



Characterizing the Typewise Top- Trading- Cycles Mechanism for Multiple- Type Housing Markets

BSE Working Paper 1341

April 2022 (Revised February 2024)

Di Feng, Bettina Klaus, Flip Klijn

bse.eu/research

Characterizing the Typewise Top-Trading-Cycles Mechanism for Multiple-Type Housing Markets*

Di Feng[†]

Bettina Klaus[‡]

Flip Klijn[§]

February 12, 2024

Abstract

We consider the generalization of the classical Shapley and Scarf housing market model (Shapley and Scarf, 1974) to so-called multiple-type housing markets (Moulin, 1995). Throughout the paper, we focus on strict preferences. When preferences are separable, the prominent solution for these markets is the typewise top-trading-cycles (tTTC) mechanism.

We first show that for lexicographic preferences, a mechanism is *unanimous* (or *onto*), *individually rational*, *strategy-proof*, and *non-bossy* if and only if it is the tTTC mechanism. Second, we obtain a corresponding characterization for separable preferences. We obtain additional characterizations when replacing [*strategy-proofness* and *non-bossiness*] with *self-enforcing group* (or *pairwise*) *strategy-proofness*. Finally, we show that for strict preferences, there is no mechanism satisfying *unanimity*, *individual rationality*, and *strategy-proofness*. We obtain further impossibility results for strict preferences based on weakening *unanimity* to *ontoness* and on extending the tTTC solution.

Our characterizations of the tTTC mechanism constitute the first characterizations of an extension of the prominent top-trading-cycles (TTC) mechanism to multiple-type housing markets.

Keywords: multiple-type housing markets; *strategy-proofness*; *non-bossiness*; *self-enforcing pairwise strategy-proofness*; top-trading-cycles (TTC) mechanism; market design.

JEL codes: C78; D47.

*We thank the editor and two reviewers for their valuable comments and suggestions. Di Feng and Bettina Klaus gratefully acknowledge financial support from the Swiss National Science Foundation (SNSF) through Project 100018_192583. Flip Klijn gratefully acknowledges financial support from AGAUR–Generalitat de Catalunya (2021-SGR-00416) and the Spanish Agencia Estatal de Investigación (MCIN/ AEI /10.13039/501100011033) through grant PID2020-114251GB-I00 and the Severo Ochoa Programme for Centres of Excellence in R&D (Barcelona School of Economics CEX2019-000915-S).

[†]Department of Finance, Dongbei University of Finance and Economics, Jianshan Jie 217, 116025, Shahekou, Dalian, P.R. China; *e-mail*: dfengdufe@outlook.com.

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

[§]Institute for Economic Analysis (CSIC) and Barcelona School of Economics, Campus UAB, 08193 Bellaterra (Barcelona), Spain; *e-mail*: flip.klijn@iae.csic.es.

1 Introduction

In many applied matching problems, indivisible goods that are in unit demand have to be assigned without monetary transfers. One of the most prominent such problems is modeled by classical Shapley-Scarf housing markets (Shapley and Scarf, 1974). Shapley and Scarf (1974) consider an exchange economy in which each agent owns an indivisible object (say, a house); each agent has preferences over houses and wishes to consume exactly one house. The objective of the market designer then is to reallocate houses among agents. When preferences are strict, Shapley and Scarf (1974) show that the strict core (defined by a weak blocking notion) has remarkable features: it is non-empty,¹ and can be easily calculated by the so-called top-trading-cycles (TTC) algorithm (due to David Gale). Moreover, the TTC mechanism that assigns the unique strict core allocation satisfies important incentive properties, *strategy-proofness* (Roth, 1982) as well as the stronger property of *group strategy-proofness* (Bird, 1984). Furthermore, Ma (1994) and Svensson (1999) show that the TTC mechanism is the unique mechanism satisfying *Pareto efficiency*, *individual rationality*, and *strategy-proofness*. Throughout the paper, we focus on strict preferences.

However, more general problems of exchanging indivisible objects that are in multi-unit demand are known to be very difficult. In this paper, we consider an extension of the classical Shapley-Scarf housing markets by allowing a specific kind of multi-unit demand: multiple-type housing markets, to use the language of Moulin (1995).² In this model, objects are of different types (say, houses, cars, etc.) and agents initially own and wish to consume exactly one object of each type. A familiar example for most readers would be the situation of students' enrollment at many universities where courses are taught in small groups and in multiple sessions (Klaus, 2008). Furthermore, for term paper presentations during a course, students may want to exchange their assigned topics and dates (Mackin and Xia, 2016); hospitals may want to improve their surgery schedule for surgeons by swapping surgery staff, operating rooms, and dates (Huh et al., 2013); and in cloud computing (Ghodsi et al., 2011, 2012) and 5G network slicing (Peng et al., 2015; Bag et al., 2019; Han et al., 2019), there may be several types of resources that agents require, including CPU, memory, and storage.

This model is firstly studied by Konishi et al. (2001). Their results are mainly negative: they show that even if we further restrict preferences to be strict and additively separable, the strict core may still be empty. Moreover, there exists no mechanism that is *Pareto efficient*, *individually rational*, and *strategy-proof*.

Despite their negative results, for (strictly) separable preferences, Wako (2005) suggests an alternative solution concept to the strict core by first decomposing a multiple-type housing

¹Roth and Postlewaite (1977) show that the strict core is single-valued.

²There are many other resource allocation models with multi-unit demand, such as Pápai (2001, 2007) and Manjunath and Westkamp (2021).

market into typewise submarkets and second, determining the strict core in each submarket. Wako (2005) calls this unique outcome the commoditywise competitive allocation and shows that it is implementable in (self-enforcing) coalition-proof Nash equilibria but not in strong Nash equilibria.³

Based on Wako's result, we investigate the mechanism that always assigns the commoditywise competitive allocation; since this allocation can be obtained by using the TTC algorithm for each object type, we refer to it as the *typewise TTC (tTTC) mechanism*. Although the tTTC mechanism is not *Pareto efficient*, it does have many desirable properties: it is *individually rational*, *strategy-proof*, and *second-best incentive compatible*, i.e., it is *not Pareto dominated* by any other *strategy-proof* mechanism (Klaus, 2008). In view of these positive results, one may wonder whether the tTTC mechanism can be characterized by weakening *Pareto efficiency* and strengthening *strategy-proofness*.

For Shapley-Scarf housing markets with strict preferences, a characterization along these lines is provided by Takamiya (2001): he shows that the TTC mechanism is the only mechanism satisfying *unanimity*, *individual rationality*, and *group strategy-proofness*.⁴ Based on Takamiya's result, one could now conjecture that this characterization of the TTC mechanism for Shapley-Scarf housing markets can be carried over to the tTTC mechanism for multiple-type housing markets. That conjecture is almost true; however, we need to weaken *group strategy-proofness* to *strategy-proofness* and *non-bossiness*.⁵ In other words, inspired by Takamiya's result for Shapley-Scarf housing markets, we show that, remarkably, the tTTC mechanism is the only mechanism satisfying *unanimity* (or *onteness*), *individual rationality*, *strategy-proofness*, and *non-bossiness* (see Theorems 1 and 2 for lexicographic and separable preferences, respectively). We obtain additional characterizations when replacing [*strategy-proofness* and *non-bossiness*] with *self-enforcing group (or pairwise) strategy-proofness* (Corollaries 1 and 2).

Our characterizations of the tTTC mechanism constitute the first characterizations of an extension of the prominent top-trading-cycles (TTC) mechanism to multiple-type housing markets. Furthermore, our results suggest that when preferences are separable, the tTTC mechanism is outstanding; first, because some efficiency in the form of *unanimity* is preserved (even if full *Pareto efficiency* cannot be reached), and second, because of its incentive robustness in the form of *strategy-proofness*, *non-bossiness*, and *self-enforcing group (pairwise) strategy-proofness* (even if full *group strategy-proofness* cannot be reached). Moreover, we also provide several impossibility results (Theorems 3 and 4, Corollaries 3 and 4) for strict (but otherwise unrestricted)

³However, (1) the commoditywise competitive allocation may be *Pareto inefficient*; and (2) the mechanism that always assigns this allocation is *not group strategy-proof* (see Wako, 2005, Section 6, for details).

⁴In fact, Takamiya's characterization is based on *onteness*, a weakening of *unanimity*. However, in the presence of *group strategy-proofness*, *onteness* coincides with *unanimity*.

⁵When preferences are strict but otherwise unrestricted, the combination of *strategy-proofness* and *non-bossiness* is equivalent to *group strategy-proofness*. Example 1 shows that this is not true for separable preferences.

preferences:

- there is no mechanism satisfying *unanimity*, *individual rationality*, and *strategy-proofness* (Theorem 3);
- there is no mechanism satisfying *ontoness*, *individual rationality*, *strategy-proofness*, and *non-bossiness* (Corollary 3);
- there is no *individually rational* and *strategy-proof* mechanism that extends the tTTC mechanism from lexicographic (separable) preferences to strict preferences (Theorem 4); and
- there is no *strategy-proof* and *non-bossy* mechanism that extends the tTTC mechanism from lexicographic (separable) preferences to strict preferences (Corollary 4).

The rest of the paper is organized as follows. In the following section, Section 2, we introduce multiple-type housing markets, mechanisms and their properties, and the tTTC mechanism. We state our results in Section 3. In Subsection 3.1, we first show that for lexicographic preferences, a mechanism is *unanimous* (or *onto*), *individually rational*, *strategy-proof*, and *non-bossy* if and only if it is the tTTC mechanism (Theorem 1). In Subsection 3.2, using Theorem 1, we obtain a corresponding characterization for separable preferences (Theorem 2). We would like to emphasize that the proof strategy to use lexicographic preferences as a “stepping stone” to obtain a corresponding result for separable preferences is, to the best of our knowledge, new. In Subsections 3.1 and 3.2 we obtain additional characterizations when replacing [*strategy-proofness* and *non-bossiness*] with *self-enforcing group (or pairwise) strategy-proofness* (Corollaries 1 and 2). In Subsection 3.3 and Appendix F, we finally show several impossibility results (Theorems 3 and 4, Corollaries 3 and 4). Section 4 concludes with a discussion of our results and how they relate to the literature.

2 The model

Multiple-type housing markets

We consider a barter economy without monetary transfers formed by n agents and $n \times m$ indivisible objects. Let $N = \{1, \dots, n\}$ be a finite *set of agents*. A nonempty subset of agents $S \subseteq N$ is a *coalition*. We assume that there exist $m \geq 1$ (*distinct*) *types of indivisible objects* and n (*distinct*) *indivisible objects of each type*. We denote the *set of types* by $T = \{1, \dots, m\}$. Note that for $m = 1$ our model equals the classical Shapley-Scarf housing market model (Shapley and Scarf, 1974).

Each agent $i \in N$ is endowed with exactly one object of each type $t \in T$, denoted by o_i^t . Hence, each agent i 's endowment is a list $o_i = (o_i^1, \dots, o_i^m)$. The set of type- t objects is $O^t = \{o_1^t, \dots, o_n^t\}$, and the set of all objects is $O = \{o_1^1, o_1^2, \dots, o_n^1, o_n^2, \dots, o_n^m\}$. In particular, $|O| = n \times m$.

For each $i \in N$, an allotment x_i assigns one object of each type to agent i , i.e., x_i is a list $x_i = (x_i^1, \dots, x_i^m) \in \Pi_{t \in T} O^t$, where $x_i^t \in O^t$ is agent i 's type- t allotment. We assume that each agent i has complete, antisymmetric, and transitive preferences R_i over all possible allotments, i.e., R_i is a linear order over $\Pi_{t \in T} O^t$.⁶ For two allotments x_i and y_i , x_i is weakly preferred to y_i if $x_i R_i y_i$, and x_i is strictly preferred to y_i if $[x_i R_i y_i \text{ and not } y_i R_i x_i]$, denoted by $x_i P_i y_i$. Finally, since preferences over allotments are strict, agent i is indifferent between x_i and y_i only if $x_i = y_i$. We denote preferences as ordered lists, e.g., $R_i : x_i, y_i, z_i$ instead of $x_i P_i y_i P_i z_i$. The set of all preferences is denoted by $\mathcal{R}$, which we will also refer to as the strict preference domain.

A preference profile specifies preferences for all agents and is denoted by a list $R = (R_1, \dots, R_n) \in \mathcal{R}^N$. We use the standard notation $R_{-i} = (R_1, \dots, R_{i-1}, R_{i+1}, \dots, R_n)$ to denote the list of all agents' preferences, except for agent i 's preferences. Furthermore, for each coalition S we define $R_S = (R_i)_{i \in S}$ and $R_{-S} = (R_i)_{i \in N \setminus S}$ to be the lists of preferences of the members of coalitions S and $N \setminus S$, respectively.

In addition to the domain of strict preferences, we consider two preference subdomains based on agents' "marginal preferences": assume that for each $i \in N$ and for each type $t \in T$, agent i has complete, antisymmetric, and transitive preferences R_i^t over the set of type- t objects O^t . We refer to R_i^t as agent i 's type- t marginal preferences, and denote by $\mathcal{R}^t$ the set of all type- t marginal preferences. Then, we can define the following two preference domains.

(Strictly) Separable preferences. Agent i 's preferences $R_i \in \mathcal{R}$ are separable if for each $t \in T$ there exist type- t marginal preferences $R_i^t \in \mathcal{R}^t$ such that for any two allotments x_i and y_i ,

$$\text{if for all } t \in T, x_i^t R_i^t y_i^t, \text{ then } x_i R_i y_i.$$

$\mathcal{R}_s$ denotes the domain of separable preferences.

Before defining our next preference domain, we introduce some notation. We use a bijective function $\pi_i : T \rightarrow T$ to order types according to agent i 's "(subjective) importance," with $\pi_i(1)$ being the most important and $\pi_i(m)$ being the least important object type. We denote π_i as an ordered list of types, e.g., by $\pi_i = (2, 3, 1)$, we mean that $\pi_i(1) = 2$, $\pi_i(2) = 3$, and $\pi_i(3) = 1$. For each agent $i \in N$ and each allotment $x_i = (x_i^1, \dots, x_i^m)$, $x_i^{\pi_i} = (x_i^{\pi_i(1)}, \dots, x_i^{\pi_i(m)})$ denotes the allotment after rearranging it with respect to the object-type importance order π_i .

⁶Preferences R_i are complete if for any two allotments x_i, y_i , $x_i R_i y_i$ or $y_i R_i x_i$; they are antisymmetric if $x_i R_i y_i$ and $y_i R_i x_i$ imply $x_i = y_i$; and they are transitive if for any three allotments x_i, y_i, z_i , $x_i R_i y_i$ and $y_i R_i z_i$ imply $x_i R_i z_i$.

(Separably) Lexicographic preferences. Agent i 's preferences $R_i \in \mathcal{R}$ are *(separably) lexicographical* if they are separable with type- t marginal preferences $(R_i^t)_{t \in T}$ and there exists an object-type importance order $\pi_i : T \rightarrow T$ such that for any two allotments x_i and y_i ,

if $x_i^{\pi_i(1)} P_i^{\pi_i(1)} y_i^{\pi_i(1)}$ or
if there exists a positive integer $k \leq m - 1$ such that
 $x_i^{\pi_i(1)} = y_i^{\pi_i(1)}, \dots, x_i^{\pi_i(k)} = y_i^{\pi_i(k)}$, and $x_i^{\pi_i(k+1)} P_i^{\pi_i(k+1)} y_i^{\pi_i(k+1)}$,
then $x_i P_i y_i$.

