

Payoff equivalence of efficient mechanisms in large matching markets

YEON-KOO CHE

Department of Economics, Columbia University

OLIVIER TERCIEUX

Department of Economics, Paris School of Economics

We study Pareto efficient mechanisms in matching markets when the number of agents is large and individual preferences are randomly drawn from a class of distributions, allowing for both common and idiosyncratic shocks. We provide a broad set of circumstances under which, as the market grows large, all Pareto efficient mechanisms—including top trading cycles (with an arbitrary ownership structure), serial dictatorship (with an arbitrary serial order), and their randomized variants—produce a distribution of agent utilities that in the limit coincides with the utilitarian upper bound. This implies that Pareto efficient mechanisms are uniformly asymptotically payoff equivalent “up to the renaming of agents.” Hence, when the conditions of our model are met, policy makers need not discriminate among Pareto efficient mechanisms based on the aggregate payoff distribution of participants.

KEYWORDS. Large matching markets, Pareto efficiency, payoff equivalence.

JEL CLASSIFICATION. C70, D47, D61, D63.

1. INTRODUCTION

Assigning indivisible resources without monetary transfers is an important problem in modern market design; applications range from the allocation of public housing, public school seats, employment contracts, and branch postings to the assignment of human organs to transplant patients. A basic desideratum in designing such a market is Pareto efficiency. If a mechanism is not Pareto efficient, a surplus can be generated and distributed in a way that benefits (at least weakly) all participants, suggesting clear room for improvement.

Yeon-Koo Che: yeonkooche@gmail.com

Olivier Tercieux: tercieux@pse.ens.fr

We are grateful to Ludovic Lelièvre, Charles Maurin, and Xingye Wu for their research assistance and to Eduardo Azevedo, Itai Ashlagi, Julien Combe, Olivier Compte, Tadashi Hashimoto, Yash Kanoria, Fuhito Kojima, Scott Kominers, SangMok Lee, Bobak Pakzad-Hurson, Debraj Ray, Rajiv Sethi, seminar participants at Columbia, NYU, Wisconsin, Maryland, University College London, and Yale, and attendees of the PSE Market Design Conference, Warwick Micro Theory Conference, KAEA Conference, NYC Market Design Workshop, and WCU Market Design Conference at Yonsei University for helpful comments. Both authors are grateful for the support from the Global Research Network program through the Ministry of Education of the Republic of Korea and the National Research Foundation of Korea (Grant NRF-2016S1A2A2912564).

In a centralized matching market, achieving Pareto efficiency is often not difficult. A number of mechanisms are known to produce efficiency, often satisfying additional desirable properties in terms of incentives and (ex ante) fairness.¹ Rather, the challenge is often *which mechanism to choose from among many Pareto efficient mechanisms*.

This issue is important because alternative Pareto efficient mechanisms often treat individual participants differently (often dramatically so). For instance, in serial dictatorship, individuals are allowed to choose objects, one at a time, according to some serial order; the first dictator (the first individual in the serial order) could very well select the object that is regarded by everyone as the best, while the last dictator (the final individual in the order) may have to settle for the object regarded by everyone as the worst. Without monetary transfers to compensate for the loss borne by the latter, the apparent conflict of interests leaves little hope for consensus in terms of selecting from alternative Pareto efficient mechanisms. Ideally, the selection must be based on some measure of aggregate welfare of participants. For instance, if one Pareto efficient mechanism yields a significantly higher utilitarian welfare level or a much more equal payoff distribution than others, that would constitute an important rationale for favoring such a mechanism.

A similar concern arises when the designer has additional policy considerations, such as “affirmative treatment” of some target group (say, identified based on their socioeconomic background). Specifically, the designer may select an efficient mechanism based on such additional goals (for instance, by elevating the target agents’ serial orders in a serial dictatorship). Any such adjustment obviously impacts the welfare of the participants at the individual level, but do these adjustment impact the total welfare of the agents or their aggregate payoff distribution? If so, how? If accommodating additional social objectives were to entail a significant loss of utilitarian welfare or to produce a significant distributive impact, this would call into question the merit of the policy intervention. This type of policy consideration requires one to evaluate the payoff consequences of alternative Pareto efficient mechanisms. Unfortunately, little is known about how alternative mechanisms perform in this regard.

The purpose of the current paper is to fill this gap while providing useful insights on practical market design in the process. To make progress, we add some structure to the model. First, we consider markets that are “large” in terms of the number of participants as well as in the number of object types. Large markets are clearly relevant in many settings. For instance, in the US National Resident Matching Program, approximately 20,000 medical applicants participate in filling the positions of 3000–4000 programs each year. In New York City, approximately 90,000 students apply to over 700 high school programs each year. Second, we assume that the agents’ preferences are randomly generated according to some reasonable distribution. Specifically, we consider a

¹Mechanisms such as (deterministic or random) serial dictatorship produce efficiency without regard for existing property rights; top trading cycles mechanisms achieve efficiency by allowing agents to trade preexisting property rights or priorities (Shapley and Scarf 1974, Abdulkadiroğlu and Sönmez 2003). These mechanisms satisfy strategy-proofness and can satisfy ex ante “equal treatment of equals” when the serial order and initial ownership are drawn at random. Efficiency may also be achieved by allowing agents to purchase the objects using “fake money” in an artificial market, as envisioned by Hylland and Zeckhauser (1979).

model in which each agent's utility from an object depends on a common component (i.e., a portion that does not vary across agents) and on an idiosyncratic component that is drawn at random independently (and thus varies across agents). Studying the limit properties of a large market with preferences randomly generated in this way provides a framework for answering our questions.

Our main finding is that all Pareto efficient mechanisms yield aggregate payoffs, or utilitarian welfare, that converge uniformly to the same limit—more precisely, the utilitarian optimum—as the economy grows large (in the sense described above). This result implies that in large economies, alternative efficient mechanisms become virtually indistinguishable in terms of the aggregate payoff distribution of the participants. In other words, agents' payoffs are asymptotically equivalent across different efficient mechanisms, up to the “renaming” of agents. This result implies that there is no reason to favor one efficient mechanism over another. From a policy perspective, this means that a Pareto efficient allocation favoring or prioritizing a certain group of individuals does not significantly harm utilitarian welfare or significantly alter the distribution of payoffs in a large market.

Importantly, our equivalence holds in terms of the distribution of ordinal ranks enjoyed by the participants, making the result robust to the particular specification of cardinal utilities assumed. The result is also robust to the institutional details (e.g., the property rights and priorities enjoyed by some agents), which makes the result readily applicable to many practical market design problems. Finally, we consider a real-world school choice environment to test the applicability of our results to realistic market settings. While such environments are two-sided, the literature often focuses on one-sided ex post efficient mechanisms (as we do in this paper) where only students are welfare-relevant entities. In addition, several cities actually use such mechanisms to assign students to schools.² Hence, we compare alternative efficient mechanisms using field data (as well as simulated data) from the New York City school choice program. As we see, the comparison supports our equivalence result.

The present paper contributes to several strands of literature. First, our equivalence result is closely related to, and complements, the equivalence result among a class of efficient mechanisms established by Abdulkadiroğlu and Sönmez (1998) and its extensions (Pathak and Sethuraman 2011, Carroll 2014, Lee and Sethuraman 2011, and Bade 2016). As we discuss in detail, this equivalence result holds only in the absence of prior ownership or priority rights, i.e., when participants are treated ex ante symmetrically via fair lotteries. By contrast, our equivalence result holds *across* arbitrary priority or ownership structures, as long as the market is sufficiently large. This generality makes our equivalence result applicable to many real-world settings where there are often priority considerations for participants (as is the case with the New York City school choice program).

Second, our result contributes to the literature on large matching markets, particularly those with a large number of object types and random preferences; see

²San Francisco and New Orleans use (or have used) the top trading cycles mechanism to assign students to schools. Random serial dictatorship, a Pareto efficient mechanism, was used in NYC high school choice program for students who were unassigned after the main round (where DA is used).

Immorlica and Mahdian (2005), Kojima and Pathak (2009), Lee (2017), Knuth (1997), Pittel (1989), Lee and Yariv (2017), Ashlagi et al. (2017), and Che and Tercieux (2015a).³ The first three papers are largely concerned with the incentive issues arising from the deferred acceptance (henceforth, DA) mechanisms of Gale and Lloyd (1962). The last five papers are concerned with the ranks of the partners achieved by agents on two sides of a market under DA. We focus on the payoffs enjoyed by agents under *efficient* mechanisms. In a preference environment closer to ours, Lee and Yariv (2017) show that *stable* mechanisms also yield the utilitarian upper bound in a large market limit. As we show below via simulation and data analysis, efficient mechanisms tend to converge much faster than do stable mechanisms, and the magnitude of the difference can be considerable for realistic market sizes. Further, our convergence result is robust, holding even for unbalanced markets, whereas their result is not, as implied by Ashlagi et al. (2017). Most importantly, the uniform equivalence of efficient mechanisms (possibly employing different priority structures) established in the current paper is quite striking and has no analogues in the existing literature.

Methodologically, the current paper utilizes a framework developed in random graph theory; see Dawande et al. (2001), for instance. In particular, the proof method is similar to the way Lee (2017) exploits the implications of the stability of agents on two sides in a suitably defined random graph; as will become clear, our method exploits the implications of Pareto efficiency for an appropriately constructed random graph.

The rest of the paper is organized as follows. Section 2 introduces the model. Section 3 presents our main theorem. Section 4 sketches the proof. The implications of our result are discussed in Section 5. Section 6 presents evidence based on the NYC school choice data.

2. SETUP

We consider a model in which a finite set of agents are assigned a finite set of objects, at most one object for each agent. Because our analysis examines the limit of a sequence of finite economies, it is convenient to index the economy by its size n . An n -economy $E^n = (I^n, O^n)$ consists of *agents* I^n and *object types* O^n , where $|I^n| = n$. For much of the analysis, we suppress the superscript n for notational simplicity.

The object types can be interpreted as schools or housing units. Each object type o has $q_o \geq 1$ *copies* or *quotas*. Because our model allows for $q_o = 1$ for all $o \in O^n$,

³Another strand of literature studying large matching markets considers a large number of agents matched with a finite number of object types (or firms/schools) on the other side; see Abdulkadiroğlu et al. (2015b), Che and Kojima (2010), Kojima and Manea (2010), Azevedo and Leshno (2016), and Che et al. (2015), among others. The assumption of a finite number of object types enables one to use a continuum economy as a limit benchmark in these models. At the same time, this feature makes the analysis and the resulting insights quite different. The two strands of large matching market models capture issues that are relevant in different real-world settings and are thus complementary. The latter model is more appropriate for situations in which there are a relatively small number of institutions, each with a large number of positions to fill. School choice in some districts, such as Boston Public Schools, could be a suitable application because only a handful of schools enroll hundreds of students each. The former model is descriptive of settings in which there are numerous participants on both sides of the market. Medical matching and school choice in some districts, such as the New York Public Schools, would fit this description.

one-to-one matching is a special case of our model. We assume that total quantity is $Q^n = \sum_{o \in O^n} q_o = n$. In addition, we assume that the number of copies of each object is uniformly bounded, i.e., there is $\bar{q} \geq 1$ such that $q_o \leq \bar{q}$ for all $o \in O^n$ and all n . The assumption that $Q^n = n$ is made only for convenience: as long as it grows at order n , our results hold. In particular, as will become clear, our argument holds even in cases in which the market is unbalanced. Similarly, the assumption that the number of copies of each object is uniformly bounded is not necessary as long as it grows at a sufficiently low rate.⁴

Throughout, we consider a general class of random preferences that allow for a positive correlation among agents' preferences on the objects. Specifically, each agent $i \in I^n$ receives the utility from obtaining object type $o \in O^n$,

$$U_i(o) = U(u_o, \xi_{i,o}),$$

where u_o is a common value, and the *idiosyncratic shock* $\xi_{i,o}$ is a random variable drawn independently and identically from $[0, 1]$ according to a uniform distribution.⁵

We further assume that the function $U(\cdot, \cdot)$ takes values in $\mathbb{R}_+$, is strictly increasing, and is continuous in both arguments. The utility of remaining unmatched is assumed to be 0 so that all agents find all objects acceptable.⁶ The symmetry of $U(\cdot, \cdot)$ can be seen as a normalization in the scaling of individual utilities, which also implies that the highest possible utility and the lowest possible utility are identical across all agents. The symmetry assumption serves to discipline interpersonal comparison of utilities. Further, as we discuss in Section 5, our core findings are robust to the rescaling of individual utilities. We assume that the agents' common value for object type $o \in O$, u_o takes an arbitrary value in $[0, 1]$ in an n -economy, and its population distribution is given by a cumulative distribution function (CDF),

$$X^n(u) = \frac{\sum_{o \in O^n: u_o \leq u} q_o}{n},$$

which denotes the fraction of the *objects* whose common value is less than or equal to u , and by another CDF,

$$Y^n(u) = \frac{|\{o \in O^n \mid u_o \leq u\}|}{n},$$

which describes the fraction of the *object types* whose common value is no greater than u . Because $q_o \geq 1$ for each $o \in O$, it follows that $X^n(\cdot) \geq Y^n(\cdot)$.

⁴As is clear from footnote 38, we can allow $\bar{q} = O(n/\log(n))$.

⁵This assumption entails no loss of generality as long as the distribution of idiosyncratic shocks is atomless. Suppose for instance, $\xi_{i,o}$ is stochastically dependent on u_o , given by a distribution function $F(\xi_{i,o}|u_o)$. One can define a new idiosyncratic shock $\epsilon_{i,o} = F(\xi_{i,o}|u_o)$, and redefine a utility function to be $\tilde{u}(u_o, \epsilon_{i,o}) = u(u_o, F^{-1}(\epsilon_{i,o}|u_o))$. Observe that $\epsilon_{i,o}$ is now independent of u_o . This change of variables works if F is atomless (in fact, even if $\xi_{i,o}$ is unbounded as long as u is bounded). Now, if we also require that for any $\epsilon > 0$, $\inf_o \Pr\{\xi_{i,o} > 1 - \epsilon\}$ does not vanish as n grows, our argument goes through (assuming F^{-1} does not fall too fast in u_o to ensure $\tilde{u}$ is strictly increasing in u_o).

