according to $\succ_x$ to form such an agent-object blocking pair.

If $\varphi_i(R, q) = \emptyset$, then by *weak non-wastefulness*, we have (s2) there exists $l \in N'$ such that $\varphi_l(R, q) = x$ and $i \succ_x l$. By *weak consistency*, $\varphi_i(R_{\{i, l\}}, \tilde{q}) = \emptyset$ (where $\tilde{q}_z = q_z - |\{j \in N' \setminus \{i, l\} : \varphi_j(R, q) = z\}|$ for all $z \in O$).

$\varphi_j(R, q) = z\}$ for all $z \in O$). By *weak non-wastefulness*, $\varphi_l(R_{\{i,l\}}, \tilde{q}) = x$. Hence, $\varphi(R_{\{i,l\}}, \tilde{q})$ is not stable under $\succ$, which contradicts Lemma 6 (a).

Thus, $\varphi_i(R, q) \neq \emptyset$ and $R_i \notin \hat{\mathcal{R}}$. Let $R_i^x \in \mathcal{R}^x$. By *strategy-proofness*, $\varphi_i((R_i^x, R_{-i}), q) = \emptyset$. Note that $|\hat{N}(R_i^x, R_{-i})| = k$ and by the induction hypothesis, $\varphi((R_i^x, R_{-i}), q)$ is stable under $\succ$.

For (s2) there exists $l \in N'$ such that $\varphi_l(R, q) = x$ and $i \succ_x l$. Hence, by $\varphi_i((R_i^x, R_{-i}), q) = \emptyset$ and stability, $\varphi_l((R_i^x, R_{-i}), q) \neq x$. For (s1) $|\{j \in N' : \varphi_j(R, q) = x\}| < q_x$, $\varphi_i((R_i^x, R_{-i}), q) = \emptyset$ and *weak non-wastefulness* imply $|\{j \in N' : \varphi_j((R_i^x, R_{-i}), q) = x\}| = q_x$. Hence, in both cases (s1) and (s2) there exists $j \in N' \setminus \{i\}$ such that $\varphi_j((R_i^x, R_{-i}), q) = x \neq \varphi_j(R, q)$.

Thus, stability, $\varphi_j((R_i^x, R_{-i}), q) = x$, and $\varphi_i((R_i^x, R_{-i}), q) = \emptyset$ imply $j \succ_x i$. Because i is the agent ranked highest according to $\succ_x$ to form an agent-object blocking pair, $\varphi_j(R, q) P_j x$ and $R_j \notin \hat{\mathcal{R}}$.

Step 1: Let $\varphi_j(R, q) = y \in O$ and $R_j^y \in \mathcal{R}^y$. By *strategy-proofness*, $\varphi_j((R_j^y, R_{-j}), q) = y$. Recall that $\varphi_j((R_i^x, R_{-i}), q) = x$. Then, by *strategy-proofness* and *individual rationality*, $\varphi_j((R_i^x, R_j^y, R_{-i,j}), q) = \emptyset$.

Note that both $|\hat{N}(R_j^y, R_{-j})| \leq k$ and $|\hat{N}(R_i^x, R_j^y, R_{-i,j})| \leq k$ and the induction hypothesis applies to both profiles. By *weak non-wastefulness*, there exists $h \in N'$ such that $\varphi_h((R_i^x, R_j^y, R_{-i,j}), q) = y \neq \varphi_h((R_j^y, R_{-j}), q)$. By stability and $\varphi_j((R_i^x, R_j^y, R_{-i,j}), q) = \emptyset$, we have $h \succ_y j$. If $y P_h \varphi_h((R_j^y, R_{-j}), q)$, then by $h \succ_y j$ and $\varphi_j((R_j^y, R_{-j}), q) = y$, $\varphi((R_j^y, R_{-j}), q)$ is not stable under $\succ$, a contradiction.

Thus, $\varphi_h((R_j^y, R_{-j}), q) P_h y$ and $R_h \notin \hat{\mathcal{R}}$.