$\mathcal{R}_l$ denotes the *domain of lexicographic preferences*.

Note that $R_i \in \mathcal{R}_l$ can be represented by a $m + 1$ -tuple $R_i = (R_i^1, \dots, R_i^m, \pi_i) = ((R_i^t)_{t \in T}, \pi_i)$, or a strict ordering of all objects,⁷ i.e., R_i lists first all $\pi(1)$ objects (according to $R_i^{\pi(1)}$), then all $\pi(2)$ objects (according to $R_i^{\pi(2)}$), and so on. We provide a simple illustration in Example 1.

Note that if $m > 1$,

$$\mathcal{R}_l \subsetneq \mathcal{R}_s \subsetneq \mathcal{R}.$$

An *allocation* x partitions the set of all objects O into agents' allotments, i.e., $x = \{x_1, \dots, x_n\}$ is such that for each $t \in T$, $\cup_{i \in N} x_i^t = O^t$ and for each pair $i \neq j$, $x_i^t \neq x_j^t$. For simplicity, sometimes we will restate an allocation as a list $x = (x_1, \dots, x_n)$. The *set of all allocations* is denoted by X , and the *endowment allocation* is denoted by $e = (o_1, \dots, o_n)$.

We assume that when facing an allocation x , there are no consumption externalities and each agent $i \in N$ only cares about his own allotment x_i . Hence, each agent i 's preferences over allocations X are essentially equivalent to his preferences over allotments $\Pi_{t \in T} O^t$. With some abuse of notation, we use notation R_i to denote an agent i 's preferences over allotments as well as his preferences over allocations, i.e., for each agent $i \in N$ and for any two allocations $x, y \in X$, $x R_i y$ if and only if $x_i R_i y_i$.⁸

A *(multiple-type housing) market* is a triple (N, e, R) . When no confusion is possible about the set of agents N and the endowment allocation e , we denote market (N, e, R) by R . Thus, the domain of strict preference profiles $\mathcal{R}^N$ also denotes the *set of all markets*.

Mechanisms and properties

Note that the following definitions and results for the domain of strict preference profiles $\mathcal{R}^N$ can be formulated for the domain of separable preference profiles $\mathcal{R}_s^N$ or the domain of lexicographic preference profiles $\mathcal{R}_l^N$.

⁷See Feng and Klaus (2022, Remark 1) for details.

⁸Note that when extending strict preferences over allotments to preferences over allocations without consumption externalities, strictness is lost because an agent is indifferent between any two allocations where he gets the same allotment.

A *mechanism* (on $\mathcal{R}^N$) is a function $f : \mathcal{R}^N \rightarrow X$ that assigns to each market $R \in \mathcal{R}^N$ an allocation $f(R) \in X$, and

- for each $i \in N$, $f_i(R)$ is *agent i 's allotment*;
- for each $i \in N$ and each $t \in T$, $f_i^t(R)$ is *agent i 's type- t allotment*

under mechanism f at R .

We next introduce and discuss some well-known properties for allocations and mechanisms. Let $R \in \mathcal{R}^N$.

First we consider a voluntary participation condition for an allocation x to be implementable without causing agents any harm: no agent will be worse off than at his endowment.

Definition 1 (Individual rationality).

An allocation $x \in X$ is *individually rational* if for each agent $i \in N$, $x_i R_i o_i$. A mechanism on $\mathcal{R}^N$ is *individually rational* if for each market, it assigns an individually rational allocation.⁹

Next, we consider two well-known efficiency criteria.

Definition 2 (Pareto efficiency).

An allocation $y \in X$ *Pareto dominates* allocation $x \in X$ if for each agent $i \in N$, $y_i R_i x_i$, and for at least one agent $j \in N$, $y_j P_j x_j$. An allocation $x \in X$ is *Pareto efficient* if there is no allocation $y \in X$ that Pareto dominates it. A mechanism on $\mathcal{R}^N$ is *Pareto efficient* if for each market, it assigns a Pareto efficient allocation.

Definition 3 (Unanimity).

An allocation $x \in X$ is *unanimously best* if for each agent $i \in N$ and each allocation $y \in X$, we have $x R_i y$.¹⁰ A mechanism on $\mathcal{R}^N$ is *unanimous* if for each market, it assigns the unanimously best allocation whenever it exists.

If a unanimously best allocation exists for $R \in \mathcal{R}^N$, then that allocation is the only *Pareto efficient* allocation for R . Hence, *Pareto efficiency* implies *unanimity*.

Next, we introduce a weaker condition than *unanimity* that guarantees that no allocation is a priori excluded.

Definition 4 (Onto).

A mechanism on $\mathcal{R}^N$ is *onto* if each allocation is assigned to some markets. In other words, a mechanism is *onto* if it is an onto function.

⁹Note that even if preferences R_i are lexicographic, $x_i R_i o_i$ does *not* imply that for each type $t \in T$, $x_i^t R_i^t o_i^t$. For instance, in Example 1, $(H_2, C_2) R_1 (H_1, C_1)$ but $C_1 R_1^C C_2$. The stronger requirement that for each type $t \in T$, $x_i^t R_i^t o_i^t$, which we call *marginal individual rationality*, is introduced for lexicographic preferences in Definition 12 (Appendix B).

¹⁰Since all preferences are strict, the set of unanimously best allocations is empty or single-valued.

It is immediate that *unanimity* implies *ontoness* (see also Lemma 1).

The next two properties are incentive or invariance properties that model that no agent can benefit from misrepresenting his preferences or alter other agents' allotments.

Definition 5 (Strategy-proofness).

A mechanism f on $\mathcal{R}^N$ is *strategy-proof* if for each $R \in \mathcal{R}^N$, each agent $i \in N$, and each preference relation $R'_i \in \mathcal{R}$, $f_i(R_i, R_{-i}) R_i f_i(R'_i, R_{-i})$, i.e., *agent i cannot manipulate mechanism f at R via R'_i .*

Next, we consider a well-known property for mechanisms that restricts each agent's influence: no agent can change other agents' allotments without changing his own allotment.

Definition 6 (Non-bossiness).

A mechanism f on $\mathcal{R}^N$ is *non-bossy* if for each $R \in \mathcal{R}^N$, each agent $i \in N$, and each $R'_i \in \mathcal{R}$, $f_i(R_i, R_{-i}) = f_i(R'_i, R_{-i})$ implies $f(R_i, R_{-i}) = f(R'_i, R_{-i})$.

We already mentioned that *unanimity* implies *ontoness*. We next show that, in the presence of *strategy-proofness* and *non-bossiness*, *ontoness* implies *unanimity*.

Lemma 1.

- (a) *If a mechanism on $\mathcal{R}^N$ is unanimous, then it is onto.*
- (b) *If a mechanism on $\mathcal{R}^N$ is strategy-proof, non-bossy, and onto, then it is unanimous.*

The proof of Lemma 1 is relegated to Appendix A.

The next property models that no coalition can benefit from misrepresenting their preferences.

Definition 7 (Group strategy-proofness).

A mechanism f on $\mathcal{R}^N$ is *group strategy-proof* if for each $R \in \mathcal{R}^N$, there is no coalition $S \subseteq N$ and no preference list $R'_S = (R'_i)_{i \in S} \in \mathcal{R}^S$ such that for each $i \in S$, $f_i(R'_S, R_{-S}) R_i f_i(R, R_{-S})$, and for some $j \in S$, $f_j(R'_S, R_{-S}) P_j f_j(R, R_{-S})$, i.e., *coalition S cannot manipulate mechanism f at R via R'_S .*

Group-strategy-proofness implies *strategy-proofness* and *non-bossiness*. In fact, using the arguments of the proof of Pápai (2000, Lemma 1), it is easy to see that, when preferences are (strict and) unrestricted (i.e., the domain is $\mathcal{R}^N$), the combination of *strategy-proofness* and *non-bossiness* coincides with *group strategy-proofness*. However, on some smaller domains of preference profiles (e.g., $\mathcal{R}_s^N / \mathcal{R}_l^N$), *group strategy-proofness* can be stronger than the combination of *strategy-proofness* and *non-bossiness* (see Example 1).

Finally, we introduce a strategic robustness property that is stronger than *strategy-proofness* and weaker than *group strategy-proofness*; *group strategy-proofness* is weakened by requiring that coalitional manipulations are “self-enforcing.”

Definition 8 (Self-enforcing group (pairwise) strategy-proofness).

A coalition $S \subseteq N$ can manipulate mechanism f on $\mathcal{R}^N$ at $R \in \mathcal{R}^N$ in a *self-enforcing manner* if S can manipulate f at R via some deviation R'_S such that for each $V \subsetneq S$ and each $\ell \in V$, $f_\ell(R'_S, R_{N \setminus S}) R_\ell f_\ell(R_V, R'_{S \setminus V}, R_{N \setminus S})$.¹¹ A mechanism f is *self-enforcing group strategy-proof* if no coalition can manipulate f in a self-enforcing manner at any $R \in \mathcal{R}^N$. A mechanism f is *self-enforcing pairwise strategy-proof* if no coalition S with $|S| \leq 2$ can manipulate f in a self-enforcing manner at any $R \in \mathcal{R}^N$.

Serizawa (2006), Alva (2017), and Biró et al. (2022a) introduce and analyze related (but slightly different) notions of *self-enforcing pairwise strategy-proofness* for various economic models. Similarly to Alva (2017, Proposition 1) and Biró et al. (2022a, Proposition 11), we show that *strategy-proofness* and *non-bossiness* are equivalent to *self-enforcing pairwise strategy-proofness* as well as *self-enforcing group strategy-proofness*. Thus, as the combination with *strategy-proofness* reveals, the invariance property *non-bossiness* embodies an important aspect of group incentive robustness.

Lemma 2.

The following statements for a mechanism f on $\mathcal{R}^N$ are equivalent.

- (i) f is strategy-proof and non-bossy;
- (ii) f is self-enforcing pairwise strategy-proof;
- (iii) f is self-enforcing group strategy-proof.

When preferences are (strict and) unrestricted (i.e., the domain is $\mathcal{R}^N$), Lemma 2 follows from the equivalence of *group strategy-proofness* with [*strategy-proofness* and *non-bossiness*]. However, recall that definitions and results in this section also apply to the domain of separable preference profiles $\mathcal{R}_s^N$ and the domain of lexicographic preference profiles $\mathcal{R}_l^N$. On these smaller domains of preference profiles, Lemma 2 establishes new results. The proof of Lemma 2 is relegated to Appendix A.

We next focus on the domain of separable preference profiles $\mathcal{R}_s^N$ (the domain of lexicographic preference profiles $\mathcal{R}_l^N$, respectively) and extend Gale's famous top-trading-cycles (TTC) algorithm to multiple-type housing markets.

Definition 9 (The type- t top-trading-cycles (TTC) algorithm).

Consider a market (N, e, R) such that $R \in \mathcal{R}_s^N$. For each type $t \in T$, let $(N, e^t, R^t) = (N, (o_1^t, \dots, o_n^t), (R_1^t, \dots, R_n^t))$ be its *associated type- t submarket*.

¹¹Intuitively, self-enforcement of a coalitional S 's manipulation means that no sub-coalition V of S has an incentive to revert to their original preference reports while the remainder of the coalition $S \setminus V$ continues to misreport.

For each type t , we define the top-trading-cycles (TTC) allocation for the type- t submarket as follows.

Input. A type- t submarket (N, e^t, R^t) .

Step 1. Let $N_1 := N$ and $O_1^t := O^t$. We construct a directed graph with the set of nodes $N_1 \cup O_1^t$. For each agent $i \in N_1$, there is an edge from the agent to his most preferred type- t object in O_1^t according to R_i^t . For each edge (i, o) we say that agent i points to type- t object o . For each type- t object $o \in O_1^t$, there is an edge from the object to its owner.

A *trading cycle* is a directed cycle in the graph. Given the finite number of nodes, at least one trading cycle exists. We assign to each agent in a trading cycle the type- t object he points to and remove all trading cycle agents and type- t objects. If there are some agents (and hence objects) left, we continue with the next step. Otherwise we stop.

Step k . Let N_k be the set of agents that remain after Step $k - 1$ and O_k^t be the set of type- t objects that remain after Step $k - 1$. We construct a directed graph with the set of nodes $N_k \cup O_k^t$. For each agent $i \in N_k$, there is an edge from the agent to his most preferred type- t object in O_k^t according to R_i^t . For each type- t object $o \in O_k^t$, there is an edge from the object to its owner. At least one trading cycle exists and we assign to each agent in a trading cycle the type- t object he points to and remove all trading cycle agents and objects. If there are some agents (and hence objects) left, we continue with the next step. Otherwise we stop.

Output. The type- t TTC algorithm terminates when each agent in N is assigned an object in O^t , which takes at most n steps. We denote the object in O^t that agent $i \in N$ obtains in the type- t TTC algorithm by $TTC_i^t(R^t)$ and the final type- t allocation by $TTC^t(R^t)$.

Definition 10 (tTTC allocations and the tTTC mechanism).

The *typewise top-trading-cycles (tTTC) allocation*, $tTTC(R)$, is the collection of all type- t TTC allocations, i.e., for each $R \in \mathcal{R}_s^N$,

$$tTTC(R) = ((TTC_1^1(R^1), \dots, TTC_1^m(R^m)) , \dots , (TTC_n^1(R^1), \dots, TTC_n^m(R^m))).$$

The *tTTC mechanism* (introduced by Wako, 2005) assigns to each market $R \in \mathcal{R}_s^N$ its tTTC allocation.

Shapley-Scarf housing market results

As mentioned before, for $m = 1$ our model equals the classical Shapley-Scarf housing market model (Shapley and Scarf, 1974) and the tTTC mechanism reduces to the standard TTC mechanism. The Shapley-Scarf housing market (with strict preferences) results that are pertinent for our analysis of multiple-type housing markets are the following.

Result 1 (Bird, 1984).

Let $m = 1$. The TTC mechanism on $\mathcal{R}^N$ is group strategy-proof.

Recall that *group strategy-proofness* implies *strategy-proofness* and *non-bossiness*. Thus, Result 1 also implies that the TTC mechanism is *non-bossy* (Miyagawa, 2002, explicitly shows this). Also recall that when preferences are (strict and) unrestricted, the combination of *strategy-proofness* and *non-bossiness* coincides with *group strategy-proofness*. Recently, Alva (2017) identifies other preference domain properties such that this equivalence holds.

Result 2 (Pápai, 2000; Takamiya, 2001; Alva, 2017).

Let $m = 1$. A mechanism on $\mathcal{R}^N$ is strategy-proof and non-bossy if and only if it is group strategy-proof.

Result 3 (Ma, 1994; Svensson, 1999).

Let $m = 1$. A mechanism on $\mathcal{R}^N$ is Pareto efficient, individually rational, and strategy-proof if and only if it is the TTC mechanism.

Result 4 (Takamiya, 2001).

Let $m = 1$. A mechanism on $\mathcal{R}^N$ is onto, individually rational, strategy-proof, and non-bossy if and only if it is the TTC mechanism.

Extension of existing Shapley-Scarf housing market results to multiple-type housing markets

The results in the previous subsection imply that for Shapley-Scarf housing markets, the TTC mechanism on $\mathcal{R}^N$ satisfies

- *Pareto efficiency* and hence *unanimity* and *ontoness*;
- *individual rationality*; and
- *group strategy-proofness* and hence *strategy-proofness* and *non-bossiness*.

The tTTC mechanism (which is defined on $\mathcal{R}_s^N$ and $\mathcal{R}_l^N$) inherits most of these properties, except for *Pareto efficiency* and *group strategy-proofness*. Hence, TTC Results 1, 2, and 3 do not extend to the tTTC mechanism when more than one object type is allocated.