⁶This feature does not play a crucial role in our result, which holds as long as a linear fraction of objects is acceptable to all agents.

We assume that these CDFs converge to well defined limits, X and Y , in the Lévy metric. To be precise, for any two distributions F and G , consider their distance measured in the Lévy metric:

$$L(F, G) := \inf\{\delta > 0 \mid F(z - \delta) - \delta \leq G(z) \leq F(z + \delta) + \delta, \forall z \in \mathbb{R}_+\}.$$

According to this measure, any two distributions are regarded as being close to each other as long as they are uniformly close at all points of continuity.⁷ It follows that the limit distributions X and Y are nondecreasing and satisfy $X(0) = 0$, $X(1) = 1$, and $X(\cdot) - Y(\cdot) \geq 0$. We assume that X has (at most) finite jumps. We allow X and Y to be fairly general, allowing for atoms.

Several special cases of this model are of interest. The first is a *finite-tier model*. In this model, the object types are partitioned into finite tiers, $\{O_1^n, \dots, O_K^n\}$, where $\bigcup_{k \in K} O_k^n = O^n$ and $O_k^n \cap O_j^n = \emptyset$. (With a slight abuse of notation, the largest cardinality K also denotes the set of indexes.) In this model, the CDFs X^n and Y^n are step functions with finite steps. This model offers a good approximation of situations in which the objects have clear tiers, such as schools classified into different categories or regions, or houses existing in clearly distinguishable tiers. A further special case is when $K = 1$, in which the support of the common value is degenerate and agents' ordinal preferences are drawn independently and identically distributed (i.i.d.) uniformly. Knuth (1997), Pittel (1989), and Ashlagi et al. (2017) employ such a model.

Another special case is the *full-support model* in which the limit distribution Y is strictly increasing in its support. This model is very similar to Lee (2017) and Lee and Yariv (2017), who also consider random preferences that consist of common and idiosyncratic terms. One difference is that their framework assumes that the common component of the payoff is also drawn randomly from a positive interval. Our model assumes common values to be arbitrary, but they can be interpreted as realizations of random draws (drawn according to the CDF Y). Viewed in this way, the full-support model is comparable to Lee's (2017), except that the current model also allows for atoms in the distribution of Y .

Unless otherwise specified, we are referring to a general model that nests these two as special cases. Fix an n -economy. We consider a class of matching mechanisms that are Pareto efficient. A *matching* μ in an n -economy is a mapping $\mu : I \rightarrow O \cup \{\emptyset\}$ such that $|\mu^{-1}(o)| \leq q_o$ for all $o \in O$, with the interpretation that agent i with $\mu(i) = \emptyset$ is unmatched. Let M denote the set of all matchings. All these objects depend on n , although their dependence is suppressed for notational simplicity.

In practice, the matching chosen by the designer depends on the realized preferences of the agents as well as on other features of the economy. For instance, if the objects O are institutions or individuals, their preferences over their matching partners typically influences the chosen matching. Alternatively, one may wish the matching to respect the existing rights that individuals may have over the objects; for instance, if the objects are housing, some units may be occupied by existing tenants who have priority over these units. Likewise, a school choice matching may favor students whose siblings

⁷Here, convergence of CDFs in the Lévy metric is equivalent to weak convergence.

already attend the school or those living nearby. Some of these factors may depend on the features not captured by their idiosyncratic component. Our model is completely general in this regard.

Specifically, we collect all assignment-relevant variables, call its generic realization a *state*, and denote it by $\omega = (\{\xi_{i,o}\}_{i \in I, o \in O}, \theta)$, where $\{\xi_{i,o}\}_{i \in I, o \in O}$ is the realized profile of the idiosyncratic component of payoffs, and θ is the realization of all other variables (e.g., agents' priorities, tie-breaking rules, etc.) that influence the matching, and we let Ω denote the set of all possible states. We make no assumption on how θ is drawn and how its realized value affects the outcome. The generality on the θ contrasts the current model with the others, many of which tend to impose a particular random structure on the agents' priorities (or objects' preferences). For instance, Ashlagi et al. (2017) assume that these priorities are drawn i.i.d. uniformly, and Lee (2017) and Lee and Yariv (2017) assume that they consist of random common shocks with full support and i.i.d. idiosyncratic shocks. Our results do not require any such assumptions.

A *matching mechanism* is a function that maps a state in Ω to a matching in M . With a slight abuse of notation, we use $\mu = \{\mu_\omega(i)\}_{\omega \in \Omega, i \in I}$ to denote a matching mechanism, which selects a matching $\mu_\omega(\cdot)$ in state ω . Let $\mathcal{M}$ denote the set of all matching mechanisms. For convenience, we often suppress the dependence of the matching mechanism on ω .

A matching $\mu \in M$ is Pareto efficient if there is no other matching $\mu' \in M$ such that $U_i(\mu'(i)) \geq U_i(\mu(i))$ for all $i \in I$ and $U_i(\mu'(i)) > U_i(\mu(i))$ for some $i \in I$. A matching mechanism $\mu \in \mathcal{M}$ is Pareto efficient if, for each state $\omega \in \Omega$, the matching it induces, i.e., $\mu_\omega(\cdot)$, is Pareto efficient. Let $\mathcal{M}_n^*$ denote the set of all Pareto efficient mechanisms in the n -economy.⁸

3. PAYOFF EQUIVALENCE OF PARETO EFFICIENT MECHANISMS

We first define an upper bound for the utilitarian welfare—the highest possible level of total surplus that can be realized under any matching mechanism. To this end, suppose that every agent is assigned an object and realizes the highest possible idiosyncratic payoff. Because the common values of the objects are distributed according to X^n , the resulting (normalized) utilitarian welfare is $\int_0^1 U(u, 1) dX^n(u)$. This obviously yields the upper bound for the utilitarian welfare in the n -economy. We consider its limit, the *limit utilitarian upper bound*:

$$U^* := \int_0^1 U(u, 1) dX(u).$$

The payoff distribution of an economy, whether a finite n -economy or its limit, can be represented by a distribution function, i.e., a nondecreasing right-continuous

⁸While Pareto efficiency is a property of a matching (i.e., for a given profile of preferences), our result pertains to a property of Pareto efficiency that holds probabilistically across realized profiles of preferences. This requires us to focus on a Pareto efficient “mechanism” that selects a Pareto efficient matching for each profile of preferences. Equivalently, a Pareto efficient mechanism induces a *random matching* if we were to think of the matching induced by an efficient mechanism as a random variable. Our results can be interpreted either way.

function F mapping from $[0, U(1, 1)]$ to $[0, 1]$. The quantity $F(z)$ is interpreted as the fraction of agents who realize payoffs no greater than z . We let F^μ denote the payoff distribution induced by mechanism μ . In particular, let

$$F^*(v) := X(U^{-1}(v; 1))$$

denote the distribution of payoffs attaining the utilitarian upper bound U^* , i.e., when each agent achieves the idiosyncratic value of 1. We are now in a position to state our main theorem.

THEOREM 1. *Recall that $\mathcal{M}_n^*$ is the set of Pareto efficient mechanisms in the n -economy. Then*

$$\sup_{\mu^n \in \mathcal{M}_n^*} L(F^{\mu^n}, F^*) \xrightarrow{P} 0.^9$$

In words, the theorem states that the distance (in the Lévy metric) between a payoff distribution resulting from *every* Pareto efficient mechanism and that of the utilitarian upper bound vanishes uniformly in probability as $n \rightarrow \infty$. More precisely, assuming the distribution F^* is continuous, the statement is as follows. Fix any $\epsilon, \delta > 0$. Then, with a probability of at least $1 - \delta$, the proportion of agents enjoying any payoff u or higher under any Pareto efficient mechanism is within ϵ of the proportion of agents enjoying a payoff of u or higher under the utilitarian upper bound for sufficiently large n . It is remarkable that the rate of convergence is “uniform” with respect to the entire class of Pareto efficient mechanisms.

The following corollary is immediate.

COROLLARY 1. *We have*

$$\inf_{\mu^n \in \mathcal{M}_n^*} \frac{\sum_{i \in I} U_i(\mu^n(i))}{|I|} \xrightarrow{P} U^*.^{10}$$

The theorem also implies that alternative Pareto efficient mechanisms become payoff equivalent uniformly as the market grows in size, that is, “up to the renaming of the agents”:

COROLLARY 2. *We have*

$$\sup_{\mu^n, \tilde{\mu}^n \in \mathcal{M}_n^*} L(F^{\mu^n}, F^{\tilde{\mu}^n}) \xrightarrow{P} 0.$$

⁹We say $Z_n \xrightarrow{P} z$, or Z_n *converges in probability* to z , where both Z_n and z are real-valued random variables, if for any $\epsilon > 0$, $\delta > 0$, there exists $N \in \mathbb{N}$ such that for all $n > N$, we have

$$\Pr\{|Z_n - z| > \epsilon\} < \delta.$$

¹⁰For a given matching, the value of the normalized sum of payoffs only depends on the CDF induced by the matching. The corollary can be generalized to any social welfare function that only depends on the induced CDF and is continuous with respect to the Lévy metric.

These results suggest that as long as agents are ex ante symmetric in their preferences, there is little ground to favor one Pareto efficient mechanism over another in terms of the total welfare of participants or aggregate payoff distribution, at least in a large economy.¹¹ This has important implications for market design. As we already noted, designers often face extra constraints arising from the existing rights or priorities of some participants over some objects, or there may be a need to treat some target group of participants affirmatively. In addition, there is a concern that accommodating such constraints or needs may sacrifice utilitarian welfare or adversely impact the aggregate distribution of payoffs. Our result implies that accommodating such constraints does not entail any significant loss in these terms in a large economy, as long as Pareto efficiency is maintained.

REMARK 1 (Approximate ex ante efficiency). Recall that our formalism allows for “probabilistic” mechanisms to the extent that “state” ω may include random variables (e.g., lotteries). Any such mechanism is Pareto efficient according to our definition as long as its realized outcome is ex post Pareto efficient. The fact that all such mechanisms attain the utilitarian upper bound in the limit economy implies that all such mechanisms are also ex ante Pareto efficient in an asymptotic sense: as the market grows large, the proportion of agents whose ex ante welfare can be improved upon by more than $\epsilon > 0$ from any ex ante Pareto dominating reallocation (if it ever exists) vanishes in probability for any $\epsilon > 0$.¹²

4. SKETCH OF THE PROOF

Here, we sketch the proof of [Theorem 1](#), which is contained in [Appendix A](#). For our current purposes, assume $X(\cdot)$ is degenerate with a single common value u^0 and that $X(\cdot) = Y(\cdot)$. In other words, the agents have only idiosyncratic payoffs and the matching is one-to-one. As we see in [Appendix A](#), the same proof argument works for the general case (with some care).

To begin, fix an arbitrary Pareto efficient mechanism $\tilde{\mu}$. We first invoke the fact that any Pareto efficient matching can be implemented by a serial dictatorship (SD)¹³ with a suitably chosen serial order (see [Abdulkadiroğlu and Sönmez 1998](#)). Let $\tilde{f}$ be the serial order, namely, a function that maps each agent in I to his serial order in $\{1, \dots, n\}$ that

¹¹Our methodology for proving equivalence rests on Pareto efficient mechanisms attaining the utilitarian upper bound in the large market. At the same time, the two results are logically separate, so outside our domain, one may occur without the other. To what extent our equivalence result generalizes beyond our domain remains an open question.

¹²This notion is similar to asymptotic efficiency defined by [Che and Tercieux \(2015a\)](#) (although it is an ex post notion). An ordinal notion weaker than ex ante Pareto efficiency is ordinal efficiency (see [Bogomolnaia and Moulin 2001](#)); a similar approximate efficiency holds with that notion as well. [Liu and Pycia \(2016\)](#) obtain a similar result using ordinal efficiency but with a different large market asymptotic. We further discuss the relationship to this paper in [Section 5](#).

¹³A serial dictatorship mechanism specifies an order over individuals and then lets the first individual—according to the specified ordering—receive his favorite object; the next individual receives his favorite item of the remaining objects, etc.

implements $\tilde{\mu}$ under a serial dictatorship. Because $\tilde{\mu}$ induces a Pareto efficient matching that depends on the state, the required serial order $\tilde{f}$ is random.

Next, for arbitrarily small $\epsilon, \delta > 0$, define the random set

$$\bar{I} := \{i \in I \mid U_i(\tilde{\mu}(i)) \leq U(u^0, 1 - \epsilon) \text{ and } \tilde{f}(i) \leq (1 - \delta)n\}.$$

The set $\bar{I}$ consists of agents who are within the $1 - \delta$ top percentile in terms of their serial order $\tilde{f}$ but fail to achieve payoff ϵ -close to the highest possible payoff.¹⁴ Because $\epsilon, \delta > 0$ are arbitrary, for the proof, it suffices to show that

$$\frac{|\bar{I}|}{n} \xrightarrow{p} 0.$$

To prove this, we exploit a result in random graph theory. It is thus worth introducing the relevant random graph model. A *bipartite graph* G consists of vertices $V_1 \cup V_2$ and edges $E \subset V_1 \times V_2$ across V_1 and V_2 (with no possible edges within vertices in each side). An *independent set* is $\bar{V}_1 \times \bar{V}_2$, where $\bar{V}_1 \subseteq V_1$ and $\bar{V}_2 \subseteq V_2$ for which no element in $\bar{V}_1 \times \bar{V}_2$ is an edge of G . A random bipartite graph $B = (V_1 \cup V_2, p)$, $p \in (0, 1)$, is a bipartite graph with vertices $V_1 \cup V_2$ in which each pair $(v_1, v_2) \in V_1 \times V_2$ is linked by an edge with probability p (independently of edges created for all other pairs). The following result provides the crucial step for our result.¹⁵

LEMMA 1 (Dawande et al. 2001). *Consider a random bipartite graph $B = (V_1 \cup V_2, p)$, where $0 < p < 1$ is a constant, and for each $i \in \{1, 2\}$, $|V_1| = n$, and $|V_2| = m = O(n)$. There is $\kappa > 0$ such that*

$$\Pr[\exists \text{ an independent set } \bar{V}_1 \times \bar{V}_2 \text{ with } \min\{|\bar{V}_1|, |\bar{V}_2|\} \geq \kappa \ln(n)] \rightarrow 0 \quad \text{as } n \rightarrow \infty.$$

This result implies that with probability converging to 1, for every independent set, at least one side of that set vanishes in relative size as $n \rightarrow \infty$.

