

Serial dictatorship: The unique optimal allocation rule when information is endogenous

SOPHIE BADE

Department of Economics, Royal Holloway College, London and Max Planck Institut, Bonn

The study of matching problems typically assumes that agents precisely know their preferences over the goods to be assigned. Within applied contexts, this assumption stands out as particularly counterfactual. Parents typically do invest a large amount of time and resources to find the best school for their children; doctors run costly tests to establish the best kidney for a given patient. In this paper, I introduce the assumption of endogenous information acquisition into otherwise standard house allocation problems. I find that there is a unique ex ante Pareto-optimal, strategy-proof, and nonbossy allocation mechanism: serial dictatorship. This stands in sharp contrast to the very large set of such mechanisms for house allocation problems without endogenous information acquisition.

KEYWORDS. Serial dictatorship, house allocation, endogenous information.

JEL CLASSIFICATION. C78.

1. INTRODUCTION

Many allocation problems of indivisible goods have to be solved without explicit markets. For some such goods, be it school slots or kidneys, the use of markets to determine allocations is perceived as immoral or repugnant. In many cases markets are explicitly forbidden. A prospering subfield of mechanism design questions how to best allocate such objects to recipients; many mechanisms that are optimal according to a host of different criteria have been found. These mechanisms have usually been designed for the case of agents precisely knowing their preferences over the goods to be assigned. However, this assumption seems counterfactual in many of the areas in which such mechanisms are used. Parents typically invest a significant amount of time on school choice; doctors need to run costly tests on kidneys to figure out which would be best for a given patient.

This paper sets out to study the allocative properties of mechanisms in conjunction with their impact on the agents' incentives to acquire information. To this end, I modify the standard model of house allocation problems in which a set of agents needs to be matched to a set of equally many objects, henceforth called houses, allowing for costly

Sophie Bade: sophie.bade@rhul.ac.uk

Thanks to Martin Hellwig, Michael Mandler, and Philipp Weinschenk, and audiences at the European Network on Matching conference in Brussels, the Jahrestagung of the Verein fuer Socialpolitik in Frankfurt, the Paris–Cologne meeting, the University of Bonn, and at Games 2012 in Istanbul.

Copyright © 2015 Sophie Bade. Licensed under the [Creative Commons Attribution-NonCommercial License 3.0](https://creativecommons.org/licenses/by-nc/3.0/). Available at <http://econtheory.org>.

DOI: 10.3982/TE1335

information acquisition on these houses. The goal is to characterize the set of strategy-proof, nonbossy, and Pareto-optimal mechanisms in this environment.

Over the years, various classes of such mechanisms have been identified for the standard case of known preferences. [Pycia and Ünver \(2013\)](#) and [Bade \(2014\)](#) characterize the very large set of all such mechanisms. Lots of room remains to impose additional requirements to select among these mechanisms. The case of housing problems with endogenous information acquisition differs sharply. In that case, there is a unique strategy-proof, nonbossy, and ex ante Pareto-optimal mechanism: serial dictatorship. The following example illustrates the outstanding role of serial dictatorship.

EXAMPLE 1. Two agents called 1 and 2 start out owning two houses, k and g , respectively; this initial allocation only changes if both agents agree to exchange houses.¹ In an environment without endogenous learning, this mechanism is strategy-proof, nonbossy, and Pareto optimal. To see that this mechanism can be Pareto dominated when agents have a choice to learn, consider the following setup. Neither agent knows whether he values house k at 8 or at 0: the agents' valuations of this house are independent draws from a distribution according to which the two possible values are equally likely. Both agents value house g at 2. Agent 1 has to pay 0.8 to learn his value of house k ; learning is free for agent 2. Both agents need to announce simultaneously whether or not they would like to swap.

Now let us consider agent 1's decision problem. If he does not learn the value of house k , he prefers to keep it (expected value of 4 vs. 2, the known value of g). If he learns the value, he prefers to swap houses if and only if he values house k at 0. Agent 2, in turn, is willing to swap with a probability of $\frac{1}{2}$.² If agent 1 learns his value of house k , he obtains an expected utility of $\frac{1}{2} \times 8 + \frac{1}{2}(\frac{1}{2} \times 2 + \frac{1}{2} \times 0) - 0.8 = 3.7$, with the last term reflecting agent 1's cost of learning. So agent 1 prefers to keep house k without learning, implying that in equilibrium agent 2 is stuck with house g , yielding an ex ante utility profile of (4, 2).

The serial dictatorship with agent 1 as the first dictator ex ante Pareto-dominates the given mechanism. For agent 1 as the first dictator, it is worthwhile to learn the value of house k and to choose it if and only if he finds it of high value (expected utility: $\frac{1}{2} \times 8 + \frac{1}{2} \times 2 - 0.8 = 4.2$). So agent 2 is matched to his ex ante preferred house with a probability of $\frac{1}{2}$. The profile of ex ante utilities is (4.2, 3). $\diamond$

This example shows that some of the bedrock of matching theory starts to crumble if one allows for endogenous information acquisition. Both mechanisms described in the example—the top trading cycles mechanism and serial dictatorship—are Pareto optimal, strategy-proof, and nonbossy in an environment without endogenous information acquisition. Either one of these mechanisms traces out the full set of Pareto-optimal

¹This mechanism is Gale's top trading cycles mechanism for two agents and two houses; a formal definition can be found in [Section 2.2](#).

²Since learning is costless, agent 2 will be willing to exchange with agent 1 if and only if he values house k at 8, which happens with probability $\frac{1}{2}$.

matchings when one allows for all possible orderings of dictators or for all possible initial allocations, respectively. With endogenous information acquisition, something very different happens. In that case, serial dictatorship may ex ante Pareto-dominate Gale's top trading cycles mechanism as shown in [Example 1](#). The main result of the paper significantly generalizes this observation. I show that for any nonbossy and strategy-proof mechanism that is not a serial dictatorship, one can find a housing problem and a (path-dependent) serial dictatorship, such that the serial dictatorship strictly ex ante Pareto-dominates the named mechanism in the given housing problem. Conversely, serial dictatorships are never dominated in this way.

The essential difference between the two mechanisms in [Example 1](#) is that the strong incentives for learning under serial dictatorship are dampened under top trading cycles. While agent 1's knowledge of the value of house k is always useful under serial dictatorship, the same knowledge is irrelevant in half of all cases under the alternative mechanism. Serial dictatorship stands out as the only mechanism that always combines optimal learning incentives with optimal allocation incentives. It is well known that serial dictatorship sets the "right" incentives for allocations: it belongs to the set of strategy-proof, nonbossy and Pareto-optimal mechanisms. What distinguishes serial dictatorship is that it is the only mechanism in this set that also sets the right incentives for information acquisition: given that any agent knows his exact choice set when he decides to learn, no information is ever wastefully acquired. This paper gives two variants of the uniqueness statement on serial dictatorship pertaining to the case of sequential and simultaneous learning, Theorems 1 and 2.

The set of information structures considered in this article is constrained in two ways: first, the agents' preferences are independent draws, implying that agents never wish to delegate their choices to better informed agents. Second, in line with the literature on standard housing problems, a no-indifference condition ensures that at least some mechanisms work optimally. With these two constraints in place, we can be sure that the suboptimality of mechanisms other than serial dictatorship is due to the agents' ability to acquire information. Serial dictatorship might outperform other mechanisms in matching problems with indifferences or correlated preferences. Still, the domination of serial dictatorship in the present article cannot be attributed to such arguments as these classical "trouble-makers" have been ruled out. The uniqueness result of the article is indeed driven by the assumption of endogenous information acquisition.

The compromise between the quality of information acquisition and of allocations is one of the main themes of the growing literature on mechanism design with endogenous information acquisition. Mechanisms are often characterized in terms of an optimal trade-off between informative and allocative efficiency. [Gerardi and Yariv \(2008\)](#) and [Bergemann and Välimäki \(2006\)](#), respectively, illustrate this trade-off in voting and auctions environments. This trade-off is relevant in the present paper: simple serial dictatorship is the one mechanism under which allocative and informative efficiency coexist. For any other mechanism, we have to face the trade-off between the two kinds of efficiency. The optimality of sequential learning is another theme of the literature on mechanisms with endogenous learning that is echoed in the present paper. [Gershkov and Szentes \(2009\)](#) as well as [Smorodinsky and Tennenholtz \(2006\)](#) present

voting models in which the voters' optimal acquisition of information is sequential. Similarly, for auctions, [Compte and Jehiel \(2007\)](#) find that ascending price auctions can dominate sealed bid auctions in terms of expected welfare. In this vein the present paper shows the unique optimality of *sequential* simple serial dictatorship when allowing for *any* sequence of information elicitation.

In [Section 2](#), I provide formal definitions of the housing problems and mechanisms under study. There I define [Example 2](#) to argue that sequential elicitation procedures might outperform simultaneous ones in the present context. With all the relevant terminology in hand, I state the two main results of the article, Theorems 1 and 2, in [Section 3](#). The proof of these two theorems revolves around three examples: [Example 2](#), which is presented in [Section 2](#); the introductory [Example 1](#), which is revisited in [Section 4](#); [Example 5](#), which is presented in [Section 4](#). To extend the arguments gleaned from these examples to the case of large housing problems, I rely on [Pycia and Ünver's \(2013\)](#) and [Bade's \(2014\)](#) "trading and braiding" mechanisms ([Section 6](#)). The proof of the two results is contained in [Section 7](#). The presentation of trading and braiding mechanisms and the proof are preceded by [Section 5](#) which sheds light on possible extensions and limits of the unique ex ante Pareto optimality of serial dictatorship.

2. THE MODEL

2.1 Agents, houses, and values

Fix two sets of agents $I = \{1, \dots, n\}$, and houses H with equally many elements ($|H| = n$) and generic elements $i, j \in N$ and $h, d, g, k \in H$. There is a finite state space Ω that consists of profiles of values $\omega := (\omega_h^i)_{h \in H, i \in I}$, where ω_h^i is the value that agent i assigns to house h and $\omega^i := (\omega_h^i)_{h \in H}$ is the vector of agent i 's valuations. Denote the set of all partitions of Ω by $\mathcal{P}$. The state $\omega \in \Omega$ is drawn from the probability distribution π , with $\pi(\omega) > 0$ for all $\omega \in \Omega$. The prior π is common knowledge among the designer and all agents.

A vector $c := (c^i)_{i \in N}$ of cost functions $c^i: \mathcal{P} \rightarrow \mathbb{R}_0^+ \cup \{\infty\}$ describes the agents' learning technologies, where $c^i(P)$ is agent i 's nonnegative (and possibly infinite) cost to learn P . Staying ignorant is free in the sense that $c^i(\{\Omega, \emptyset\}) = 0$ holds for all i . Adopt the understanding that $\hat{\omega}_h^i$ not only denotes agent i 's value of house h in state $\hat{\omega}$, but also the event $\{\omega \mid \omega_h^i = \hat{\omega}_h^i\}$ that agent i values house h at $\hat{\omega}_h^i$. Define ζ^i as the algebra on Ω that is generated by all events $\hat{\omega}_h^i$. It is assumed that no agent can learn anything about any other agent's preferences: formally, $c^i(P) = \infty$ holds for all $P \not\subset \zeta^i$. An agent who has acquired the partition P knows the event $P(\omega)$ at state ω . The partition according to which agent i knows his value for each of the houses is called $\bar{P}^i$.³

To ensure comparability of the present model to standard housing models, I impose two further assumptions. First, the agents' preferences are drawn independently; formally, $\pi(E^i \cap E^j) = \pi(E^i)\pi(E^j)$ holds for all $E^i \in \zeta^i$ and $E^j \in \zeta^j$. The assumption implies that agent i 's posterior value of a house does not change if he finds out what some other agent knows. To see this, define $\bar{\omega}_h^i(E)$ as agent i 's expected value of house h when he

³So $\bar{P}^i$ is the finest partition P with the feature $P \subset \zeta^i$, $\bar{P}^i(\omega) = \omega^i$ holds for all $\omega \in \Omega$.

knows event E . Observe that $\bar{\omega}_h^i(E) = \bar{\omega}_h^i(E \cap G)$ holds when $E \in \zeta^i$ and $G = \bigcap_{j \neq i} G^j$ with $G^j \in \zeta^j$, since

$$\bar{\omega}_h^i(E \cap G) = \frac{\sum_{\omega_h^i \subset E} \pi(\omega_h^i \cap G) \omega_h^i}{\pi(E \cap G)} = \frac{\sum_{\omega_h^i \subset E} \pi(\omega_h^i) \pi(G) \omega_h^i}{\pi(E) \pi(G)} = \bar{\omega}_h^i(E),$$

where the crucial equality follows from the independence of any event ω_h^i and G .⁴ Without the assumption of independence, some agents might find it beneficial to delegate their decision; there would also be scope for signaling.