Step 2: Let $\varphi_h((R_j^y, R_{-j}), q) = z \in O$. As in Step 1, we use *strategy-proofness* to replace R_h by R_h^z in the problems $((R_j^y, R_{-j}), q)$ and $((R_i^x, R_j^y, R_{-i,j}), q)$. Then again we find a new agent h' with $R_{h'} \notin \hat{\mathcal{R}}$, etc.

Because the set of agents is finite, continuing along the lines of Steps 1 and 2 will ultimately lead to a contradiction. $\square$

The proof that for all $N' \in \mathcal{N}$ and $(R, q) \in \bigcup_{N' \in \mathcal{N}} \mathcal{R}^{N'} \times \mathcal{Q}$, $\varphi(R, q) = DA^\succ(R, q)$ is similar to the proof of Lemma 4 (which only uses *strategy-proofness*).

SUPPLEMENTARY APPENDIX (not intended for publication)

D Discussion of Kojima and Manea (2010)

We presented two characterizations of the class of DA-mechanisms with responsive priorities, i.e., any of these mechanisms determines the outcome by solely using the priority orderings over individual agents (see Remark 1). In contrast, Kojima and Manea (2010) characterize the class of DA-mechanisms with so-called acceptant substitutable priorities: a larger class of mechanisms than the class of responsive DA-mechanisms that is based on priorities that are determined by a choice function that reflects substitutability in preferences over sets of agents (see also Hatfield and Milgrom, 2005).

Formally, Kojima and Manea (2010) define a *priority for a type* $x \in O$ with capacity q_x as a (choice) correspondence $C_{q_x} : 2^N \rightarrow 2^N$ satisfying for all $M \subseteq N$, $C_{q_x}(M) \subseteq M$ and $|C_{q_x}(M)| \leq q_x$. C_{q_x} is *substitutable* if for all $M' \subseteq M \subseteq N$, $C_{q_x}(M) \cap M' \subseteq C_{q_x}(M')$. C_{q_x} is *acceptant* if for all $M \subseteq N$, $|C_{q_x}(M)| = \min\{q_x, |M|\}$. Clearly, taking a linear order $\succ_x$ over agents and defining C_{q_x} by choosing the $\min\{q_x, |M|\}$ best agents in M according to $\succ_x$ defines an acceptant substitutable priority. This particular class of priorities coincides with the class of (acceptant) responsive priorities employed in this paper (see Remark 1). Note that if $q_x = 1$, then the class of acceptant substitutable priorities coincides with our class of (acceptant) responsive priorities.

Since in our model resources can change, acceptant substitutable priorities for a type $x \in O$ in our model are modeled by a profile $C_x = (C_{q_x})_{1 \leq q_x \leq \bar{q}_x}$. We refer to Kojima and Manea (2010) for the definition of the DA-mechanism based on a acceptant substitutable priority structure—with our extended notion of acceptant substitutable priorities it is straightforward how this definition extends to our setup which allows for changing capacities and agents.¹³

Non-wastefulness as used by Kojima and Manea (2010)¹⁴ implies *weak non-wastefulness*, and *individual rationality*. Note that Kojima and Manea's (2010) *non-wastefulness* condition is equivalent to *individual rationality* and *non-wastefulness*. Furthermore, *non-wastefulness* is equivalent to the absence of agent-object pairs $(i, x) \in N \cup \{\emptyset\}$ such that $x P_i \varphi_i(R, q)$ and (s1) $|\{j \in N : \varphi_j(R, q) = x\}| < q_x$. Thus, *non-wastefulness* already incorporates an important part of stability. Note that *non-wastefulness* already eliminates priority mechanisms and linear programming mechanisms (these mechanisms do satisfy *weak non-wastefulness*, but *not non-wastefulness*).