To prove our result, it therefore suffices to show that $\bar{I}$ forms a vanishing side of an independent set in an appropriately defined random graph. Consider a random bipartite graph consisting of I on one side and O on the other side where an edge is created between $i \in I$ and $o \in O$ if and only if $\xi_{i,o} > 1 - \epsilon$. Let

$$\bar{O} := \{o \in O \mid \tilde{f}(\tilde{\mu}(o)) \geq (1 - \delta)n\}$$

be the (random) set of objects that are assigned to the agents who are at the bottom δ percentile in terms of the serial order $\tilde{f}$. See Figure 1 for a graphical representation of the construction, where the set I is ordered according to (a realization of) the serial order.

We first observe that the (random) subgraph $\bar{I} \times \bar{O}$ is an independent set. To see this, suppose to the contrary: there is an edge between an agent $i \in \bar{I}$ and an object $o \in \bar{O}$ in

¹⁴Strictly speaking, we should focus on individuals receiving payoffs lower than $U(u^0, 1 - \epsilon)$. However, given that the utility functions are continuous, there is little loss in focusing our attention on agents receiving less than $U(u^0, 1 - \epsilon)$. This point is made clear in the proof.

¹⁵The original statement by Dawande et al. (2001) assumes that $|V_1| = |V_2| = n$. It is easily verified that their arguments also apply under our more general assumptions.

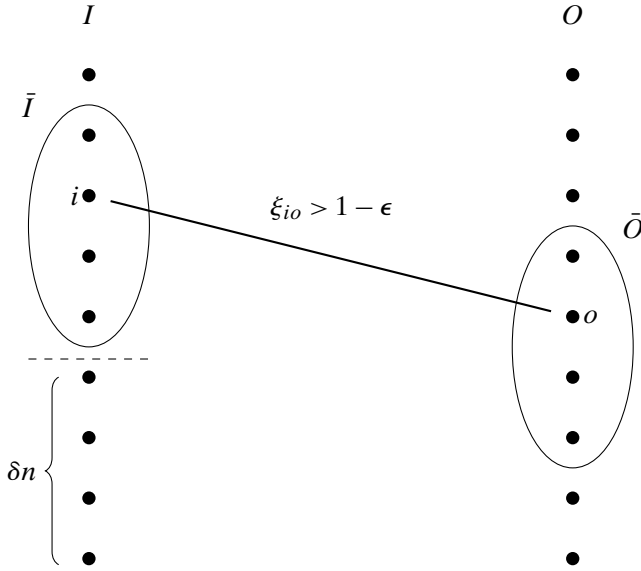


FIGURE 1. Illustration of a random graph and sets $\bar{I}$ and $\bar{O}$.

some state ω (as illustrated in Figure 1). By construction of $\bar{I}$, agent $i \in \bar{I}$ must realize less than $1 - \epsilon$ of idiosyncratic payoff from $\tilde{\mu}_\omega(i)$. However, the fact that there is an edge between i and o means that i would gain more than $1 - \epsilon$ in an idiosyncratic payoff from o . Thus, agent i must prefer o to his match $\tilde{\mu}_\omega(i)$. However, the fact that $o \in \bar{O}$ means that o is not yet claimed and is thus available when agent i (who is within the top $1 - \delta$ of serial order $\tilde{f}_\omega$) picks $\tilde{\mu}_\omega(i)$. This is a contradiction, proving that $\bar{I} \times \bar{O}$ is an independent set.

Next we observe that $|\bar{O}| \geq \delta n$, meaning that $\bar{O}$ never vanishes in probability. Lemma 1 then implies that set $\bar{I}/n$ must vanish in probability. Importantly, this result applies uniformly to all mechanisms in $\mathcal{M}^*$: If we define the sets $\bar{I}(\tilde{\mu})$ and $\bar{O}(\tilde{\mu})$ for each $\tilde{\mu} \in \mathcal{M}^*$ as above, for each $\tilde{\mu} \in \mathcal{M}^*$, $\bar{I}(\tilde{\mu}) \times \bar{O}(\tilde{\mu})$ forms an independent set of the *same* random graph. This explains the uniform convergence.

REMARK 2. If the mechanism $\tilde{\mu}$ were a serial dictatorship with a “deterministic” serial order f , a simple direct argument would prove the result. First, let us note that we can think of each agent as drawing his preferences “along the algorithm,” i.e., he draws his preferences for the stage when it is his turn to make a choice. Obviously, the distribution of i ’s preferences is not affected by the choices of agents ahead of that agent in the serial order. Fix any arbitrary $\epsilon, \delta > 0$ and let E_i be the event that at agent i ’s turn to make a choice, there remains at least one object o such that $U_i(o) \geq U(u^0, 1 - \epsilon)$. Then all agents except those in the bottom δ -percentile serial orders enjoy idiosyncratic payoffs ϵ -close to the upper bound with probability:

$$\begin{aligned} & \Pr\{U_i(\tilde{\mu}(i)) \geq U(u^0, 1 - \epsilon) \text{ for all } i \text{ with } f(i) < (1 - \delta)n\} \\ & \geq \Pr\left\{\bigcap_{i \in I: f(i) < (1 - \delta)n} E_i\right\} \end{aligned}$$

$$\begin{aligned}
 &= 1 - \Pr \left\{ \bigcup_{i \in I: f(i) < (1-\delta)n} E_i^c \right\} \\
 &\geq 1 - (1-\delta)n(1-\epsilon)^{\delta n} \rightarrow 1 \quad \text{as } n \rightarrow \infty.
 \end{aligned}$$

However, this argument does not work for an arbitrary Pareto efficient mechanism. For a general Pareto efficient mechanism, the serial order implementing the mechanism is, in general, not independent of the agents' preferences (which is required in the last inequality of the above string).¹⁶ Our general proof using random graph theory avoids this difficulty and allows us to obtain a uniform convergence result.

REMARK 3 (Rate of convergence). To evaluate an asymptotic result like [Theorem 1](#), a natural question is how large the market has to be for the result to have “significant bite.” One way to answer that question is to consider the rate of convergence. Specifically, we may ask how fast the probability that $\sup_{\mu^n \in \mathcal{M}_n^*} L(F^{\mu^n}, F^*)$ falls below some small $\delta > 0$ converges to 1. For a meaningful analysis of the convergence rate of interest, we assume $X^n = X$ for any n .¹⁷ As can be seen from the above sketch of the proof, the rate of convergence of the above probability is the same as the rate at which the probability that there is an independent set where the size of both sides is linear in n goes to 0. In [Appendix B](#), we show that this rate is “arbitrarily close” to $(1/n!)^2$. This gives a precise sense in which the convergence stated in [Theorem 1](#) is very fast. The simulation result displayed in [Figure 2](#) supports this conclusion. It shows that the welfare performance of alternative efficient mechanisms is within 1% of each other and is 10% points above that of DA, even for market with size $n = 1000$.

REMARK 4 (Absolute ranks). While the purpose of the current section is largely to illustrate the proof of [Theorem 1](#), it yields as a side product an implication that is of independent interest. Specifically, the analysis implies that when payoffs are purely idiosyncratic, all agents, except for a fraction vanishing in probability, enjoy payoffs ϵ -close to the upper bound in any Pareto efficient one-to-one matching. Specifically, we can show that all agents, except for a proportion vanishing in probability, enjoy ranks in the order very close to $\log(n)$. Formally, we state the following proposition.

¹⁶This is not just for the sake of generality: this is necessary for standard mechanisms. In particular, requiring strategy-proofness need not eliminate such dependence. Indeed, the top trading cycles mechanism (TTC) is Pareto efficient and strategy-proof, but if one implements TTC via serial dictatorship, the serial order must depend on the agents' (reported) preferences. For instance, consider an economy with three individuals and three objects. Assume that individual k owns o_k for each $k = 1, 2, 3$. Assume that 1 and 3 rank o_2 first, while both 1 and 3 rank the object they initially own (o_1 and o_3 , respectively) as their least preferred object. It is easy to show that if 2 ranks o_1 first, then 1 and 2 are involved in a trade under the top trading cycles mechanism; so serial dictatorship, which replicates this outcome, must place 3 in the last position of the serial order. Symmetrically, if 2 ranks o_3 first, then 1 must be the last in the serial order.

¹⁷Alternatively, we could assume that each object draws a common value according to distribution X and have X^n be the empirical distribution of common values. In that case, we know by the Dvoretzky–Kiefer–Wolfowitz inequality that $L(X^n, X)$ converges in probability to 0 at the very fast rate of $\exp(-n)$.

PROPOSITION 1. *Assume that for each n , $X^n(\cdot)$ is degenerate with a single common value u^0 . Let $h : \mathbb{N} \rightarrow \mathbb{R}$ be any function such that $h(n) = \omega(\log(n))$ and $h(n) = o(n)$.¹⁸ Let $\tilde{I}_\mu := \{i \in I \mid \text{rank}_\mu(i) \leq h(n)\}$. Then,*

$$\inf_{\mu \in \mathcal{M}_n^*} \frac{|\tilde{I}_\mu|}{n} \xrightarrow{P} 1,$$

where $\mathcal{M}_n^*$ is the set of Pareto efficient matchings in the n -economy with $\bar{q}_o = 1$ for all $o \in O^n$.

See [Appendix C](#) for the proof.

[Frieze and Pittel \(1995\)](#) and [Knuth \(1996\)](#) established a similar property but under two specific Pareto efficient mechanisms: random serial dictatorship and top trading cycles. [Proposition 1](#) establishes that the property holds under *all* Pareto efficient mechanisms. While this result seems closely related to [Theorem 1](#) and the argument sketched earlier, its proof is not an immediate corollary of [Theorem 1](#). First, in our proof, the probability of adding an edge in our random graph must go to 0 as n increases, which requires us to develop a new result on the size of independent sets in random graphs (see [Appendix B](#)). Second, when we relate cardinal payoffs to ordinal rankings, we have to provide a nontrivial bound on the lower tail of a binomial distribution, which requires the use of a large deviation theory (see [Appendix C](#)).

REMARK 5 (Bounded support of the idiosyncratic payoffs). A key assumption used for [Theorem 1](#) is that the upper bound on the support of the distribution of the idiosyncratic shocks does not depend on the market size. An interpretation could be that one's payoff depends on the ordinal preference rank of the object one receives.¹⁹ Nevertheless, if one takes the cardinal utility perspective seriously, it may be reasonable for the support of idiosyncratic payoffs to expand as the economy grows large. It turns out that [Theorem 1](#) is robust to expanding support. [Appendix B](#) shows that our result remains valid as long as the upper bound of the idiosyncratic payoff is $o(n/\log(n))$.

5. DISCUSSION

We now discuss the implications of our results and the roles played by the assumptions we made.

Robustness to cardinalization and ordinal equivalence

Our welfare measure is cardinal in nature, and hence is susceptible to rescaling of utilities. Scaling of individual utilities would matter at a superficial level, for instance, if we scale up the utilities of some group of agents and keep the others the same, efficient

¹⁸The Bachmann–Landau notation $\omega(\cdot)$ (not to be confused with our notation on “state”) is defined as follows. A function $h(n)$ is in $\omega(f(n))$, or simply $h(n) = \omega(f(n))$, if $|h(n)| \geq k|f(n)|$ for every positive number k .

¹⁹See [Section 5](#) for further discussion of this interpretation.

mechanisms that treat them differently would result in different aggregate utilities.²⁰ Nevertheless, there are important senses in which our insights are robust to ex ante heterogeneity.

First, regardless of how individual utilities are (re)scaled, in all efficient mechanisms, the fraction of agents who receive an arbitrarily high *idiosyncratic payoff* converges to 1 as the market grows large. Second, on more normative grounds, given that we are using the utilitarian criterion (which assigns equal weight to individual welfare), the symmetry assumption essentially reflects the idea of treating agents identically. Finally and most importantly, many institutions assess the performance of a mechanism based on the distributions of “relative ranks” that agents achieve, namely, the ordinal rank of the object obtained in each agent’s ranking divided by the number of objects allocated.²¹ Such a measure is clearly invariant to the scaling of the individual utilities. Our large market equivalence result implies equivalence in this measure.

More generally, in the original environment, our result implies equivalence of efficient mechanisms in this ordinal welfare sense. To this end, for each $u_0 \in [0, 1]$, let

$$\rho(u_0) := 1 - \left(X(u_0) + \int_{u_0}^1 \Pr\{\xi_{io} : U(u, \xi_{io}) < U(u_0, 1)\} dX(u) \right)$$

denote the (expected) proportion of objects that give higher utility than an object with common value u_0 and idiosyncratic value of 1. Now define a CDF $Z^* : [0, 1] \rightarrow [0, 1]$, given by $Z^*(r) := 1 - X(\rho^{-1}(r))$, which gives the upper bound of the fraction of agents who can achieve a (normalized) rank of r or better. [Theorem 1](#) implies the equivalence in ordinal rankings.