Second, to avoid the difficulties that arise in housing problems with indifferences, I assume that any agent i who is faced with the nonstrategic problem of choosing a house from some subset $S \subset H$ has a unique optimal plan of action that consists of a partition P together with a choice function $C: \Omega \rightarrow S$ that prescribes a unique choice from S for every cell of P , so $C(\omega) = C(\omega')$ for $\omega \in P(\omega')$. The condition is satisfied if $\bar{\omega}_h^i(E) \neq \bar{\omega}_g^i(E)$ holds for any $h \neq g$ and any E that is an element of a partition P with $c^i(P) < \infty$ and if for every $S \subset H$, there is a unique P that maximizes $\sum_{E \in P} \max_{h \in S} \bar{\omega}_h^i(E) - c^i(P)$.⁵ The two assumptions of independence and no indifference imply that the following results are indeed driven by the novel assumption of endogenous information acquisition. The outstanding role of serial dictatorship cannot be attributed to an appearance of weak or correlated preferences under the guise of endogenous information acquisition. The two assumptions are discussed at length in Section 5.

The vector $\mathcal{H} = (I, H, \Omega, \pi, c)$ of sets of agents I and houses H , a state space Ω , a probability distribution π on Ω , and cost functions c that all satisfy the assumptions discussed above constitutes a *housing problem (with endogenous information acquisition)*.

A *matching* $\mu: I \rightarrow H$. A *submatching* $\sigma: I_\sigma \rightarrow H_\sigma$ is a bijection with $I_\sigma \subset I$ and $H_\sigma \subset H$. The set of all submatchings that are not matchings is denoted by $\bar{\mathcal{M}}$. For any particular submatching σ , the sets of unmatched agents and houses are denoted by $\bar{I}_\sigma$ and $\bar{H}_\sigma$. The house assigned to agent i under the submatching σ is $\sigma(i)$. Submatchings σ are also interpreted as sets, where a pair (i, h) belongs to the set σ if and only if $\sigma(i) = h$ under the interpretation of σ as a function.

An *outcome function* $f: \Omega \rightarrow \mathcal{M} \times \mathcal{P}^n$ maps any state ω to a matching $\mu[\omega] \in \mathcal{M}$ and profile of information partitions $(P^i[\omega])_{i \in I}$. Outcome functions describe the different matchings achieved and the learning undertaken at all states ω . At ω , agent i knows the event $P^i\omega$ when he acquires the partition $P^i[\omega]$ as prescribed by the outcome function f . The ex ante utility U^i of agent i associated with a given outcome function $f: \Omega \rightarrow \mathcal{M} \times \mathcal{P}^n$ is defined as

$$U^i(f) = \sum_{\omega \in \Omega} \pi(\omega) (\omega_{\mu[\omega](i)}^i - c^i(P^i[\omega])).$$

One outcome function f is said to (ex ante Pareto) *dominate* another outcome function f' if $U^i(f) \geq U^i(f')$ holds for all $i \in I$ and if $U^j(f) > U^j(f')$ holds for some $j \in I$.

⁴Note that any $E \in \zeta^i$ can be represented as the union of events $\omega_h^i \subset E$.

⁵Since $P \not\subset \zeta^i$ implies $c^i(P) = \infty$, agent i 's unique utility maximizing partition P must be ζ^i -measurable.

2.2 Standard housing problems

A housing problem $\mathcal{H} = (I, H, \Omega, \pi, c)$ is a *standard housing problem* if Ω is a singleton. Dropping I and H , and omitting π and c , which are irrelevant when Ω is a singleton, I denote a standard housing problem by ω , the profile of preferences (that is known to occur). In a standard housing problem, an agent i has a unique optimal plan of action for every choice set $S \subset H$ if and only if his preference over any two different houses is strict ($\omega_h^i \neq \omega_g^i$ for all $h \neq g$ and i). So in the subset of standard housing problems, the no-indifference condition of the present article is equivalent to the standard condition of strict preferences. The condition of independently drawn preferences is trivially satisfied in standard housing problems. The set of all standard housing problems is denoted by $\Theta := \{\omega \mid \omega_h^i \neq \omega_g^i \text{ for all } h \neq g \text{ and } i\}$. An outcome function $f: \Omega \rightarrow \mathcal{M} \times \mathcal{P}^n$ for a standard housing problem maps the only state $\omega \in \Omega$ to a matching $\mu[\omega] \in \mathcal{M}$ and the trivial partition $\{\emptyset, \Omega\}$ for every agent. Within the set of standard housing problems Θ , any outcome for a particular problem ω can, consequently, be identified with the matching $\mu[\omega]$.

A (*direct*) *mechanism* is a function $\varphi: \Theta \rightarrow \mathcal{M}$ mapping profiles of preferences $\omega \in \Theta$ to matchings $\varphi(\omega) \in \mathcal{M}$. Such a mechanism is considered strategy-proof if the truthful revelation of preferences is a weakly dominant strategy. It is nonbossy (as defined by Satterthwaite and Sonnenschein 1981) if an agent can only change the allocation of some other agent if he also changes his own allocation. This implies that any misreport of preferences that does not change the agent's own assignment does not change anyone else's assignment. The mechanism φ is considered Pareto optimal if $\varphi(\omega)$ is Pareto optimal for any ω . The following three canonical mechanisms are strategy-proof, nonbossy, and Pareto optimal.

According to a *simple serial dictatorship*, one agent, the first dictator, is matched to the best house out of H according to his stated preferences.⁶ Next, another agent, the second dictator, is matched to his most preferred house out of the remainder, and so forth, until all houses are matched. I denote a simple serial dictatorship as a direct mechanism by $\delta: \Theta \rightarrow \mathcal{M}$. The simple serial dictatorship in which agent i is the i th dictator is called δ^* . The reason for the qualifier “simple” arises since *path-dependent serial dictatorships*, denoted by $\gamma: \Theta \rightarrow \mathcal{M}$, also play a role in the present paper. This type of serial dictatorship generalizes simple serial dictatorships insofar as that the identity of any current dictator is allowed to depend on all preceding dictators' choices.⁷

Gale's top trading cycles, the third canonical mechanism,⁸ starts out with a matching μ called the initial endowment. Each agent points to the agent who has been endowed with the house he likes best according to his stated preferences. At least one cycle forms. All agents in such cycles are assigned the houses that they point to. The procedure is repeated with the remaining houses and agents until all houses are assigned.⁹

⁶Simple serial dictatorship has been characterized by Svensson (1999).

⁷Path-dependent serial dictatorships were introduced by Pápai (2001) under the name of sequential dictatorship.

⁸This mechanism was first defined by Shapley and Scarf (1974), who attribute it to David Gale.

⁹These three mechanisms are well defined when any agent has a unique most preferred house in any set of houses, as is the case for any $\omega \in \Theta$. If we allow for indifferences, the mechanisms cease to be well defined.

2.3 Dynamic direct revelation mechanisms

In this section, I define the grand set of mechanisms considered in the present article together with a list of canonical examples. Let me first argue that the sequence of preference announcements matters in mechanisms with endogenous information acquisition.

EXAMPLE 2. Two different dynamic versions of the serial dictatorship δ^* stand out: the designer might either simultaneously elicit the preferences of all agents; alternatively, the designer might elicit the preferences of all agents in order of their index i and thereby allow each agent to tailor his information acquisition to his actual choice set. To see that this difference matters, let $\mathcal{H}^b = (I, H, \Omega, \pi, c)$ with $H = \{d, g, k\}$ and three a priori identical agents. Each agent assigns value 8 or 0 (with probability $\frac{1}{2}$) to house d . The values of houses g and k are known to be 5 and 2, respectively. Assume that it costs each agent $c = 0.1$ to learn his type. If the designer simultaneously elicits preferences, it is worthwhile for agents 1 and 2 to learn their type. However, if the designer elicits preferences sequentially, then agent 2 will only learn his type if agent 1 did not choose house d . The sequential mechanism ex ante Pareto-dominates the mechanism of simultaneous elicitation, as the second dictator will not spend the cost $c = 0.1$ when learning is of no consequence to his decision. $\diamond$

In a *dynamic (direct) mechanism* the designer can fix any order of the agents' announcements. A rooted tree t , called a *c-tree*, describes the agents' communication to the designer. The initial node of a c-tree is labeled with the first agent to declare a preference. The next agent to declare a preference is allowed to depend on the declaration of the prior agent(s); branches terminate when all agents have declared their type. The designer can freely choose the sequencing of announcements as well as the information sets on the c-tree t . An agent's information set on a c-tree determines what he knows about the preceding announcements when it is his turn to reveal his type to the designer. The dynamic mechanism induced by the c-tree t and the direct mechanism φ is denoted as $\langle \varphi, t \rangle$.

Applying the dynamic mechanism $\langle \varphi, t \rangle$ to a housing problem $\mathcal{H}$, one obtains the extensive form game $\langle \varphi, t \rangle(\mathcal{H})$. This game starts with a chance node in which nature draws the state ω from π . Agents get to declare their preferences in the order determined by the c-tree t . Any agent gets to choose an information partition right before the node in which he declares his preference. The information sets in the extensive form game reflect the privacy of learning as well as the revelations implied by the c-tree t . Agent i 's utility in an end node is calculated as the difference between his value of the house he is assigned and the learning cost he incurred on the path to the node.

Of course, in many contexts, sequential learning might be impractical. This is the case when learning takes up much time or when there is a large number of agents. I therefore also study the class of mechanisms in which the designer simultaneously elicits all preferences. Formally a mechanism $\langle \varphi, t^s \rangle$ is defined as a *simultaneous (direct) mechanism* where t^s is the c-tree according to which no agent knows anything about the other agents' announcements when he announces his own preferences.

A *sequential simple serial dictatorship* or *3S dictatorship* is defined as the dynamic direct revelation mechanism $\langle \delta, t^\delta \rangle$, where t^δ is the c-tree, according to which any dictator knows the preference announcement of all preceding dictators when it is his turn to announce his preferences. When considering the simple serial dictatorship δ^* (where agent i is the i th dictator), I let $t^{\delta^*} = t^*$. Analogously a dynamic direct revelation mechanism is a *sequential path-dependent serial dictatorship* $\langle \gamma, t^\gamma \rangle$, where t^γ is such that agents publicly announce their preferences in the sequence in which they become dictators.