In the characterizations obtained by Kojima and Manea (2010), next to *non-wastefulness*, two monotonicity properties are employed: *individually rational (IR) monotonicity* and *weak Maskin*

¹³Notice that any acceptant priority structure satisfies the *law of demand*: for all $x \in O$ and all $M' \subseteq N' \subseteq N$, $|C_x(M')| \leq |C_x(N')|$. Hence, by Hatfield and Milgrom (2005, Theorem 11), all acceptant substitutable DA-mechanisms satisfy *strategy-proofness*.

¹⁴**Kojima and Manea (2010) Non-Wastefulness:** For all $(R, q) \in \mathcal{R}^N \times \mathcal{Q}$, all $x \in O_+(q) \cup \{\emptyset\}$ with $|\{j \in N : \varphi_j(R, q) = x\}| < q_x$, and all $i \in N$, $\varphi_i(R, q) R_i x$.

monotonicity. We first define both monotonicity properties (using the equivalent “unilateral” definitions).¹⁵

Given $i \in N$ and $R_i \in \mathcal{R}$, a strategy $\bar{R}_i \in \mathcal{R}$ is an *individually rational (IR) monotonic transformation* of R_i at $\varphi_i(R, q)$ if for all $x \in O$, $x \bar{P}_i \varphi_i(R, q)$ and $x \bar{P}_i \emptyset$ imply $x P_i \varphi_i(R, q)$.

Individually Rational (IR) Monotonicity: For all $(R, q) \in \mathcal{R}^N \times \mathcal{Q}$, all $i \in N$, and all $\bar{R}_i \in \mathcal{R}_i$, if $\bar{R}_i$ is an IR monotonic transformation of R_i at $\varphi_i(R, q)$ and $\bar{R} = (\bar{R}_i, R_{-i})$, then for all $j \in N$, $\varphi_j(\bar{R}, q) \bar{R}_j \varphi_j(R, q)$.

Given $i \in N$ and $R_i \in \mathcal{R}$, a strategy $\bar{R}_i \in \mathcal{R}$ is a *monotonic transformation* of R_i at $\varphi_i(R, q)$ if for all $x \in O \cup \{\emptyset\}$, $x \bar{P}_i \varphi_i(R, q)$ implies $x P_i \varphi_i(R, q)$.

Weak Maskin Monotonicity: For all $(R, q) \in \mathcal{R}^N \times \mathcal{Q}$, all $i \in N$, and all $\bar{R}_i \in \mathcal{R}_i$, if $\bar{R}_i$ is a Maskin monotonic transformation of R_i at $\varphi_i(R, q)$ and $\bar{R} = (\bar{R}_i, R_{-i})$, then for all $j \in N$, $\varphi_j(\bar{R}, q) \bar{R}_j \varphi_j(R, q)$.

By Kojima and Manea (2010, Theorem 1) all DA-mechanisms with acceptant substitutable priorities satisfy *non-wastefulness* and *individually rational monotonicity*. It is easily seen that the class of priority mechanisms satisfies neither *non-wastefulness* nor *individually rational monotonicity*. The same is true for the class of linear programming mechanisms.

By Kojima and Manea (2010, Theorem 2) all DA-mechanisms with acceptant substitutable priorities satisfy *non-wastefulness*, *population monotonicity*, and *weak Maskin monotonicity*. It is easily seen that the class of priority mechanisms satisfies neither *non-wastefulness* nor *weak Maskin monotonicity*. The same is true for the class of linear programming mechanisms.

Our characterizations of responsive DA-mechanisms are based on simple and intuitive properties except for *strategy-proofness* (and again, this property makes the difference). In the characterizations of Kojima and Manea (2010) relatively stronger axioms are used and in each of their characterizations, both the classes of priority mechanisms and linear programming mechanisms violate at least two properties. Our results further support the use of responsive DA-mechanism as a practical solution in real-life matching markets where incentive compatibility is important.

Next, we present an example of an acceptant substitutable priority that is not responsive.