Proposition 1. *The tTTC mechanism on $\mathcal{R}_s^N$ ($\mathcal{R}_l^N$, respectively) satisfies unanimity, ontoness, individual rationality, strategy-proofness, non-bossiness, and self-enforcing group (pairwise) strategy-proofness.*

The proof of Proposition 1 is relegated to Appendix A.

Example 1 below shows that the tTTC mechanism on $\mathcal{R}_s^N / \mathcal{R}_l^N$, is neither *Pareto efficient* nor *group strategy-proof*.

Example 1 (*tTTC is neither Pareto efficient nor group strategy-proof*).

Consider the market with $N = \{1, 2\}$, $T = \{H(ouse), C(ar)\}$, $O = \{H_1, H_2, C_1, C_2\}$, and where each agent i 's endowment is (H_i, C_i) . The preference profile $R \in \mathcal{R}_i^N$ is as follows:¹²

$$R_1 : H_2, \mathbf{H}_1, \mathbf{C}_1, C_2,$$

$$R_2 : C_1, \mathbf{C}_2, \mathbf{H}_2, H_1.$$

Thus, agent 1, who primarily cares for houses, would like to trade houses but not cars and agent 2, who primarily cares about cars, would like to trade cars but not houses. One easily verifies that $tTTC(R) = ((H_1, C_1), (H_2, C_2))$, the no-trade allocation. However, note that since preferences are lexicographic, both agents would be strictly better off if they traded cars and houses. Thus, allocation $((H_2, C_2), (H_1, C_1))$ Pareto dominates $tTTC(R)$. Hence, $tTTC$ is not *Pareto efficient*. Furthermore, assume that both agents (mis)report their preferences as follows:

$$R'_1 : H_2, \mathbf{H}_1, C_2, \mathbf{C}_1,$$

$$R'_2 : C_1, \mathbf{C}_2, H_1, \mathbf{H}_2.$$

Then, $tTTC(R') = ((H_2, C_2), (H_1, C_1))$, making both agents better off compared to $tTTC(R)$. Hence, $tTTC$ is not *group strategy-proof*. Finally, note that

$$tTTC_1(R_1, R'_2) = (H_2, C_1) P_1 (H_2, C_2) = tTTC_1(R')$$

and

$$tTTC_2(R'_1, R_2) = (H_2, C_1) P_2 (H_1, C_1) = tTTC_2(R'),$$

and hence R' is not a manipulation in a self-enforcing manner; the tTTC mechanism does not violate *self-enforcing group (pairwise) strategy-proofness* (Proposition 1). Finally, we note that since tTTC is also *onto* (Proposition 1), it follows that *self-enforcing group strategy-proofness* and *ontoness* do not imply *Pareto efficiency* (which is known to be implied by *group strategy-proofness* and *ontoness*). $\diamond$

While Example 1 shows that tTTC mechanism is not *Pareto efficient*, Klaus (2008) shows that it is *second-best incentive compatible*, i.e., there exists no other *strategy-proof* mechanism that *Pareto dominates* the tTTC mechanism. At the end of her paper, Klaus (2008) presents a mechanism for classical housing markets that is different from the TTC mechanism and satisfies *individual rationality*, *second-best incentive compatibility*, and *strategy-proofness*. This mechanism can be extended to multiple-type housing markets by applying it typewise; thus, the tTTC mechanism is not the unique mechanism that satisfies these properties.

¹²In all examples we indicate endowments in boldface.

Example 1 also shows that the tTTC mechanism does not satisfy the three properties that are used in Result 3. Is there another mechanism that does satisfy the three properties? The following result gives an answer in the negative: there is no mechanism that satisfies *Pareto efficiency*, *individual rationality*, and *strategy-proofness*, neither on the domain of separable preference profiles nor on the domain of lexicographic preference profiles.

Result 5 (Impossible trinity).

- (a) *Let $m > 1$. There is no mechanism on $\mathcal{R}_s^N$ that is Pareto efficient, individually rational, and strategy-proof (Konishi et al., 2001, Proposition 4.1).*
- (b) *Let $m > 1$. There is no mechanism on $\mathcal{R}_l^N$ that is Pareto efficient, individually rational, and strategy-proof (Sikdar et al., 2017, Theorem 2).*

Result 5 implies that there is no other mechanism that does better than the tTTC mechanism by satisfying the three properties on either the domain of separable preference profiles or the domain of lexicographic preference profiles. However, the tTTC mechanism on $\mathcal{R}_s^N$ ($\mathcal{R}_l^N$, respectively) does satisfy all the properties used in Result 4. In the next section we answer the question if Takamiya's characterization of the TTC mechanism for Shapley-Scarf housing markets can be extended to characterize the tTTC mechanism for multiple-type housing markets.

Finally, Proposition 1 and Example 1 also demonstrate that the equivalence of *strategy-proofness* and *non-bossiness* with *group strategy-proofness* (Result 2) does not extend to multiple-type housing markets with separable or lexicographic preferences (because *strategy-proofness* and *non-bossiness* do not imply *group strategy-proofness*).

3 Characterizing the tTTC mechanism

From now on, we focus on the multiple-type extension of the Shapley-Scarf housing market model as introduced by Moulin (1995) with more than 1 agent and more than 1 type, i.e., $|N| = n > 1$ and $|T| = m > 1$.¹³

3.1 Characterizing the tTTC mechanism for lexicographic preferences

We first show that Takamiya's result (Takamiya, 2001, Corollary 4.16) can indeed be extended to characterize the tTTC mechanism for lexicographic preferences.

¹³One agent multiple-type housing market problems are rather trivial since no trade occurs and for just one object type, we are back to the Shapley-Scarf housing market model.

Theorem 1. *A mechanism on $\mathcal{R}_i^N$ is*

- *unanimous (or onto),*
- *individually rational,*
- *strategy-proof, and*
- *non-bossy*

if and only if it is the tTTC mechanism.

From Proposition 1 it follows that the tTTC mechanism satisfies *unanimity* (or *ontoness*), *individual rationality*, *strategy-proofness*, and *non-bossiness*. Next, we explain the uniqueness part of the proof; the full proof that there is no other mechanism that satisfies the above properties is relegated to Appendices B and C.

First, we establish several auxiliary results for a mechanism f satisfying the properties of Theorem 1 (Appendix B): invariance of f under (Maskin) monotonic transformations (Lemma 3) and *marginal individual rationality* (Lemma 4). Next, we assume that a mechanism f that is not equal to the tTTC mechanism, but has the same properties, exists (Appendix C). We then obtain a contradiction via a well-constructed sequence of preference profiles (by using the lexicographic nature of preferences).

Lemma 2 implies the following corollary.

Corollary 1. *A mechanism on $\mathcal{R}_i^N$ is*

- *unanimous (or onto),*
- *individually rational, and*
- *self-enforcing group (or pairwise) strategy-proof*

if and only if it is the tTTC mechanism.

Note that even if one does not consider the domain of lexicographic preference profiles as an interesting or relevant preference profile domain for multiple-type housing markets, Theorem 1 serves as an important stepping stone to establish the corresponding characterization of the tTTC mechanism for separable preferences, see Subsection 3.2. To the best of our knowledge, the technical tool of “lifting up” a result from lexicographic preferences to separable preferences is used here for the first time.

We establish the logical independence of the properties in Theorem 1 (Corollary 1) in Appendix D.

3.2 Characterizing the tTTC mechanism for separable preferences

Note that for lexicographic preferences, under the tTTC mechanism, the importance order of types plays no role because the allocation of each type only depends on the agents' marginal preferences of each type, i.e., for each market R and type t , $tTTC^t(R) = TTC(R_1^t, \dots, R_n^t)$. Thus, one could conjecture that Theorem 1 also holds for separable preferences. This conjecture is correct.

Theorem 2. *A mechanism on $\mathcal{R}_s^N$ is*

- *unanimous (or onto),*
- *individually rational,*
- *strategy-proof, and*
- *non-bossy*

if and only if it is the tTTC mechanism.

From Proposition 1 it follows that the tTTC mechanism on $\mathcal{R}_s^N$ satisfies *unanimity* (or *ontoness*), *individual rationality*, *strategy-proofness*, and *non-bossiness*. Next, we explain the uniqueness part of the proof; the full proof that there is no other mechanism that satisfies the above properties is relegated to Appendix E.

The uniqueness part of the proof works as follows. We assume that a mechanism is *unanimous* (or *onto*), *individually rational*, *strategy-proof*, and *non-bossy*. By Theorem 1, we know that if all agents happen to have lexicographic preferences, then the tTTC allocation is assigned. Next, we consider a preference profile such that only one agent has separable and non-lexicographic preferences. We show that for this agent, if he (mis)reports lexicographic preferences without changing his marginal preferences, then he must receive the same allotment. According to Theorem 1, the allotment (in fact, the whole allocation) then equals the tTTC allotment (allocation). Hence, f assigns the tTTC allocation if all but one agent have lexicographic preferences. By applying this preference replacement argument, one by one, for all other agents, we conclude that f equals the tTTC mechanism on the domain of separable preference profiles.

Lemma 2 implies the following corollary.

Corollary 2. *A mechanism on $\mathcal{R}_s^N$ is*

- *unanimous (or onto),*
- *individually rational, and*
- *self-enforcing group (or pairwise) strategy-proof*

if and only if it is the *tTTC* mechanism.

The examples in Appendix D are well-defined on the domain of separable preference profiles and establish the logical independence of the properties in Theorem 2 (Corollary 2).

Having established Theorem 2 on the domain of separable preference profiles, a natural question that emerges is whether it is a maximal domain for the existence of a mechanism that satisfies all properties, i.e., where all properties are compatible. While we do not have an answer to this question and leave it for future research, we note that there are economically interesting domains that contain non-separable preferences and on which there exists a mechanism that satisfies all properties. For instance, consider the market with $N = \{1, 2, 3\}$, $T = \{H(ouse), C(ar)\}$, $O = \{H_1, H_2, H_3, C_1, C_2, C_3\}$, and where each agent i 's endowment is (H_i, C_i) . Let agents have (and report) preferences where they primarily care about houses. However, each agent's preferences over cars is allowed to depend on the house he receives. For instance,

$$R_1 : (H_2, C_2), (H_2, C_1), (H_2, C_3), (H_3, C_3), (H_3, C_2), (H_3, C_1), (H_1, C_1), (H_1, C_2), (H_1, C_3).$$

Obviously, agent 1's preferences R_1 are not separable. Sikdar et al. (2017, Theorem 1) implies that on this domain Sikdar et al.'s (2017) mTTC mechanism satisfies all four properties from our Theorem 2.

3.3 Impossibility results for strict preferences

Note that for $m > 1$ the tTTC mechanism is not well-defined for strict preferences since for non-separable preferences, marginal type preferences cannot be derived. Then, a natural question is if there exists an extension of the tTTC mechanism to the domain of strict preference profiles that satisfies our properties. First, observe that the impossibility trinity result (Result 5) implies that for strict preferences, no mechanism satisfies *Pareto efficiency*, *individual rationality*, and *strategy-proofness*. Our next result shows that weakening *Pareto efficiency* to *unanimity* cannot resolve this impossibility.

Theorem 3. *Let $m > 1$. Then, no mechanism on $\mathcal{R}^N$ is*

- *unanimous,*
- *individually rational, and*
- *strategy-proof.*

Proof. Without loss of generality, let $m = 2$. Suppose that there is a mechanism $f : \mathcal{R}^N \rightarrow X$ that is *unanimous*, *individually rational*, and *strategy-proof*. Let $x, y \in X \setminus \{e\}$ be such that at

x agents 1 and 2 swap their endowments of type 2, i.e.,

$$\begin{aligned} x_1 &= (o_1^1, o_2^2, o_1^3, o_1^4, \dots, o_1^m), \\ x_2 &= (o_2^1, o_1^2, o_2^3, o_2^4, \dots, o_2^m), \\ \text{and for each } i &= 3, \dots, n, \quad x_i = o_i \end{aligned}$$

and at y agents 1 and 2 swap their endowments of type 1, i.e.,

$$\begin{aligned} y_1 &= (o_2^1, o_1^2, o_1^3, o_1^4, \dots, o_1^m), \\ y_2 &= (o_1^1, o_2^2, o_2^3, o_2^4, \dots, o_2^m), \\ \text{and for each } i &= 3, \dots, n, \quad y_i = o_i. \end{aligned}$$

Obviously, $x \neq y$.

Let $R \in \mathcal{R}^N$ be such that agents 1 and 2 prefer only their allotments at x and y to their endowments, they disagree on which allocation is the better one, and each other agent ranks his endowments highest, i.e.,

$$\begin{aligned} R_1 &: x_1, y_1, o_1, \dots, \\ R_2 &: y_2, x_2, o_2, \dots, \\ \text{and for each } i &= 3, \dots, n, \quad R_i : o_i, \dots \end{aligned}$$

Note that $R \in \mathcal{R}^N \setminus \mathcal{R}_s^N$. There are only three *individually rational* allocations at R : x , y , and e .

Let

- $R'_1 : x_1, o_1, \dots,$
- $R'_2 : y_2, o_2, \dots,$
- $R''_1 : y_1, o_1, \dots,$ and
- $R''_2 : x_2, o_2, \dots$

Suppose that $f(R) = e$. Then, by *unanimity* of f , $f(R'_2, R_{-2}) = x$, which implies that agent 2 has an incentive to misreport R'_2 at R ; contradicting *strategy-proofness* of f . Therefore, $f(R) \in \{x, y\}$.

Suppose that $f(R) = x$. Then, by *strategy-proofness* of f , $f_2(R'_2, R_{-2}) \neq y_2$ and hence by *individual rationality* of f , $f(R'_2, R_{-2}) = e$. However, by *unanimity* of f , $f(R'_1, R'_2, R_{-\{1,2\}}) = y$, which implies that agent 1 has an incentive to misreport R'_1 at (R'_2, R_{-2}) ; contradicting *strategy-proofness* of f .

Suppose that $f(R) = y$. Then, by *strategy-proofness* of f , $f_1(R'_1, R_{-1}) \neq x_1$ and hence, by *individual rationality* of f , $f(R'_1, R_{-1}) = e$. However, by *unanimity* of f , $f(R'_1, R'_2, R_{-\{1,2\}}) = x$, which implies that agent 2 has an incentive to misreport R'_2 at (R'_1, R_{-1}) ; contradicting *strategy-proofness* of f . $\square$

Examples 2, 3, and 4 in Appendix D are well-defined on the domain of strict preference profiles and establish the logical independence of the corresponding properties in Theorem 3.

Our next impossibility result is established by weakening *unanimity* to *ontoness* and by adding *non-bossiness*.

Corollary 3. *Let $m > 1$. Then, no mechanism on $\mathcal{R}^N$ is*

- *onto,*
- *individually rational,*
- *strategy-proof, and*
- *non-bossy.*

Note that Lemma 2 implies that *strategy-proofness* and *non-bossiness* in Corollary 3 can be replaced by *self-enforcing group* (or *pairwise*) *strategy-proofness*. In fact, on the domain of strict preference profiles, *strategy-proofness* and *non-bossiness* imply *group strategy-proofness*.

Proof. Lemma 1 (b) together with Theorem 3 implies Corollary 3.