2.4 Equilibria and implementation

A (mixed) strategy profile in $\langle \varphi, t \rangle(\mathcal{H})$ is considered an *equilibrium* if it is a perfect Bayesian equilibrium and if agents truthfully announce their types in the sense that any agent i reveals his (true) ex post preferences $\bar{\omega}^i$ to the designer. In the standard case there exists at most one equilibrium. In that case, each agent knows his ranking ω^i and the question is just whether telling it is a best reply. My next example demonstrates that matching mechanisms with endogenous information acquisition might have multiple equilibria.

EXAMPLE 3. Consider a housing problem $\mathcal{H} = (I, H, \Omega, \pi, c)$ as follows: $n = 2$, $H = \{k, g\}$, and Ω has four equiprobable states. Agent 1's valuation of house k might be either 8 or 0; he is sure to value house g at 3. Conversely, agent 2's valuation of house g might be either 8 or 0; he is sure to value house k at 3. It costs each agent 0.1 to find out his preference. Let φ be Gale's top trading cycles mechanism where agent 1 starts out owning house k . The game $\langle \varphi, t^s \rangle(\mathcal{H})$ (in which both agents need to announce their rankings simultaneously) has two equilibria. According to the first, neither agent learns anything and always points to the house he was endowed with. According to the other, both agents learn their true values and point to the house they find to be of higher value. Note that in either one of these equilibria, the agents tell the truth. $\diamond$

Every strategy profile in the game $\langle \varphi, t \rangle(\mathcal{H})$ is associated with an outcome function $f: \Omega \rightarrow \mathcal{M} \times \mathcal{P}^n$ in the sense that the matching $\mu[\omega]$ and the set of partitions $(P^i[\omega])_{i \in I}$ (so $f(\omega) = (\mu[\omega], (P^i[\omega])_{i \in I})$) obtain at state ω when agents follow the strategy profile. A mechanism $\langle \varphi, t \rangle$ is said to *implement* a vector of ex ante utilities $(U^1(f); \dots; U^n(f))$ in the housing problem $\mathcal{H}$ if $\langle \varphi, t \rangle(\mathcal{H})$ has an equilibrium strategy profile that is associated with the outcome function f . If *all* utility vectors implemented by $\langle \varphi, t \rangle(\mathcal{H})$ dominate *all* utility vectors implemented by a different dynamic direct revelation mechanism $\langle \varphi', t' \rangle$ in $\mathcal{H}$, then $\langle \varphi, t \rangle$ is said to (ex ante Pareto) *dominate* $\langle \varphi', t' \rangle$ at $\mathcal{H}$, which is denoted by $\langle \varphi, t \rangle(\mathcal{H}) \succ^* \langle \varphi', t' \rangle(\mathcal{H})$. I say that a mechanism $\langle \varphi, t \rangle$ is (ex ante) *Pareto optimal* in a set of mechanisms if this set contains no alternative mechanism $\langle \varphi', t' \rangle$ such that $\langle \varphi', t' \rangle(\mathcal{H}) \succ^* \langle \varphi, t \rangle(\mathcal{H})$ holds for some housing problems $\mathcal{H}$.

Note that the set of Pareto-optimal mechanisms might be empty. This is the case if for every $\langle \varphi, t \rangle$, there exists an alternative mechanism $\langle \varphi', t' \rangle$ and a housing problem $\mathcal{H}$ such that $\langle \varphi', t' \rangle(\mathcal{H}) \succ^* \langle \varphi, t \rangle(\mathcal{H})$. Restricted to the set of standard housing problems Θ ,

the present notion of a Pareto-optimal mechanism coincides with the standard notion given in [Section 2.2](#). To see this, consider a direct mechanism $\varphi: \Theta \rightarrow \mathcal{M}$ that is Pareto optimal according to the notion just defined.¹⁰ This implies that there is no alternative mechanism φ' and no housing problem ω such that the matching $\varphi'(\omega)$ Pareto-dominates the matching $\varphi(\omega)$. But φ' might be any mechanism, including a constant one that maps all ω to the same matching μ . So φ is Pareto optimal according to the notion defined here if and only if there exists no profile of preferences ω such that $\varphi(\omega)$ is dominated by some matching μ . In sum, φ satisfies the standard definition of a Pareto-optimal mechanism.

3. THE UNIQUENESS OF SERIAL DICTATORSHIP

It is the goal of this article to characterize all strategy-proof and nonbossy mechanisms φ with c-trees t such that the dynamic direct revelation mechanism $\langle \varphi, t \rangle$ is ex ante Pareto optimal. The next two theorems show that simple serial dictatorship is the only such mechanism whether one allows for all dynamic direct revelation mechanisms or only for the simultaneous ones.

THEOREM 1. *A mechanism $\langle \varphi^\circ, t^\circ \rangle$ is Pareto optimal in the set of all mechanisms $\langle \varphi, t \rangle$ with φ nonbossy and strategy-proof if and only if $\langle \varphi^\circ, t^\circ \rangle$ is a 3S dictatorship.*

A very similar observation holds when one restricts attention to the set of simultaneous matching mechanisms.

THEOREM 2. *A simultaneous mechanism $\langle \varphi^\circ, t^s \rangle$ is Pareto optimal in the set of all simultaneous mechanisms $\langle \varphi, t^s \rangle$ with φ nonbossy and strategy-proof if and only if φ° is a simple serial dictatorship.*

Serial dictatorships have another outstanding welfare property: any mechanism that is not a simple serial dictatorship is dominated by a path-dependent serial dictatorship in some housing problem. Since this observation holds for dynamic as well as simultaneous mechanisms, I state only one remark that covers both cases.

REMARK 1. Fix any mechanism in the set of dynamic (simultaneous), strategy-proof, nonbossy direct revelation mechanisms that are not a 3S (simultaneous simple serial) dictatorship. There exists a housing problem in which this mechanism is dominated by a dynamic (simultaneous) path-dependent serial dictatorship.

So for any strategy-proof and nonbossy φ that is not a simple serial dictatorship and any c-tree t , there exist path-dependent serial dictatorships γ, γ' and housing problems $\mathcal{H}, \mathcal{H}'$ such that $\langle \gamma, t^\gamma \rangle$ dominates $\langle \varphi, t \rangle$ at $\mathcal{H}$ and $\langle \gamma', t^s \rangle$ dominates $\langle \varphi, t^s \rangle$ at $\mathcal{H}'$. In the [Appendix](#), I show that [Remark 1](#) cannot be strengthened by replacing path-dependent dictatorships with simple serial dictatorships. The next section contains two examples with $n \leq 3$ that do not just illustrate the two theorems; they serve as the backbone of the proof.

¹⁰I omit t here, which does not matter given that no agent can learn in a standard problem.

4. SMALL HOUSING PROBLEMS

With just two agents there is only one strategy-proof and Pareto-optimal mechanism other than serial dictatorship: Gale's top trading cycles mechanism. Let us revisit [Example 1](#) to see that the “only if” part of [Theorems 1 and 2](#) as well as [Remark 1](#) hold for $n = 2$.

EXAMPLE 4. Let $n = 2$ and $H = \{g, k\}$. Fix $\langle \varphi, t \rangle$ with φ Gale's top trading cycles mechanism in which agents 1 and 2 start out owning house k and g , respectively. According to t , at least one agent, say agent 1, has to declare his preferences before knowing the preference of the other. Reconsider the housing problem defined in [Example 1](#), which can now be defined succinctly as $\mathcal{H}^a = (I, H, \Omega, \pi, c)$ with $n = 2$, $H = \{k, g\}$, $\pi(\omega_k^i = 8) = \pi(\omega_k^i = 0) = \frac{1}{2}$, and $\pi(\omega_g^i = 2) = 1$ for $i = 1, 2$, $c^1(\bar{P}^1) = 0.8$, and $c^2(\bar{P}^2) = 0$. As argued in the [Introduction](#), agent 1's costs outweigh his benefit of learning; since agent 1 ex ante prefers house k , there is no exchange in equilibrium, yielding the ex ante utility profile $(4, 2)$. Conversely, if agent 1 is the first dictator, learning is worthwhile for him; in this case, agent 2 has a chance to obtain his ex ante preferred house g and the ex ante utility profile implemented by $\langle \delta^*, t^* \rangle(\mathcal{H})$ is $(4.2; 3)$. $\diamond$

There was only one reference to the sequence of announcements: according to t , (at least) one agent has to announce his preferences before knowing the preferences of the other. As this holds for simultaneous and sequential versions of the mechanism and as the outcome of serial dictatorship depends on only one announcement when there are just two agents, the above scenario shows the “only if” part of [Theorems 1 and 2](#) for $n = 2$. Since any simple serial dictatorship is path-dependent, [Example 4](#) also proves [Remark 1](#) for the case that $n = 2$.

To see that the “if” part of [Theorem 1](#) holds when $n = 2$, I fix an arbitrary housing problem $\mathcal{H}$ with $H = \{g, k\}$. I first show that agent 1 (weakly) prefers any equilibrium of $\langle \delta^*, t^* \rangle(\mathcal{H})$ to any equilibrium of any other mechanism $\langle \varphi, t \rangle(\mathcal{H})$. As the first dictator, agent 1 obtains the utility

$$U^* := \max_{P \subset \zeta^1} \left(\sum_{E \in P} \pi(E) \max_{h \in \{g, k\}} \bar{\omega}_h^1(E) - c^1(P) \right) \geq \max\{\bar{\omega}_g^1(\Omega), \bar{\omega}_k^1(\Omega)\},$$

where the lower bound represents the utility of not learning and then choosing the ex ante preferred house.

Now fix an arbitrary strategy for agent 2 in $\langle \varphi, t \rangle(\mathcal{H})$. Observe that this strategy might determine whether agent 1 gets to choose from $\{g, k\}$ or whether he is matched to g or k . Abusing notation, I denote the event that agent 1 gets to choose from S given agent 2's fixed equilibrium strategy by S as well. Any such event must be an element of ζ^2 , given that these are the only events on which agent 2 can condition his strategy.¹¹ Agent 2's strategy implies a distribution ρ over agent 1's choice sets $S \in \{\{g, k\}, \{g\}, \{k\}\}$.

¹¹It was assumed that $c^2(P) = \infty$ for $P \not\subset \zeta^2$.

Since the agents' preferences are independently drawn, the expected value that agent 1 assigns to any house does not vary with his knowledge of any event in ζ^2 . Under $\langle \varphi, t \rangle$, agent 1 might have to declare (and therefore learn) his preferences before he knows whether he has any choice ($S = \{g, k\}$) or not ($S = \{g\}$ or $\{k\}$). Since agent 1's utility can only increase when he may condition his choice to learn on the event S , the expression $\rho(\{g, k\})U^* + \rho(\{g\})\bar{\omega}_g^1(\Omega) + \rho(\{k\})\bar{\omega}_k^1(\Omega) \leq U^*$ yields an upper bound on agent 1's expected utility in $\langle \varphi, t \rangle(\mathcal{H})$ for the fixed strategy of agent 2.

