

Stability in repeated matching markets

CE LIU

Department of Economics, Michigan State University

This paper develops a framework for studying repeated matching markets. The model departs from the Gale–Shapley matching model by having a fixed set of long-lived players (firms) match with a new generation of short-lived players (workers) in every period. I define history-dependent and self-enforcing matching processes in this repeated matching environment and characterize the firms' payoffs. Firms fall into one of two categories: some firms must obtain the same payoff as they would in static stable matchings, and this holds at every patience level; meanwhile, repetition and history dependence can enlarge the set of sustainable payoffs for the other firms, provided that the firms are sufficiently patient. In large matching markets with correlated preferences, the first kind of firms corresponds to “elite” firms that make up at most a vanishingly small fraction of the market. The vast majority of firms fall into the second category.

KEYWORDS. Gale–Shapley, matching, repeated games, stability.

JEL CLASSIFICATION. C71, C72, C73, C78, D47.

1. INTRODUCTION

College admission, hospital—resident matching, and entry-level hiring are ongoing matching processes that take place every year. One side of these markets—the firms—is long-lived players, whereas the other side—namely the students, residents, or workers—participates in the matching process on only a few occasions, sometimes only once. However, much of the theoretical analysis of matching environments treats both sides of the market as short-lived players, ignoring the possibility for dynamic incentives that could be used as a carrot and stick to motivate the long-lived players.

To understand the scope of these dynamic incentives, I consider a two-sided one-to-many matching model with long-lived *firms* and short-lived *workers*. In each period, a new generation of workers enters the market and lives for one period. The stage game is the canonical one-to-many matching market à la Gale and Shapley (1962) among the firms and workers currently in the market. I define a stability notion—self-enforcing matching process—that generalizes static stability to this repeated environment. Specifically, a matching process is a complete contingent plan that specifies a current stage-game matching as a function of past histories. A blocking coalition in the stage game can

Ce Liu: celiu@msu.edu

I am grateful to Nageeb Ali for his continuing guidance and encouragement. The paper has also benefited from the thoughtful and constructive feedback of three referees. I also thank Christopher Chambers, Henrique de Oliveira, Jon Eguia, Simone Galperti, Ed Green, Fuhito Kojima, Maciej Kotowski, Xiao Lin, Elliot Lipnowski, Zizhen Ma, Arijit Mukherjee, Esteban Peralta, Ran Shorrer, Joel Sobel, Qinggong Wu, Hanzhe Zhang, and audience members at various seminars and workshops for helpful comments.

comprise a firm and a set of workers who also find this deviation profitable, but unlike in the theory of static matching, firms care about continuation play and are unwilling to deviate if doing so leads to unattractive future outcomes. A matching process is said to be self-enforcing if it is immune to not only unilateral deviations by firms or workers, but also sequential blocking coalitions that can be chained together by a firm over a possibly infinite horizon.

The goal of this paper is to investigate what can be sustained through self-enforcing matching processes. I find that firms fall into either one of two categories. The first kind of firms must obtain the same payoff as they would in static stable matchings, and this holds regardless of the firms' patience level. By contrast, the other kind of firms can obtain payoffs that are distinct from what can be sustained in static stable matchings; in fact, when patience is sufficiently high, the only restriction is for these firms to obtain payoffs that are higher than their respective minmax values.¹ Finally, I show that in large matching markets with random and correlated preferences, the first kind of firms make up at most a vanishingly small fraction of the market.

Let us illustrate the effect of history dependence through a stylized example based on the matching market between hospitals and medical students. Suppose that there are three firms (hospitals): f_1 and f_2 are urban hospitals while f_r is rural. Firms are long-lived players, each with two hiring slots to fill every year. On the other side of the market, five representative workers (students) $w_1, \dots, w_5$ enter the market looking for residency jobs each year. Workers are short-lived players in this market.