An alternative proof, based on Konishi et al. (2001), is as follows. Since the domain is $\mathcal{R}^N$, *strategy-proofness* and *non-bossiness* imply *group strategy-proofness*. Moreover, *group strategy-proofness* and *ontoness* imply *Pareto efficiency*. Thus, Konishi et al. (2001, Proposition 4.1) implies Corollary 3. $\square$

Examples 2, 3, and 4 in Appendix D are well-defined on the domain of strict preference profiles and establish the logical independence of *ontoness*, *individual rationality*, and *strategy-proofness* in Corollary 3. The *non-bossiness* example, Example 5, in Appendix D can be extended to strict preferences for $m = 1$; for $m > 1$ and with separable preferences, the mechanism is extended by applying it typewise to all object types. The latter extension method does not work for strict preferences and the independence of *non-bossiness* from the other properties in Corollary 3 is an **open problem** for $m > 1$.

We provide two further impossibility results that are based on extending the tTTC mechanism from lexicographic (separable) preferences to strict preferences in Appendix F.

Finally, an interesting question is whether, besides the domain of strict preferences, there is a natural and economically interesting preference domain that subsumes separable preferences; and if so, whether such preference domains have enough structure to establish again either possibility results (as in the previous section) or impossibility results (as in this section and Appendix F). We believe that these important questions are beyond the scope of the current paper, and we leave them for future research.

4 Discussion

Shapley-Scarf housing markets

Our results (Theorem 1 and Theorem 2) can be compared to Takamiya (2001, Corollary 4.16) for Shapley-Scarf housing markets. In the proof of Theorem 1 we make explicit use of the steps used by the TTC algorithm to compute the TTC allocation. In contrast, Takamiya's proof is not based on the TTC algorithm. Instead, his proof is based on *strict core-stability*,¹⁴ i.e., the absence of weak blocking coalitions and profitable coalitional deviations. His proof consists of two steps: (1) *strict core-stability* implies *group strategy-proofness* and (2) *group strategy-proofness* and *ontoness* together imply *Pareto efficiency*. Since the tTTC mechanism neither satisfies *Pareto efficiency* nor *group strategy-proofness*, our results and proof strategy are logically independent. Moreover, Takamiya's proof strategy cannot be extended to multiple-type housing markets because weak blocking coalitions and profitable coalitional deviations need not coincide (see Feng and Klaus, 2022, for details).

Furthermore, comparing the classical TTC characterization by Ma (1994) with that of Takamiya (2001) yields the following result. For Shapley-Scarf housing markets, an *individually rational* and *strategy-proof* mechanism is *Pareto efficient* if and only if it is *unanimous* and *non-bossy*. However, this result does not extend to multiple-type housing markets, as illustrated in Example 1, which shows that *tTTC* is not *Pareto efficient* (recall that there, the no-trade allocation $tTTC(R) = ((H_1, C_1), (H_2, C_2))$ is Pareto dominated by the full-trade allocation $((H_2, C_2), (H_1, C_1))$).

Object allocation problems with multi-demand and without ownership

Our results can be compared to Monte and Tumennasan (2015) and Pápai (2001) for object allocation problems with multi-demand and without ownership, i.e., agents can consume more than one object, and the set of objects is a social endowment.

While Monte and Tumennasan (2015) still assume that objects are of different types and agents can only consume one object of each type, Pápai (2001) imposes no consumption restriction.¹⁵ Although both models are slightly different, their characterization results are similar: the only mechanisms satisfying *Pareto efficiency*, *strategy-proofness*, and *non-bossiness* are sequential dictatorships. Clearly, if agents, like in our model, have property rights, sequential dictatorship mechanisms will not satisfy *individual rationality*. Thus, their characterization results imply an impossibility result for our model, in line with our Theorem 3; however, note that our efficiency notion in Theorem 3, *unanimity*, is weaker than *Pareto efficiency*.

¹⁴A mechanism is *strictly core-stable* if it always assigns a strict core allocation.

¹⁵In Pápai (2001), agents can consume any set of objects, and their preferences are linear orders over all sets of objects.

Object allocation problems with multi-demand and with ownership

Finally, we compare our results (Theorems 1 and 2) to Pápai (2003).

Similarly to Pápai (2001), Pápai (2003) considers a more general model of allocating objects to the set of agents who can consume any set of objects. In contrast to Pápai (2001), each object now is owned by an agent and each agent has strict preferences over all objects, and his preferences over sets of objects are monotonically responsive to these “objects-preferences.”¹⁶ In our model, we impose more structure by assuming that (i) the set of objects is partitioned into sets of exogenously given types and (ii) each agent owns and wishes to consume one object of each type.

Pápai (2003) considers *strategy-proofness* and *non-bossiness* (as we do) and she introduces two additional (non-standard) properties: *trade sovereignty* and *strong individual rationality*. *Trade sovereignty* requires that every feasible allocation that consists of “admissible transactions” should be realized at some preference profile; it allows for trade restrictions and some objects never being traded and is hence weaker than *ontoness* (for details see Pápai, 2003). *Strong individual rationality* requires that for each agent and all preference relations with the same objects-preferences as the agent has, *individual rationality* holds (for details see Pápai, 2003). Note that *strong individual rationality* is stronger than *individual rationality*. For instance, if agent 1’s endowment is (H_1, C_1) , and his objects-preferences are $R_1 : H_2, H_1, C_1, C_2$, then allotment (H_2, C_2) is not *strongly individually rational*.¹⁷

Pápai (2003) shows that the set of mechanisms satisfying *trade sovereignty*, *strong individual rationality*, *strategy-proofness*, and *non-bossiness* coincides with the set of segmented trading cycle mechanisms. In this class of mechanisms, all objects are (endogenously) decomposed into different segments that can be expressed as the components of a trading possibility graph (which can express trading restrictions that can even mean that certain objects cannot be traded). Agents can own at most one object per segment and the TTC algorithm is then executed separately for each segment. The set of segmented trading cycle mechanisms is large and, for our model, would include the tTTC mechanism, the no-trade mechanism, and many segmented trading cycles mechanisms with restricted trades.

The tTTC mechanism is a specific segmented trading cycle mechanism in the sense that all

¹⁶Formally, let O be a finite set of objects. A preference relation $\succeq$ over all non-empty sets of objects is *monotonically responsive* if (i) it is monotonic, i.e., for any two non-empty subsets of objects, $A, B \subseteq O$, $A \subseteq B$ implies that $B \succeq A$; and (ii) *responsive*, i.e., there exists a strict “objects-preference relation” over all objects, R , such that for any two distinct objects $o, o' \in O$, and a subset of objects $A \subseteq O \setminus \{o, o'\}$, $o P o'$ implies that $\{o\} \cup A \succ \{o'\} \cup A$. In our model, since agents’ allotments have a fixed number of objects, monotonicity of preferences over sets of objects plays no role. Furthermore, given our constraint that each agent needs to receive an object of each type, responsiveness corresponds to separability.

¹⁷Let $\tilde{\succ}_1 : (H_2, C_1), (H_1, C_1), (H_2, C_2), (H_1, C_2)$ and $\hat{\succ}_1 : (H_2, C_1), (H_2, C_2), (H_1, C_1), (H_1, C_2)$. Note that both preferences are responsive to R_1 . We see that $(H_2, C_2) \hat{\succ}_1 (H_1, C_1)$ but $(H_1, C_1) \tilde{\succ}_1 (H_2, C_2)$. Thus, (H_2, C_2) is *individually rational* at $\hat{\succ}_1$ but *not individually rational* at $\tilde{\succ}_1$.

segments are a priori determined by object types. Thus, our characterization result of the tTTC mechanism can be seen as characterizing a specific segmented trading cycle mechanism while Pápai characterizes the whole class of segmented trading cycle mechanisms. On the one hand, we weaken *strong individual rationality* to *individual rationality* but strengthen *trade sovereignty* to *ontoness*. On the other hand, we consider two different preference domains that reflect some responsiveness through separability. Therefore, while there is a close connection between our models and results, there is no direct logical relation between Pápai (2003)'s result and ours (Theorems 1 and 2).

Appendix

A Proof of Lemmas 1 and 2 and Proposition 1

Proof of Lemma 1. (a) Let f on $\mathcal{R}^N$ be *unanimous*. Fix any allocation $x \in X$. Let $R \in \mathcal{R}^N$ be a preference profile such that x is *unanimously* best under R . Then, by *unanimity* of f , $f(R) = x$. Hence, f is an *onto* function.

(b) Let f on $\mathcal{R}^N$ be *strategy-proof*, *non-bossy*, and *onto*. Let $x \in X$ and $R \in \mathcal{R}^N$ be a preference profile such that x is unanimously best under R . By *ontoness* of f , there exists a preference profile $R' \in \mathcal{R}$ such that $f(R') = x$. Let $i \in N$ and $y = f(R_i, R'_{-i})$. By *strategy-proofness* of f , we have $y_i R_i x_i$. Since x_i is agent i 's most preferred allotment, we have $y_i = x_i$. Then, by *non-bossiness* of f , we have $f(R_i, R'_{-i}) = y = x = f(R')$. By applying this argument repeatedly for all agents in $N \setminus \{i\}$, we find that $f(R) = x = f(R')$. So, f is *unanimous*. $\square$

Proof of Lemma 2. Note that (iii) $\Rightarrow$ (ii) is immediate.

We first prove (ii) $\Rightarrow$ (i). It is easy to see that *self-enforcing pairwise strategy-proofness* implies *strategy-proofness* (by considering $S = \{i\}$). We will show that *self-enforcing pairwise strategy-proofness* also implies *non-bossiness*. Suppose that a mechanism f is *self-enforcing pairwise strategy-proof* and *bossy*. Then there exist $R \in \mathcal{R}^N$, $i, j \in N$, and $R'_i \in \mathcal{R}$ such that $f_i(R) = f_i(R'_i, R_{-i})$ and $f_j(R) \neq f_j(R'_i, R_{-i})$. Let $S = \{i, j\}$ and let $R'_j = R_j$. Then, $(R'_S, R_{-S}) = (R'_i, R_{-i})$ and $f_i(R'_S, R_{-S}) R_i f_i(R)$. Moreover, since $f_j(R'_S, R_{-S}) \neq f_j(R)$, we can assume, without loss of generality, that $f_j(R'_S, R_{-S}) P_j f_j(R)$. This proves that coalition $S = \{i, j\}$ can manipulate f at R via deviation (R'_i, R'_j) . Finally, note that $f_i(R'_S, R_{-S}) R_i f_i(R_i, R'_j, R_{-\{i,j\}})$ is equivalent to $f_i(R'_S, R_{-S}) R_i f_i(R)$ and that $f_j(R'_S, R_{-S}) R_j f_j(R'_i, R_j, R_{-\{i,j\}})$ is equivalent to $f_j(R'_S, R_{-S}) R_j f_j(R'_S, R_{-S})$, both of which are satisfied. Thus, $S = \{i, j\}$ can manipulate f in a self-enforcing manner at R , which contradicts our assumption that f is *self-enforcing pairwise strategy-proof*.

Finally, we will show (i) $\Rightarrow$ (iii). Let f be a *strategy-proof* and *non-bossy* mechanism. By contradiction, assume that there exists a coalition $S \subseteq N$ that can manipulate f at $R \in$

$\mathcal{R}^N$ via deviation $R'_S \in \mathcal{R}^S$ in a self-enforcing manner. Without loss of generality, we can assume that S is such that there is no proper coalition $T \subsetneq S$ such that T can manipulate f at R via deviation R'_T in a self-enforcing manner. Let $j \in S$ such that $f_j(R'_S, R_{-S}) P_j f_j(R)$. Note that since f is *strategy-proof*, $S \setminus \{j\} \neq \emptyset$. Let $i \in S \setminus \{j\}$. Then, since deviation R'_S is self-enforcing, $f_i(R'_S, R_{-S}) R_i f_i(R_i, R'_{S \setminus \{i\}}, R_{-S})$. Since f is *strategy-proof*, this implies that $f_i(R'_S, R_{-S}) = f_i(R_i, R'_{S \setminus \{i\}}, R_{-S})$ and thus, by *non-bossiness*, $f(R'_S, R_{-S}) = f(R_i, R'_{S \setminus \{i\}}, R_{-S})$. Let $T = S \setminus \{i\}$. Then, since $j \in T$, coalition T is a manipulating coalition at R via R'_T . In fact, the deviation R'_T at R is self-enforcing. To see this, let $W \subsetneq T$ and $V = W \cup \{i\}$. Note that $V \subsetneq S$. Then, since deviation R'_S is self-enforcing, for each $\ell \in W$,

$$f_\ell(R'_T, R_{N \setminus T}) = f_\ell(R'_S, R_{N \setminus S}) R_\ell f_\ell(R_V, R'_{S \setminus V}, R_{N \setminus S}).$$

Moreover, since $V = W \cup \{i\}$ and $S = T \cup \{i\}$,

$$f(R_V, R'_{S \setminus V}, R_{N \setminus S}) = f(R_W, R_i, R'_{S \setminus (W \cup \{i\})}, R_{N \setminus S}) = f(R_W, R'_{T \setminus W}, R_{N \setminus T}).$$

Hence, it follows that for each $\ell \in W$,

$$f_\ell(R'_T, R_{N \setminus T}) R_\ell f_\ell(R_W, R'_{T \setminus W}, R_{N \setminus T}).$$

Therefore, $T \subsetneq S$ can manipulate f at R via deviation R'_T in a self-enforcing manner. This contradicts the minimality of S . $\square$

Proof of Proposition 1. It is straightforward to check that the tTTC mechanism on $\mathcal{R}_s^N$ is *individually rational* and *unanimous* (and hence *onto*).

We next show that the tTTC mechanism on $\mathcal{R}_s^N$ inherits *strategy-proofness* from the TTC mechanism. Let $R \in \mathcal{R}_s^N$, $i \in N$, and $\hat{R}_i \in \mathcal{R}_s$ with marginal preferences $(\hat{R}_i^1, \dots, \hat{R}_i^m)$. By the definition and *strategy-proofness* of the TTC mechanism, for each $t \in T$, $tTTC_i^t(R) = TTC_i^t(R^t) R_i^t TTC_i^t(\hat{R}_i^t, R_{-i}^t) = tTTC_i^t(\hat{R}_i, R_{-i})$. Then, by the separability of preferences, we have $tTTC_i(R) R_i tTTC_i(\hat{R}_i, R_{-i})$ and the tTTC mechanism is *strategy-proof*.

Finally, to show that the tTTC mechanism on $\mathcal{R}_s^N$ is *non-bossy*, let $R \in \mathcal{R}_s^N$, $i \in N$, and $\hat{R}_i \in \mathcal{R}_s$, with marginal preferences $(\hat{R}_i^1, \dots, \hat{R}_i^m)$, be such that $tTTC_i(R) = tTTC_i(\hat{R}_i, R_{-i})$. Thus, for each $t \in T$, $tTTC_i^t(R) = tTTC_i^t(\hat{R}_i, R_{-i})$. Moreover, by definition of the tTTC mechanism, we have for each $t \in T$, $tTTC_i^t(R) = TTC_i^t(R^t)$ and $tTTC_i^t(\hat{R}_i, R_{-i}) = TTC_i^t(\hat{R}_i^t, R_{-i}^t)$. Thus, for each $t \in T$, $TTC_i(R^t) = TTC_i(\hat{R}_i^t, R_{-i}^t)$, and since the TTC mechanism is *non-bossy*, we have that for each $t \in T$, $TTC(R^t) = TTC(\hat{R}_i^t, R_{-i}^t)$. Then, for each $t \in T$, $tTTC^t(R) = tTTC^t(\hat{R}_i, R_{-i})$. Thus, $tTTC(R) = tTTC(\hat{R}_i, R_{-i})$ and the tTTC mechanism is *non-bossy*.