Consequently, for $\langle \varphi, t \rangle(\mathcal{H})$ to have an equilibrium that Pareto-dominates the equilibria of $\langle \delta^*, t^* \rangle(\mathcal{H})$, agent 1 must obtain the utility U^* under the equilibrium of $\langle \varphi, t \rangle(\mathcal{H})$. However, the no-indifference condition implies that

$$\max_{P \subset \zeta^1} \sum_{E \in P} \pi(E) \max_{h \in \{g, k\}} \bar{\omega}_h^1(E) - c^1(P)$$

is uniquely maximized by a partition P^* and a function $\mu[\cdot](1): \Omega \rightarrow \{g, k\}$ that maps every state ω to a house $\mu[\omega](1)$ for agent 1. This uniqueness implies that agent 1's matches must also be described by $\mu[\cdot](1)$ for agent 1 to obtain U^* under the equilibrium of $\langle \varphi, t \rangle(\mathcal{H})$. Since there are only two agents, there is no leeway with respect to agent 2's matches. For agent 1 to obtain utility U^* in $\langle \varphi, t \rangle(\mathcal{H})$, agent 2 must—in every state ω —be matched with the house $\mu[\omega](2) \in \{g, k\}$ that is not $\mu[\omega](1)$. Since the function $\mu[\cdot](2)$ also describes agent 2's matches under the equilibrium of $\langle \delta^*, t^* \rangle(\mathcal{H})$, we can conclude that $\langle \varphi, t \rangle(\mathcal{H})$ cannot have an equilibrium that Pareto-dominates the unique equilibrium of $\langle \delta^*, t^* \rangle(\mathcal{H})$.

The “if” part of [Theorem 2](#) follows from the same arguments as the announcement of a single agent (the first dictator) determines the outcome of a serial dictatorship with just two agents, so $\langle \delta^*, t^* \rangle$ and $\langle \delta^*, t^s \rangle$ are identical for $n = 2$. The treatment of the case that $n = 2$ already contains most arguments of the proof for any n . However, when $n = 2$, all serial dictatorships are simple. Therefore, some example with $n > 2$ is in order to preview the arguments pertaining to path-dependent serial dictatorships. Next I show that path-dependent serial dictatorships can be dominated by other path-dependent serial dictatorships.

EXAMPLE 5. Take $n = 3$ and $H = \{g, k, d\}$. Consider the sequential path-dependent serial dictatorship $\langle \gamma, t^\gamma \rangle$ with agent 1 as the first dictator. If he chooses g , then agent 2 gets to choose from $\{k, d\}$; otherwise, agent 3 becomes the next dictator. Define the housing problem $\mathcal{H}^c$ such that agent 1's utility vector for the three houses is either $(2, 1, 0)$ or $(0, 2, 1)$, each with probability $\frac{1}{2}$.¹² Agent 1 faces a cost of 0.1 to learn his type. The utility vectors of agents 2 and 3 are known to be $\omega^2 = (10, 2, 0)$ and $\omega^3 = (2, 10, 0)$, respectively. The unique equilibrium of $\langle \gamma, t^\gamma \rangle(\mathcal{H}^c)$ yields the vector $(1.9; 1; 1)$ of expected utilities to agents 1, 2, and 3.

Now consider the alternative sequential path-dependent serial dictatorship $\langle \gamma', t^{\gamma'} \rangle$ that also starts with agent 1 as the first dictator, but then continues with 3 as the next

¹²Note that the a correlation of values of houses does not conflict with the independence assumption, which only requires that the preferences of different agents are drawn independently.

dictator if agent 1 chooses g and with agent 2 otherwise. The vector of ex ante utilities implemented by $\langle \gamma', t^{\gamma'} \rangle(\mathcal{H}^c)$ is $(1.9; 5; 5)$. The crucial difference between $\langle \gamma, t^{\gamma} \rangle(\mathcal{H}^c)$ and $\langle \gamma', t^{\gamma'} \rangle(\mathcal{H}^c)$ is that under the latter, agents 2 and 3 get to choose from sets where they face a utility differential of 10. Conversely, under $\langle \gamma, t^{\gamma} \rangle(\mathcal{H}^c)$, agents 2 and 3 get to choose only in situations of relatively minor relevance: when either one is called to choose, he faces a utility differential of just 2.

Since there is only one agent in $\mathcal{H}^c$ who has any information to acquire, the timing of announcements does not matter in [Example 5](#); the equilibrium sets of the games $\langle \varphi, t \rangle(\mathcal{H}^c)$ and $\langle \varphi, t^s \rangle(\mathcal{H}^c)$ are identical for any t . Any path-dependent serial dictatorship with just three agents that is not a simple serial dictatorship is—up to renaming—identical to γ . In sum, the example shows that the “only if” part of [Theorems 1](#) and [2](#) as well as [Remark 1](#) are true when one only considers the case of path-dependent serial dictatorships and $n = 3$. $\diamond$

5. LIMITATIONS AND EXTENSIONS

The “if” part of [Theorem 1](#) states that 3S dictatorship is Pareto optimal. The “only if” part together with [Remark 1](#) states that for each mechanism $\langle \varphi, t \rangle$ that is not a 3S dictatorship, there exists a housing problem $\mathcal{H}$ and a path-dependent serial dictatorship such that the latter dominates $\langle \varphi, t \rangle$ at the housing problem $\mathcal{H}$. So the “if” part could be strengthened by showing that it holds for a yet larger domain of housing problems than the one defined in [Section 2](#). Conversely the “only if” part (together with [Remark 1](#)) could be strengthened by showing that it holds on a subdomain. Here I show that one cannot enlarge the domain by much and have the “if” part continue to hold. On the other hand, the “only if” part also holds on a much smaller domain.

The domain of housing problems could be enlarged by dropping the condition of no indifference or of independently drawn preferences, or by allowing for there to be more or less houses than agents. Dropping the no-indifference condition is problematic since serial dictatorships cease to be well defined when we allow for indifferences: the truthful revelation of preferences need not imply unique choices from sets of houses. To circumvent this problem, I modify the notion of (path-dependent and simple) serial dictatorship in housing problems with indifferences. Letting m be the number of some dictator’s most preferred houses in the set of houses from which he is entitled to choose, I require that this dictator faces a probability of $1/m$ to be matched to any one of these m houses. Without indifferences, this notion of serial dictatorship reduces to the standard notion. The following example¹³ shows that 3S dictatorship can be dominated when we drop the no-indifference condition.

EXAMPLE 6. Consider a housing problem $\mathcal{H} = (I, H, \Omega, \pi, c)$ with $n = 2$ and $H = \{k, g\}$. Let Ω have four equiprobable states, where each agent’s utility schedule might be either $(1, 2)$ or $(2, 1)$. Agent 1’s cost of learning his own preference is 3; agent 2’s is 0. As the first dictator, agent 1 optimally stays ignorant, assigns a value of 1.5 to each house, and, consequently, obtains either one of the two houses with probability $\frac{1}{2}$. Facing a probability

¹³I would like to thank one of the referees for this example.

$\frac{1}{2}$ to be matched to either house, agent 2 also obtains a utility of 1.5. On the other hand, agent 2 as the first dictator chooses the house that gives him a utility of 2, implying a utility vector of (1.5; 2). So the serial dictatorship with agent 2 moving first dominates the serial dictatorship with agent 1 moving first. $\diamond$

Example 6 shows that serial dictatorship may be dominated by another mechanism in housing problems in which agents are indifferent between choices. Agent 1 does not have a unique optimal plan of action for the choice set H ; not learning and choosing k , and not learning and choosing g , as well as any mixture thereof are all optimal. To see that similar issues arise in the standard model with indifferences, consider a variation of **Example 6** in which agent 1's and agent 2's utility schedules are known to be (1.5, 1.5) and (1, 2), respectively. The serial dictatorship with agent 2 as the first dictator yields a utility vector of (1.5; 2) and, therefore, Pareto-dominates the other serial dictatorship, which yields a utility vector of (1.5; 1.5), given that agent 1 as the first dictator is equally likely to be matched to either one of the two—indifferent—houses.

Ehlers (2002) showed that the set of Pareto-optimal, strategy-proof, and nonbossy mechanisms is empty when the domain of preference profiles includes indifferences. So we certainly need to impose *some* condition of no-indifference on housing problems with endogenous information acquisition for such mechanisms to exist. Applied to standard housing problems, this no-indifference condition has to imply that preferences over different houses are strict.

The no-indifference condition defined in **Section 5** is not the only one that satisfies the requirement just mentioned: the strictness of preferences in standard housing problems is also ensured if no agent is indifferent between any two houses according to any ω in the support of π . However, **Example 6** shows that this alternative condition does not suffice for the existence of an ex ante Pareto-optimal, strategy-proof, and nonbossy mechanism.¹⁴ The same example furthermore shows that the alternative condition does not imply the no-indifference condition I impose. To see that the converse implication does not hold either, modify the housing problem defined in **Example 4** such that agent 1 is, with a small probability, indifferent between the two houses, and such that his preferences are, with the complementary probability, determined as in the description of **Example 4**. Assume furthermore that the acquisition of any partition according to which agent 1 learns whether he is indifferent or not has infinite cost. Keep all other aspects of **Example 4** fixed. The no-indifference condition defined in **Section 5** holds since agents 1 and 2 each have unique optimal plans when faced with the choice between house g and house k .

Finally observe that, in parallel to the standard case, the no-indifference condition holds generically. It is violated if an agent who learned some partition at a cost below infinity is ex post indifferent between two houses or if an agent is indifferent between learning two different partitions when he faces the problem to choose a house from

¹⁴By **Theorem 1**, we know that for any mechanism $\langle \varphi, t \rangle$ that is not a 3S dictatorship, we can find a housing problem and a mechanism $\langle \varphi', t' \rangle$ such that $\langle \varphi, t \rangle$ is dominated by $\langle \varphi', t' \rangle$ at the given housing problem. Now fix any 3S dictatorship $\langle \delta, t^\delta \rangle$ and allow for indifferences. Following **Example 6**, we can then construct a housing problem $\mathcal{H}$ and an alternative 3S dictatorship $\langle \delta', t^{\delta'} \rangle$ that dominates $\langle \delta, t^\delta \rangle$ at $\mathcal{H}$.

some set S . Mathematically, if we consider $\mathbb{R}^n$ to be the parameter space,¹⁵ then the parameters at which indifferences obtain are defined by a finite set of linear equalities and, thus, form a set of Lebesgue measure 0.

To see that the Pareto optimality of serial dictatorship also fails when we allow for correlated preferences, consider the following example.¹⁶

EXAMPLE 7. Consider a housing problem $\mathcal{H} = (I, H, \Omega, \pi, c)$ with $n = 2$, $H = \{g, k\}$, Ω consisting of two equiprobable states $\tilde{\omega}$ and $\hat{\omega}$ with $\tilde{\omega}^1 = (2, 0) = \hat{\omega}^2$ and $\hat{\omega}^1 = (0, 4) = \tilde{\omega}^2$, and $c^1(\bar{P}^1) = 0.1$ and $c^2(\bar{P}^2) = \infty$. So agent 1 prefers house g whenever agent 2 prefers house k and vice versa. Moreover, while learning is cheap for agent 1, it is prohibitive for agent 2. Agent 2 therefore always likes to cede the choice to agent 1. At the given housing problem with correlated preferences, the serial dictatorship with 2 as the first dictator is Pareto-dominated by the serial dictatorship with agent 1 as the first dictator (the respective vectors of ex ante utilities are $(2, 2)$ for the first kind of serial dictatorship and $(2.9, 3)$ for the second). $\diamond$

There is yet a further problem in environments with correlated preferences: even simple serial dictatorships need not have truth-telling equilibria for all c-trees t .¹⁷ In sum, we can say that the “if” part of [Theorem 1](#) (and [2](#)) fails if we drop the conditions of no-indifference and/or of independence. While it might be possible to relax these conditions somewhat, the search for the maximal domain of housing problems for which 3S dictatorship is ex ante Pareto optimal goes beyond the scope of this paper. Let me just say that the case of $|I| \neq |H|$ can easily be accommodated. If there are less agents than houses, the proof of the optimality of 3S dictatorship (and the analogous simultaneous case) goes through unchanged. If there are more agents than houses, we can only allow for some trivial changes: any optimal mechanism has to start out as a serial dictatorship. Once all (real) houses have been assigned and only dummy houses remain, we can use any mechanism.