Example 1. Let $N = \{s_1, s_2, j_1, j_2\}$ be four economists looking for a new academic position. Furthermore, s_1, s_2 are seniors with specializations 1 or 2 and j_1, j_2 are juniors with specializations 1 or 2. The intuition behind the priorities that we define is the following. An economics department z has the following preferences for hiring. If only one economist can be hired the priority ranking for hiring is $s_1 \succ_z s_2 \succ_z j_1 \succ_z j_2$. However, if two positions can be filled, the department always would like to fill both positions. So, $\bar{q}_z = 2$. Furthermore, they would like to hire the two seniors;

¹⁵It can be shown that either of the two monotonicity properties implies *truncation invariance*.

but if only one senior s_i can be hired, then they are interested in also hiring the junior j_i in the same field. To be more specific, we assume that the department's priority ranking for hiring two economists is $\{s_1, s_2\} \succ_z \{s_1, j_1\} \succ_z \{s_2, j_2\} \succ_z \{s_1, j_2\} \succ_z \{s_2, j_1\} \succ_z \{j_1, j_2\}$.

Note that the department loosely speaking has lexicographic preferences. Priorities C_{q_z} based on these preferences work as follows. For $q_z = 2$ (the case $q_z = 1$ is obvious) it follows that

$$C_{q_z}(M) = \begin{cases} M & \text{if } |M| \leq 2, \\ \{s_1, s_2\} & \text{if } \{s_1, s_2\} \subsetneq M, \\ \{s_i, j_i\} & \text{if } \{s_i, j_i\} \subsetneq M, i \in \{1, 2\}. \end{cases}$$

It is easily verified that $C_z = (C_{q_z})_{1 \leq q_z \leq 2}$ is an acceptant and substitutable priority. Since $C_{q_z=2}(\{s_1, j_1, j_2\}) = \{s_1, j_1\}$ and $C_{q_z=2}(\{s_2, j_1, j_2\}) = \{s_2, j_2\}$, $C_{q_z=2}$ is not responsive. $\diamond$

E Independence

Note that below any mechanism satisfying *weak non-wastefulness* also satisfies *non-wastefulness* and *weak non-wastefulness* may be replaced by *non-wastefulness* without altering the independence of the axioms in any of our characterizations.

E.1 Independence of Properties in Theorem 1

The following example shows that on the domain of all problems, *two-agent consistent conflict resolution* is independent from all other properties in Theorem 1.

Example 2. The following mechanism f , which is not a responsive DA-mechanism, satisfies *unavailable type invariance*, *individual rationality*, *weak non-wastefulness*, *truncation invariance*, and *strategy-proofness*. Let $N = \{1, 2, 3\}$, $O = \{x, y\}$, and $\bar{q}_x = 2$ and $\bar{q}_y = 1$. Furthermore, $\succ_x: 1 \ 2 \ 3$, $\succ'_x: 1 \ 3 \ 2$, and $\succ_y: 1 \ 2 \ 3$. Let $\succ = (\succ_x, \succ_y)$ and $\succ' = (\succ'_x, \succ_y)$. Then, for each problem $(R, q) \in \mathcal{R}^N \times \mathcal{Q}$,

$$f(R, q) = \begin{cases} DA^{\succ'}(R, q) & \text{if } q_x = 2 \text{ and } x \text{ is agent 1's favorite object in } O_+(q) \text{ and} \\ DA^{\succ}(R, q) & \text{otherwise.} \end{cases}$$

It is easy to see that f satisfies *unavailable type invariance*, *individual rationality*, *weak non-wastefulness*, *truncation invariance*, and *strategy-proofness*.

Not two-agent consistent conflict resolving: In Example 2, let $R \in \mathcal{R}^N$ be such that $y P_1 x P_1 \emptyset$ and $R_2 = R_3 = R^x$. Let $q \in \mathcal{Q}$ be such that $q_x = 2$ and $q_y = 0$. Then $\varphi_2(R, q) = \emptyset$ and $\varphi_3(R, q) = x$ whereas $\varphi_2(R, 1_{xy}) = x$ and $\varphi_3(R, 1_{xy}) = \emptyset$.