Since the tTTC mechanism on $\mathcal{R}_s^N$ is *strategy-proof* and *non-bossy*, by Lemma 2, it also satisfies *self-enforcing group (pairwise) strategy-proof*. $\square$

B Auxiliary properties and results

In this appendix, we introduce auxiliary properties and obtain results that are key for the proof of Theorem 1 (Appendix C). While some of the results below can also be proven for separable preferences, we focus on lexicographic preferences because Theorem 1 deals with lexicographic preferences.

We introduce the well-known property of (*Maskin*) *monotonicity*, which requires that if an allocation is chosen, then that allocation will still be chosen if each agent shifts it up in his preferences.

Let $i \in N$. Given preferences $R_i \in \mathcal{R}_i$ and an allotment x_i , let $L(x_i, R_i) = \{y_i \in \Pi_{t \in T} O^t \mid x_i R_i y_i\}$ be the *lower contour set* of R_i at x_i . Preference relation $R'_i \in \mathcal{R}_i$ is a *monotonic transformation* of R_i at x_i if $L(x_i, R_i) \subseteq L(x_i, R'_i)$. Similarly, given a preference profile $R \in \mathcal{R}_i^N$ and an allocation x , a preference profile $R' \in \mathcal{R}_i^N$ is a *monotonic transformation* of R at x if for each $i \in N$, R'_i is a monotonic transformation of R_i at x_i .

Definition 11 (Monotonicity).

A mechanism f on $\mathcal{R}_i^N$ is *monotonic* if for each $R \in \mathcal{R}_i^N$ and for each monotonic transformation $R' \in \mathcal{R}_i^N$ of R at $f(R)$, we have $f(R') = f(R)$.

We show that *strategy-proofness* and *non-bossiness* imply *monotonicity*.

Lemma 3. *If a mechanism on $\mathcal{R}_i^N$ is strategy-proof and non-bossy, then it is monotonic.*

Proof. The proof is a straightforward extension of Takamiya (2001, Theorem 4.12) and Pápai (2001, Lemma 1). Suppose mechanism f on $\mathcal{R}_i^N$ is *strategy-proof* and *non-bossy*. Let $R \in \mathcal{R}_i^N$ and let $x = f(R)$. Let $R' \in \mathcal{R}_i^N$ be a monotonic transformation of R at x . Let $i \in N$ and $y = f(R'_i, R_{-i})$. By *strategy-proofness* of f , we have $x_i R_i y_i$, which implies that $y_i \in L(x_i, R_i) \subseteq L(x_i, R'_i)$. However, by *strategy-proofness* of f , we also have $y_i R'_i x_i$. Thus, since $y_i \in L(x_i, R'_i)$, $x_i = y_i$. Then, by *non-bossiness* of f , we have $x = y$. By applying this argument sequentially for all agents in $N \setminus \{i\}$, we find that $f(R) = x = f(R')$. $\square$

The converse of Lemma 3 is not true: lexicographic preferences are not rich enough to satisfy Alva's (2017) preference domain richness condition *two-point connectedness*.

Next, we introduce a “marginal version” of monotonic preference transformations. Let $i \in N$. Given preferences $R_i \in \mathcal{R}_i$ and an allotment x_i , for each type t , consider the associated marginal preferences R_i^t and marginal allotment x_i^t . Let $L(x_i^t, R_i^t) = \{y_i^t \in O^t \mid x_i^t R_i^t y_i^t\}$ be the lower contour set of R_i^t at x_i^t . Marginal preference relation $\hat{R}_i^t$ is a monotonic transformation of R_i^t at x_i^t if $L(x_i^t, R_i^t) \subseteq L(x_i^t, \hat{R}_i^t)$.

Fact 1. Let x_i be an allotment. Let $R_i, \hat{R}_i$ be lexicographic preferences such that (1) $\pi_i = \hat{\pi}_i$ and (2) for each $t \in T$, $\hat{R}_i^t$ is a monotonic transformation of R_i^t at x_i^t . Then, $\hat{R}_i$ is a monotonic transformation of R_i at x_i .

Proof. We show that $L(x_i, R_i) \subseteq L(x_i, \hat{R}_i)$. Let $y_i \in L(x_i, R_i)$ with $y_i \neq x_i$. Then, $x_i P_i y_i$. Restate y_i and x_i as $y_i^{\pi_i} = (y_i^{\pi_i(1)}, \dots, y_i^{\pi_i(m)})$ and $x_i^{\pi_i} = (x_i^{\pi_i(1)}, \dots, x_i^{\pi_i(m)})$, respectively. Let k be the first type for which x_i and y_i assign different objects, i.e., for all $l < k$, $y_i^{\pi_i(l)} = x_i^{\pi_i(l)}$ and $y_i^{\pi_i(k)} \neq x_i^{\pi_i(k)}$. Since $x_i P_i y_i$ and preferences are lexicographic, we have $x_i^{\pi_i(k)} P_i^{\pi_i(k)} y_i^{\pi_i(k)}$. Thus, $y_i^{\pi_i(k)} \in L(x_i^{\pi_i(k)}, R_i^{\pi_i(k)}) \subseteq L(x_i^{\pi_i(k)}, \hat{R}_i^{\pi_i(k)})$, which implies that $x_i^{\pi_i(k)} \hat{P}_i^{\pi_i(k)} y_i^{\pi_i(k)}$. Then, since $\pi_i = \hat{\pi}_i$, $x_i \hat{P}_i y_i$, i.e., $y_i \in L(x_i, \hat{R}_i)$. $\square$

Therefore, by *monotonicity*, if an agent receives an allotment and shifts each of its objects up in the marginal preferences (without changing his importance order), he still receives that allotment and the allotments of the other agents do not change either.

Next, for lexicographic preferences, we introduce a new property, *marginal individual rationality*, which is a stronger property than *individual rationality*.

Definition 12 (Marginal individual rationality).

A mechanism f on $\mathcal{R}_i^N$ is *marginally individually rational* if for each $R \in \mathcal{R}_i^N$, each $i \in N$, and each $t \in T$, $f_i^t(R) R_i^t o_i^t$.

Lemma 4. *If a mechanism on $\mathcal{R}_i^N$ is unanimous, individually rational, strategy-proof, and non-bossy, then it is marginally individually rational.*

Proof. Suppose mechanism f on $\mathcal{R}_i^N$ is *unanimous, individually rational, strategy-proof, non-bossy*, and not *marginally individually rational*, i.e., there exist a preference profile $R \in \mathcal{R}_i^N$, an agent $i \in N$, and a type $t \in T$ such that $o_i^t P_i^t f_i^t(R)$. Then, by *individual rationality* of f , we know that $t \neq \pi_i(1)$.

Let $x \equiv f(R)$. Consider a preference profile $\hat{R} \in \mathcal{R}_i^N$ such that

for agent i ,

- $\hat{R}_i^t : o_i^t, x_i^t, \dots$,
- for each $\tau \in T \setminus \{t\}$, $\hat{R}_i^\tau : x_i^\tau, \dots$, and
- $\hat{\pi}_i = \pi_i$;

and for each agent $j \in N \setminus \{i\}$,

- for each $\tau \in T$, $\hat{R}_j^\tau : x_j^\tau, \dots$, and
- $\hat{\pi}_j = \pi_j$.

Note that, by Fact 1, $\hat{R}$ is a monotonic transformation of R at x . By Lemma 3, f is *monotonic*. Thus, $f(\hat{R}) = x$.

Next, consider a preference profile $(\bar{R}_i, \hat{R}_{-i}) \in \mathcal{R}_i^N$, where $\bar{R}_i$ is such that

- for each $\tau \in T$, $\bar{R}_i^\tau = \hat{R}_i^\tau$, and
- $\bar{\pi}_i(1) = t$.

Note that $\bar{R}_i$ can be interpreted as a linear order over all objects such that $\bar{R}_i : o_i^t, \dots$, i.e., object o_i^t is the most preferred object.

Let $y \equiv f(\bar{R}_i, \hat{R}_{-i})$. By *individual rationality* of f , $y_i^t = o_i^t$. Thus, $y_i \neq x_i$. By *strategy-proofness* of f , $x_i = f(\hat{R}_i, \hat{R}_{-i}) \hat{P}_i f(\bar{R}_i, \hat{R}_{-i}) = y_i$. Since agent i gains in type t by misreporting at $\hat{R}_i$ (i.e., $y_i^t = o_i^t \hat{P}_i^t f_i^t(\hat{R}) = x_i^t$), he must lose in some other more important type according to $\hat{\pi}_i$. That is, there is a type $t' \neq t$ such that (1) $\hat{\pi}_i^{-1}(t') < \hat{\pi}_i^{-1}(t)$ and (2) $x_i^{t'} \hat{P}_i^{t'} y_i^{t'}$. In particular, $x_i^{t'} \neq y_i^{t'}$.

Next, consider a preference profile $\bar{R} \equiv (\bar{R}_i, \bar{R}_{-i})$ such that

for each agent $j \in N \setminus \{i\}$,

- $\bar{R}_j^t : y_j^t, \dots$,
- for each $\tau \in T \setminus \{t\}$, $\bar{R}_j^\tau = \hat{R}_j^\tau$, and
- $\bar{\pi}_j = \hat{\pi}_j$.

Note that the only relevant difference between $\bar{R}$ and $(\bar{R}_i, \hat{R}_{-i})$ is that under $\bar{R}$, each agent $j \neq i$ positions y_j^t as his most preferred type- t object. Thus, $\bar{R}$ is a monotonic transformation of $(\bar{R}_i, \hat{R}_{-i})$ at y . Therefore, by *monotonicity* of f , $f(\bar{R}) = y$.

However, under $\bar{R}$, for each agent $k \in N$, his most preferred allotment is $z_k = (x_k^1, \dots, x_k^{t-1}, y_k^t, x_k^{t+1}, \dots, x_k^m)$. Note that $z = (z_k)_{k \in N} \in X$ is an allocation because z is a mixture of y (for type t) and x (for other types). Thus, by *unanimity* of f , $f(\bar{R}) = z$. So, $y = z$. However, for type t' , $z_i^{t'} = x_i^{t'} \neq y_i^{t'}$, a contradiction. $\square$

C Proof of Theorem 1: uniqueness

Proof of Theorem 1: uniqueness. Suppose that there is a mechanism $f : \mathcal{R}_i^N \rightarrow X$, different from the $tTTC$ mechanism, that satisfies the properties listed in Theorem 1 (by Lemma 1, *onteness* and *unanimity* can be used interchangeably). Then, there is a market R such that $y \equiv f(R) \neq tTTC(R) \equiv x$. In particular, there is a type t such that $(y_1^t, \dots, y_n^t) \neq (x_1^t, \dots, x_n^t)$.

By Lemma 3, both mechanisms, f and $tTTC$, are *monotonic*. By Lemma 4, both mechanisms, f and $tTTC$, are *marginally individually rational*. Since both mechanisms are *marginally individually rational*, for each $i \in N$ and each $\tau \in T$, $y_i^\tau R_i^\tau o_i^\tau$ and $x_i^\tau R_i^\tau o_i^\tau$. So, we can define a preference profile $\hat{R} \in \mathcal{R}_i^N$ such that

for each agent $i \in N$,

- $\hat{R}_i^t : \begin{cases} x_i^t, y_i^t, o_i^t, \dots & \text{if } x_i^t R_i^t y_i^t \\ y_i^t, x_i^t, o_i^t, \dots & \text{if } y_i^t R_i^t x_i^t \end{cases}$
- for each $\tau \in T \setminus \{t\}$, $\hat{R}_i^\tau : y_i^\tau, o_i^\tau, \dots$, and
- $\hat{\pi}_i = \pi_i$.

Note that, by Fact 1, $\hat{R}$ is a monotonic transformation of R at y . Since f is *monotonic*, $f(\hat{R}) = y$. Furthermore, since $\hat{R}^t$ is a monotonic transformation of R^t at x^t , *monotonicity* of the *TTC* mechanism implies $tTTC^t(\hat{R}) = TTC(\hat{R}^t) = x^t$.

Next, consider a preference profile $\bar{R} \in \mathcal{R}_I^N$ such that

for each agent $i \in N$,

- $\bar{R}_i^t : x_i^t, o_i^t, \dots$,
- for each $\tau \in T \setminus \{t\}$ $\bar{R}_i^\tau = \hat{R}_i^\tau$, and
- $\bar{\pi}_i = \pi_i$.

Note that the only relevant difference between $\bar{R}$ and $\hat{R}$ is that under $\bar{R}$, each agent $i \in N$ positions x_i^t as his most preferred type- t object and his endowment o_i^t as his second preferred.

Under $\bar{R}$, each agent i 's most preferred allotment is $z_i \equiv (y_i^1, \dots, y_i^{t-1}, x_i^t, y_i^{t+1}, \dots, y_i^m)$. Note that $z = (z_i)_{i \in N} \in X$ is an allocation because z is a mixture of x (for type t) and y (for other types). Thus, by *unanimity* of f , $f(\bar{R}) = z$.

Recall that since $(x_1^t, \dots, x_n^t) = tTTC^t(\hat{R}) = TTC(\hat{R}^t)$, $(x_1^t, \dots, x_n^t)$ is obtained by applying the *TTC* algorithm to preference profile $\hat{R}^t$. For each $i \in N$, let s_i be the step of the *TTC* algorithm at which agent i receives object x_i^t . Without loss of generality, assume that if $i < i'$ then $s_i \leq s_{i'}$.

Next, we will show that $f(\hat{R}) = z$ by using that $f(\bar{R}) = z$ and replacing, step-by-step, each $\bar{R}_i$ with $\hat{R}_i$. More specifically, we will replace the individual preferences in the order $n, n-1, \dots, 1$.

We first show that $f(\bar{R}_{-n}, \hat{R}_n) = z$. Suppose $x_n^t \hat{R}_n^t y_n^t$. Then, $(\bar{R}_{-n}, \hat{R}_n)$ is a monotonic transformation of $\bar{R}$ at z . By *monotonicity* of f , $f(\bar{R}_{-n}, \hat{R}_n) = f(\bar{R}) = z$.

Now suppose $y_n^t \hat{P}_n^t x_n^t$. Let $\tau \in T$ such that $\pi_n(\tau) = 1 < \pi_n(t)$ (if $\pi_n(t) = 1$, then skip this step). Since f is *strategy-proof*, preferences are lexicographic, and τ is the most important type for agent n , we have $f_n^\tau(\bar{R}_{-n}, \hat{R}_n) \hat{R}_n^\tau f_n^\tau(\bar{R})$. Since $\tau \neq t$, $f_n^\tau(\bar{R}) = z_n^\tau = y_n^\tau$ and $f_n^\tau(\bar{R}_{-n}, \hat{R}_n) \hat{R}_n^\tau y_n^\tau$.

Since $\tau \neq t$, it follows from the definition of $\hat{R}_n^\tau$ that y_n^τ is the best type- τ object. So, $f_n^\tau(\bar{R}_{-n}, \hat{R}_n) = y_n^\tau$. Now one can, sequentially, from more to less important types, apply similar arguments to show that

$$\text{for each type } t' \in T \text{ with } \pi_n(t') < \pi_n(t), f_n^{t'}(\bar{R}_{-n}, \hat{R}_n) = y_n^{t'} = f_n^{t'}(\bar{R}). \quad (1)$$

Since f is *marginally individually rational*, $f_n^t(\bar{R}_{-n}, \hat{R}_n) \in \{x_n^t, y_n^t, o_n^t\}$. Suppose $f_n^t(\bar{R}_{-n}, \hat{R}_n) = o_n^t$ and $o_n^t \neq x_n^t$. Then, $f_n^t(\bar{R}) = z_n^t = x_n^t \hat{P}_n^t o_n^t = f_n^t(\bar{R}_{-n}, \hat{R}_n)$, which together with (1) would contradict the *strategy-proofness* of f . Hence, $f_n^t(\bar{R}_{-n}, \hat{R}_n) \in \{x_n^t, y_n^t\}$.