Without stating a result on the minimal domain on which the “only if” part of [Theorems 1](#) and [2](#) and [Remark 1](#) are valid, let me argue that these results also hold on a much smaller domain of housing problems. To this end restrict the domain of housing problems as follows. First, any house has at most two different values for each agent, so $\omega_h^i \in \{\bar{\omega}_h^i, \underline{\omega}_h^i\}$ holds for all i, h and all $\omega \in \Omega$. Second, for each agent i , there is at most one nontrivial partition $\tilde{P}^i \neq \{\emptyset, \Omega\}$ such that $c^i(\tilde{P}^i) < \infty$. So there does not need to be much uncertainty in the housing problem: it suffices for each agent i to know that house h is either of some high value $\bar{\omega}_h^i$ or some low value $\underline{\omega}_h^i$. Moreover, just one informational choice per agent, either learn $\tilde{P}^i$ or remain ignorant ($\{\Omega, \emptyset\}$), is sufficient. Since these two additional restrictions are satisfied by all examples used to prove the “only if” part of [Theorems 1](#) and [2](#) as well as [Remark 1](#), these results are also valid for the restricted domain.

¹⁵For a fixed number of states in the space Ω , the following parameters need to be determined to describe a housing problem: a value $\omega_h^i \in \mathbb{R}$ for each agent, each house, and each state, a probability vector $\pi \in \Delta^{|\Omega|}$ on the state space, and a finite set of cost values $c^i(P^i) \in \mathbb{R} \cup \{\infty\}$, assigning a cost for every partition $P^i \subset \zeta^i$ on Ω for every agent i .

¹⁶I would like to thank one of the referees for this example.

¹⁷An example is available on request.

6. TRADING AND BRAIDING MECHANISMS

The set of all strategy-proof, nonbossy, and Pareto-optimal direct revelation mechanisms φ has been characterized by Pycia and Ünver (2013) and Bade (2014) as the set of trading and braiding mechanisms. In trading and braiding mechanisms, just as in Gale's top trading cycles mechanism, there is an initial allocation of all houses to the agents, and assignments are then determined through trade in cycles. Trading and braiding mechanisms generalize Gale's top trading cycles mechanism in three ways: First of all, agents can own more than one house before they leave with their assignment. Once an owner of multiple houses leaves the mechanism, his as of yet unmatched houses are passed on to the remaining agents via a fixed inheritance rule. Ownership of multiple houses was introduced by Pápai (2000). Second, in addition to ownership there is a new form of control called *brokerage*. Brokerage was introduced by Pycia and Ünver (2013). Third, a trading and braiding mechanism might terminate with a braid. Braids are designed to match three agents and houses with the goal to avoid one particular matching. Braids were introduced by Bade (2014).

A trading and braiding mechanism is defined using a set of control rights functions. A *control rights function* at some submatching σ $c_\sigma: \overline{H}_\sigma \rightarrow \overline{I}_\sigma \times \{o, b\}$ assigns control rights over any unmatched house to some unmatched agent and specifies a type of control. If $c_\sigma(h) = (i, x)$, then agent i *controls* house h at σ . If $x = o$, then i *owns* h ; if $x = b$ he *brokers* h . Control rights functions satisfy the following three criteria:

- (C1) If more than one house is brokered, then there are exactly three houses and they are brokered by three different agents.
- (C2) If exactly one house is brokered then there are at least two owners.
- (C3) No broker owns a house.

A *general control rights structure* c maps a set of submatchings σ to control rights functions c_σ . For now just assume that c is defined for sufficiently many submatchings to ensure that the following algorithm is well defined for any fixed ω .

Initialize with $r = 1$, $\sigma_1 = \emptyset$

Round r : Only consider the remaining houses and agents $\overline{H}_{\sigma_r}$ and $\overline{N}_{\sigma_r}$.

Braiding: If more than one house is brokered under c_{σ_r} let B be the braid defined (below) by the avoidance matching ν with $c_{\sigma_r}(\nu(i)) = (i, b)$. Terminate the process with the matching $\sigma_r \cup B(\overline{\omega})$ where $\overline{\omega}$ is the restriction of ω to $\overline{H}_{\sigma_r}$ and $\overline{I}_{\sigma_r}$. If not, go on to the next step.

Pointing: Each house points to the agent who controls it, so $h \in \overline{H}_{\sigma_r}$ points to $i \in \overline{I}_{\sigma_r}$ with $c_{\sigma_r}(h) = (i, \cdot)$. Each owner points to his most preferred house. Each broker points to his most preferred owned house.

Cycles: Select at least one cycle. Define σ° such that $\sigma^\circ(i) := h$ if i points to h in one of the selected cycles.

Continuation: Define $\sigma_{r+1} := \sigma_r \cup \sigma^\circ$. If σ_{r+1} is a matching terminate the process with σ_{r+1} . If not, continue with round $r + 1$.

A submatching σ is *reachable under c* at ω if some round of a trading and braiding process can start with σ . A submatching σ is *c -relevant* if it is reachable under c at some ω .¹⁸ A submatching σ' is a *direct c -successor of* some c -relevant σ if there exists a profile of preferences ω such that σ is reachable under $c(\omega)$ and σ' arises out of matching a single cycle at σ . A *control rights structure c* maps any c -relevant submatching σ to a control rights function c_σ and satisfies requirements (C4), (C5), and (C6).

Fix a c -relevant submatching σ° together with a direct c -successor σ .

(C4) If $i \notin N_\sigma$ owns h at σ° then i owns h at σ .

(C5) If at least two owners at σ° remain unmatched at σ and if $i \notin N_\sigma$ brokers h at σ° then i brokers h at σ .

(C6) If i owns h at σ° and σ , and if $i' \notin N_\sigma$ brokers h' at σ° but not at σ , then i owns h' at σ and i' owns h at $\sigma \cup \{(i, h')\}$.

The *braid B* is a Pareto optimal, strategy-proof, and nonbossy mechanism for a problem with exactly three houses and three agents.¹⁹ It is fully defined through an *avoidance matching ν* . Matchings $B(\omega)$ are chosen to avoid matching i to $\nu(i)$. For any ω let $\text{PO}(\omega)$ be the set of Pareto optima μ . If $\min_{\mu \in \text{PO}(\omega)} |\{i : \mu(i) = \nu(i)\}|$ is attained at a unique μ^* then let $B(\omega) = \mu^*$. If not, at least two agents must rank some house $h^* = \nu(i^*)$ at the top and the pair (i^*, h^*) is decisive in the following sense. If only one agent $j \neq i^*$ ranks h^* at the top then $B(\omega)$ is the unique minimizer that matches j to h^* . If both agents $i \neq i^*$ rank h^* at the top, then $B(\omega)$ is agent i^* 's preferred minimizer.

The trading and braiding mechanism defined by the control rights structure c is also denoted c and $c(\omega)$ is the outcome of the trading and braiding mechanism c at the profile of preferences ω . Any Pareto optimal, strategy-proof, and nonbossy mechanism has a unique representation as a trading and braiding mechanism and any trading and braiding mechanism has these three properties.

The canonical mechanisms introduced at the end of Section 2.2 can now be represented as special cases of trading and braiding mechanisms. Path-dependent serial dictatorships require that for each c -relevant σ , there exists an i_σ such that $c_\sigma(h) = (i_\sigma, o)$ for all $h \in \overline{H}_\sigma$, and simple serial dictatorships are special cases of path-dependent serial dictatorships with $i_\sigma = i_{\sigma'}$ if $|\sigma| = |\sigma'|$; a control rights structure c defines Gale's top trading cycles mechanism if there exists a matching μ such that $c_\sigma(h) = (\mu^{-1}(h), o)$ for all $h \in H$.

¹⁸Consider a control rights structure c with three agents $\{1, 2, 3\}$ and 4 houses $\{e, g, k, h'\}$, where agent 1 starts out owning house e and g and agent 2 starts out owning the remaining houses. Suppose ω is such that $\omega_e^1 \geq \omega_h^1$ and $\omega_{h'}^2 \geq \omega_h^2$ holds for all $h \in H$. Then $\{(1, e)\}$, $\{(2, h')\}$ and $\{(1, e), (2, h')\}$ are examples of submatchings that are reachable under $c(\omega)$. The submatching $\{(1, g)\}$ is c -relevant since agent 1 could appropriate house g , but not it is not reachable under $c(\omega)$ given that 1 prefers e to g . The submatching $\{(3, e)\}$ is not c -relevant since 3 does not own any house at the start of the mechanism.

¹⁹Bade (2014) defines braids for housing problems with three houses and at least as many agents. I state a simpler definition here, since the present paper is only concerned with housing problems with equally many agents and houses.

7. PROOFS

To prove that 3S dictatorship cannot be ex ante Pareto-dominated (the “if” part of [Theorem 1](#)), suppose the 3S dictatorship $\langle \delta^*, t^* \rangle$ was dominated by some $\langle \varphi, t \rangle$ at some housing problem $\mathcal{H}$. Under $\langle \delta^*, t^* \rangle(\mathcal{H})$, agent 1 obtains the ex ante utility²⁰

$$\max_{P \subset \zeta^1} \left(\sum_{E \in P} \pi(E) \max_{h \in H} \bar{w}_h^1(E) - c^1(P) \right).$$

Pick an equilibrium of $\langle \varphi, t \rangle(\mathcal{H})$ and fix the strategies of agents $\{2, \dots, n\}$ to the ones prescribed by that equilibrium. Agent 1’s problem then consists of announcing preferences that determine choices from a set $\{H^1, \dots, H^L\}$ of choice sets that are possible according to all other agents’ strategies. The strategies of agents $2, \dots, n$ imply a distribution ρ on $\{H^1, \dots, H^L\}$. Let H^l not only denote a choice set for agent 1, but also the event that agent 1 gets to choose from H^l . Since the strategies of agents 2 through n determine the set H^l that agent 1 gets to choose from and since any agent i can only condition his strategy on events in ζ^i , any event H^l can be represented as $\bigcap_{i=2}^n E_i^l$ for some events $E_i^l \in \zeta^i$ for all $i = 2, \dots, n$.

Agent 1 may have to announce (and learn) his preferences before he knows which choice set he is facing. Since agent 1’s utility can only increase, as he gets to choose a separate information partition for every choice set H^l , his ex ante utility in the fixed equilibrium of $\langle \varphi, t \rangle(\mathcal{H})$ cannot be higher than

$$\begin{aligned} & \sum_{l=1}^L \rho(H^l) \max_{P \subset \zeta^1} \left(\sum_{E \in P} \pi(E|H^l) \max_{h \in H^l} \bar{w}_h^1(E \cap H^l) - c^1(P) \right) \\ &= \sum_{l=1}^L \rho(H^l) \max_{P \subset \zeta^1} \left(\sum_{E \in P} \pi(E) \max_{h \in H^l} \bar{w}_h^1(E) - c^1(P) \right) \\ &\leq \max_{P \subset \zeta^1} \left(\sum_{E \in P} \pi(E) \max_{h \in H} \bar{w}_h^1(E) - c^1(P) \right). \end{aligned}$$

The equality holds since all agents’ preferences are drawn independently, which, in turn, implies that $\bar{w}_h^1(E) = \bar{w}_h^1(E \cap H^l)$ holds for all $h \in H$ and all $1 \leq l \leq L$. The inequality holds since the maximum in some set $S \subset \mathbb{R}$ cannot be smaller than the maximum in any subset of S . So we can conclude that agent 1’s utility in the fixed equilibrium of $\langle \varphi, t \rangle(\mathcal{H})$ is no higher than his utility as the first dictator.