COROLLARY 3 (Equivalence in ordinal rankings). *For each $r \in [0, 1]$, let $Z^{\mu^n}(r)$ denote a CDF measuring the fraction of population in an n -economy who achieve normalized rank of r or better under mechanism μ^n . Then*

$$\sup_{\mu^n \in \mathcal{M}^*} L(Z^{\mu^n}, Z^*) \xrightarrow{P} 0 \quad \text{and} \quad \sup_{\mu^n, \tilde{\mu}^n \in \mathcal{M}^*} L(Z^{\mu^n}, Z^{\tilde{\mu}^n}) \xrightarrow{P} 0.$$

Role of ex ante symmetry

While our model is robust to the rescaling of cardinal utilities in the sense discussed above, ex ante symmetry remains a crucial assumption necessary for our equivalence. To illustrate this, suppose there are two object types H and L , each with quota $\frac{1}{2}n$ such that $u_H \geq u_L = 0$, and there are two agent types, 1 and 2, with their utilities given by

	$u^1(u_o, \xi_{io})$	$u^2(u_o, \xi_{io})$
$o = H$	$1 + \xi_{io}$	ξ_{io}
$o = L$	ξ_{io}	ξ_{io}

²⁰Imagine two serial dictatorship mechanisms that yield systematically different serial orders based on group membership.
²¹[Featherstone \(2015\)](#) is the first paper studying “rank efficient mechanisms,” i.e., mechanisms with rank distributions that are not stochastically dominated.

Assume there are as many type 1 agents as type 2 agents. Alternative Pareto efficient mechanisms do not yield the same utilitarian welfare even in the limit. Indeed, if the serial dictatorship is run with the serial order chosen so that type 1 agents precede all type 2 agents, then (given that agents attain idiosyncratic payoffs arbitrarily close to the upper bound of 1) the limit welfare will converge to $(\frac{1}{2})(2) + (\frac{1}{2})(1) = 1.5$ with high probability. But if the serial order is chosen so that type 2 agents precede all type 1 agents, then the limit welfare will converge to $(\frac{1}{2})(1) + (\frac{1}{2})(1) = 1$ with high probability. More importantly, nonequivalence persists even in ordinal welfare: the normalized ranks of type 1 agents converge (in probability) to 0 in the former mechanism, but they converge to $\frac{1}{2}$ in the latter mechanism (recall they all prefer tier 1 objects to tier 2 objects), while those of type 2 agents converge to 0 in both mechanisms.

Role of preference diversity

The utilitarian efficiency of Pareto efficient allocations stated in [Corollary 1](#) is remarkable and surprising given the fact that monetary transfers are not allowed. One interpretation is that a large market makes utilities virtually transferable by creating “rich” opportunities for agents to trade on idiosyncratic payoffs. In other words, objects that are uniformly valued by the participants can be transferred from one set of agents to another set without entailing much loss in terms of the idiosyncratic payoffs. In this sense, a large market can act as a “substitute” for monetary transfers.

At the same time, the role preference diversity plays for our equivalence result should not be exaggerated. It is indeed true that as the market grows large, it becomes increasingly *feasible* for agents to match with objects that yield very high idiosyncratic payoffs for them.²² Yet, this by itself does not imply that *all* Pareto efficient mechanisms match agents in that way.

First, rich preference diversity does not mean that no conflicts of interests or competition exist in our model. On the contrary, the presence of common payoff shocks entails significant conflicts of interests, as would be present in many real-world matching markets. To see this, consider a (very) special case of our environment in which there are two “tiers” of objects, all agents agreeing (i.e., irrespective of the realization of idiosyncratic payoffs) that objects in the first tier are better than those in the second tier. In such a case, agents have clear conflicting rankings over ex post efficient mechanisms (a group of agents may be favored in getting the objects in the first tier for some ex post efficient mechanism, while some other group would be favored in others).

Second, such conflicts of interests could easily entail significant loss of welfare for agents and yield distinct utilitarian welfares across different Pareto efficient mechanisms. This is indeed the case if the matching is two-sided, i.e., the objects’ welfare is

²²This is actually implied by the well known Erdős–Rényi theorem. To see this, as before, one can represent our matching model by a random bipartite graph in which the nodes on both sides represent agents and objects, respectively, and an edge between an agent and an object is generated whenever $\xi_{io} > 1 - \epsilon$ for any arbitrary (but fixed) ϵ . The Erdős–Rényi theorem implies the existence of a perfect matching in large random graphs. Adapted to our random graph, this means that, with high probability, all agents can be assigned objects that yield within ϵ of the maximal idiosyncratic payoffs as $n \rightarrow 1$.

included in Pareto efficiency (think of them as “schools”). The Gale and Shapley’s deferred acceptance algorithm is a Pareto efficient mechanism in such a setting. But as long as the market is unbalanced (imagine that there are more agents than objects, and the difference grows linearly in n), we know from [Ashlagi et al. \(2017\)](#) and [Che and Tercieux \(2015a\)](#) that agents on the long side of the market compete in DA to such an extent that they suffer a significant loss of welfare even when there are no common shocks. The “two tiers” market described above effectively involves the same kind of imbalance (with respect to the top tier objects), so Pareto efficient mechanisms are *not* utilitarian efficient, even though it becomes increasingly feasible to find a matching that yields high idiosyncratic payoffs for all agents on both sides as the market grows.²³ This also means that large-market payoff equivalence does not hold when the matching is two-sided.²⁴

Therefore, it is rather remarkable that, in the one-sided matching market, conflicts of interests do not give rise to a significant loss in utilitarian welfare or a significant difference in welfare across alternative efficient mechanisms. Also striking, and not anticipated from the modeling structure at all, is the uniformity in the convergence rates in which *all* efficient mechanisms become equivalent as the market grows. The uniform convergence result appears particularly important for the payoff equivalence to *matter* for market sizes that are of practical relevance, as we seen in our simulation below (see [Figures 2 and 3](#)) and in the data set based on the New York City high school assignment (see [Section 6](#)).

Connection with the existing equivalence results

The current equivalence result is reminiscent of a similar equivalence result obtained by [Abdulkadiroğlu and Sönmez \(1998\)](#) between two well known mechanisms—random serial dictatorship (RSD) and TTC with random ownership—and of the large market equivalence result obtained by [Che and Kojima \(2010\)](#) between random serial dictatorship and a probabilistic serial mechanism as well as their extensions ([Pathak and Sethuraman 2011](#), [Carroll 2014](#), [Lee and Sethuraman 2011](#), [Bade 2016](#), and [Liu and Pycia 2016](#)). While these results consider arbitrary preferences on the agents, they assume *ex ante* symmetric random priorities with respect to the objects. By contrast, our equivalence result does not impose any structure on the priorities on the object side and in fact holds across different priority structures. That *ex ante* symmetric random priorities are not required for our equivalence is particularly important in practice, since, in many settings such as school choice and medical matching, applicants are typically treated differently by schools or hospitals through priorities, preference rankings, or interview decisions. It also implies that serial dictatorship mechanisms with priorities chosen to “minimize” utilitarian welfare and “maximize” it all converge uniformly in terms of the utilitarian

²³To see this, a random graph in [footnote 22](#) can be modified so that an edge between an agent and an object is created if and only if the payoffs of the agent *and* the object are both very high. Then the Erdős–Rényi’s result applied to this random graph establishes the feasibility of a matching that realizes high idiosyncratic payoffs.

²⁴Serial dictatorship is obviously Pareto efficient even in a two-sided economy, and [Theorem 1](#) implies that the agents enjoy close to the upper bound in this allocation, whereas they do not in the DA allocation, another Pareto efficient allocation.

welfare attained. This has no analogue in the existing equivalence results. At the same time, our equivalence result requires a certain structure on the agents' preferences (to consist of common values and i.i.d. idiosyncratic shocks), and holds only in the limit as the number of agents and objects (types) becomes large, whereas the equivalence results mentioned above holds for any finite economy. Ultimately, our result complements the existing focus on (ex post) Pareto efficient mechanisms that treat agents symmetrically in terms of tie-breaking or ex ante assignment of property rights. One implication of our result is that the fairness achieved by symmetric treatment of agents does not compromise welfare significantly, neither in terms of utilitarian welfare nor of the payoff distribution among agents.

Efficiency versus stability

One may wonder to what extent the payoff performances of efficient mechanisms are due to *efficiency*. Indeed, a similar payoff performance is known to arise from a stable matching in some large market settings. Knuth (1997) and Pittel (1989), among others, show that if the agents' ordinal preferences are drawn i.i.d. uniformly and the market is balanced, the aggregate welfare of the agents under a stable matching approaches the utilitarian upper bound.²⁵ Lee and Yariv (2017) show a similar result in a large balanced market in which the agents on both sides have random preferences that consist of common and idiosyncratic shocks, and the common shock has full support.

Figure 2 shows the normalized utilitarian welfare performance of the stable matching in comparison with three different efficient mechanisms in a simulated assignment problem. In this simulation, agents' preferences are given by $U(u_o, \xi_{io}) = u_o + \xi_{io}$ and their priorities are given by school utilities $V(v_i, \eta_{io}) = v_i + \eta_{io}$, where u_o , ξ_{io} , v_i , and η_{io} are all drawn independently and uniformly from $[0, 1]$. This specification allows for a reasonable amount of correlation in students' preferences and their priorities, an environment conducive to high welfare performance for a stable matching. The three efficient mechanisms we consider are as follows. Top trading cycles (TTC), proposed by Abdulkadiroğlu and Sönmez (2003), yields a Pareto efficient allocation by executing trades among agents in multiple rounds.²⁶ For the two remaining efficient mechanisms, we randomly draw a serial order of agents 100 times and select the best- and worst-performing assignments in terms of utilitarian welfare (these mechanisms are named SDmax and SDmin, respectively). As pointed out earlier, any efficient assignment corresponds to SD with some serial order, so the two SDs encompass a large range of welfare performances achievable by efficient mechanisms. The objective is to see the range of variations in utilitarian welfare associated with different Pareto efficient mechanisms,

²⁵Specifically, they show that the (preference) rank of the objects agents enjoy converges to $\log(n)$ on average, which means that the idiosyncratic payoffs are on the order of $1 - \log(n)/n$ on average. Because the common values are degenerate in their environment, the result follows.

²⁶In each round, each applicant points to his most preferred object among those available, each object points to the applicant with the highest priority for the object among those available, and the applicants associated with a cycle (which must exist due to the finite number of participants) are assigned the objects they point to and exit the market, along with the seats they are assigned. The same process is repeated with the remaining participants, until all participants are exhausted.

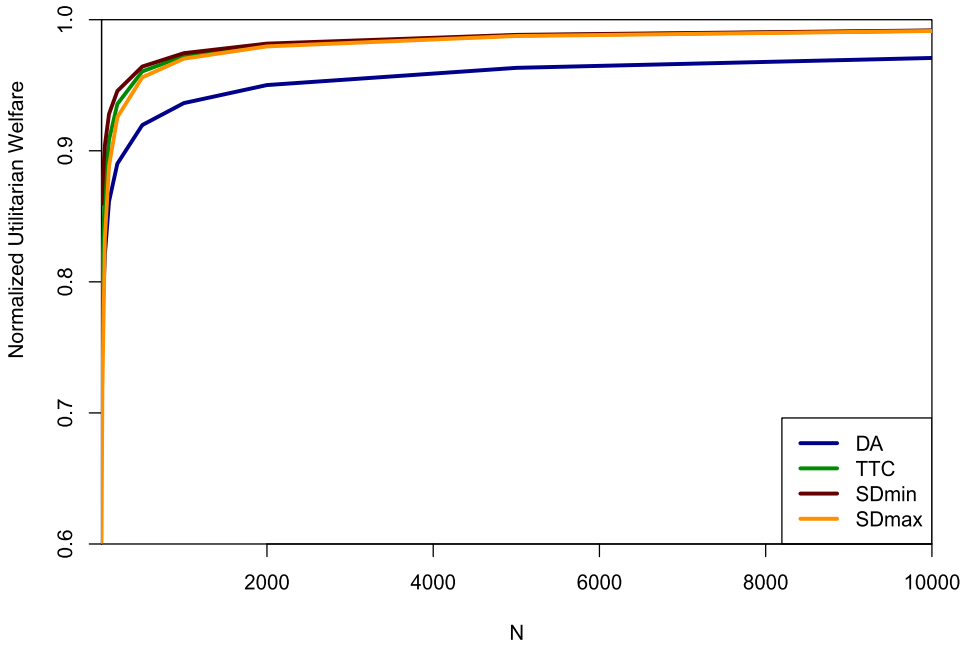


FIGURE 2. Normalized utilitarian welfare under alternative mechanisms.

and compare it with the DA that selects the student-optimal stable assignment for the agents.²⁷

For each mechanism, we normalize the utilitarian welfare by subtracting the common value component u_o ; essentially, we focus on the idiosyncratic utility realized by the mechanism. (This is without loss since the common value u_o affects all agents identically.) Specifically, we average the normalized utilitarian welfare over 50 iterations of the preference and priority draws $\{(u_o, \xi_{io}, v_i, \eta_{io})\}_{(i,o)}$.²⁸

All mechanisms, including DA, perform well and improve as the market grows large. These results are in line with our main findings (in particular, [Corollary 1](#)) and with [Lee and Yariv \(2017\)](#). What is not implied by these results and can be learned from the simulations is the speed of convergence and the uniformity across mechanisms. In this regard, two observations can be made. The first is the degree of payoff proximity of alternative Pareto efficient mechanisms. For market size $n = 2000$, all efficient mechanisms—even the SDmin—realize more than 98% of the highest possible surplus. The uniformity

²⁷Reporting the results for DA serves an additional purpose. The difference between DA and efficient mechanisms makes it clear that the similarity across efficient mechanisms is not driven by the correlation in preferences that are assumed in the simulation. See [footnote 36](#) for a related point.