Suppose that $f_n^t(\bar{R}_{-n}, \hat{R}_n) = y_n^t$. By the definition of the TTC algorithm, x_n^t is agent n 's most preferred type- t object among the remaining objects at Step s_n of the TTC algorithm at preference profile $\hat{R}^t$. Therefore, object y_n^t is removed (i.e., assigned to some agent) at some Step $s^* < s_n$ of the TTC algorithm at preference profile $\hat{R}^t$.

Let C be the trading cycle of the TTC algorithm at preference profile $\hat{R}^t$ that contains y_n^t . Suppose C only contains one agent, say $j \neq n$. Then, among all objects present at Step s^* , agent j most prefers his own endowment, i.e., $o_j^t = y_n^t$. Hence, $x_j^t = tTTC_j^t(\hat{R}) = TTC_j(\hat{R}^t) = y_n^t = o_j^t$. So, by definition of $\bar{R}$, we have that at $(\bar{R}_{-n}, \hat{R}_n)$ agent j 's marginal preferences of type t are given by $\bar{R}_j^t : o_j^t, \dots$. By *marginal individual rationality* of f , $f_j^t(\bar{R}_{-n}, \hat{R}_n) = y_n^t$, which contradicts $f_n^t(\bar{R}_{-n}, \hat{R}_n) = y_n^t$.

Hence, C consists of agents $i_1, i_2, \dots, i_K$ (with $K \geq 2$) and type- t objects $o_{i_1}^t, \dots, o_{i_K}^t$ such that $n \notin \{i_1, \dots, i_K\}$ and $y_n^t \in \{o_{i_1}^t, \dots, o_{i_K}^t\}$. Without loss of generality, the cycle C is ordered $(i_1, i_2, \dots, i_K)$. Note that at $(\bar{R}_{-n}, \hat{R}_n)$, for each $i_k \in \{i_1, \dots, i_K\}$, agent i_k 's marginal preferences of type t are $\bar{R}_{i_k}^t : o_{i_{k+1}}^t (= x_{i_k}^t), o_{i_k}^t, \dots$ (modulo K). Without loss of generality, assume that $y_n^t = o_{i_1}^t$. It follows from $f_n^t(\bar{R}_{-n}, \hat{R}_n) = y_n^t$ and *marginal individual rationality* of f that $f_{i_K}^t(\bar{R}_{-n}, \hat{R}_n) = o_{i_K}^t$. Subsequently, for each agent $i_k \in \{i_2, \dots, i_K\}$, $f_{i_k}^t(\bar{R}_{-n}, \hat{R}_n) = o_{i_k}^t$. Therefore, $f_{i_1}^t(\bar{R}_{-n}, \hat{R}_n) \neq o_{i_2}^t$. Moreover, $f_{i_1}^t(\bar{R}_{-n}, \hat{R}_n) \neq o_{i_1}^t$ because $f_n^t(\bar{R}_{-n}, \hat{R}_n) = y_n^t = o_{i_1}^t$. Thus, $o_{i_1}^t \bar{P}_{i_1} f_{i_1}^t(\bar{R}_{-n}, \hat{R}_n)$, which violates *marginal individual rationality* of f . Therefore, $f_n^t(\bar{R}_{-n}, \hat{R}_n) \neq y_n^t$. Hence,

$$f_n^t(\bar{R}_{-n}, \hat{R}_n) = x_n^t = f_n^t(\bar{R}). \quad (2)$$

Having established (1) and (2), one can use arguments similar to those for (1) to show that

$$\text{for each type } t' \in T \text{ with } \pi_n(t') > \pi_n(t), f_n^{t'}(\bar{R}_{-n}, \hat{R}_n) = y_n^{t'} = f_n^{t'}(\bar{R}). \quad (3)$$

From (1), (2), and (3) it follows that for each type $\tau \in T$, $f_n^\tau(\bar{R}_{-n}, \hat{R}_n) = f_n^\tau(\bar{R})$. Hence, $f_n(\bar{R}_{-n}, \hat{R}_n) = f_n(\bar{R})$. By *non-bossiness* of f , $f(\bar{R}_{-n}, \hat{R}_n) = f(\bar{R}) = z$.

By applying repeatedly the same arguments for agents $i = n-1, \dots, 1$, we can sequentially replace each $\bar{R}_i$ with $\hat{R}_i$, and conclude that $f(\hat{R}) = f(\bar{R}) = z$. However, since $(y_1^t, \dots, y_n^t) \neq (x_1^t, \dots, x_n^t)$, there exists an agent j such that $y_j^t \neq x_j^t$. Hence, $f_j^t(\hat{R}) = y_j^t \neq x_j^t = z_j^t$, a contradiction. $\square$

D Independence of properties in Theorem 1

The following examples establish the logical independence of the properties in Theorem 1 (Corollary 1) on $\mathcal{R}_i^N$. We label the examples by the property/properties that is/are not satisfied.

Example 2 (*Ontoness and unanimity*).

The no-trade mechanism that always assigns the endowment allocation to each market is *individually rational*, (*group*) *strategy-proof*, and *non-bossy*, but neither *onto* nor *unanimous*. $\diamond$

The no-trade mechanism in Example 2 is well-defined on $\mathcal{R}_l^N$, $\mathcal{R}_s^N$, and $\mathcal{R}^N$.

Example 3 (*Individual rationality*).

By ignoring property rights that are established via the endowments, we can easily adjust the well-known mechanism of serial dictatorship to our setting: based on an ordering of agents, we let agents sequentially choose their allotments. Serial dictatorship mechanisms have been shown in various resource allocation models to satisfy *Pareto efficiency* (and hence *ontoness* and *unanimity*), *strategy-proofness*, and *non-bossiness*; since property rights are ignored, they violate *individual rationality* (e.g., see Monte and Tumennasan, 2015, Theorem 1). $\diamond$

The serial dictatorship mechanism in Example 3 is well-defined on $\mathcal{R}_l^N$, $\mathcal{R}_s^N$, and $\mathcal{R}^N$.

Example 4 (*Strategy-proofness*).

We adapt so-called Multiple-Serial-IR mechanisms introduced by Biró et al. (2022b) for their circulation model to our multiple-type housing markets model. A Multiple-Serial-IR mechanism is determined by a fixed order of the agents. At any preference profile and following the order, the mechanism lets each agent pick his most preferred allotment from the available objects such that this choice together with previous agents' choices is compatible with an *individually rational* allocation. Formally,

Input. An order $\delta = (i_1, \dots, i_n)$ of the agents and a multiple-type housing market $R \in \mathcal{R}_l^N$.

Step 0. Let $Y(0)$ be the set of individually rational allocations in X .

Step 1. Let Y_1 be the set of agent i_1 's allotments that are compatible with some allocation in $Y(0)$, i.e., Y_1 consists of all $y_{i_1} \in \Pi_{t \in T} O^t$ for which there exists an allocation $x \in Y(0)$ such that $x_{i_1} = y_{i_1}$.

Let $y_{i_1}^*$ be agent i_1 's most preferred allotment in Y_1 , i.e., for each $y_{i_1} \in Y_1$, $y_{i_1}^* R_i y_{i_1}$.

Let $Y(1) \subseteq Y(0)$ be the set of allocations in $Y(0)$ that are compatible with $y_{i_1}^*$, i.e., $Y(1)$ consists of all $x \in Y(0)$ with $x_{i_1} = y_{i_1}^*$.

Step $k = 2, \dots, n$. Let Y_k be the set of agent i_k 's allotments that are compatible with some allocation in $Y(k-1)$.

Let $y_{i_k}^*$ be agent i_k 's most preferred allotment in Y_k .

Let $Y(k) \subseteq Y(k-1)$ be the set of allocations in $Y(k-1)$ that are compatible with $y_{i_k}^*$.

Output. The allocation of the Multiple-Serial-IR mechanism associated with δ at R is $MSIR(\delta, R) \equiv (y_1^*, y_2^*, \dots, y_n^*)$.

Given an order δ , the associated Multiple-Serial-IR mechanism Δ assigns to each market R the allocation $\Delta(R) \equiv MSIR(\delta, R)$.

Biró et al. (2022b) show that Multiple-Serial-IR mechanisms are *individually rational* and *Pareto efficient*.

Next, we show that Multiple-Serial-IR mechanisms are *non-bossy*. Let $\delta = (i_1, \dots, i_n)$ be an order of the agents and let Δ denote the associated Multiple-Serial-IR mechanism.

Let $R \in \mathcal{R}_i^N$, $i \in N$, and $R'_i \in \mathcal{R}_i$. Let $R' \equiv (R'_i, R_{-i})$, $x \equiv \Delta(R)$, and $y \equiv \Delta(R')$. Assume $y_i = x_i$. We show that $y = x$.

Let $i_k \equiv i$. Since $y_i = x_i$ and for each $\ell = 2, \dots, k-1, k+1, \dots, n$, $R'_{i_\ell} = R_{i_\ell}$, agent i_1 's choice at Step 1 under R' is restricted in the same way as agent i_1 's choice at Step 1 under R . Thus, since $R'_{i_1} = R_{i_1}$, we have $y_{i_1} = x_{i_1}$. Similar arguments show that for each $\ell = 2, \dots, k-1, k+1, \dots, n$, $y_{i_\ell} = x_{i_\ell}$. Hence, Δ is non-bossy.

In the context of multiple-type housing markets, Konishi et al. (2001) show that there is no mechanism that is *Pareto efficient*, *individually rational*, and *strategy-proof*. Since Multiple-Serial-IR mechanisms are *Pareto efficient* and *individually rational*, they are not *strategy-proof*. We include a simple illustrative example for $n = 2$ agents and $m = 2$ types for completeness.

Let $N = \{1, 2\}$ and $T = \{H(ouse), C(ar)\}$. For each $i \in N$, let (H_i, C_i) be agent i 's endowment. Let $R \in \mathcal{R}_i^N$ be given by

$$\mathbf{R}_1 : H_2, \mathbf{H}_1, C_2, \mathbf{C}_1,$$

$$\mathbf{R}_2 : H_1, \mathbf{H}_2, \mathbf{C}_2, C_1.$$

Consider the Multiple-Serial-IR mechanism Δ induced by $\delta = (1, 2)$, i.e., agent 1 moves first (note that since there are only two agents, when agent 1 picks his allotment, the final allocation is completely determined). Since allocation $x \equiv ((H_2, C_2), (H_1, C_1))$ is *individually rational* at R and $x_1 = (H_2, C_2)$ is agent 1's most preferred allotment, $\Delta(R) = x$.

Next, consider $\mathbf{R}'_2 : \mathbf{C}_2, C_1, H_1, \mathbf{H}_2$. Note that at (R_1, R'_2) , only $y \equiv ((H_2, C_1), (H_1, C_2))$ and e are *individually rational*. Thus, agent 1 can only pick y_1 or o_1 . Since $y_1 R_1 o_1$, agent 1 picks y_1 and hence $\Delta(R_1, R'_2) = y$. Finally, we see that $y_2 R_2 x_2$, which implies that agent 2 has an incentive to misreport R'_2 at R . Hence, the Multiple-Serial-IR mechanism induced by $\delta = (1, 2)$ is not *strategy-proof*. $\diamond$

The mechanism in Example 4 is well-defined on $\mathcal{R}_i^N$, $\mathcal{R}_s^N$, and $\mathcal{R}^N$.

Note that if $n = 2$, then any mechanism is *non-bossy*. Thus, for our last independence example, we assume $n > 2$.

Example 5 (Non-bossiness).

We first provide an example of a mechanism for $n = 3$ and $m = 1$. Let $N = \{1, 2, 3\}$ and $T = \{H(ouse)\}$. Let $R \in \mathcal{R}^N$. We say that agents 1 and 3 are *in conflict* if H_2 is the most preferred object for both R_1 and R_3 . Similarly, we say that agents 1 and 2 are *in conflict* if H_3 is the most preferred object for both R_1 and R_2 . Let mechanism f be defined as follows: for each $R \in \mathcal{R}^N$,

- (a) if agents 1 and 2 are in conflict, then (i) transform R_2 to $\bar{R}_2$ by dropping H_3 to the bottom, i.e., $\bar{R}_2 : \dots, H_3$, while keeping the relative order of H_1 and H_2 , and (ii) set $f(R) \equiv TTC(R_1, \bar{R}_2, R_3)$;
- (b) if agents 1 and 3 are in conflict, then (i) transform R_3 to $\bar{R}_3$ by dropping H_2 to the bottom, i.e., $\bar{R}_3 : \dots, H_2$, while keeping the relative order of H_1 and H_3 , and (ii) set $f(R) \equiv TTC(R_1, R_2, \bar{R}_3)$;
- (c) if agent 1 is not in conflict with either agent 2 or agent 3, then $f(R) \equiv TTC(R)$.

It is easy to verify that f is *individually rational* and *unanimous*. We prove that f is *strategy-proof* in Appendix D.1. To see that f is *bossy*, let R be such that

$$\mathbf{R}_1 : H_3, \mathbf{H}_1, H_2,$$

$$\mathbf{R}_2 : H_3, \mathbf{H}_2, H_1,$$

$$\mathbf{R}_3 : H_2, \mathbf{H}_3, H_1.$$

Then, since agents 1 and 2 are in conflict, we have $\bar{\mathbf{R}}_2 : \mathbf{H}_2, H_1, H_3$ and $f(R) = TTC(\bar{\mathbf{R}}_2, R_{-2})$. In particular, for each $i = 1, 2, 3$, $f_i(R) = H_i$. Next consider $\mathbf{R}'_1 : \mathbf{H}_1, \dots$. Then, $f(R'_1, R_{-1}) = TTC(R'_1, R_{-1})$. In particular, $f_1(R'_1, R_{-1}) = H_1$, $f_2(R'_1, R_{-1}) = H_3$, and $f_3(R'_1, R_{-1}) = H_2$. Therefore, $f_1(R'_1, R_{-1}) = H_1 = f_1(R)$, $f_2(R'_1, R_{-1}) = H_3 \neq H_2 = f_2(R)$, and $f_3(R'_1, R_{-1}) = H_2 \neq H_3 = f_3(R)$. Hence, f is *bossy* (and not *Pareto efficient*).

Next, we extend mechanism f from $n = 3$ to any $n > 3$. Let $n > 3$ and recall that $m = 1$. An object $o \in O$ is *acceptable* for agent $i \in N$ if $o R_i H_i$. Let mechanism g be defined as follows: for each $R \in \mathcal{R}^N$,

Case (A) if some agent $i \in \{4, \dots, n\}$ finds some object different from his endowment acceptable, then set $g(R) \equiv TTC(R)$;

Case (B) if each agent $i \in \{4, \dots, n\}$ only finds his own endowment acceptable, then

- let $N' \equiv \{1, 2, 3\}$ and for each $i \in N'$, let $g_i(R) \equiv f_i(R_{|N'})$ where $R_{|N'}$ denotes the preferences of agents in N' restricted to $\{H_1, H_2, H_3\}$;
- for each agent $i \in \{4, \dots, n\}$, $g_i(R) \equiv H_i$.

Since f and TTC are *individually rational* and *unanimous*, g is *individually rational* and *unanimous*. Since f is *bossy*, g is *bossy* as well.