The left panel of Table 1 shows the firms' stage-game utilities. For this example, we shall assume that firms have additively separable utilities from matched workers and derive 0 from unfilled positions. Observe that w_5 is every firm's least preferred worker yielding a payoff of 1; on the contrary, the maximum payoff a firm can obtain from any matching is 9. Workers' preferences over firms are in the right panel of Table 1. Each worker prefers to work for any company over unemployment. Observe that the firm f_r (which represents the rural hospital) is the worst firm for all workers.

As illustrated in Figure 1, there are two stable matchings in the stage game: $m_{\mathcal{W}}$ is the worker-optimal stable matching while $m_{\mathcal{F}}$ is optimal for firms. Both $m_{\mathcal{F}}$ and $m_{\mathcal{W}}$ match

TABLE 1. Example: preferences.

$u_f(w)$						$\succ$			
	w_1	w_2	w_3	w_4	w_5	w_1	f_2	f_1	f_r
f_1	5	4	3	2	1	w_2	f_2	f_1	f_r
f_2	2	4	5	3	1	w_3	f_1	f_2	f_r
f_r	2	5	3	4	1	w_4	f_1	f_2	f_r
						w_5	f_1	f_2	f_r

¹In particular, firms may be able to obtain payoff profiles that are on their efficient frontier. Note that from the firms' perspective, the firm-proposing stable matching is ordinally efficient among all *stable* matchings; however, it may not be on the firms' efficient frontier because all the firms may be better off from an unstable matching.

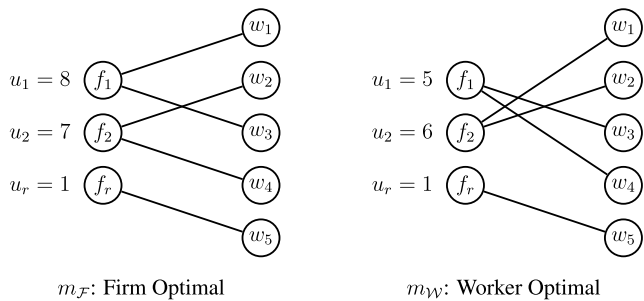


FIGURE 1. Stable matchings.

f_r with the worst worker w_5 while leaving its other position unfilled. The matching m_0 in Figure 2 matches f_r with w_2 , but m_0 is unstable: f_1 and w_2 will form a blocking pair.

Now suppose how players match in the future can be based on how they matched in the past. Consider the following “triggering” matching process μ^0 : firms and workers match according to m_0 on the path of play; if any firm has deviated in the past, players will instead match according to the worker-optimal stage-game matching m_W . Note that this is the familiar idea of Nash reversion, but the stage game is a cooperative game.

The stage-game matching m_0 is played in every period in the matching process μ^0 . A one-shot deviation principle, established later in Lemma 1, shows that μ^0 is also self-enforcing when firm patience is high. Lemma 1 shows that a matching process is self-enforcing if and only if two requirements are satisfied at every history of the market:

- No worker wishes to unilaterally leave her matched firm.
- No firm finds it profitable to conduct a one-shot deviation with a group of workers who also find this deviation profitable.

These requirements are met at every off-path history of μ_0 : m_W is a static stable matching, so by construction, there are no profitable deviations. At every on-path history (where m_0 is played), no workers want to unilaterally leave their firm since everyone prefers employment. By following μ^0 , f_1 receives 6 in every period. If f_1 were to deviate with any worker, it would receive no more than 9 in the current period and 5 in all future periods. As long as f_1 ’s discount factor δ satisfies $\delta > 3/4$, it would not find such

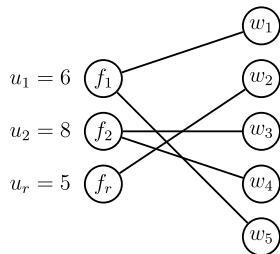


FIGURE 2. Matching m_0 : An unstable matching.

TABLE 2. Workers share identical preference.