²⁸For SDmax and SDmin, we have 50 preference draws of $\nu := \{(u_o, \xi_{io})\}_{(i,o)}$. For each preference draw, we obtain the maximum or minimum normalized utilitarian welfares among 100 draws of serial orders. We then average the welfare over 50 preference draws. Specifically, let $\bar{U}(\nu; f)$ denote the normalized utilitarian welfare for a profile ν of preferences and a serial order f . Then $\text{SDmin} := \frac{1}{50} \sum_{\nu \in \tilde{\Omega}} (\min_{f \in \tilde{F}} \bar{U}(\nu; f))$, $\text{SDmax} := \frac{1}{50} \sum_{\nu \in \tilde{\Omega}} (\max_{f \in \tilde{F}} \bar{U}(\nu; f))$, where $\tilde{F}$ is the set of 100 random samples of serial orders, and $\tilde{\Omega}$ is the set of 50 random samples of agents' preferences.

of convergences across efficient mechanisms is indeed remarkable, for they become essentially indistinguishable for any $n \geq 1000$.²⁹ Second, although DA performs well in utilitarian efficiency, there is a clear difference relative to the efficient mechanisms. For $n = 2000$, there is a 3% point difference relative to SDmin, and the tangible difference of at least 2% of points remains even for $n = 10,000$. This suggests that the welfare converges more slowly to the utilitarian upper bound under DA than under efficient mechanisms. Indeed, the difference becomes more pronounced as the common value component of agents' preferences becomes more important.³⁰

Robustness to the large market asymptotics

We have assumed that the number of object types grows linearly in the market size while the quota for each object type is bounded at $\bar{q}$, as has been assumed by several authors (e.g., Ashlagi et al. 2017, Che and Tercieux 2015a, Immorlica and Mahdian 2005, Kojima and Pathak 2009, Knuth 1997, and Pittel 1989). In fact, our result continues to hold even when the number of object types grows much more slowly in the market size. Specifically, the number of object types can grow as slow as on the order $\log(n)$, and the quota/copies of each object type can grow as fast as on the order $n/\log(n)$. At the same time, the equivalence does not extend if the quota grows faster, e.g., linearly in the market size, as has been assumed by a few authors (Abdulkadiroğlu et al. 2015b, Che and Kojima 2010, Azevedo and Leshno 2016, and Liu and Pycia 2016). To see this, suppose there are two object types, say o_1 and o_2 , each with quota $n/2$. Assume further that all agents prefer o_1 objects to o_2 objects, regardless of their idiosyncratic shocks. If a serial dictatorship is run with the serial order given by the realized value of $\xi_{i,o_1} - \xi_{i,o_2}$ (so that the agents with high values go first), the expected (normalized) utilitarian welfare converges to $\frac{1}{2}U(u_1, \frac{3}{4}) + \frac{1}{2}U(u_2, \frac{3}{4})$ as $n \rightarrow \infty$. By contrast, if the serial dictatorship is run with the opposite serial order (namely, in the descending order of $\xi_{i,o_2} - \xi_{i,o_1}$), the expected (normalized) utilitarian welfare converges to $\frac{1}{2}U(u_1, \frac{1}{4}) + \frac{1}{2}U(u_2, \frac{1}{4})$ as $n \rightarrow \infty$. This example shows that our equivalence result is distinct from, and is thus not implied by, the equivalence found by these authors.

Ultimately, which asymptotics is valid depends on the applications. According to our simulations below, the asymptotics allowed here covers a very large range. Indeed, we conducted simulations with $k = 10, 20, 30, 40, 50$ schools, each with a quota of 100 and $100k$ in total students, where the students' preferences and priorities are generated

²⁹Note that the current equivalence result holds across different priorities and therefore is distinct from—not implied by—the equivalence result by Abdulkadiroğlu and Sönmez (1998) and its extensions. Indeed, for very small n , the difference between SDmin and SDmax is appreciable, suggesting that the classical equivalence result is not applicable here. But even for a modestly large n , the difference between the two mechanisms vanishes.

³⁰Simulations demonstrating this point are available from the authors as well as in Che and Tercieux (2015b). The intuition follows from Ashlagi et al. (2017) and Che and Tercieux (2015a). Namely, the limit result of Lee and Yariv (2017) does not extend to a situation in which students compete for scarce seats at good schools. As the common value component becomes more important (e.g., the support of u_o increases), agents compete more vigorously for objects with high common value. Such a competition entails significant welfare loss in a stable mechanism, but not under Pareto efficient mechanisms.

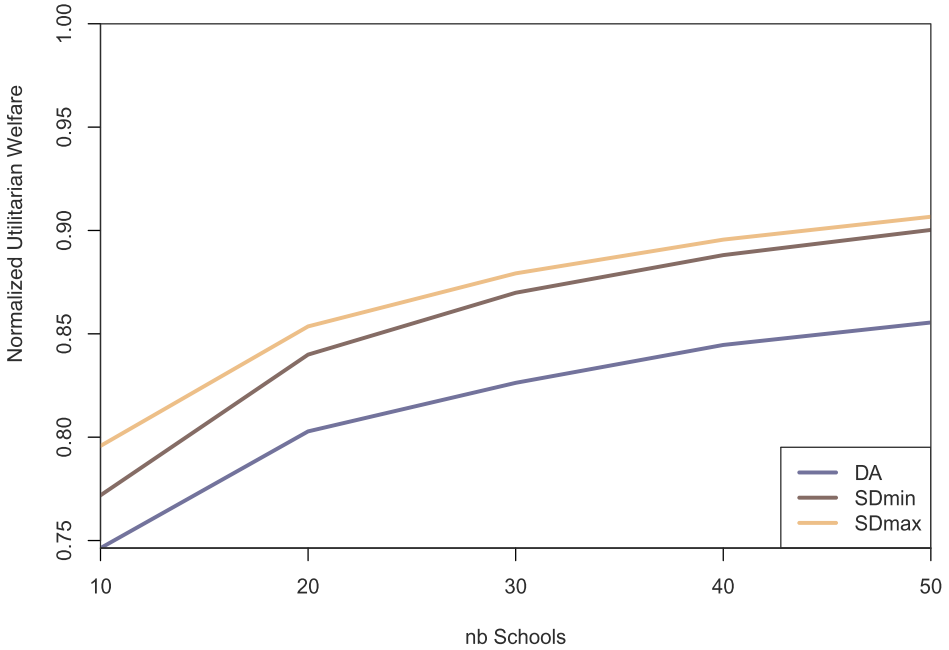


FIGURE 3. Normalized utilitarian welfare for schools across alternative school quotas.

based on the earlier simulations. This specification allows for a reasonable amount of correlation in students' preferences and their priorities, as would be natural in a typical school choice setting. In addition, this covers typical school choice settings in which the quota (the number seats at a given school) is typically in the 100s, and the number of schools/programs can be in the 10s.³¹ Figure 3 shows the utilitarian welfare of alternative mechanisms as a percentage of its upper bound. The difference between SDmax and SDmin captures the degree of possible payoff inequivalence across Pareto efficient mechanisms. As can be seen, the difference between two mechanisms is less than 1.5% for 20 schools (a very realistic number for schools) and less than 1% for 30 schools. By contrast, the difference between SDmax and DA is more than 5% and does not shrink even for $k = 50$.

6. EVIDENCE FROM NYC SCHOOL CHOICE

What do our results imply for realistic markets? We study this question based on the choice data supplied by the New York City Department of Education. In New York City, approximately 90,000 students (mostly in the 8th grade) are assigned to over 700 public high school programs through an annual centralized matching process. We focus on the main round (round 2) of assignment. In that round, each student submits a rank ordered list (ROL) of up to 12 programs, and each program ranks applicants who listed it on their

³¹In the New York City high school assignment, the number of programs is around 750, and the number of students is close to 80,000. In the Boston school system, the number schools is 18 for about 4000 students willing to enter grade 9.

ROIs according to its priority criteria, which depend on the types of the program.³² The priorities are coarse for many programs, and a single (uniform) lottery is used to break ties for all programs. Given the ROIs and priorities, a student-proposing DA algorithm is used to generate an assignment.

We use the 2009–2010 choice data to calibrate the assignments that would arise under DA and three Pareto efficient mechanisms: TTC, random serial dictatorship (RSD), and versions of SDmin and SDmax.³³ For DA and TTC, we run the mechanism using a random serial order as a tie-breaker. We then iterate 100 times and obtain the average distribution of ranks enjoyed by the participants, namely the average number of students achieving each rank $r = 1, 2, \dots, 12$. For RSD, we again generate 100 random draws of serial order and run SD with each serial order, and obtain an average distribution of ranks. For SDmax (SDmin), we select a serial order from among the 100 draws to minimize (maximize) the total sum of ranks enjoyed by the agents.³⁴ We iterate this procedure 100 times and obtain the average distribution of ranks enjoyed by the participants under each mechanism. Here again, the purpose is to see the range of variations in distribution of ranks associated with different Pareto efficient mechanisms, and to compare it with the DA. These mechanisms differ in the way that they treat the participants, so there is no a priori expectation for the relation of the distributions.

Before proceeding, a couple of remarks are in order. First, following the existing literature, we assume that the observed ROIs of the applicants represent their truthful preference ranking of top programs. This assumption is not entirely innocuous because the strategy-proofness of DA does not apply when the applicants' ROIs are truncated (see Haeringer and Klijn 2009). Nevertheless, approximately 80% of the participants did not fill out their ROIs, suggesting that truncation was not a binding constraint (see Abdulkadiroğlu et al. 2009 and Abdulkadiroğlu et al. 2015a for the same assumption). Second, under the current DA algorithm, programs do not specify priorities for students unless they rank them in their ROIs. To calibrate TTC, we assume that programs assign lower priorities to students who do not rank them than to those who do rank them. For our purpose, this does not appear to pose a serious problem.³⁵

Table 1 and Figure 4 present the distribution of preference ranks achieved by the applicants under alternative efficient mechanisms, using DA as a control mechanism. They exhibit a striking resemblance in the rank distribution across alternative Pareto efficient mechanisms, and a noticeable difference between them and the DA outcome.

³²The programs are categorized into several types in terms of admissions method: screened, limited unscreened, unscreened, ed-op, zoned, and audition. See Che and Tercieux (2015a) for a detailed description of the data and the institutional details.

³³RSD is a serial dictatorship where the applicants' serial orders are determined at random.

³⁴In this definition, the number of unassigned students is not taken into account. An alternative would be to assign rank 13 to unassigned students and define SDmax and SDmin accordingly. The figures we obtain in this section are almost the same under this alternative.

³⁵First of all, the Pareto efficiency of TTC is not affected by this feature. Second, the TTC outcome is likely to lie between those of SDmax and SDmin, which appear to be close to each other according to our finding below.

	DA	TTC	SDmin	SDmax	RSD
1	35,200.87 (53.67)	38,090.25 (36.58)	37,632.46 (59.06)	37,674.23 (62.59)	37,657.08 (51.03)
2	14,006.80 (53.01)	13,256.99 (46.08)	13,272.93 (66.06)	13,348.43 (64.78)	13,307.03 (61.02)
3	8168.72 (41.93)	7157.68 (41.24)	7075.81 (53.59)	7115.60 (54.77)	7103.74 (51.60)
4	4882.67 (35.32)	4025.68 (31.32)	3974.64 (38.37)	4007.42 (47.09)	3983.21 (41.23)
5	2976.64 (29.75)	2382.62 (25.83)	2364.11 (40.66)	2380.82 (32.90)	2374.91 (34.81)
6	1716.71 (20.81)	1347.35 (21.12)	1350.64 (27.38)	1344.07 (27.23)	1343.15 (25.16)
7	996.40 (19.27)	746.87 (17.07)	795.09 (20.54)	790.22 (23.89)	789.61 (20.66)
8	592.47 (16.46)	443.39 (12.92)	471.46 (17.12)	468.89 (17.18)	471.48 (15.55)
9	336.74 (11.8)	265.24 (11.00)	289.05 (14.20)	283.68 (13.23)	287.17 (14.51)
10	190.38 (9.00)	150.22 (8.26)	175.24 (10.22)	171.56 (10.91)	174.88 (11.12)
11	122.17 (6.34)	100.79 (6.02)	115.07 (8.93)	111.79 (8.48)	112.97 (9.16)
12	66.22 (5.41)	54.22 (4.90)	70.94 (7.15)	67.66 (7.57)	69.58 (7.64)
Unassigned	8458.21 (29.31)	9693.70 (31.31)	10,127.56 (29.77)	9950.63 (26.70)	10,040.19 (47.17)

We ran 100 iterations of each algorithm with independent draws of lotteries and focused on the average performance of each algorithm, including DA. Standard errors are given in parentheses.

TABLE 1. Ranks achieved by participants under five different algorithms.

While the source of the resemblance is not immediately clear, the difference these mechanisms exhibit relative to the DA outcome suggests that the resemblance is not driven by the special nature of the underlying preferences.^{36,37}

Recall also that the programs have intrinsic priorities in the data, and the alternative Pareto efficient mechanisms differ in the way that the programs’ priority information is used to generate the assignment. Hence, the resemblance across alternative Pareto efficient mechanisms cannot be explained by the equivalence result of [Abdulkadiroğlu and Sönmez \(1998\)](#) or its extensions by recent authors. These authors focus on an environment in which programs have no intrinsic priorities and find equivalence of efficient mechanisms that treat agents randomly in an ex ante symmetric manner. Importantly, the equivalence in these papers holds only in the ex ante sense (in terms of the lotteries the agents receive), and it does not imply that a similar rank distribution would result from different Pareto efficient mechanisms using different priority systems.

³⁶For example, if all applicants submit the same ROL of programs, the rank distribution (i.e., the fraction of agents getting their first choice, second choice, and so on) would be identical across all assignments, and there would be no difference between DA assignments and efficient assignments. Likewise, if there are no conflicts of interests, again, all agents will be assigned to their top choice under both an efficient mechanism and DA. The difference between the DA and efficient mechanisms suggests that neither scenario holds here.