In the following we consider the house allocation model, i.e., for all $x \in O$, $q_x = 1$. Therefore, instead of denoting capacity vectors, we simply denote the set of available real types, e.g., for $O' \subseteq O$, $O' \neq \emptyset$, (R, O') denotes a problem where one copy of each type in O' is available.

For any strict order π of agents in N , we denote the corresponding *serial dictatorship mechanism* by f^π ; for example, if $\pi : 1 \ 2 \ \dots \ (n-1) \ n$, then f^π works as follows: for each problem (R, O') , first agent 1 chooses his preferred object in O' , then agent 2 chooses his preferred object from the remaining objects $O' \setminus \{f_1^\pi(R, O')\}$, etc. Note that for each strict order π of N , $f^\pi = DA^{\succ^\pi}$ where $\succ^\pi$ equals the priority order where for all $x \in O$, $\succ_x^\pi = \pi$. Thus, each serial dictatorship mechanism f^π satisfies *unavailable type invariance*, *individual rationality*, *weak non-wastefulness*, *two-agent consistent conflict resolution*, *truncation invariance*, and *strategy-proofness*.

The following examples establish the independence of the properties (properties not mentioned in the examples follow easily).

Not unavailable type invariant: Let $n \geq 3$ and $\pi : 1 \ 2 \ 3 \ \dots \ (n-1) \ n$ and $\pi' : 1 \ n \ (n-1) \ \dots \ 3 \ 2$. Then, for each problem (R, O') ,

$$\varphi(R, O') = \begin{cases} f^\pi(R, O') & \text{if } A(R_1) = \emptyset \text{ and} \\ f^{\pi'}(R, O') & \text{otherwise.} \end{cases}$$

Not individually rational: Let $\pi : 1 \ 2 \ \dots \ (n-1) \ n$. For each $R_n \in \mathcal{R}$, let $\hat{R}_n$ be such that $A(\hat{R}_n) = O$ and $\hat{R}_n|O = R_n|O$. Then, for each problem (R, O') ,

$$\varphi(R, O') = f^\pi((R_{-n}, \hat{R}_n), O').$$

Not weakly non-wasteful: Fix an object $y \in O$ and $\pi : 1 \ 2 \ \dots \ (n-1) \ n$. Then, for each problem (R, O') ,

$$\varphi(R, O') = f^\pi(R, O' \setminus \{y\}).$$

Not two-agent consistent conflict resolving: Let π and π' be two distinct strict orders of agents in N . Then, for each problem (R, O') ,

$$\varphi(R, O') = \begin{cases} f^\pi(R, O') & \text{if } O' = O \text{ and} \\ f^{\pi'}(R, O') & \text{otherwise.} \end{cases}$$

Not truncation invariant: Let $N = \{1, 2, 3\}$, $O = \{x, y\}$, $\succ_x : 1 \ 2 \ 3$, $\succ'_x : 2 \ 1 \ 3$, and $\succ_y : 3 \ 1 \ 2$. Let $\succ = (\succ_x, \succ_y)$ and $\succ' = (\succ'_x, \succ_y)$. Then, for each problem (R, O') ,

$$\varphi(R, O') = \begin{cases} DA^\succ(R, O') & \text{if } \emptyset P_3 x \text{ and } x \in O' \text{ and} \\ DA^{\succ'}(R, O') & \text{otherwise.} \end{cases}$$

Let $R_1 : x \emptyset y$, $R_2 : x \emptyset y$, $R_3 : y \emptyset x$, and $R'_3 : y \ x \ \emptyset$. Let $R = (R_1, R_2, R_3)$ and $R' = (R_1, R_2, R'_3)$. Note that R_3 is a truncation of R'_3 and $\varphi_3(R, \{x, y\}) = y = \varphi_3(R', \{x, y\})$. However, $\varphi_1(R, \{x, y\}) = x$ and $\varphi_2(R', \{x, y\}) = x$; a contradiction of *truncation invariance*. Next, we show *two-agent consistent conflict resolution*, and *strategy-proofness* for this mechanism.