Due to the no-indifference condition, agent 1’s optimal choices as the first dictator imply a unique match for him for every state ω . These matches can be described by the function $\mu[\cdot](1): \Omega \rightarrow H$. For agent 1 to be at least as well off under the equilibrium of $\langle \varphi, t \rangle(\mathcal{H})$ as under $\langle \delta^*, t^* \rangle(\mathcal{H})$, the house that agent 1 is matched with under the equilibrium of $\langle \varphi, t \rangle(\mathcal{H})$ must also be described by $\mu[\cdot](1)$.

²⁰Note that agent 1 here maximizes over all partitions $P \in \mathcal{P}$ that are subsets of ζ^1 . This is without loss of generality since $c^1(P) = \infty$ holds for any partition P with $P \not\subset \zeta^1$.

Fixing a house h^* that agent 1 chooses for some state ω as the first dictator, compare agent 2's utility in the event $E := \{\omega \mid \mu[\omega](1) = h^*\}$ under serial dictatorship and under the equilibrium of $\langle \varphi, t \rangle(\mathcal{H})$. Since agent 1 can only base his decision on a partition $P \subset \zeta^1$, the event E that agent 1 picks h^* is an element of ζ^1 . The independence assumption implies that knowing this event E has no impact on the assessment of any event that is relevant for the other agents' decisions; formally, $\pi(\omega_h^i | E) = \pi(\omega_h^i)$ holds for all $i \in \{2, \dots, n\}$ and all $h \in H$. Under serial dictatorship, agent 2 gets to choose from the set $H \setminus \{h^*\}$. He (weakly) prefers this choice to any other mechanism that matches the houses $H \setminus \{h^*\}$ to the agents $\{2, \dots, n\}$. This preference follows the same arguments given for agent 1's preference to be the first dictator. Since h^* was chosen arbitrarily, these observations hold for any possible choice by agent 1 as the first dictator.

We can conclude that conditioning on agent 1 being at least as well off as the first dictator in $\langle \varphi, t \rangle(\mathcal{H})$, agent 2 cannot be made any better off under $\langle \varphi, t \rangle(\mathcal{H})$ than under $\langle \delta^*, t^* \rangle(\mathcal{H})$. The claim then follows by an inductive application of these arguments to all consecutive dictators. The “if” part of [Theorem 2](#) can be shown using a minor modification of the above arguments. I postpone this proof to the [Appendix](#).

The proof of the “only if” part of [Theorem 1](#) together with [Remark 1](#) starts with the observation that for any mechanisms $\langle \varphi^\circ, t^\circ \rangle$ to be Pareto optimal in the sets of dynamic mechanisms, the direct revelation mechanism φ° must itself be Pareto optimal. Otherwise $\langle \varphi^\circ, t^\circ \rangle$ is dominated by some constant mechanism at some standard housing problem ω . [Pycia and Ünver's \(2013\)](#) and [Bade's \(2014\)](#) characterization then implies that φ° can be represented as a unique trading and braiding mechanism c . I subdivide the set of dynamic direct mechanisms $\langle c, t \rangle$ that are not 3S dictatorships into three categories: (I) c is not a path-dependent serial dictatorship; (II) c is a path-dependent serial dictatorship without being a simple serial dictatorship; (III) c is a simple serial dictatorship δ , but t is not equal to t^δ . [Lemmas 1, 2, and 3](#) then show that the “only if” part of [Theorem 1](#) holds restricted to mechanisms belonging to categories I, II, and III, respectively. All proofs can be found in the [Appendix](#).

LEMMA 1. *Fix any trading and braiding mechanism c that is not a path-dependent serial dictatorship and fix any c -tree t . There exists a housing problem $\mathcal{H}^A$ and a simple serial dictatorship δ such that $\langle c, t \rangle$ is dominated by $\langle \delta, t^\delta \rangle$ at $\mathcal{H}^A$.*

The proof of this lemma for the case of $n = 2$ is contained in [Example 1](#), which demonstrates the domination of Gale's top trading cycles by a serial dictatorship in some housing problems with two agents. To extend this idea of proof to [Lemma 1](#) (for any n), situations similar to Gale's top trading cycles with just two agents need to be identified in any trading and braiding mechanism that is not a path-dependent serial dictatorship.

Fix any c that is not a path-dependent serial dictatorship. By the definition of trading and braiding mechanisms there must exist a c -relevant submatching σ^* with either two owners or two brokers. Say agents 1 and 2 own houses k and g respectively, or say that agent 1 brokers g and agent 2 brokers k . Restricted to agents 1 and 2, and houses g and k , let $\mathcal{H}^A$ be identical to the problem defined in [Example 1](#). The preferences of all other agents are known. Define a matching μ such that $\mu(1) = k$, $\mu(2) = g$, and $\sigma^* \subset \mu$.

Assume that any agent $i \neq 1, 2$ prefers his match under μ to all other houses. Note that σ^* is reached in any equilibrium of $\langle c, t \rangle(\mathcal{H}^A)$. Once σ^* is reached, we face a housing problem that is strategically identical to [Example 1](#); consequently, the mechanism $\langle c, t \rangle$ is dominated by a 3S dictatorship at $\mathcal{H}^A$. The next two lemmas state the “only if” part of [Theorem 1](#) and [Remark 1](#) for the categories II and III.

LEMMA 2. *Fix any path-dependent serial dictatorship γ that is not a serial dictatorship and fix any c -tree t . There exist a housing problem $\mathcal{H}^C$ and a path-dependent serial dictatorship γ' such that $\langle \gamma, t \rangle$ is dominated by $\langle \gamma', t' \rangle$ at $\mathcal{H}^C$.*

LEMMA 3. *Fix a serial dictatorship δ together with a c -tree $t \neq t^\delta$. There exists a housing problem $\mathcal{H}^B$ such that the sequential serial dictatorship $\langle \delta, t \rangle$ is dominated by the 3S dictatorship $\langle \delta, t^\delta \rangle$ at $\mathcal{H}^B$.*

Lemmas 2 and 3 were proven by [Examples 5](#) and [2](#), respectively, for the case of $n = 3$. The proof of these lemmas for larger n consists of embedding these examples into housing problems with n houses and agents. Jointly Lemmas 2 and 3 constitute the proof of the “only if” part of [Theorem 1](#) and the relevant portion of [Remark 1](#): they show that for any conceivable deviation from 3S dictatorship, there exists a housing problem such that some path-dependent serial dictatorship dominates that mechanism at this housing problem.

Some tweaking of the preceding proof suffices to show the “only if” part of [Theorem 2](#) (and the relevant portion of [Remark 1](#)). The reason is that the sequentiality of announcements matters neither for the arguments brought forward with respect to [Examples 4](#) and [5](#) nor for their embedding in larger housing problems. The housing problems $\mathcal{H}$ chosen to prove Lemmas 1 and 2 are defined such that at most one agent learns in any of the equilibria that are relevant for the proofs. But any equilibrium of a game $\langle \varphi, t \rangle(\mathcal{H})$ in which at most one agent learns is an equilibrium of the game $\langle \varphi, t^s \rangle(\mathcal{H})$, implying that some minor translation work is needed to make the proof of [Theorem 1](#) applicable to [Theorem 2](#); the details can be found in the [Appendix](#).

8. CONCLUSION

If one allows for endogenous information acquisition in housing problems, simple serial dictatorships stand out from the large set of strategy-proof, nonbossy, and Pareto-optimal mechanisms. Whether one looks at mechanisms that dynamically elicit preferences or only at the subset of mechanisms in which preferences are elicited simultaneously, simple serial dictatorships are the only ex ante Pareto-optimal mechanisms.

Within the set of strategy-proof and nonbossy mechanisms, serial dictatorships are unique in the sense that they always provide optimal learning incentives. In a 3S dictatorship, each agent knows his choice set when it is his turn to learn and choose. Agents can, therefore, perfectly tailor their learning to fit the questions at stake. [Example 1](#) shows that this is not the case for Gale’s top trading cycles mechanism with just two agents: in this case, one agent, say agent 1, needs to decide what to learn when he

only knows the distribution over his possible choice sets. That example was constructed such that this agent avoids learning and, therefore, refuses any exchange. In addition, agent 2 would rather award dictator rights to agent 1 to get agent 1 to learn, than to keep the house he is initially endowed with. The main argument of the proof was that any strategy-proof, nonbossy, and Pareto-optimal direct choice mechanism that is not a serial dictatorship in a sense embeds Gale's top trading cycles mechanism with just two agents.

Abdulkadiroğlu and Sönmez (1998) show that serial dictatorship with the order of dictators randomly drawn from a uniform distribution is equivalent to Gale's top trading cycles mechanism with the endowment drawn from a uniform distribution. To see that this equivalence result does not hold for the case of endogenous information acquisition, reconsider the housing problem $\mathcal{H}^a$ defined and discussed in Example 1. Suppose each agent is assigned each of the roles (first or second dictator and owner of g or k , respectively) with probability $\frac{1}{2}$. The agents get to know the role they are assigned before they decide whether to acquire information about the houses. The expected utility profiles for the two different serial dictatorships are $(3, 5)$ and $(4.2, 3)$, whereas for Gale's top trading cycles mechanism, they are $(3, 5)$ and $(4, 2)$. So the vectors of expected utilities for serial dictatorship with a random order of dictators and for Gale's top trading cycles mechanism with random endowments are $\frac{1}{2}(3; 5) + \frac{1}{2}(4.2; 3) = (3.6; 4)$ and $\frac{1}{2}(3; 5) + \frac{1}{2}(4; 2) = (3.5; 3.5)$. We can conclude that randomization of serial dictatorship Pareto-dominates the randomization of Gale's top trading cycles in the present example.

There is another question relating to random matching mechanisms: could sequential serial dictatorships be replaced by random matching mechanisms in Remark 1? Example 5 suggests that this is not possible. Recall the arguments brought forward to show that agent 1 would have to be the first dictator in any mechanism that dominates the path-dependent serial dictatorship of Example 5. Applying the same arguments to the present case, we obtain that agent 1 would have to obtain the same matches as he does as the first dictator in any dominating random matching mechanism. So the randomization can only concern agents 2 and 3. But it is preferable that agents 2 and 3 are not randomly assigned to choose from the set that remains after agent 1 appropriated a house. For each agent, there is one set in which his choice matters much to him (high utility differential between the remaining houses) and another set in which his choice matters less (low utility differential between the remaining houses).

It could also be interesting to study ex ante Pareto optimality without the assumption of endogenous learning. This study could be couched in a version of the model considered here with $c^i(P) = 0$ for any partition P containing only elements $E \in \zeta^i$, meaning that agents face no cost of learning their own types. Observe that the cost of learning in Example 5 played no role, so the argument that any path-dependent serial dictatorship is dominated by another path-dependent serial dictatorship for some housing problem $\mathcal{H}$ also applies to this special case.

Example 5 can be reinterpreted as an illustration of conflict between bossiness and ex ante Pareto optimality. To see this, change the mechanism γ defined in that example to a very similar type of bossy serial dictatorship in which the identity of the second dictator does not depend on whether agent 1 chooses house g or k , but rather on whether

he ranks house d at the bottom or not. Say agent 2 becomes the second dictator if and only if agent 1 ranks house d at the bottom. This is a bossy mechanism, since agent 1's assignment does not change when announcing either $(2, 1, 0)$ or $(2, 0, 1)$. However, the assignments to the following two dictators will vary with agent 1's announcement if their preferences are aligned. Now observe that for $\mathcal{H}^c$, the housing problem given in [Example 5](#), this bossy mechanism is essentially identical to the path-dependent serial dictatorship defined there: $\mathcal{H}^c$ is defined such that agent 1 chooses house g if and only if he ranks house d lowest. Consequently, for $\mathcal{H}^c$ the given form of bossy serial dictatorship is dominated by the alternative path-dependent serial dictatorship γ' defined in the same example. This is but one example; it is not known whether ex ante Pareto optimality generally conflicts with bossiness.