$u_f(w)$						$\succ$			
	w_1	w_2	w_3	w_4	w_5	w_1	f_1	f_2	f_r
f_1	5	4	3	2	1	w_2	f_1	f_2	f_r
f_2	2	4	5	3	1	w_3	f_1	f_2	f_r
f_r	2	5	3	4	1	w_4	f_1	f_2	f_r
						w_5	f_1	f_2	f_r

one-shot deviations profitable. A similar argument rules out any profitable one-shot deviations involving f_2 . In light of Lemma 1, μ^0 is a self-enforcing matching process for $\delta > 3/4$.

We have argued so far that it is possible to use history dependence to expand the set of stable outcomes. The next example shows that certain preference configurations can severely limit this possibility.

Consider the market in Table 2: the only difference from Table 1 is that now all workers share a common preference ranking $f_1 \succ f_2 \succ f_r$ over firms. The stage game has a unique stable matching m^* , as depicted in Figure 3. We argue below that no self-enforcing matching process can sustain any matching other than m^* no matter how patient the firms are.

To see why, first observe that as the workers’ favorite firm, f_1 finds its favorite workers $\{w_1, w_2\}$ available in every future generation: whenever f_1 is not matched to $\{w_1, w_2\}$, it can always poach them. Since $\{w_1, w_2\}$ also happens to give f_1 the highest possible stage-game payoff, it is impossible to punish or reward f_1 through continuation value. Essentially, f_1 ’s “minmax” payoff is the same as its maximum payoff, so it is impossible to motivate f_1 through dynamic enforcement. As a result, f_1 acts like a short-lived player and will always match with $\{w_1, w_2\}$ after every history.

Since $\{f_1, w_1, w_2\}$ are always matched together, they are essentially inactive. This makes f_2 the workers’ favorite among active firms. Meanwhile, f_2 finds its favorite active workers, $\{w_3, w_4\}$, always available in every future generation. The only way to credibly remove w_3 or w_4 from f_2 is to match them with f_1 ; otherwise f_2 can simply poach them back. But this is impossible since f_1 is always occupied by $\{w_1, w_2\}$ at every history. Without changes in continuation value, f_2 also behaves myopically and matches with

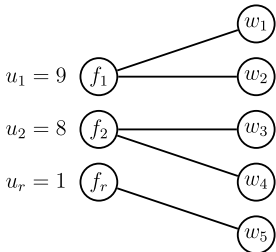


FIGURE 3. Matching m^* : The unique stable matching.

$\{w_3, w_4\}$ at every history. A similar “peeling” argument along workers’ shared preference list ensures that f_r is always matched with w_5 in every self-enforcing matching process.

None of the arguments so far has involved the firms’ patience, so for all $0 < \delta < 1$, the only self-enforcing matching process is the one where m^* is played after every history, and the market functions like a one-shot interaction. It is also worth noting that the uniqueness of the static stable matching is not responsible for the collapse of dynamic enforcement. In fact, if the firms share a common ranking over workers, the market will have a unique static stable matching. But in this case, it is possible to sustain other stage-game matchings in a self-enforcing matching process: I provide an example that illustrates this possibility in Appendix A.6.

In the examples above, repeated interaction had starkly different implications for what can be sustained in a matching market. The key difference between these two settings is the size of the top coalition sequence. A firm and a group of representative workers form a *top coalition* in the stage game if they are mutual favorites. The top coalition sequence is identified by iteratively finding and removing new top coalitions in the stage game until no more top coalitions can be found. In the first example, the top coalition sequence is empty; by contrast, all players are in the top coalition sequence in the second example. The results in this paper generalize these observations.