Next, we show that g is *strategy-proof*. First, we verify that no agent $i \in \{4, \dots, n\}$ can profitably misreport his preferences. If R is in case (A), then a misreport by agent i that creates another profile in case (A) does not lead to a more preferred allotment because TTC is *strategy-proof*; a misreport that creates a profile in case (B) assigns endowment H_i to agent i . In either

case, the misreport does not yield a more preferred allotment for agent i . If R is in case (B), then each agent $i \in \{4, \dots, n\}$ obtains his most preferred object (his own endowment) and hence cannot gain by misreporting his preferences.

Second, no agent in $\{1, 2, 3\}$ can “move” R from case (A) to case (B) nor from case (B) to case (A). If R is in case (A), no agent in $\{1, 2, 3\}$ can profitably misreport his preferences because TTC is *strategy-proof*. If R is in case (B), no agent in $\{1, 2, 3\}$ can profitably misreport his preferences because f is *strategy-proof*. Hence, g is *strategy-proof*.

Finally, we extend mechanism g from Shapley-Scarf housing markets to multiple-type housing markets with lexicographic (or separable) preferences by applying it typewise to all object types. Let h be the mechanism that assigns the objects of each type t according to g . Then, h is *unanimous* (and hence *onto*), *individually rational*, and *strategy-proof*, but *bossy*. $\diamond$

The mechanism in Example 5 is well-defined on $\mathcal{R}_i^N$ and $\mathcal{R}_s^N$ (but not on $\mathcal{R}^N$).

D.1 Proof of *strategy-proofness* in Example 5

We show that mechanism f on $\mathcal{R}^N$ defined in Example 5 for $n = 3$ and $m = 1$ is *strategy-proof*.

Proof. Let $R \in \mathcal{R}^N$. We consider three cases.

Case 1. Preferences of agent 1 are $R_1 : H_1, \dots$

By *individual rationality* of f , $f_1(R) = H_1$ and since this is his most preferred object, agent 1 cannot gain by misreporting his preferences.

Let R'_2 be some misreport of agent 2. Since neither agents 1 and 2 nor agents 1 and 3 are in conflict at R or at (R_1, R'_2, R_3) , mechanism f yields the corresponding TTC allocations at R and (R_1, R'_2, R_3) . Hence, by *strategy-proofness* of TTC , agent 2 does not have a profitable deviation at R . Similarly, agent 3 does not have a profitable deviation at R .

Case 2. Preferences of agent 1 are $R_1 : H_2, H_1, H_3$. (Since agents 2 and 3 play a symmetric role in the definition of f , similar symmetric arguments work for preferences $R_1 : H_3, H_1, H_2$.) Agents 1 and 2 are not in conflict. Hence, by *strategy-proofness* of TTC , agent 2 does not have a profitable deviation at R .

Next, we verify that agent 1 does not have a profitable deviation at R .

Case 2.a. Preferences of agent 2 are $R_2 : H_2, \dots$

Note that by *individual rationality* of f we have $f_2(R) = H_2$. So, $f_1(R) = H_1$ and by reporting other preferences, agent 1 still cannot obtain H_2 . So, agent 1 does not have a profitable deviation at R .

Case 2.b. Preferences of agent 2 are $R_2 : H_3, H_2, H_1$ and preferences of agent 3 are $R_3 : H_2, H_3, H_1$.

Agents 1 and 3 are in conflict and one easily verifies that $f(R)$ is the no-trade allocation. In particular, agent 1 receives his endowment H_1 at R . Obviously, by reporting preferences

$R'_1 : H_1, \dots$, agent 1 still obtains H_1 . Any other misreport of agent 1's preferences yields the no-trade allocation. So, agent 1 does not have a profitable deviation at R .

Case 2.c. Preferences of agent 2 are $R_2 : H_1, \dots$ or preferences of agent 2 are $R_2 : H_3, H_1, H_2$ or [preferences of agent 2 are $R_2 : H_3, H_2, H_1$ and preferences of agent 3 are $R_3 : H_1, H_2, H_3$, $R_3 : H_1, H_3, H_2$, or $R_3 : H_2, H_1, H_3$].

It is easy but cumbersome to verify that $f_1(R) = H_2$, i.e., agent 1 receives his most preferred object H_2 .¹⁸ So, agent 1 does not have a profitable deviation at R .

Case 2.d. Preferences of agent 2 are $R_2 : H_3, H_2, H_1$ and preferences of agent 3 are $R_3 : H_3, \dots$. By *individual rationality* of f we have that for all preferences R'_1 (in particular for $R'_1 = R_1$) of agent 1, $f_3(R'_1, R_2, R_3) = H_3$ and thus, also $f_2(R'_1, R_2, R_3) = H_2$. Hence, $f_1(R'_1, R_2, R_3) = H_1$, i.e., independently of his reported preferences, agent 1 receives his own object H_1 . So, agent 1 does not have a profitable deviation at R .

Finally, we verify that agent 3 does not have a profitable deviation at R .

Case 2.I. Preferences of agent 3 are $R_3 : H_3, \dots$.

By *individual rationality* of f , $f_3(R) = H_3$ and since this is his most preferred object, agent 3 cannot gain by misreporting his preferences.

Case 2.II. Preferences of agent 3 are $R_3 : H_1, \dots$.

Agents 1 and 3 are not in conflict and by *strategy-proofness* of TTC , agent 3 does not have a profitable deviation at R .

Case 2.III. Preferences of agent 3 are $R_3 : H_2, H_3, H_1$.

Agents 1 and 3 are in conflict and one easily verifies that $f_3(R) = H_3$. Any possible profitable misreport of preferences by agent 3 requires that H_2 is acceptable and appears in second position. Hence, the only possible candidate for a profitable deviation is $R'_3 : H_1, H_2, H_3$. However, if $R_2 : H_1, \dots$ or $R_2 : H_2, \dots$, then $f_3(R_1, R_2, R'_3) = H_3$; and if $R_2 : H_3, \dots$, then $f_3(R_1, R_2, R'_3) = H_1$. So, agent 3 does not have a profitable deviation at R .

Case 2.IV.i. Preferences of agent 3 are $R_3 : H_2, H_1, H_3$ and preferences of agent 2 are $R_2 : H_1, \dots$ or $R_2 : H_2, \dots$.

Agents 1 and 3 are in conflict and for any possible deviation R'_3 , $f_3(R_1, R_2, R'_3) = H_3 = f_3(R)$. Hence, agent 3 does not have a profitable deviation at R .

Case 2.IV.ii. Preferences of agent 3 are $R_3 : H_2, H_1, H_3$ and preferences of agent 2 are $R_2 : H_3, \dots$.

Agents 1 and 3 are in conflict and one easily verifies that $f_3(R) = H_1$. Any possible profitable misreport of preferences by agent 3 requires that H_2 is acceptable and appears in second position. Hence, the only possible candidate for a profitable deviation is $R'_3 : H_1, H_2, H_3$. However, $f_3(R_1, R_2, R'_3) = H_1$. So, agent 3 does not have a profitable deviation at R .

¹⁸Note that for $R_3 : H_2, \dots$, agents 1 and 3 are in conflict and hence, agent 3 cannot receive H_2 .

Case 3. Preferences of agent 1 are $R_1 : H_2, H_3, H_1$. (Since agents 2 and 3 play a symmetric role in the definition of f , similar symmetric arguments work for preferences $R_1 : H_3, H_2, H_1$.) Agents 1 and 2 are not in conflict. Hence, by *strategy-proofness* of TTC , agent 2 does not have a profitable deviation at R .

Next, we verify that agent 1 does not have a profitable deviation at R .

Case 3.a. Preferences of agent 2 are $R_2 : H_1, \dots$

One immediately verifies that $f_1(R) = H_2$, which is his most preferred object. So, agent 1 does not have a profitable deviation at R .

Case 3.b. Preferences of agent 2 are $R_2 : H_2, \dots$ and preferences of agent 3 are $R_3 : H_1, \dots$ or $R_3 : H_2, H_1, H_3$.

By *individual rationality* of f we have that for any preference relation R'_1 (in particular for $R'_1 = R_1$) of agent 1, $f_2(R'_1, R_2, R_3) = H_2$. Hence, $f_1(R'_1, R_2, R_3) \neq H_2$, i.e., agent 1 does not obtain his most preferred object. Since $f_1(R) = H_3$ is his second most preferred object, we conclude that agent 1 does not have a profitable deviation at R .

Case 3.c. Preferences of agent 2 are $R_2 : H_2, \dots$ and preferences of agent 3 are $R_3 : H_3, \dots$ or $R_3 : H_2, H_3, H_1$;

or

Case 3.d. Preferences of agent 2 are $R_2 : H_3, H_2, H_1$ and preferences of agent 3 are $R_3 : H_3, \dots$ or $R_3 : H_2, H_3, H_1$.

In cases 3.c and 3.d, we have that for any possible deviation R'_1 , $f_1(R'_1, R_2, R_3) = H_1 = f_1(R)$. Hence, agent 1 does not have a profitable deviation at R .

Case 3.e. Preferences of agent 2 are $R_2 : H_3, \dots$ and preferences of agent 3 are $R_3 : H_1, \dots$;

or

Case 3.f. Preferences of agent 2 are $R_2 : H_3, H_2, H_1$ and preferences of agent 3 are $R_3 : H_2, H_1, H_3$;

or

Case 3.g. Preferences of agent 2 are $R_2 : H_3, H_1, H_2$ and preferences of agent 3 are $R_3 : H_2, \dots$;

or

Case 3.h. Preferences of agent 2 are $R_2 : H_3, H_1, H_2$ and preferences of agent 3 are $R_3 : H_3, \dots$

In cases 3.e, 3.f, 3.g, and 3.h, $f_1(R) = H_2$, i.e., agent 1 receives his most preferred object H_2 .¹⁹ So, agent 1 does not have a profitable deviation at R .

Finally, we verify that agent 3 does not have a profitable deviation at R . Cases 3.I, 3.II, and 3.III below are as 2.I, 2.II, and 2.III. There is a small difference between cases 2.IV and 3.IV.

Case 3.I. Preferences of agent 3 are $R_3 : H_3, \dots$

By *individual rationality* of f , $f_3(R) = H_3$ and since this is his most preferred object, agent 3 cannot gain by misreporting his preferences.

¹⁹Note that for $R_3 : H_2, \dots$, agents 1 and 3 are in conflict and hence, agent 3 cannot receive H_2 .

Case 3.II. Preferences of agent 3 are $R_3 : H_1, \dots$

Agents 1 and 3 are not in conflict and by *strategy-proofness* of TTC , agent 3 does not have a profitable deviation at R .

Case 3.III. Preferences of agent 3 are $R_3 : H_2, H_3, H_1$.

Agents 1 and 3 are in conflict and one easily verifies that $f_3(R) = H_3$. Any possible profitable misreport of preferences by agent 3 requires that H_2 is acceptable and appears in second position. Hence, the only possible candidate for a profitable deviation is $R'_3 : H_1, H_2, H_3$. However, if $R_2 : H_1, \dots$, then $f_3(R_1, R_2, R'_3) = H_3$; and if $R_2 : H_3, \dots$ or $R_2 : H_2, \dots$, then $f_3(R_1, R_2, R'_3) = H_1$. So, agent 3 does not have a profitable deviation at R .

Case 3.IV.i. Preferences of agent 3 are $R_3 : H_2, H_1, H_3$ and preferences of agent 2 are $R_2 : H_1, \dots$

Agents 1 and 3 are in conflict and for any possible deviation R'_3 , $f_3(R_1, R_2, R'_3) = H_3 = f_3(R)$. Hence, agent 3 does not have a profitable deviation at R .

Case 3.IV.ii. Preferences of agent 3 are $R_3 : H_2, H_1, H_3$ and preferences of agent 2 are $R_2 : H_2, \dots$ or $R_2 : H_3, \dots$

Agents 1 and 3 are in conflict and one easily verifies that $f_3(R) = H_1$. Any possible profitable misreport of preferences by agent 3 requires that H_2 is acceptable and appears in second position. Hence, the only possible candidate for a profitable deviation is $R'_3 : H_1, H_2, H_3$. However, $f_3(R_1, R_2, R'_3) = H_1$. So, agent 3 does not have a profitable deviation at R . $\square$

E Proof of Theorem 2: uniqueness

Proof of Theorem 2: uniqueness. Suppose that mechanism $f : \mathcal{R}_s^N \rightarrow X$ satisfies the properties listed in Theorem 2 (by Lemma 1, *ontoness* and *unanimity* can be used interchangeably). We will show that for each $R \in \mathcal{R}_s^N$, $f(R) = tTTC(R)$. We introduce the following notation. For any agent $i \in N$ and any two separable preferences $R_i, \bar{R}_i \in \mathcal{R}_s$, we write $R_i \sim \bar{R}_i$ if they induce the same marginal preferences, i.e., for each $t \in T$, $R_i^t = \bar{R}_i^t$.

Let $R \in \mathcal{R}_s^N$ such that each agent has lexicographic preferences, i.e., $R \in \mathcal{R}_l^N$. Since the restriction of f to $\mathcal{R}_l^N$ satisfies the properties listed in Theorem 1, it immediately follows from Theorem 1 that $f(R) = tTTC(R)$.

Let $R \in \mathcal{R}_s^N$ such that only one agent does not have lexicographic preferences. We can assume, without loss of generality, that $R_1 \in \mathcal{R}_s \setminus \mathcal{R}_l$ and for each agent $j \neq 1$, $R_j \in \mathcal{R}_l$. Let $y \equiv f(R)$.

For each $t \in T$, define $R'_1(t) \in \mathcal{R}_l$ such that $R'_1(t) \sim R_1$ and the most important type of $R'_1(t)$ is type t . Since $R_1 \sim R'_1(1) \sim R'_1(2) \sim \dots \sim R'_1(m)$, it follows from the definition of $tTTC$ that $x \equiv tTTC(R) = tTTC(R'_1(1), R_{-1}) = tTTC(R'_1(2), R_{-1}) = \dots = tTTC(R'_1(m), R_{-1})$. We will show that $y = x$.

Let $t \in T$. From the case where each agent has lexicographic preferences, it follows that $f(R'_1(t), R_{-1}) = tTTC(R'_1(t), R_{-1}) = x$. By *strategy-proofness* of f when moving from $(R'_1(t), R_{-1})$ to (R_1, R_{-1}) , $x_1 = f_1(R'_1(t), R_{-1}) R'_1(t) f_1(R_1, R_{-1}) = y_1$. Then, since $R'_1(t) \sim R_1$ and $R'_1(t)$ is a lexicographic preference relation where t is the most important type, $x_1^t R_1^t y_1^t$.

Since for each $t \in T$, $x_1^t R_1^t y_1^t$ and since $R_1 \in \mathcal{R}_s$, we have $x_1 R_1 y_1$. By *strategy-proofness* of f when moving from (R_1, R_{-1}) to $(R'_1(t), R_{-1})$, we have that $y_1 = f_1(R_1, R_{-1}) R_1 f_1(R'_1(t), R_{-1}) = x_1$. Hence, $x_1 = y_1$. By *non-bossiness* of f , we have that $y = f(R_1, R_{-1}) = f(R'_1(t), R_{-1}) = x$.

Let $R \in \mathcal{R}_s^N$ such that exactly two agents do not have lexicographic preferences. We can assume, without loss of generality, that $R_1, R_2 \in \mathcal{R}_s \setminus \mathcal{R}_l$ and for each agent $j \neq 1, 2$, $R_j \in \mathcal{R}_l$. Let $y \equiv f(R)$.