Finally, let me say that a relaxation of the restrictions I imposed on the domain of housing problems might lead to a wealth of interesting results on matching with endogenous information acquisition. One stylized fact about matching mechanisms used in practice is that they often do not allow the participants to submit complete rankings; instead, they only permit short lists of a few top choices. Maybe such mechanisms fare better than classical matching mechanisms in housing problems in which agents are indifferent over all “ex ante unknown” objects. So a theory of matching with endogenous information acquisition and ex ante indifference could explain why such mechanisms are so common in practical applications. Alternatively, a theory of matching problems under endogenous information acquisition with correlated values could serve to analyze (and possibly optimize) the institutions through which agents share information before submitting their preferences to matching mechanisms.

APPENDIX

For any fixed trading and braiding mechanism c and c -relevant submatching σ^* , the mechanism c' defined through $c'_\sigma = c_{\sigma^* \cup \sigma}$ for all submatchings σ such that $\sigma^* \cup \sigma$ is c -relevant is called the *submechanism* of c at σ^* . This mechanism matches the agents in $\bar{I}_{\sigma^*}$ to the houses in $\bar{H}_{\sigma^*}$.

PROOF OF LEMMA 1. Since c is not a path dependent serial dictatorship one can fix a minimal c -relevant σ^* such that there exists no single agent who owns all houses $\bar{H}_{\sigma^*}$ at σ^* . Assume first that $c_{\sigma^*}(k) = (1, o)$ and $c_{\sigma^*}(g) = (2, o)$ holds for houses $k, g \in \bar{H}_{\sigma^*}$, and that agent 1 has to announce his preferences before he learns the announcement of agent 2 according to the c -tree t .

Define the housing problem $\mathcal{H}^A$ as follows. Fix a matching μ such that $\mu(1) = k$, $\mu(2) = g$, and $\sigma^* \subset \mu$. The preferences of all agents $i \notin \{1, 2\}$ are known. For any agent $i \neq 1, 2$ and any $h \neq \mu(i)$ assume $\omega_{\mu(i)}^i = 1 > \omega_h^i$. The values $0 > \omega_h^i$ that agents $i = 1, 2$ assign to houses h other than g and k are known. Restricted to agents 1, 2 and houses k, g , the housing problem $\mathcal{H}^A$ is identical to the housing problem $\mathcal{H}^a$ as defined in [Example 4](#).

No agent learns in the unique equilibrium of $\langle c, t \rangle(\mathcal{H}^A)$. To see this, I show first that μ obtains when all agents reveal their ex ante preferences. Since $\omega_{\sigma^*(i)}^i = 1 \geq \omega_h^i$ holds

for all $h \in H$ and all $i \in I_{\sigma^*}$, and since σ^* is c -relevant, it must be reached.²¹ Once σ^* is reached, agents 1 and 2 are endowed with houses k and g , respectively. Since agent 1 points to house k at σ^* the submatching $\sigma^* \cup \{(1, k)\}$ is reached next. By (C4), agent 2 continues to own house g at $\sigma^* \cup \{(1, k)\}$. The submatching $\sigma^* \cup \{(1, k), (2, g)\}$ is reached since agent 2 points to house g . The Pareto optimality of the submechanism of c at $\sigma^* \cup \{(1, k), (2, g)\}$, implies that all remaining agents are matched in accordance with μ .

Consider any agent $i \in I_{\sigma^*}$, fix the strategies of all other agents in I_{σ^*} to truth-telling, and fix the strategies of the agents in $\bar{I}_{\sigma^*}$ arbitrarily. To see that telling the truth is a best reply for agent i , observe that the submatching σ^* is reached if i tells the truth and all other agents follow the fixed strategy profile. So if agent i tells the truth, he obtains house $\sigma^*(i)$, his most preferred house among all houses in H . In sum, truth-telling is a best response for all agents in I_{σ^*} no matter what we assume about the strategies of the agents in $\bar{I}_{\sigma^*}$. Given that no $i \in I_{\sigma^*}$ has the choice to learn, the submatching σ^* must be reached in any equilibrium of $\langle c, t \rangle(\mathcal{H}^A)$.

Next consider the submechanism of c at σ^* . The strategic situation faced by agents 1 and 2 in that submechanism is nearly identical to the one they face in the top trading cycles mechanism constructed in [Example 4](#). The only difference is that they have some additional strategies in the present mechanism: pointing to houses other than g or k . However, all these additional strategies are dominated for any outcome of learning given that according to $\mathcal{H}^A$, $\omega_k^i, \omega_g^i > \omega_h^i$ holds for $i = 1, 2$, all $h \in H \setminus \{g, k\}$, and all $\omega \in \Omega$. The fact that not learning is the unique equilibrium in [Example 4](#) implies that conditioning on all agents in I_{σ^*} telling the truth, agents 1 and 2 best respond by not learning and truthfully revealing their ex ante preferences, so the submatching $\sigma^* \cup \{(1, k), (2, g)\}$ must be reached in any equilibrium.

The submechanism of c at $\sigma^* \cup \{(1, k), (2, g)\}$ is a trading and braiding mechanism and, therefore, is strategy-proof. This implies that the truthful revelation of their known preferences is a best reply for any agent in $\bar{I}_{\sigma^*} \setminus \{1, 2\}$, given that the agents in $I_{\sigma^*} \cup \{1, 2\}$ follow the best reply strategies described so far. In sum, we can conclude that not learning and telling the truth is the unique equilibrium of $\langle c, t \rangle(\mathcal{H}^A)$. So the profile of ex ante utilities implemented by $\langle c, t \rangle$ in $\mathcal{H}^A$ is $(4; 2; 1; \dots; 1)$.

If there are no two different owners at σ^* , there must be two brokers, say $c_{\sigma^*}(g) = (1, b)$ and $c_{\sigma^*}(k) = (2, b)$. With two brokers at σ^* , the submechanism prescribed by c at σ^* must be a braid. Without loss of generality (w.l.o.g.) we have $\bar{I}_{\sigma^*} = \{1, 2, 3\}$ and $\bar{H}_{\sigma^*} = \{g, k, e\}$. Given that 3 prefers $\mu(3) = e$ to g, k and given that 1 and 2 prefer g and k to e , agent 3 is matched with $\mu(3) = e$ under $\langle c, t \rangle$. The situation faced by 1 and 2 is strategically identical to the situation they face in [Example 4](#). So also in this case $\langle c, t \rangle$ implements the profile of ex ante utilities $(4; 2; 1; \dots; 1)$ in $\mathcal{H}^A$.

Now consider the serial dictatorship $\langle \delta^*, t^* \rangle(\mathcal{H}^A)$. Since agents 1 and 2 are the first two dictators under δ^* , and since they prefer houses k and g to all other houses in any

²¹For σ^* to be c -relevant, one agent $i_{\emptyset} \in I_{\sigma^*}$ must be the initial dictator. Since this initial agent prefers $\sigma^*(i_{\emptyset})$ to all other houses, he appropriates $\sigma^*(i_{\emptyset})$. The c -relevance of σ^* then requires that at the submatching $\{(i_{\emptyset}, \sigma^*(i_{\emptyset}))\}$, an agent $i \in I_{\sigma^*} \setminus \{i_{\emptyset}\}$ turns into the next dictator. The fact that σ^* is reached follows by induction.

state ω , their equilibrium behavior is the same as in the serial dictatorship discussed in [Example 4](#). The submechanism after the assignment of 1 and 2 is a serial dictatorship that matches each of the remaining agents i with $\mu(i)$, since each i of these agents prefers $\mu(i)$ to all remaining houses. So the unique profile of expected utilities implemented by $\langle \delta^*, t^* \rangle(\mathcal{H}^A)$ is $(4.2; 3; 1; \dots; 1)$, which Pareto-dominates the unique outcome of $\langle c, t \rangle(\mathcal{H}^A)$. $\square$

PROOF OF LEMMA 2. Fix a path-dependent serial dictatorship $\tilde{\gamma} = c$ that is not a simple serial dictatorship. Let σ^* be a minimal c -relevant submatching such that the dictator at the following submatching depends on the choice of the current dictator. Formally, assume that there exist $g, k, d \in \overline{H}_{\sigma^*}$ such that $c_{\sigma^*}(h) = (1, o)$ for all $h \in \overline{H}_{\sigma^*}$, $c_{\tilde{\sigma}}(h) = (2, o)$ for all $h \in \overline{H}_{\tilde{\sigma}}$ where $\tilde{\sigma} = \sigma^* \cup \{(1, g)\}$, and $c_{\hat{\sigma}}(h) = (3, o)$ for all $h \in \overline{H}_{\hat{\sigma}}$, where $\hat{\sigma} = \sigma^* \cup \{(1, k)\}$.²²

Define the housing problem $\mathcal{H}^C$ as follows. Fix a submatching σ' such that $\sigma^* \subset \sigma'$, $I_{\sigma'} = I \setminus \{1, 2, 3\}$, and $H_{\sigma'} = H \setminus \{g, k, d\}$. The preferences of all agents $i \neq 1, 2, 3$ are known with $\omega_{\sigma'(i)}^i = 1 \geq \omega_h^i$ holding for all $h \in H$. The values $0 > \omega_h^i$ that agents $i = 1, 2, 3$ assign to houses h other than g, k , and d are known. Restricted to agents 1, 2, 3 and houses g, k, d , the housing problem $\mathcal{H}^C$ is identical to the housing problem $\mathcal{H}^c$ as defined in [Example 5](#).

The proof that the truth-telling strategy profile according to which agent 1 learns is the only equilibrium is nearly identical to its counterpart for the preceding lemma. All agents but agent 1 know their preferences and, therefore, have a unique truth-telling strategy. Following the arguments in the preceding proof (all agents in I_{σ^*} best respond by telling the truth), σ^* must be reached in any equilibrium of $\langle \tilde{\gamma}, t \rangle(\mathcal{H}^C)$. At σ^* , agent 1 becomes the dictator. His choice problem is nearly identical to that in $\langle \gamma, t^\gamma \rangle(\mathcal{H}^c)$: in addition to $\{g, k, d\}$ (his choice set in $\langle \gamma, t^\gamma \rangle(\mathcal{H}^c)$), there might be some more houses in $\overline{H}_{\sigma^*}$; however, agent 1 strictly prefers g, k , and d to any of these additional houses for any state $\omega \in \Omega$. So agent 1's optimal behavior in $\langle \tilde{\gamma}, t \rangle(\mathcal{H}^C)$ is implied by his optimal behavior in $\langle \gamma, t^\gamma \rangle(\mathcal{H}^c)$ in [Example 5](#): agent 1 is better off learning his preferences than not. The preferences of all remaining agents are known; therefore, truth-telling is a best reply for them in the submechanism following agent 1's choice. The ex ante utilities of the agents are $(1.9; 1; \dots; 1)$.