For *two-agent consistent conflict resolution*, consider (R, O') and (R, O'') . Since the allocation of y always follows the same priority order, *two-agent consistent conflict resolution* could only be violated for x . But then $x \in O'$ and $x \in O''$ and $\varphi(R, O') = \varphi(R, O'')$.

For *strategy-proofness*, note that agents 1 and 2 cannot change the priority structure by reporting a false preference relation. Consider agent 3 and a problem (R, O') . Obviously, if $|O'| = 1$, then agent 3 cannot profitably manipulate by reporting a false preference relation. Let $O' = \{x, y\}$. Now if agent 3's first choice is y or $\emptyset$ (or agent 3 receives his first choice), then agent 3 receives his first choice under (R, O') and agent 3 cannot profitably manipulate. Let agent 3's first choice be x .

If $R_3 : x \emptyset y$ and agent 3 does not receive his first choice, then $\varphi_3(R, \{x, y\}) = \emptyset$ and agent 3 can only change the priority structure by reporting a preference relation R'_3 with $\emptyset P'_3 x$. But then by *individual rationality*, $\varphi_3((R_{-3}, R'_3), \{x, y\}) \neq x$ and agent 3 cannot profitably manipulate by reporting a false preference relation.

If $R_3 : x y \emptyset$ and agent 3 does not receive his first choice, then $\varphi_3(R, \{x, y\}) = y$. Now the same argument as above establishes that agent 3 can never receive x by reporting a false preference relation.

Not *strategy-proof*: Let $\succ$ be a priority structure. Then, the responsive DA-mechanism based on the object-optimal matching that is obtained by using Gale and Shapley's (1962) object-proposing deferred-acceptance mechanism satisfies all properties except *strategy-proofness*.

E.2 Independence of Properties in Theorem 4

In the following we consider the house allocation model, i.e., for all $x \in O$, $q_x = 1$. Therefore, instead of denoting capacity vectors, we simply denote the set of available real objects, e.g., for $O' \subseteq O$, $O' \neq \emptyset$, (R, O') denotes a problem where one copy of each type in O' is available.

For any strict order π of agents in N , we denote the corresponding *serial dictatorship mechanism* by f^π ; for example, if $\pi : 1 \ 2 \ \dots \ (n-1) \ n$, then f^π works as follows: for each problem (R, O') such that $R \in R^{N'}$, first agent $\min N'$ chooses his preferred object in O' , then agent $\min N' \setminus \{\min N'\}$ chooses his preferred object from the remaining objects $O' \setminus \{f_1^\pi(R, O')\}$, etc. Note that for each strict order π of N , $f^\pi = DA^{\succ^\pi}$ where $\succ^\pi$ equals the priority order where for all $x \in O$, $\succ_x^\pi = \pi$. Thus, each serial dictatorship mechanism f^π satisfies *unassigned type invariance*, *individual rationality*, *weak non-wastefulness*, *weak consistency*, *truncation invariance*, and *strategy-proofness*.

The following examples establish the independence of the properties (properties not mentioned in the examples follow easily).

Not *unassigned type invariant*: Ehlers and Klaus (2007, Example 1) introduce a mechanism that violates *unassigned type invariance*, but satisfies *efficiency* and *consistency* (and the other properties of Theorem 4).