³⁷The average sum of ranks for TTC is smaller than that under SDmax. This should not necessarily come as a surprise since SDmax minimizes the average sum of ranks over “only” 100 draws of serial orders. Recall that TTC uses the intrinsic priorities from the data, supplemented by the random draws in case of ties, whereas the SD only uses random draws. Hence, it is not unlikely that the (100) random draws used for SD do not include the one replicating TTC. Importantly, the differences between SDmax, SDmin, and TTC are small relative to the DA outcome.

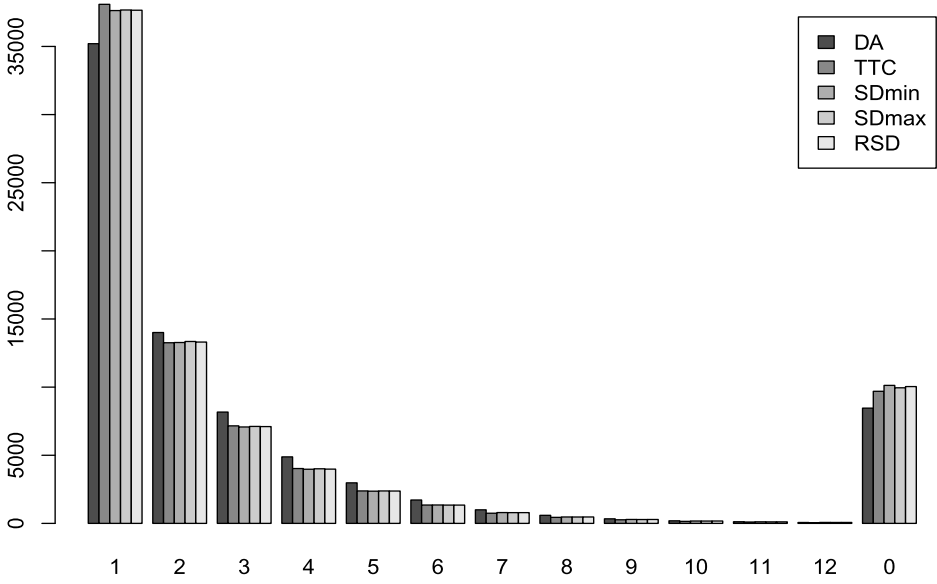


FIGURE 4. Rank distribution under alternative mechanisms (averaged across 100 iterations).

APPENDIX A: PROOF OF THEOREM 1

A.1 Preliminaries

For an n -economy and for each $u \in [0, 1]$, let $O_{\geq u}^n := \{o \in O^n \mid u_o \geq u\}$ and $O_{\leq u}^n := \{o \in O^n \mid u_o \leq u\}$ denote the set of object types that yield the common value no less than u and the set of object types that yield the common value no greater than u , respectively. The numbers of objects with types in $O_{\geq u}^n$ and $O_{\leq u}^n$ are, respectively, denoted by $Q_{\geq u}^n$ and $Q_{\leq u}^n$. For notational simplicity, we suppress the dependence of these sets on n , with the exception of X^n .

Now consider any Pareto efficient mechanism $\mu \in \mathcal{M}^*$. By a well known result (e.g., [Abdulkadiroğlu and Sönmez 1998](#)), any Pareto efficient matching can be equivalently implemented by a serial dictatorship mechanism with a suitably chosen serial order. Let SD^{f_μ} be the serial dictatorship mechanism where, for each state ω , a serial order $f_\mu(\omega) : I \rightarrow I$, a bijective mapping, is chosen so as to implement $\mu_\omega(\cdot)$. That is, for each state $\omega \in \Omega$, the serial order f_μ is chosen so that $SD_\omega^{f_\mu(\omega)}(i) = \mu_\omega(i)$ for each $i \in I$. Since the matching μ arising from the mechanism depends on the random state ω , the serial order f implementing μ is a random variable. In the sequel, we study a Pareto efficient matching mechanism μ via the associated SD^{f_μ} . To avoid clutter, we now suppress the dependence of f on μ .

Given an n -economy, for any Pareto efficient mechanism μ and the associated serial order f , let

$$I_{\geq u}(\mu) := \{i \in I \mid f(i) \leq Q_{\geq u}\}$$

be the set of agents who have a serial order within the total supply of objects whose common values are at least u (in the equivalent serial dictatorship implementation).

For any ϵ , the set

$$\bar{I}_{\geq u}^{\epsilon}(\mu) = \{i \in I_{\geq u}(\mu) \mid U_i(\text{SD}^f(i)) \leq U(u, 1 - \epsilon)\}$$

consists of the agents who realize payoff no greater than $U(u, 1 - \epsilon)$ while having a serial order within $Q_{\geq u}$. The following lemma is crucial for the main result.

LEMMA 2. *For any $\epsilon, \gamma > 0$,*

$$\Pr\left[\exists \mu \in \mathcal{M}^* \text{ and } u \text{ such that } \frac{|\bar{I}_{\geq u}^{\epsilon}(\mu)|}{|I|} \geq \gamma\right] \rightarrow 0$$

as $n \rightarrow \infty$.

PROOF. Fix any $\epsilon > 0$ and $\gamma > 0$. We first build a random bipartite graph on $I \cup O$ where an edge (i, o) is added if and only if $\xi_{i,o} > 1 - \epsilon$.

Now choose any $\delta \in (0, 1)$. For each $\mu \in \mathcal{M}^*$ and u , define random sets $I_{\geq u}^{\delta}(\mu) := \{i \in I \mid f(i) \leq (1 - \delta)Q_{\geq u}\}$, $\bar{I}_{\geq u}^{\epsilon, \delta}(\mu) := \{i \in I_{\geq u}^{\delta}(\mu) \mid U_i(\text{SD}^f(i)) \leq U(u, 1 - \epsilon)\}$, and

$$\bar{O}_{\geq u}^{\delta}(\mu) := \{o \in O_{\geq u} \mid \exists i \in \mu^{-1}(o) \text{ s.t. } f(i) > (1 - \delta)Q_{\geq u}\},$$

which consist of object types in $O_{\geq u}$ assigned to the agents with serial order worse than $(1 - \delta)Q_{\geq u}$.

We argue that the set $\bar{I}_{\geq u}^{\epsilon, \delta}(\mu) \times \bar{O}_{\geq u}^{\delta}(\mu)$ must be an independent set of the random bipartite graph on $I \cup O$. To prove this, suppose otherwise. Then there would exist an edge $(i, o) \in \bar{I}_{\geq u}^{\epsilon, \delta}(\mu) \times \bar{O}_{\geq u}^{\delta}(\mu)$. Then

$$U_i(o) > U(u, 1 - \epsilon) \geq U_i(\text{SD}^f(i)),$$

where the strict inequality holds since $\xi_{i,o} > 1 - \epsilon$ (i.e., (i, o) is an edge), $o \in O_{\geq u}$, and since $U(\cdot, \cdot)$ is monotonic (in particular, strictly increasing in idiosyncratic component). The weak inequality holds because $i \in \bar{I}_{\geq u}^{\epsilon, \delta}$. In addition, we must have

$$f(i) \leq (1 - \delta)Q_{\geq u} < f(i') \quad \text{for some } i' \in \mu^{-1}(o),$$

where the first inequality comes from the fact that $i \in \bar{I}_{\geq u}^{\epsilon, \delta}(\mu) \subset I_{\geq u}^{\delta}(\mu)$, while the second inequality from the fact that $o \in \bar{O}_{\geq u}^{\delta}(\mu)$. Thus, this means that when i becomes the dictator under SD^f , object o is still available. But $U_i(o) > U_i(\text{SD}^f(i))$ means that i chooses an object worse than o , which yields a contradiction.

Since $\bar{I}_{\geq u}^{\epsilon, \delta}(\mu) \times \bar{O}_{\geq u}^{\delta}(\mu)$ is an independent set for each $\mu \in \mathcal{M}^*$ and $u \in [0, 1]$, and since $|I| = n$, applying Lemma 1, we have that, for any $\gamma' > 0$,

$$\Pr[\exists \mu \in \mathcal{M}^* \text{ and } u \in [0, 1] \text{ s.t. } \min\{|\bar{I}_{\geq u}^{\epsilon, \delta}(\mu)|, |\bar{O}_{\geq u}^{\delta}(\mu)|\} \geq \gamma' n] \rightarrow 0 \quad (1)$$

as n goes to infinity.

Fix any $\gamma' > 0$. Recall that $|\bar{I}_{\geq u}^{\epsilon, \delta}(\mu)| \leq |I_{\geq u}^{\delta}(\mu)| \leq (1 - \delta)Q_{\geq u}$ and $|\bar{O}_{\geq u}^{\delta}(\mu)| \geq \delta Q_{\geq u}/\bar{q}$. Hence, if $|\bar{I}_{\geq u}^{\epsilon, \delta}(\mu)| \geq \gamma'n$, then we must have $|\bar{O}_{\geq u}^{\delta}(\mu)| \geq \delta\gamma'n/((1 - \delta)\bar{q})$, where recall $\bar{q}$ is the upper bound for the number of copies for each object type.

Hence, as $n \rightarrow \infty$,

$$\begin{aligned} & \Pr[\exists \mu \in \mathcal{M}^* \text{ and } u \text{ s.t. } |\bar{I}_{\geq u}^{\epsilon, \delta}(\mu)| \geq \gamma'n] \\ &= \Pr\left[\exists \mu \in \mathcal{M}^* \text{ and } u \text{ s.t. } |\bar{I}_{\geq u}^{\epsilon, \delta}(\mu)| \geq \gamma'n \text{ and } |\bar{O}_{\geq u}^{\delta}(\mu)| \geq \frac{\delta\gamma'}{(1 - \delta)\bar{q}}n\right] \\ &\leq \Pr\left[\exists \mu \in \mathcal{M}^* \text{ and } u \text{ s.t. } \min\{|\bar{I}_{\geq u}^{\epsilon, \delta}(\mu)|, |\bar{O}_{\geq u}^{\delta}(\mu)|\} \geq \min\left\{1, \frac{\delta}{(1 - \delta)\bar{q}}\right\}\gamma'n\right] \\ &\rightarrow 0, \end{aligned}$$

where the convergence follows from (1).³⁸

Finally, by construction, $|\bar{I}_{\geq u}^{\epsilon, \delta}(\mu)| \geq |\bar{I}_{\geq u}^{\epsilon}(\mu)| - \delta Q_{\geq u} - 1$. Since $Q_{\geq u} \leq n$, we get that

$$\frac{|\bar{I}_{\geq u}^{\epsilon, \delta}(\mu)|}{|I|} \geq \frac{|\bar{I}_{\geq u}^{\epsilon}(\mu)|}{|I|} - \delta - \frac{1}{|I|}$$

for each $\mu \in \mathcal{M}^*$. Hence, it follows that

$$\begin{aligned} & \Pr\left[\exists \mu \in \mathcal{M}^* \text{ and } u \text{ s.t. } \frac{|\bar{I}_{\geq u}^{\epsilon}(\mu)|}{|I|} \geq \gamma' + \delta\right] \\ &\leq \Pr\left[\exists \mu \in \mathcal{M}^* \text{ and } u \text{ s.t. } \frac{|\bar{I}_{\geq u}^{\epsilon, \delta}(\mu)|}{|I|} \geq \gamma' + \frac{1}{|I|}\right] \rightarrow 0. \end{aligned}$$

Set δ and γ' such that $\delta + \gamma' = \gamma$. Then

$$\Pr\left[\exists \mu \in \mathcal{M}^* \text{ and } u \text{ s.t. } \frac{|\bar{I}_{\geq u}^{\epsilon}(\mu)|}{|I|} \geq \gamma\right] \rightarrow 0.$$

□

We are now ready to prove [Theorem 1](#).

A.2 Proof of [Theorem 1](#)

To prove the statement, we show that the payoff distributions induced by Pareto efficient mechanisms converge to F^* in the sense defined earlier.

Fix any $\epsilon > 0$. We show that, as $n \rightarrow \infty$,

$$\Pr\left[\sup_{\mu \in \mathcal{M}^*} \sup_z \max\{F^{\mu}(z - \epsilon) - F^*(z), F^*(z) - F^{\mu}(z + \epsilon)\} \geq \epsilon\right] \rightarrow 0,$$

³⁸Here we use the assumption that $\bar{q}$ does not increase in n . If $\bar{q}$ increases in n at the rate of $O(n/\log(n))$, then one can check that the lower bound in the above equation is $\omega(\log(n))$. Using [Lemma 1](#), one can show that [Lemma 2](#)—and thus [Theorem 1](#)—holds with this lower bound.

where F^* and F^μ are, respectively, the CDF of the payoff induced by the limit utilitarian upper bound and the CDF of the payoffs induced by mechanism μ in $\mathcal{M}^*$.

Let

$$J^\mu(z) := \{i \in I \mid U_i(\mu(i)) \leq z\}$$

denote the set of agents enjoying payoff of at most z under matching mechanism μ . Obviously, $F^\mu(z) = |J^\mu(z)|/n$. Let $u(z)$ be such that $U(u(z), 1) = z$ for each $z \in \hat{Z} := [U(0, 1), U(1, 1)]$. (This is well defined since $U(\cdot, 1)$ is continuous and strictly increasing.) Note that the function $u : \hat{Z} \rightarrow [0, 1]$ so defined is continuous and increasing. Let us fix ϵ' such that for any common value $u \leq u(z) + \epsilon'$, we have $U(u, 1) \leq z + \epsilon$ for each $z \in \hat{Z} := [U(0, 1), U(1, 1)]$. Note that this is well defined since $U(u(z), 1) = z$ and $U(\cdot, 1)$ is continuous and strictly increasing. Further observe that ϵ' is strictly positive. Clearly, for any $z \in \hat{Z}$, any agent matched with an object having common value no greater than $u(z) + \epsilon'$ must be in $J^\mu(z + \epsilon)$. This means that $|J^\mu(z + \epsilon)| \geq Q_{\leq u(z) + \epsilon'} = X^n(u(z) + \epsilon')n$ for all $\mu \in \mathcal{M}^*$. By definition, for each z , $F^*(z) = X(u(z))$.