Just as in [Example 5](#), it would be a Pareto improvement for agents 2 and 3 to “switch.” So the alternative path-dependent serial dictatorship $\langle \tilde{\gamma}', t^{\tilde{\gamma}'} \rangle$, where $\tilde{\gamma}$ and $\tilde{\gamma}'$ are identical except that agents 2 and 3 switch roles (formally, $\tilde{\gamma}' = c'$ with $c_\sigma(h) = (2, o) \Rightarrow c'_\sigma(h) = (3, o)$, $c_\sigma(h) = (3, o) \Rightarrow c'_\sigma(h) = (2, o)$, $c_\sigma(h) = c'_\sigma(h)$ for all other h), ex ante Pareto-dominates the given mechanism $\langle \tilde{\gamma}, t \rangle$ at the housing problem $\mathcal{H}^C$ and the game $\langle \tilde{\gamma}', t^{\tilde{\gamma}'} \rangle(\mathcal{H}^C)$ yields the expected utilities $(1.9; 5, 5.1; \dots; 1)$ to all players. $\square$

PROOF OF LEMMA 3. Consider the set of all relevant submatchings σ for which t prescribes that the agents in I_σ announce their preferences in the order in which they become dictators. Let σ^* be maximal in this set. This implies that if the first I_{σ^*} agents'

²²I omit the definition of the second component b_σ since in a path-dependent serial dictatorship, we have that $b_\sigma(h) = o$ for all reachable σ . Since agents 2 and 3 have a choice, there must be at least one house other than g and k .

declarations lead to the submatching σ^* , then t prescribes that some agent i , say 2, who is not the dictator at σ^* must according to t reveal his preference before learning the announcement of the dictator at σ^* , say 1. Since both 1 and 2 need to declare their preferences, their choices must be followed by at least one more agent, say 3. Define the housing problem $\mathcal{H}^B$ like $\mathcal{H}^C$ in the preceding proof with the one exception that restricted to agents 1, 2, 3 and houses k, g, d , the housing problem $\mathcal{H}^B$ is identical to the housing problem $\mathcal{H}^b$ as defined in [Example 2](#).

Given the parallel setup of the preceding and the current mechanism, truth-telling is a best reply for all agents in I_{σ^*} in $\langle \delta, t \rangle(\mathcal{H}^B)$. This implies that σ^* must be reached in any equilibrium of $\langle \delta, t \rangle(\mathcal{H}^B)$. The problem faced by agents 1, 2, and 3 at σ^* is nearly identical to that in [Example 2](#), the only difference being the availability of some inferior houses $\overline{H}_{\sigma^*} \setminus \{g, k, d\}$. Agents 1 and 2 will, therefore, both learn their value of house d in the unique equilibrium of $\langle \delta, t \rangle(\mathcal{H}^B)$.

To see that $\langle \delta^*, t^* \rangle$ dominates $\langle \delta, t \rangle$ at $\mathcal{H}^B$ observe that, just as in [Example 2](#), agent 2 will only acquire information in $\langle \delta^*, t^* \rangle(\mathcal{H}^B)$ when it is relevant to his decision. So agent 2 will only become informed when agent 1 does not choose house d . The equilibria of $\langle \delta, t \rangle(\mathcal{H}^B)$ and $\langle \delta^*, t^* \rangle(\mathcal{H}^B)$ induce the identical mappings from states ω to matchings μ . The equilibrium outcome functions differ only in one respect: there are some states ω under which agent 2 acquires information in the unique equilibrium of $\langle \delta, t \rangle(\mathcal{H}^B)$ but does not do so in the unique equilibrium of $\langle \delta^*, t^* \rangle(\mathcal{H}^B)$. So all agents but agent 2 obtain the same ex ante utility in the respective equilibria of $\langle \delta, t \rangle(\mathcal{H}^B)$ and $\langle \delta^*, t^* \rangle(\mathcal{H}^B)$. The 3S dictatorship $\langle \delta^*, t^* \rangle$ Pareto-dominates $\langle \delta, t \rangle$ at $\mathcal{H}^B$ since agent 2's expected cost of information acquisition is lower in the equilibrium of the former. $\square$

PROOF OF [THEOREM 2](#) AND THE “SIMULTANEOUS” PART OF [REMARK 1](#). This proof consists in a few amendments of the proof of [Theorem 1](#). The solution concept, which requires that all agents' ex post preferences lead to unique choices of houses from sets, remains well defined: by assumption, any (possible) ex post preferences of any agent strictly rank any two different houses. However, simple serial dictatorship might not have a unique equilibrium when all agents are forced to learn simultaneously. Some agents might have multiple optimal partitions, given all other agents' learning choices. To show that $\langle \delta^*, t^s \rangle$ cannot be dominated by any $\langle \varphi, t^s \rangle$ at any housing problem $\mathcal{H}$, I show that a *selected* equilibrium of $\langle \delta^*, t^s \rangle(\mathcal{H})$ cannot be dominated by any of the equilibria of $\langle \varphi, t^s \rangle(\mathcal{H})$.

This equilibrium is selected as follows. Let i be the first dictator, who strictly prefers some equilibrium of $\langle \delta^*, t^s \rangle(\mathcal{H})$ to another. Discard all equilibria that are inferior according to i 's preference.²³ If only one equilibrium survives, terminate the process; if not, use the preferences of the next dictator, who strictly ranks any two of the remaining equilibria to reduce the set yet further. Continue this process until only a single equilibrium survives or until all agents are indifferent between all surviving equilibria. The application of the proof of the “if” part of [Theorem 1](#) to the selected equilibrium yields the

²³This first dictator i cannot be agent 1, since the no-indifference condition implies that he has a unique optimal plan of choice from the set H .

desired result: simultaneous simple serial dictatorship is Pareto optimal in the set of all simultaneous matching mechanisms.

To prove the “only if” part of [Theorem 2](#) fix any $\langle c, t^s \rangle$ such that c is not a path-dependent serial dictatorship and define $\mathcal{H}^A$ as in the proof of [Lemma 1](#). Since only one agent learns in the unique equilibrium of $\langle \delta^*, t^s \rangle(\mathcal{H}^A)$, this strategy profile is also the unique equilibrium of $\langle \delta^*, t^s \rangle(\mathcal{H}^A)$. This, together with the observation that $\langle \delta^*, t^s \rangle$ dominates $\langle c, t^s \rangle$ at $\mathcal{H}^A$, implies that $\langle \delta^*, t^s \rangle$ dominates $\langle c, t^s \rangle$ at $\mathcal{H}^A$. The proof of [Lemma 2](#), which shows that any path-dependent but not simple serial dictatorship is dominated by a 3S dictatorship at some housing problem, directly applies to the simultaneous case. The reason is that only one agent has the choice to learn in $\mathcal{H}^C$, the housing problem constructed to prove [Lemma 2](#), so $\langle c, t \rangle(\mathcal{H}^C) = \langle c, t^s \rangle(\mathcal{H}^C)$ holds for all mechanisms $\langle c, t \rangle$. In sum, we have that for any c not a simple serial dictatorship, there exists a housing problem $\mathcal{H}$ and a path-dependent serial dictatorship γ such that $\langle c, t^s \rangle$ is dominated by $\langle \gamma, t^s \rangle$ at $\mathcal{H}$, proving the “only if” part of [Theorem 2](#) as well as the part of [Remark 1](#) that pertains to simultaneous mechanisms. $\square$

Let me finally reconsider the strengthening of [Remark 1](#) to the claim that the dominating mechanism can always be a simple serial dictatorship. To this end, suppose that the path-dependent serial dictatorship $\langle \gamma, t^\gamma \rangle$ as defined in [Example 5](#) was dominated by a 3S dictatorship $\langle \delta, t^\delta \rangle$ in some housing problem $\mathcal{H}$. Following the arguments in the proof that serial dictatorship cannot be dominated ([Section 7](#)), agent 1 would have to be the first dictator under δ to keep his utility at least as high as under γ .²⁴

Now consider the case with agent 2 as the second dictator under δ . In set $\{k, d\}$, agent 2 will choose just as he does under $\langle \gamma, t^\gamma \rangle$. If he replicates the outcome of $\langle \gamma, t^\gamma \rangle(\mathcal{H})$ for the other two choice sets $\{g, k\}$ and $\{g, d\}$, then $\langle \delta, t^\delta \rangle(\mathcal{H})$ leads to exactly the same matchings as $\langle \gamma, t^\gamma \rangle(\mathcal{H})$ and can, therefore, not be dominating. So agent 2's choices in these two sets must lead to different matchings. Agent 3 would consequently sometimes obtain a house he would not have chosen. Given that the agents' preferences are independently drawn, agent 3 would be made worse off. The alternative 3S dictatorship with agent 3 as the second dictator can be ruled out mutatis mutandis. In sum, there is no housing problem such that the path-dependent serial dictatorship $\langle \gamma, t^\gamma \rangle$ is dominated by a simple serial dictatorship $\langle \delta, t^\delta \rangle$ at that housing problem. [Remark 1](#) cannot be strengthened in the conjectured way in the case of dynamic mechanisms. The same applies to the case of simultaneous mechanisms given that all of the above arguments continue to hold if we replace t^γ and t^δ by t^s throughout the preceding paragraph.

REFERENCES

Abdulkadiroğlu, Atila and Tayfun Sönmez (1998), “Random serial dictatorship and the core from random endowments in house allocation problems.” *Econometrica*, 66, 689–701. [404]

²⁴Actually if the housing problem was such that agent 1's most preferred house in all possible states is the other two agents' least preferred house in all states, agent 1 could also be made the second or third dictator. But in that case, the serial dictatorship with agent 1 as the first dictator, keeping the order between the other two fixed, yields the same outcome in the housing problem.

Bade, Sophie (2014), “Pareto optimal, strategy proof, and non-bossy matching mechanisms.” Report, Royal Holloway, University of London. [386, 388, 399, 400, 402]

Bergemann, Dirk and Juuso Välimäki (2006), “Information in mechanism design.” In *Advances in Economics and Econometrics: Theory and Applications, Ninth World Congress of the Econometric Society*, Vol. 1 (Richard Blundell, Whitney Newey, and Torsten Persson, eds.), 186–221, Cambridge University Press, Cambridge. [387]

Compte, Olivier and Philippe Jehiel (2007), “Auctions and information acquisition: Sealed bid or dynamic formats?” *RAND Journal of Economics*, 38, 355–372. [388]

Ehlers, Lars (2002), “Coalitional strategy-proof house allocation.” *Journal of Economic Theory*, 105, 298–317. [397]

Gerardi, Dino and Leeat Yariv (2008), “Information acquisition in committees.” *Games and Economic Behavior*, 62, 436–459. [387]

Gershkov, Alex and Balázs Szentes (2009), “Optimal voting schemes with costly information acquisition.” *Journal of Economic Theory*, 144, 36–68. [387]

Pápai, Szilvia (2000), “Strategyproof assignment by hierarchical exchange.” *Econometrica*, 68, 1403–1433. [399]

Pápai, Szilvia (2001), “Strategyproof and nonbossy multiple assignments.” *Journal of Public Economic Theory*, 3, 257–271. [390]

Pycia, Marek and M. Utku Ünver (2013), “Incentive compatible allocation and exchange of discrete resources.” Working paper, UCLA. [386, 388, 399, 402]

Satterthwaite, Mark A. and Hugo Sonnenschein (1981), “Strategy-proof allocation mechanisms at differentiable points.” *Review of Economic Studies*, 48, 587–597. [390]

Shapley, Lloyd and Herbert Scarf (1974), “On cores and indivisibility.” *Journal of Mathematical Economics*, 1, 23–37. [390]

Smorodinsky, Rann and Moshe Tennenholtz (2006), “Overcoming free riding in multi-party computations—The anonymous case.” *Games and Economic Behavior*, 55, 385–406. [387]

Svensson, Lars-Gunnar (1999), “Strategy-proof allocation of indivisible goods.” *Social Choice and Welfare*, 16, 557–567. [390]