Section 2 first introduces the repeated matching market and the top coalition sequence. Theorem 1 then shows that regardless of firms’ patience, players in the top coalition sequence always match in the same way as they do in static stable matchings. Theorem 2 complements Theorem 1 and proves a folk theorem for players outside of the top coalition sequence, so with sufficient patience, they may obtain matches that are unattainable in static stable matchings. Theorem 1 has a simple intuition. If a firm is in a top coalition, its minmax payoff is equal to its maximum stage-game payoff, so it cannot be motivated through dynamic enforcement. In standard repeated games, this would only arise under very strong assumptions on payoffs, but in two-sided matching markets, this occurs naturally with top coalitions. As a result, at every history, these firms must always match with their top coalition workers, so a top coalition can be treated as “inactive” players and removed from the stage game. Applying this argument iteratively yields Theorem 1. Theorem 2 is a folk theorem for the remaining “active” players in the reduced game and follows from the standard arguments in Fudenberg, Kreps, and Maskin (1990).

Section 3 builds on the repeated matching model but allows the worker population in each period to be drawn randomly. Preferences are correlated: both the firms and the workers can be divided into quality classes; players always prefer matches from a higher quality class, but within the same class, preference is heterogeneous and random. In this setting, I will call the best class of firms *elite firms* if the size of this quality class is vanishingly small relative to the size of the best class of workers. Theorem 3 shows that for every fixed discount factor, as the market size grows, all elite firms (should they exist) must obtain almost their maximum payoffs at every history. In other words, elite firms are untouchable in large matching markets. Elite firms, however, make up only a vanishingly small fraction of the market by definition. Theorem 4 shows that as long as a firm quality class makes up a nonvanishing fraction of the market, the range of payoffs

that can be sustained in self-enforcing matching processes will be non-vanishing as the market size grows large. This contrasts with elite firms, whose range of possible payoffs becomes degenerate in large markets. Perhaps surprisingly, Theorem 4 also applies to the top quality class as long as its size is not vanishing (in this case, there would be no elite firms in the market).

To understand the intuition of Theorem 3, suppose that f^* is an elite firm and that the value of each worker is bounded between $[0, 1]$. As the market size grows, the workers who are worth at least $(1 - \epsilon)$ to f^* will likely far outnumber the total hiring slots at elite firms. Since f^* faces competition only from other elite firms, it is almost guaranteed to fill all its hiring slots with these workers. As the market size grows large, this payoff guarantee will approach f^* 's payoff upper bound despite the randomness in preference realizations. One challenge in proving Theorem 4 is that in large matching markets, players obtain approximately efficient payoffs from all static stable matchings, so these matchings cannot be used to punish deviating firms.² Instead, I show that a variant of the worker-proposing serial dictatorship can punish firms effectively even in large matching markets.

Related literature This paper is related to several different lines of research. First, my paper is part of a large and active literature on community enforcement, which studies how repeated interactions can lead to desirable outcomes that are not sustainable in one-shot interactions. See, for example, Kandori (1992), Ellison (1994), Wolitzky (2013), Ali and Miller (2016), Acemoglu and Wolitzky (2020, 2021), and Deb, Sugaya, and Wolitzky (2020). The main difference between my paper and the existing literature is that I focus on two-sided and one-to-many matching environments, which may contain top coalition players who cannot be motivated dynamically even when patience is high. To quantify the impact of these untouchable players, I build on techniques from the large matching market literature to obtain asymptotic characterizations of their relative size in the market.

Second, this paper is also related to the literature on dynamic matching. Du and Livne (2016) and Doval (2022) consider the existence of self-enforcing arrangements in a setting where matching is one-to-one, and players leave the market permanently once matched.³ Another strand of this literature investigates self-enforcing arrangements in matching markets where the links among long-lived players can be revised over time. See, for example, Corbae, Temzelides, and Wright (2003), Damiano and Lam (2005), Kurino (2020), Newton and Sawa (2015), Kadam and Kotowski (2018a,b), Kotowski (2020), and Altınok (2021). The main difference in the current paper is that I study a setting where a fixed set of long-lived players match with multiple short-lived players in every period. As a result of this difference, while dynamic incentives typically impede either stability or efficiency in the existing literature, in my paper, they are used as a carrot and stick to enforce more stable outcomes.

²See Pittel (1989, 1992) and Lee (2016) for large-market results on this point; Ashlagi, Kanoria, and Leshno (2017) show that this is true even in small matching markets.