Then

$$\begin{aligned} & \Pr \left[\sup_{\mu \in \mathcal{M}^*} \sup_z \left(-\frac{|J^\mu(z + \epsilon)|}{|I|} + F^*(z) \right) \geq \epsilon \right] \\ &= \Pr \left[\sup_{\mu \in \mathcal{M}^*} \sup_{z \in \hat{Z}} \left(-\frac{|J^\mu(z + \epsilon)|}{|I|} + F^*(z) \right) \geq \epsilon \right] \\ &\leq \Pr \left[\sup_{z \in \hat{Z}} (-X^n(u(z) + \epsilon') + X(u(z))) \geq \epsilon \right] \\ &= \Pr \left[\sup_u (-X^n(u + \epsilon') + X(u)) \geq \epsilon \right] \\ &\leq \Pr \left[\sup_u (-X^n(u + \min\{\epsilon, \epsilon'\}) + X(u)) \geq \min\{\epsilon, \epsilon'\} \right] \\ &\rightarrow 0 \end{aligned} \tag{2}$$

as $n \rightarrow \infty$. The first equality comes from the fact that for any $z < U(0, 1)$, $F^*(z) = 0$; the first inequality is by definition of ϵ' and the convergence follows since $L(X^n, X) \rightarrow 0$ as $n \rightarrow \infty$.

For the next part, recall that $I_{\geq u}(\mu) := \{i \in I \mid f(i) \leq Q_{\geq u}\}$ and $\tilde{I}_{\geq u}^\epsilon(\mu) := \{i \in I_{\geq u}(\mu) \mid U_i(\text{SD}^f(i)) \leq U(u, 1 - \epsilon)\}$, where SD^f is the SD rule implementing μ . Let $I_{\leq u}(\mu) := \{i \in I \mid f(i) \leq Q_{\leq u}\}$. We also extend the function u such that $u(z) = 0$ for any $z \in [U(0, 0), U(0, 1)]$.

Fix ϵ'' such that for any common value $u \geq u(z) - \epsilon''$, we have $z - \epsilon \leq U(u, 1 - \epsilon'')$ for all $z \in [U(0, 0), U(1, 1)]$. Note that this is well defined since $U(u(z), 1) \geq z$ and U is continuous and strictly increasing in both components. In addition, $\epsilon'' > 0$. Observe that $U(\mu(i)) \leq z - \epsilon$ implies $i \in I_{\leq u(z) - \epsilon''} \cup \tilde{I}_{\geq u(z) - \epsilon''}^{\epsilon''}(\mu)$. Hence,

$$\frac{|J^\mu(z - \epsilon)|}{|I|} \leq \frac{|I_{\leq u(z) - \epsilon''}|}{|I|} + \frac{|\tilde{I}_{\geq u(z) - \epsilon''}^{\epsilon''}(\mu)|}{|I|}.$$

We obtain

$$\begin{aligned}
& \Pr \left[\sup_{\mu \in \mathcal{M}^*} \sup_z \left(\frac{|J^\mu(z - \epsilon)|}{|I|} - F^*(z) \right) \geq \epsilon \right] \\
& \leq \Pr \left[\sup_{\mu \in \mathcal{M}^*} \sup_z \left(\frac{|I_{\leq u(z) - \epsilon''}|}{|I|} + \frac{|\tilde{I}_{\geq u(z) - \epsilon''}^{\epsilon''}(\mu)|}{|I|} - F^*(z) \right) \geq \epsilon \right] \\
& \leq \Pr \left[\sup_{\mu \in \mathcal{M}^*} \sup_z (X^n(u(z) - \epsilon'') - X(u(z))) \geq \frac{\epsilon}{2} \right] + \Pr \left[\sup_{\mu \in \mathcal{M}^*} \sup_z \frac{|\tilde{I}_{\geq u(z) - \epsilon''}^{\epsilon''}(\mu)|}{|I|} \geq \frac{\epsilon}{2} \right] \\
& = \Pr \left[\sup_u (X^n(u - \epsilon'') - X(u)) \geq \frac{\epsilon}{2} \right] + \Pr \left[\sup_{\mu \in \mathcal{M}^*} \sup_z \frac{|\tilde{I}_{\geq u(z) - \epsilon''}^{\epsilon''}(\mu)|}{|I|} \geq \frac{\epsilon}{2} \right] \quad (3) \\
& \leq \Pr \left[\sup_u \left(X^n \left(u - \min \left\{ \epsilon'', \frac{\epsilon}{2} \right\} \right) - X(u) \right) \geq \min \left\{ \epsilon'', \frac{\epsilon}{2} \right\} \right] \\
& \quad + \Pr \left[\sup_{\mu \in \mathcal{M}^*} \sup_u \frac{|\tilde{I}_{\geq u}^{\epsilon''}(\mu)|}{|I|} \geq \frac{\epsilon}{2} \right] \\
& \rightarrow 0,
\end{aligned}$$

where the convergence holds by $L(X^n, X) \rightarrow 0$ and [Lemma 2](#).

Combining (2) and (3) and the fact that $F^\mu(z) = |J^\mu(z)|/n$, we conclude that

$$\Pr \left[\sup_{\mu \in \mathcal{M}^*} \sup_z \max \{ F^\mu(z - \epsilon) - F^*(z), -F^\mu(z + \epsilon) + F^*(z) \} \geq \epsilon \right] \rightarrow 0$$

as $n \rightarrow \infty$.

APPENDIX B: FURTHER ANALYSIS OF INDEPENDENT SETS IN RANDOM GRAPHS

We propose a class of random bipartite graphs where the only difference from [Dawande et al. \(2001\)](#) is that now the probability of linking two nodes $(v_1, v_2) \in V_1 \times V_2$ is given by $p(n)$, where $p(n)$ can go to 0 as n goes to infinity. The next result shows that as long as $p(n)$ goes to 0 at a lower rate than $\log(n)/n$, the probability that there is an independent set where the size of both sides is linear in n goes to 0 at a rate “arbitrarily close” to $(1/n!)^2$.

LEMMA 3. *Fix any function $f : \mathbb{N} \rightarrow \mathbb{R}$ such that $\lim_{n \rightarrow \infty} f(n)/n = 0$. Consider a random bipartite graph $B = (V_1 \cup V_2, p(n))$, where $0 < p(n) < 1$ satisfies $\lim_{n \rightarrow \infty} p(n)/(\log(n)/n) = \infty$ and for each $i \in \{1, 2\}$, $|V_1| = n$ and $|V_2| = m = O(n)$. For any n large enough,*

$$\Pr[\exists \text{ an independent set } \bar{V}_1 \times \bar{V}_2 \text{ with } \min\{|\bar{V}_1|, |\bar{V}_2|\} \geq \gamma n] \leq \left(\frac{1}{f(n)!} \right)^2.$$

PROOF. Fix a random bipartite graph $B = (V_1 \cup V_2, p(n))$, where $0 < p(n) < 1$ satisfies $\lim_{n \rightarrow \infty} p(n)/(\log(n)/n) = \infty$ and for each $i \in \{1, 2\}$, $|V_1| = n$ and $|V_2| = m = O(n)$. Observe that if $\bar{V}_1 \times \bar{V}_2$ is an independent set, then so is any subset of $\bar{V}_1 \times \bar{V}_2$. This has two im-

plications. The first implication is that whenever there is an independent set $\bar{V}_1 \times \bar{V}_2$ with $\min\{|\bar{V}_1|, |\bar{V}_2|\} = a(n)$, then there is a *balanced* independent set with size $a(n)$.³⁹ Let $Z_{a(n)}$ be the number of balanced independent sets $\bar{V}_1 \times \bar{V}_2$ with $|\bar{V}_1| = |\bar{V}_2| = a(n)$. Thus, we want to show that the probability of having a balanced independent set $\bar{V}_1 \times \bar{V}_2$ with $|\bar{V}_1| = |\bar{V}_2| = a(n) \geq \gamma n$ is smaller than $(1/(f(n)!))^2$ for any n large enough. In the sequel, we sometimes abuse notation and write p and a for $p(n)$ and $a(n)$, respectively. The second implication is that whenever there is a balanced independent set of size a , there must be a balanced independent set of smaller size. Otherwise stated, $\Pr\{\exists a \geq \gamma n \text{ s.t. } Z_a \geq 1\} = \Pr\{Z_{\gamma n} \geq 1\}$. Thus, we can assume without loss of generality that $a = \gamma n$ and just show that $\Pr\{Z_a \geq 1\} \leq (1/(f(n)!))^2$ for any n large enough. We have

$$\Pr\{Z_a \geq 1\} \leq \mathbb{E}(Z_a) = \binom{n}{a} \binom{m}{a} ((1-p)^a)^a \leq \left(\frac{n^a}{a!}\right) \left(\frac{m^a}{a!}\right) ((1-p)^a)^a.$$

The computation of $\mathbb{E}(Z_a)$ follows from the following argument. A set of vertices $A \cup B$, where $A \subseteq V_1$ and $B \subseteq V_2$, forms an independent set if there is no edge between every pair of vertices $v_1 \in A$ and $v_2 \in B$. Suppose that $|A| = |B| = a$. Since the probability of an edge is p , the probability that a given $A \cup B$ forms an independent set is $((1-p)^a)^a$. There are $\binom{n}{a}$ different ways to choose a subset $A \subseteq V_1$ of size a and $\binom{m}{a}$ ways to choose a subset $B \subseteq V_2$ of size a . Hence, the number of pairs of subsets (A, B) we are considering is $\binom{n}{a} \binom{m}{a}$.

We now claim that for any n large enough, $\gamma n \geq (\log(n) + \log(m))/\log(1/(1-p(n)))$. This is clear if $p(n)$ does not go to 0. So assume that $p(n) \rightarrow 0$. Because $m = O(n)$, we have that for some constant c and for any n large enough, $m \leq cn$. Hence, it is enough to show that for n large enough, $\gamma n \geq (\log(n) + \log(cn))/\log(1/(1-p(n)))$. The last inequality can simply be written as $\log(1-p(n)) \leq -(\log(n) + \log(cn))/\gamma n$. Now fix any $\rho > 0$. Using the fact that $\log(1+x)/x$ converges to 1 as $x \rightarrow 0$, we obtain that, for n large enough, $\log(1-p(n)) \leq (1-\rho)(-p(n))$. Thus, for n large enough, $\log(1-p(n)) \leq -(\log(n) + \log(cn))/\gamma n$ is implied by $(1-\rho)(-p(n)) \leq -(\log(n) + \log(cn))/\gamma n$. This condition can be written as $\gamma(1-\rho) \geq (2\log(n)/n)/p(n) + (\log(c)/n)/p(n)$. Because $\lim p(n)/(\log(n)/n) = \infty$, the right-hand side goes to 0 as n grows and so the previous inequality is satisfied for n large enough.

Hence, for n large enough,

$$(1-p)^a \leq (1-p)^{(\log(n)+\log(m))/\log(1/(1-p))} = (1-p)^{\log_{1-p}(n^{-1})} (1-p)^{\log_{1-p}(m^{-1})} = n^{-1} m^{-1}.$$

Thus, for n large enough, we get that $(n^a/a!)(m^a/a!)((1-p)^a)^a \leq (1/a!)^2$. Because $\lim f(n)/n = 0$, for any n large enough, $f(n) \leq \gamma n$ and so for n large enough, $a! = (\gamma n)! \geq f(n)!$. Thus, we obtain that for n large enough,

$$\Pr\{Z_a \geq 1\} \leq \left(\frac{1}{a!}\right)^2 \leq \left(\frac{1}{f(n)!}\right)^2. \quad \square$$

Assuming that the idiosyncratic shocks are drawn from a uniform distribution on $[0, \bar{\xi}(n)]$, in the random graph built in the proof of [Theorem 1](#), we have that $p(n) =$

³⁹An independent set $\bar{V}_1 \times \bar{V}_2$ is balanced if $|\bar{V}_1| = |\bar{V}_2|$.

$\varepsilon/\bar{\xi}(n)$. Hence, since the above result requires $\lim p(n)/(\log(n)/n) = \infty$, [Theorem 1](#) must hold as long as $\bar{\xi}(n)/(n/\log(n)) \rightarrow 0$, i.e., $\bar{\xi}(n) = o(n/\log(n))$, as claimed in [Remark 5](#).

APPENDIX C: PROOF OF [PROPOSITION 1](#)

In the sequel, assume $X^n(\cdot)$ is degenerate with a single common value u^0 and that $X^n(\cdot) = Y^n(\cdot)$, i.e., $\bar{q}_o = 1$ for all $o \in O^n$. For the proof, we focus on *relative rank*, i.e., $\text{rank}_\mu(i)/n$. Accordingly, let $g(n) := h(n)/n$. Clearly, $g(n) = \omega(\log(n)/n)$ and $\lim_{n \rightarrow \infty} g(n) = 0$. Then we can rewrite $\bar{I}_\mu = \{i \in I \mid \text{rank}_\mu(i)/n \leq g(n)\}$. We want to show that

$$\inf_{\mu \in \mathcal{M}_n^*} \frac{|\bar{I}_\mu|}{n} \xrightarrow{p} 1,$$

where $\mathcal{M}_n^*$ is the set of Pareto efficient matchings in the n -economy.

Let function $\epsilon : \mathbb{N} \rightarrow \mathbb{R}_{++}$ be defined by $\epsilon(n) := \frac{1}{4}g(n)$ for all $n \in \mathbb{N}$. Note that $\epsilon(n)$ must go to 0 and $\lim_{n \rightarrow \infty} \epsilon(n)/(\log(n)/n) = \infty$. Let $\bar{I}_\mu^\epsilon := \{i \in I \mid \xi_{i\mu(i)} \geq 1 - \epsilon(n)\}$. Using the very same random graph as in our sketch of the proof in [Section 4](#), we can use [Lemma 3](#) (since $\lim_{n \rightarrow \infty} \epsilon(n)/(\log(n)/n) = \infty$) to obtain that

$$\inf_{\mu \in \mathcal{M}_n^*} \frac{|\bar{I}_\mu^\epsilon|}{n} \xrightarrow{p} 1.$$

The following claim is then crucial for our result. It states that with probability going to 1 as $n \rightarrow \infty$, an individual i assigned to o with $\xi_{io} \geq 1 - \epsilon(n)$ gets a relative rank smaller than $g(n)$.