Not *individually rational*: Let $\pi : 1 \ 2 \ \dots \ (n-1) \ n$. For each $R_n \in \mathcal{R}$, let $\hat{R}_n$ be such that $A(\hat{R}_n) = O$ and $\hat{R}_n|O = R_n|O$. Then, for each problem (R, O') such that $R \in R^{N'}$,

$$\varphi(R, O') = \begin{cases} f^\pi((R_{-n}, \hat{R}_n), O') & \text{if } n \in N' \text{ and} \\ f^\pi(R, O') & \text{otherwise.} \end{cases}$$

Not *weakly non-wasteful*: Fix an object $y \in O$ and $\pi : 1 \ 2 \ \dots \ (n-1) \ n$. Then, for each problem (R, O') such that $R \in R^{N'}$,

$$\varphi(R, O') = f^\pi(R, O' \setminus \{y\}).$$

Not *weakly consistent*: Let $\pi : 1 \ 2 \ 3 \ \dots \ (n-1) \ n$ and $\pi' : 1 \ n \ (n-1) \ \dots \ 3 \ 2$. Then, for each problem (R, O') such that $R \in R^{N'}$,

$$\varphi(R, O') = \begin{cases} f^\pi(R, O') & \text{if } 1 \in N' \text{ and} \\ f^{\pi'}(R, O') & \text{otherwise.} \end{cases}$$

Not *strategy-proof*: Let $\succ$ be a priority structure. Then, the responsive DA-mechanism based on the object-optimal matching that is obtained by using Gale and Shapley's (1962) object-proposing deferred-acceptance mechanism satisfies all properties except *strategy-proofness*.

Récents cahiers de recherche du CIREQ
Recent Working Papers of CIREQ

Si vous désirez obtenir des exemplaires des cahiers, vous pouvez les télécharger à partir de notre site Web <http://www.cireqmontreal.com/cahiers-de-recherche>

If you wish to obtain copies of the working papers, you can download them directly from our website, <http://www.cireqmontreal.com/cahiers-de-recherche>

- 14-2011 Sprumont, Y., "Constrained-Optimal Strategy-Proof Assignment : Beyond the Groves Mechanisms", décembre 2011, 18 pages
- 01-2012 Long, N.V., V. Martinet, "Combining Rights and Welfarism : A New Approach to Intertemporal Evaluation of Social Alternatives", janvier 2012, 43 pages
- 02-2012 Atewamba, C., G. Gaudet, "Pricing of Durable Nonrenewable Natural Resources under Stochastic Investment Opportunities", novembre 2011, 24 pages
- 03-2012 Ehlers, L., "Top Trading with Fixed Tie-Breaking in Markets with Indivisible Goods", mars 2012, 25 pages
- 04-2012 Andersson, T., L. Ehlers, L.-G. Svensson, "(Minimally) ϵ -Incentive Compatible Competitive Equilibria in Economies with Indivisibilities", avril 2012, 13 pages
- 05-2012 Bossert, W., H. Peters, "Single-Plateaued Choice", mai 2012, 15 pages
- 06-2012 Bencheikroun, H., G. Martín-Herrán, "Farsight and Myopia in a Transboundary Pollution Game", juillet 2012, 19 pages
- 07-2012 Bossert, W., L. Ceriani, S.R. Chakravarty, C. D'Ambrosio, "Intertemporal Material Deprivation", juin 2012, 21 pages
- 08-2012 Bossert, W., Qi, C.X., J.A. Weymark, "Extensive Social Choice and the Measurement of Group Fitness in Biological Hierarchies", juillet 2012, 25 pages
- 09-2012 Bossert, W., K. Suzumura, "Multi-Profile Intertemporal Social Choice", juillet 2012, 20 pages
- 10-2012 Bakis, O., B. Kaymak, "On the Optimality of Progressive Income Redistribution", août 2012, 42 pages
- 11-2012 Poschke, M., "The Labor Market, the Decision to Become an Entrepreneur, and the Firm Size Distribution", août 2012, 29 pages
- 12-2012 Bossert, W., Y. Sprumont, "Strategy-proof Preference Aggregation", août 2012, 23 pages
- 13-2012 Poschke, M., "Who Becomes an Entrepreneur? Labor Market Prospects and Occupational Choice", septembre 2012, 49 pages
- 14-2012 Bencheikroun, H., G. Gaudet, H. Lohoues, "Some Effects of Asymmetries in a Common Pool Natural Resource Oligopoly", août 2012, 24 pages