For each $t \in T$, define $R'_2(t) \in \mathcal{R}_l$ such that $R'_2(t) \sim R_2$ and the most important type of $R'_2(t)$ is type t . Since $R_2 \sim R'_2(1) \sim R'_2(2) \sim \dots \sim R'_2(m)$, it follows from the definition of $tTTC$ that $x \equiv tTTC(R) = tTTC(R'_2(1), R_{-2}) = tTTC(R'_2(2), R_{-2}) = \dots = tTTC(R'_2(m), R_{-2})$. We will show that $y = x$.

Let $t \in T$. At preference profile $(R'_2(t), R_{-2})$, only agent 1 has non-lexicographic preferences. Thus, from the previous case, $f(R'_2(t), R_{-2}) = tTTC(R'_2(t), R_{-2}) = tTTC(R) = x$. By *strategy-proofness* of f when moving from $(R'_2(t), R_{-2})$ to (R_2, R_{-2}) , we have that $x_2 = f_2(R'_2(t), R_{-2}) R'_2(t) f_2(R_2, R_{-2}) = y_2$. Then, since $R'_2(t) \sim R_2$ and $R'_2(t)$ is a lexicographic preference relation where t is the most important type, $x_2^t R_2^t y_2^t$.

Since for each $t \in T$, $x_2^t R_2^t y_2^t$ and since $R_2 \in \mathcal{R}_s$, we have $x_2 R_2 y_2$. By *strategy-proofness* of f when moving from (R_2, R_{-2}) to $(R'_2(t), R_{-2})$, $y_2 = f_2(R_2, R_{-2}) R_2 f_2(R'_2(t), R_{-2}) = x_2$. Hence, $x_2 = y_2$. By *non-bossiness* of f , we have that $y = f(R_2, R_{-2}) = f(R'_2(t), R_{-2}) = x$.

We can apply repeatedly the same arguments to obtain that for each $k = 0, 1, \dots, n$ and each preference profile $R \in \mathcal{R}_s^N$ where exactly k agents have non-lexicographic preferences, $f(R) = tTTC(R)$. Thus, for each $R \in \mathcal{R}_s^N$, $f(R) = tTTC(R)$. $\square$

F Two further impossibility results for strict preferences

A mechanism $f : \mathcal{R}^N \rightarrow X$ extends the $tTTC$ mechanism from $\mathcal{R}_l^N$ ($\mathcal{R}_s^N$, respectively) to $\mathcal{R}^N$, if for each $R \in \mathcal{R}_l^N$ ($\mathcal{R}_s^N$, respectively), $f(R) = tTTC(R)$.

The following theorem captures another impossibility result.

Theorem 4. *Let $m > 1$. Then, no mechanism satisfying individual rationality and strategy-proofness extends the $tTTC$ mechanism from lexicographic (separable) preferences to strict preferences.*

Proof. Without loss of generality, let $m = 2$. Suppose that there is a mechanism $f : \mathcal{R}^N \rightarrow X$ that is *individually rational* and *strategy-proof* and that coincides with the $tTTC$ mechanism

on $\mathcal{R}_i^N$ ($\mathcal{R}_s^N$, respectively). Let $x, y \in X \setminus \{e\}$ be such that at x agents 1 and 2 swap their endowments of types 1 and 2, i.e.,

$$\begin{aligned} x_1 &= (o_2^1, o_2^2, o_1^3, o_1^4, \dots, o_1^m), \\ x_2 &= (o_1^1, o_1^2, o_2^3, o_2^4, \dots, o_2^m), \\ \text{and for each } i &= 3, \dots, n, \quad x_i = o_i \end{aligned}$$

and at y agents 1 and 2 swap their endowments of type 1, i.e.,

$$\begin{aligned} y_1 &= (o_2^1, o_1^2, o_1^3, o_1^4, \dots, o_1^m), \\ y_2 &= (o_1^1, o_2^2, o_2^3, o_2^4, \dots, o_2^m), \\ \text{and for each } i &= 3, \dots, n, \quad y_i = o_i. \end{aligned}$$

Obviously, $x \neq y$.

Next, we define lexicographic preferences for agent 1 by listing a strict ordering of all objects. At R_1 , agent 1's type order is $1, 2, \dots, m$ and the only object he prefers to one of his endowments is the type 1 endowment of agent 2, i.e.,

$$R_1 : o_2^1, o_1^1, o_3^1, \dots, o_n^1, o_1^2, \dots, o_n^2, o_1^3, \dots, o_n^3, o_1^m, \dots, o_n^m$$

At R'_1 , agent 1's type order is $1, 2, \dots, m$ and the only objects he prefers to some of his endowments are the type 1 and 2 endowments of agent 2, i.e.,

$$R'_1 : o_2^1, o_1^1, o_3^1, \dots, o_n^1, o_2^2, o_1^2, o_3^2, \dots, o_n^2, o_1^3, \dots, o_n^3, o_1^m, \dots, o_n^m.$$

Similarly, we define lexicographic preferences for agent 2 by listing a strict ordering of all objects. At R_2 , agent 2's type order is $1, 2, \dots, m$ and the only objects he prefers to some of his endowments are the type 1 and 2 endowments of agent 1, i.e.,

$$R_2 : o_1^1, o_2^1, o_3^1, \dots, o_n^1, o_1^2, o_2^2, o_3^2, \dots, o_n^2, o_1^3, \dots, o_n^3, o_1^m, \dots, o_n^m.$$

Next, we define lexicographic preferences for all remaining agents as follows. For each $i = 3, \dots, n$, agent i prefers his full endowment to all other allotments, i.e.,

$$\text{for each } t \in T, \quad R_i^t : o_i^t, \dots$$

Finally, let R'_2 be strict and *non-separable* preferences for agent 2 such that agent 2 prefers only his allotment at x to his endowment, i.e.,

$$R'_2 : x_2, o_2, \dots$$

Note that $(R'_1, R'_2, R_{N \setminus \{1,2\}}) \in \mathcal{R}^N \setminus \mathcal{R}_s^N$. There are only two *individually rational* allocations at $(R'_1, R'_2, R_{N \setminus \{1,2\}})$: e and x .

Since R is a profile of lexicographic preferences, we have $f(R) = tTTC(R) = y$. By *individual rationality* of f , $f(R_1, R'_2, R_{N \setminus \{1,2\}}) = e$. Hence, by *individual rationality* and *strategy-proofness* of f , $f(R'_1, R'_2, R_{N \setminus \{1,2\}}) = e$.

Since $(R'_1, R_2, R_{N \setminus \{1,2\}})$ is a profile of lexicographic preferences, we have $f(R'_1, R_2, R_{N \setminus \{1,2\}}) = tTTC(R'_1, R_2, R_{N \setminus \{1,2\}}) = x$. Therefore, agent 2 has an incentive to misreport R_2 at $(R'_1, R'_2, R_{N \setminus \{1,2\}})$; contradicting *strategy-proofness* of f . $\square$

The no-trade rule (Example 2, Appendix D) is *individually rational*, *strategy-proof*, and does not extend the tTTC mechanism from lexicographic (separable) preferences to strict preferences. The following mechanism that extends tTTC from lexicographic (separable) preferences to strict preferences is *individually rational* but not *strategy-proof*: the mechanism assigns the tTTC allocation on $\mathcal{R}_l^N$ ($\mathcal{R}_s^N$, respectively) and the endowment allocation on $\mathcal{R}^N \setminus \mathcal{R}_l^N$ ($\mathcal{R}^N \setminus \mathcal{R}_s^N$, respectively). The independence of *strategy-proofness* is an **open problem**.

Lemma 5. *Let $m > 1$. If a mechanism satisfies strategy-proofness, non-bossiness, and extends the tTTC mechanism from lexicographic (separable) preferences to strict preferences, then it satisfies individual rationality.*

Proof. Suppose that there is a mechanism $f : \mathcal{R}^N \rightarrow X$ that is *strategy-proof*, *non-bossy*, and that coincides with the tTTC mechanism on $\mathcal{R}_l^N$ ($\mathcal{R}_s^N$, respectively). By Lemma 3 (which can be proven for $\mathcal{R}^N$ using the same arguments), f is *monotonic*.

By contradiction, assume that f is not *individually rational*. Thus, there exists a profile $R \in \mathcal{R}^N$ and an agent $i \in N$ such that $o_i P_i f_i(R)$. Without loss of generality, let $i = 1$.

Let $x \equiv f(R)$. Let $\hat{R}_1 \in \mathcal{R}$ be such that agent 1 positions o_1 first and x_1 second, i.e.,

$$\hat{R}_1 : o_1, x_1, \dots$$

For each agent $j = 2, 3, \dots, n$, let $\hat{R}_j \in \mathcal{R}_l$ be such that agent j positions x_j first, i.e.,

$$\text{for each } t \in T, \hat{R}_j^t : x_j^t, \dots$$

It is easy to see that $\hat{R}$ is a monotonic transformation of R at x . Thus, by *monotonicity* of f , $f(\hat{R}) = x$.

Next, let $\bar{R}_1 \in \mathcal{R}_l$ be such that

$$\text{for each } t \in T, \bar{R}_1^t : o_1^t, x_1^t, \dots$$

Note that $(\bar{R}_1, \hat{R}_{-1}) \in \mathcal{R}_l^N$. Thus, $f(\bar{R}_1, \hat{R}_{-1}) = tTTC(\bar{R}_1, \hat{R}_{-1})$, and in particular, $f_1(\bar{R}_1, \hat{R}_{-1}) = o_1$. However, then $f_1(\bar{R}_1, \hat{R}_{-1}) = o_1 \hat{P}_1 x_1 = f_1(\hat{R})$ and agent 1 has an incentive to misreport $\bar{R}_1$ at $\hat{R}$, which contradicts with *strategy-proofness* of f . $\square$

Now, Theorem 4 and Lemma 5 imply the following impossibility result.

Corollary 4. *Let $m > 1$. Then, no mechanism satisfying strategy-proofness and non-bossiness extends the tTTC mechanism from lexicographic (separable) preferences to strict preferences.*

The no-trade rule (Example 2, Appendix D) is *strategy-proof* and *non-bossy*, and does not extend the tTTC mechanism from lexicographic (separable) preferences to strict preferences. The independence of *strategy-proofness* (*non-bossiness*, respectively) is an **open problem**.

References

- Alva, S. (2017). When is manipulation all about the ones and twos? Working paper, <https://samsonalva.com/PWSP.pdf>.
- Bag, T., S. Garg, Z. Shaik, and A. Mitschele-Thiel (2019). Multi-numerology based resource allocation for reducing average scheduling latencies for 5G NR wireless networks. In *2019 European Conference on Networks and Communications*, pp. 597–602.
- Bird, C. G. (1984). Group incentive compatibility in a market with indivisible goods. *Economics Letters* 14(4), 309–313.
- Biró, P., F. Klijn, and S. Pápai (2022a). Balanced exchange in a multi-unit Shapley-Scarf market. Barcelona School of Economics Working Paper 1342.
- Biró, P., F. Klijn, and S. Pápai (2022b). Serial rules in a multi-unit Shapley-Scarf market. *Games and Economic Behavior* 136, 428–453.
- Feng, D. and B. Klaus (2022). Preference revelation games and strict cores of multiple-type housing market problems. *International Journal of Economic Theory* 18(1), 61–76.
- Ghodsí, A., V. Sekar, M. Zaharia, and I. Stoica (2012). Multi-resource fair queueing for packet processing. In *Proceedings of the ACM SIGCOMM 2012 Conference on Applications, Technologies, Architectures, and Protocols for Computer Communication*, pp. 1–12.
- Ghodsí, A., M. Zaharia, B. Hindman, A. Konwinski, S. Shenker, and I. Stoica (2011). Dominant resource fairness: Fair allocation of multiple resource types. In *Proceedings of the 8th USENIX Conference on Networked Systems Design and Implementation*, pp. 323–336.
- Han, B., V. Sciancalepore, D. Feng, X. Costa-Perez, and H. D. Schotten (2019). A utility-driven multi-queue admission control solution for network slicing. In *IEEE INFOCOM 2019-IEEE Conference on Computer Communications*, pp. 55–63.
- Huh, W. T., N. Liu, and V.-A. Truong (2013). Multiresource allocation scheduling in dynamic environments. *Manufacturing & Service Operations Management* 15(2), 280–291.

- Klaus, B. (2008). The coordinate-wise core for multiple-type housing markets is second-best incentive compatible. *Journal of Mathematical Economics* 44(9-10), 919–924.
- Konishi, H., T. Quint, and J. Wako (2001). On the Shapley–Scarf economy: the case of multiple types of indivisible goods. *Journal of Mathematical Economics* 35(1), 1–15.
- Ma, J. (1994). Strategy-proofness and the strict core in a market with indivisibilities. *International Journal of Game Theory* 23(1), 75–83.
- Mackin, E. and L. Xia (2016). Allocating indivisible items in categorized domains. In *Proceedings of the Twenty-Fifth International Joint Conference on Artificial Intelligence*, pp. 359–365.
- Manjunath, V. and A. Westkamp (2021). Strategy-proof exchange under trichotomous preferences. *Journal of Economic Theory* 193, 105197.
- Miyagawa, E. (2002). Strategy-proofness and the core in house allocation problems. *Games and Economic Behavior* 38(2), 347–361.
- Monte, D. and N. Tumennasan (2015). Centralized allocation in multiple markets. *Journal of Mathematical Economics* 61, 74–85.
- Moulin, H. (1995). *Cooperative Microeconomics: A Game-Theoretic Introduction*. Princeton, NJ: Princeton University Press.
- Pápai, S. (2000). Strategyproof assignment by hierarchical exchange. *Econometrica* 68(6), 1403–1433.
- Pápai, S. (2001). Strategyproof and nonbossy multiple assignments. *Journal of Public Economic Theory* 3(3), 257–271.
- Pápai, S. (2003). Strategyproof exchange of indivisible goods. *Journal of Mathematical Economics* 39(8), 931–959.
- Pápai, S. (2007). Exchange in a general market with indivisible goods. *Journal of Economic Theory* 132(1), 208–235.
- Peng, M., C. Wang, J. Li, H. Xiang, and V. Lau (2015). Recent advances in underlay heterogeneous networks: Interference control, resource allocation, and self-organization. *IEEE Communications Surveys & Tutorials* 17(2), 700–729.
- Roth, A. E. (1982). Incentive compatibility in a market with indivisible goods. *Economics Letters* 9(2), 127–132.
- Roth, A. E. and A. Postlewaite (1977). Weak versus strong domination in a market with indivisible goods. *Journal of Mathematical Economics* 4(2), 131–137.

- Serizawa, S. (2006). Pairwise strategy-proofness and self-enforcing manipulation. *Social Choice and Welfare* 26(2), 305–331.
- Shapley, L. S. and H. Scarf (1974). On cores and indivisibility. *Journal of Mathematical Economics* 1(1), 23–37.
- Sikdar, S., S. Adah, and L. Xia (2017). Mechanism design for multi-type housing markets. In *Proceedings of the Thirty-First AAAI Conference on Artificial Intelligence*, pp. 684–690.
- Svensson, L.-G. (1999). Strategy-proof allocation of indivisible goods. *Social Choice and Welfare* 16(4), 557–567.
- Takamiya, K. (2001). Coalition strategy-proofness and monotonicity in Shapley–Scarf housing markets. *Mathematical Social Sciences* 41(2), 201–213.
- Wako, J. (2005). Coalition-proof Nash allocation in a barter game with multiple indivisible goods. *Mathematical Social Sciences* 49(2), 179–199.