CLAIM 1. *We have $\Pr\{\forall \mu \in \mathcal{M}_n^* : \bar{I}_\mu^\epsilon \subset \bar{I}_\mu\} \rightarrow 1$ as $n \rightarrow \infty$.*

PROOF. For each agent i , let X_i be the number of objects o such that $\xi_{io} \geq 1 - \epsilon(n)$. Hence, if $\xi_{i\mu(i)} \geq 1 - \epsilon(n)$ for some μ , then we must have $\text{rank}_\mu(i) \leq X_i$. So if $X_i \leq ng(n)$, then, for any μ , whenever $i \in \bar{I}_\mu^\epsilon$ ($\Leftrightarrow \xi_{i\mu(o)} \geq 1 - \epsilon(n)$), we must have $\text{rank}_\mu(i)/n \leq g(n)$, so $i \in \bar{I}_\mu$. To prove the claim, therefore, it is enough to show that, with probability approaching 1 as $n \rightarrow \infty$, $X_i \leq ng(n)$ for all $i \in I$. Note that X_i follows a binomial distribution $B(n, \epsilon(n))$ (recall that ξ_{io} follows a uniform distribution with support $[0, 1]$). In the sequel, we let $Y_i := n - X_i$. We have

$$\begin{aligned} \Pr\{\exists i \text{ with } X_i > ng(n)\} &\leq \sum_{i \in I} \Pr\{X_i > ng(n)\} \\ &= |I| \Pr\{X_i > ng(n)\} \\ &= |I| \Pr\{Y_i < n(1 - g(n))\}, \end{aligned}$$

where the inequality is by the union bound and the second equality is by definition of Y_i . We further note that $Y_i = n - X_i$ follows a binomial distribution $B(n, 1 - \epsilon(n))$. Using

the inequality provided by [Arratia and Gordon \(1989\)](#) to bound $\Pr\{Y_i < n(1 - g(n))\}$, we obtain⁴⁰

$$\begin{aligned}\Pr\{\exists i \text{ with } X_i > ng(n)\} &\leq |I| \Pr\{Y_i < n(1 - g(n))\} \\ &\leq n \exp[-nD(1 - g(n) \parallel 1 - \epsilon(n))] \\ &= n \exp\left[-n\left((1 - g(n)) \log\left(\frac{1 - g(n)}{1 - \epsilon(n)}\right) + g(n) \log\left(\frac{g(n)}{\epsilon(n)}\right)\right)\right],\end{aligned}$$

where $D(a \parallel p)$ stands for the relative entropy between Bernoulli(a) and Bernoulli(p).

Now choose $\rho > 0$ small enough so that $\log(4) - (\frac{3}{4})[1 + \rho] > \frac{1}{2}$. Using the fact that $\lim_{x \rightarrow 0} \log(1 + x)/x = 1$, we must have that, for n large enough, $\log(1 - g(n))/(1 - \epsilon(n)) = \log(1 + (\epsilon(n) - g(n))/(1 - \epsilon(n))) \geq [1 + \rho](\epsilon(n) - g(n))/(1 - \epsilon(n))$ (recall that $\epsilon(n) - g(n)/(1 - \epsilon(n)) < 0$). Thus, for n large enough, we can simplify the above expression as

$$\begin{aligned}&\frac{n}{\exp\left[n(1 - g(n)) \log \frac{1 - g(n)}{1 - \epsilon(n)} + ng(n) \log \frac{g(n)}{\epsilon(n)}\right]} \\ &\leq \frac{n}{\exp\left[n(1 - g(n)) \frac{\epsilon(n) - g(n)}{1 - \epsilon(n)} [1 + \rho] + ng(n) \log \frac{g(n)}{\epsilon(n)}\right]} \\ &= \frac{n}{\exp\left[\frac{g(n)}{\log(n)/n} \log(n) \log \frac{g(n)}{\epsilon(n)} - \frac{(g(n) - \epsilon(n))}{g(n)} \left(\frac{1 - g(n)}{1 - \epsilon(n)}\right) \frac{g(n)}{\log(n)/n} \log(n) [1 + \rho]\right]} \\ &= \frac{n}{\exp\left[\left(\log \frac{g(n)}{\epsilon(n)} - \frac{(g(n) - \epsilon(n))}{g(n)} \left(\frac{1 - g(n)}{1 - \epsilon(n)}\right) [1 + \rho]\right) \frac{g(n)}{\log(n)/n} \log(n)\right]} \\ &= \frac{n}{\exp\left[\left(\log(4) - \frac{3}{4} \left(\frac{1 - g(n)}{1 - \epsilon(n)}\right) [1 + \rho]\right) \frac{g(n)}{\log(n)/n} \log(n)\right]} \\ &\leq \frac{n}{\exp\left[\frac{1}{2} \frac{g(n)}{\log(n)/n} \log(n)\right]} \\ &= \frac{n}{\exp\left[\log\left(n^{\frac{1}{2} \frac{g(n)}{\log(n)/n}}\right)\right]} \\ &= \frac{n}{n^{\frac{1}{2} \frac{g(n)}{\log(n)/n}}} = \frac{1}{n^{\frac{1}{2} \frac{g(n)}{\log(n)/n} - 1}},\end{aligned}$$

⁴⁰The bound provided by [Arratia and Gordon \(1989\)](#) is closely related to Sanov's theorem; see Chapter 11.4 of [Cover and Thomas \(1991\)](#). We note that a standard approach to provide a bound on the binomial cumulative distribution consists in using Hoeffding's inequality. However, this bound is not fine enough for our purpose and we need to appeal to bounds provided in the large deviation theory.

where the third equality uses the assumption that $\epsilon(n) = \frac{1}{4}g(n)$, while the last inequality follows since $\log(4) - (\frac{3}{4})[1 + \rho]((1 - g(n))/(1 - \epsilon(n))) \geq \log(4) - (\frac{3}{4})[1 + \rho] > \frac{1}{2}$ (recall that $g(n) \geq \epsilon(n)$). Thus, since $g(n)/(\log(n)/n) \rightarrow \infty$, the above term must converge to 0. $\square$

Now we are ready to complete the proof of [Proposition 1](#). Fix any $\gamma > 0$. We have that

$$\Pr\left\{\forall \mu \in \mathcal{M}_n^* : \frac{|\bar{I}_\mu|}{n} > \gamma\right\} \geq \Pr\left\{\forall \mu \in \mathcal{M}_n^* : \frac{|\bar{I}_\mu^\epsilon|}{n} > \gamma \text{ and } \bar{I}_\mu^\epsilon \subset \bar{I}_\mu\right\} \rightarrow 1,$$

where the convergence result follows from [Claim 1](#) and the fact that $\inf_{\mu \in \mathcal{M}_n^*} |\bar{I}_\mu^\epsilon|/n \xrightarrow{P} 1$. We therefore conclude that

$$\inf_{\mu \in \mathcal{M}_n^*} \frac{|\bar{I}_\mu|}{n} \xrightarrow{P} 1,$$

as was to be shown.

REFERENCES

- Abdulkadiroğlu, Atila, Nikhil Agarwal, and Parag A. Pathak (2015a), “The welfare effects of coordinated assignment: Evidence from the NYC HS match.” NBER Working Paper 21046. [259]
- Abdulkadiroğlu, Atila, Yeon-Koo Che, and Yosuke Yasuda (2015b), “Expanding ‘choice’ in school choice.” *American Economic Journal: Microeconomics*, 7, 1–42. [242, 257]
- Abdulkadiroğlu, Atila, Parag A. Pathak, and Alvin E. Roth (2009), “Strategy-proofness versus efficiency in matching with indifference: Redesigning the NYC high school match.” *American Economic Review*, 99, 1954–1978. [259]
- Abdulkadiroğlu, Atila and Tayfun Sönmez (1998), “Random serial dictatorship and the core from random endowments in house allocation problems.” *Econometrica*, 66, 689–701. [241, 247, 254, 257, 260, 261]
- Abdulkadiroğlu, Atila and Tayfun Sönmez (2003), “School choice: A mechanism design approach.” *American Economic Review*, 93, 729–747. [240, 255]
- Arratia, Richard and Louis Gordon (1989), “Tutorial on large deviations for the binomial distribution.” *Bulletin of Mathematical Biology*, 51, 125–131. [268]
- Ashlagi, Itai, Yashodhan Kanoria, and Jacob D. Leshno (2017), “Unbalanced random matching markets: The stark effect of competition.” *Journal of Political Economy*, 125, 69–98. [242, 244, 245, 254, 257]
- Azevedo, Eduardo M. and Jacob D. Leshno (2016), “A supply and demand framework for two-sided matching markets.” *Journal of Political Economy*, 124, 1235–1268. [242, 257]
- Bade, Sophie (2016), “Random serial dictatorship: The one and only.” Unpublished paper, Royal Holloway College. [241, 254]

- Bogomolnaia, Anna and Hervé Moulin (2001), “A new solution to the random assignment problem.” *Journal of Economic Theory*, 100, 295–328. [247]
- Carroll, Gabriel (2014), “A general equivalence theorem for allocation of indivisible objects.” *Journal of Mathematical Economics*, 51, 163–177. [241, 254]
- Che, Yeon-Koo, Jinwoo Kim, and Fuhito Kojima (2015), “Stable matching in large economies.” Unpublished paper, Columbia University. [242]
- Che, Yeon-Koo and Fuhito Kojima (2010), “Asymptotic equivalence of probabilistic serial and random priority mechanisms.” *Econometrica*, 78, 1625–1672. [242, 254, 257]
- Che, Yeon-Koo and Olivier Tercieux (2015a), “Efficiency and stability in large matching markets.” Unpublished paper, Columbia University and Paris-Jourdan Sciences Economiques. [242, 247, 254, 257, 259]
- Che, Yeon-Koo and Olivier Tercieux (2015b), “Payoff equivalence of efficient mechanisms in large markets.” Unpublished paper, Columbia University and Paris-Jourdan Sciences Economiques. [257]
- Cover, Thomas M. and Joy A. Thomas (1991), *Elements of Information Theory*. Wiley Series in Telecommunications. Wiley, New York. [268]
- Dawande, Milind, Pinar Keskinocak, Jayashankar M. Swaminathan, and Sridhar Tayur (2001), “On bipartite and multipartite clique problems.” *Journal of Algorithms*, 41, 388–403. [242, 248, 265]
- Featherstone, Clayton (2015), “Rank efficiency: Investigating a widespread ordinal welfare criterion.” Unpublished paper, The Wharton School, University of Pennsylvania. [252]
- Frieze, Alan and Boris Pittel (1995), “Probabilistic analysis of an algorithm in the theory of markets in indivisible goods.” *The Annals of Applied Probability*, 5, 768–808. [251]
- Gale, David and S. Shapley Lloyd (1962), “College admissions and the stability of marriage.” *American Mathematical Monthly*, 69, 9–15. [242]
- Haeringer, Guillaume and Flip Klijn (2009), “Constrained school choice.” *Journal of Economic Theory*, 144, 1921–1947. [259]
- Hylland, Aanund and Richard Zeckhauser (1979), “The efficient allocation of individuals to positions.” *Journal of Political Economy*, 87, 293–314. [240]
- Immorlica, Nicole and Mohammad Mahdian (2005), “Marriage, honesty, and stability.” In *SODA '05: Proceedings of the Sixteenth Annual ACM–SIAM Symposium on Discrete Algorithms*, 53–62, Society for Industrial and Applied Mathematics, Philadelphia, Pennsylvania. [242, 257]
- Knuth, Donald E. (1996), “An exact analysis of stable allocation.” *Journal of Algorithms*, 20, 431–444. [251]
- Knuth, Donald E. (1997), *Stable Marriage and Its Relation to Other Combinatorial Problems: An Introduction to the Mathematical Analysis of Algorithms*, volume 10 of CRM

Proceedings & Lecture Notes. American Mathematical Society, Providence, Rhode Island. [242, 244, 255, 257]

Kojima, Fuhito and Mihai Manea (2010), “Incentives in the probabilistic serial mechanism.” *Journal of Economic Theory*, 145, 106–123. [242]

Kojima, Fuhito and Parag A. Pathak (2009), “Incentives and stability in large two-sided matching markets.” *American Economic Review*, 99, 608–627. [242, 257]

Lee, SangMok (2017), “Incentive compatibility of large centralized matching markets.” *Review of Economic Studies*, 84, 444–463. [242, 244, 245]

Lee, SangMok and Leeat Yariv (2017), “On the efficiency of stable matchings in large markets.” Unpublished paper, University of Pennsylvania. [242, 244, 245, 255, 256, 257]

Lee, Thiam and Jay Sethuraman (2011), “Equivalence results in the allocation of indivisible objects: A unified view.” Unpublished paper, Columbia University. [241, 254]

Liu, Qingmin and Marek Pycia (2016), “Ordinal efficiency, fairness, and incentives in large markets.” Unpublished paper, Columbia University. [247, 254, 257]

Pathak, Parag A. and Jay Sethuraman (2011), “Lotteries in student assignment: An equivalence result.” *Theoretical Economics*, 6, 1–17. [241, 254]

Pittel, Boris (1989), “The average number of stable matchings.” *SIAM Journal on Discrete Mathematics*, 2, 530–549. [242, 244, 255, 257]

Shapley, Lloyd S. and Herbert E. Scarf (1974), “On cores and indivisibility.” *Journal of Mathematical Economics*, 1, 23–37. [240]

Co-editor George J. Mailath handled this manuscript.

Manuscript received 28 February, 2017; final version accepted 18 May, 2017; available online 23 May, 2017.