³See Ünver (2010), Anderson, Ashlagi, Gamarnik, and Kanoria (2015), Baccara, Lee, and Yariv (2019), Leshno (2022), and Akbarpour, Li, and Gharan (2020) for the welfare implications of various dynamic matching algorithms in such markets.

Third, the current paper is also part of a nascent literature that combines repeated games and cooperative games. [Bernheim and Slavov \(2009\)](#) study a repeated version of Condorcet winner. [Ali and Liu \(2020\)](#) consider coalitional deviations in general repeated games where the stage game can be either a strategic-form game or a cooperative game. The solution concept in this paper builds on the full history dependence and subgame-perfection requirement in these papers. The main difference is in the form of effective coalitions: in both these papers, the effective coalitions consist of subsets of long-lived players; in the current paper, however, the effective coalitions are those that consist of a single long-lived player and multiple generations of short-lived players. [Ali and Liu \(2020\)](#) also focus on the effects of public versus secret payments, whereas the current paper focuses on matching markets without transfers. More recently, [Bardhi, Guo, and Strulovici \(2023\)](#) use a similar solution concept as [Ali and Liu \(2020\)](#) to study early career discrimination in matching markets where wages are flexible.

2. REPEATED MATCHING MARKET

In this section, I first review the benchmark static matching environment and then extend the model to repeated matching markets. I also introduce the notion of top coalition sequence and show that it determines whether or not a firm can be motivated through continuation play.

2.1 Model

Players At the beginning of each period $t = 0, 1, 2, \dots$, a new generation of workers $\mathcal{W}$ enter the market to match with a fixed set of firms $\mathcal{F}$. Firms are long-lived players who persist through time. Workers are short-lived and remain in the market for only one period, but the composition of the workers in each generation is the same. Matching is one-to-many: each firm f has $q_f > 0$ hiring slots to fill in every period.

Each worker w has a strict preference relation $\succ_w$ over the set of firms and being unmatched (being unmatched is denoted w). I write $f \succeq_w f'$ if either $f \succ_w f'$ or $f = f'$.

Each firm f has a utility function $\tilde{u}_f : 2^{\mathcal{W}} \rightarrow \mathbb{R}$ defined on all subsets of workers. Firms' utility functions are strict ($\tilde{u}_f(W) = \tilde{u}_f(W')$ only if $W = W'$) and responsive: for all $f \in \mathcal{F}$, $W \subseteq \mathcal{W}$, and $w, w' \notin W$,⁴

- $\tilde{u}_f(W \cup \{w'\}) > \tilde{u}_f(W \cup \{w\})$ if and only if $\tilde{u}_f(\{w'\}) > \tilde{u}_f(\{w\})$
- $\tilde{u}_f(W \cup \{w\}) > \tilde{u}_f(W)$ if and only if $\tilde{u}_f(\{w\}) > \tilde{u}_f(\emptyset)$.

That is, replacing a worker with someone better (or adding an acceptable worker) makes the firm f better off. Firms share a common discount factor δ and evaluate a sequence of flow utilities through exponential discounting.

⁴Firms' utility functions are defined even for groups of workers that exceed its capacity constraint: this allows for a simpler notation for utility functions. However, the concepts of stage-game matching and feasible deviation both require firms to respect their capacity constraints, which rules out the possibility for any firm f to match with more than q_f workers.

Stage game The stage game in every period is a static one-to-many coalitional matching game played between the firms and the workers who are active in that period. Formally, a stage-game matching m is a mapping defined on the set $\mathcal{F} \cup \mathcal{W}$ such that (i) for every $w \in \mathcal{W}$, $m(w) \in \mathcal{F} \cup \{w\}$, (ii) for every $f \in \mathcal{F}$, $m(f) \subseteq \mathcal{W}$ and $|m(f)| \leq q_f$, and (iii) $w \in m(f)$ if and only if $m(w) = f$. Let M denote the set of all stage-game matchings. For each $f \in \mathcal{F}$, let $u_f : M \rightarrow \mathbb{R}$ be firm f 's utility function over stage-game matchings induced from its preference over workers: $u_f(m) \equiv \tilde{u}_f(m(f))$ for all $m \in M$.

A stage-game matching m is subject to three types of deviations: (i) a deviation by a firm f , where f fires a subset of its employees and leaves those positions unfilled; (ii) a deviation by a worker w , where w leaves her employer and remains unmatched; (iii) a deviation consisting of a firm f and a subset of workers W , where f and W match together and abandon any other preexisting match partners. However, observe that the firm f firing a subset of its employees is equivalent to f deviating with the workers that remain employed by f . Therefore, it is without loss to focus only on the latter kind of deviation in addition to deviations by individual workers.

We say that a matching m is *acceptable* to worker w if $m(w) \succeq_w w$. The coalitional deviation $\{f\} \cup W$ from m is said to be *feasible* to f if $|W| \leq q_f$ and $f \succ_w m(w)$ for $w \in W \setminus m(f)$: each worker $w \in W$ is either already working for f or finds herself better off to do so. Furthermore, the coalitional deviation $\{f\} \cup W$ is said to be *profitable* for f if $\tilde{u}_f(W) > u_f(m)$. Finally, a stage-game matching m is *stable* if all workers find it acceptable, and no firm can find any coalition deviation that is both feasible and profitable.⁵

In static matching models, there is no need to specify the resulting matching outcome after a deviation. This is because in static matching environments, the profitability of a deviation does not depend on how others respond. However, in repeated matching markets where the past influences the future, we have to specify what outcome is realized after a deviation. To this end, let $[m, (f, W)] \in M$ denote the resulting stage-game matching after coalition $\{f\} \cup W$ deviates from stage-game matching m . I make the following assumption.

ASSUMPTION 1. *The stage-game matching $m' = [m, (f, W)]$ satisfies $m'(f) = W$, and $m'(f') = m(f') \setminus W$ for all $f' \neq f$.*

Assumption 1 states that members of the deviating coalition are matched together in the resulting stage-game matching; in addition, players abandoned by the deviators remain unmatched, while those untouched by the deviation remain matched as before. One can make alternative assumptions that specify how other players may further deviate within the period after the initial deviation. As long as it is possible to identify the firm that initiated the first deviation, identical results follow.⁶

⁵Note that the notion of stability outlined above is stronger than pairwise stability à la Gale–Shapley but weaker than the strong core, which would rule out deviations involving multiple firms. These different notions of stability coincide in static matching environments where preferences satisfy substitutability, but can lead to different analyses in dynamic settings. I discuss this point in more detail after introducing the dynamic stability notion in Definition 1.

⁶Assumption 1 is only needed for establishing Lemma 2 in Appendix A.1: it is always possible to identify the firm responsible for a deviation.

Finally, it is worth noting that even though the stage game is a cooperative game featuring deviations by coalitions, the firms play a much more active role than the workers: essentially, a firm can choose any worker who prefers them to their current match. At the end of this section, I discuss an alternative normal-form game with firms as the only active players, and argue that the repeated matching market can be analyzed by studying subgame perfect Nash equilibria in the corresponding repeated normal-form game.

Repeated matching market The timing in each period is as follows: at the beginning of the period, a realization $\omega \in \Omega$ is drawn from a public randomization device (by, for example, a centralized matching clearing house); based on ω and the history of past interactions, a recommended stage-game matching is created for the players who are currently in the market; firms and workers then decide whether to deviate from this recommendation, which leads to the realized stage-game matching. Note that the public randomization device is *not* intended to represent the random realization of players' preferences.⁷ Instead, it represents the ability of the matching clearing house to randomize its recommendations.

A t -period ex ante histories $h = (\omega_\tau, m_\tau)_{\tau=0}^{t-1}$ specifies a sequence of past realizations from the public randomization device and matching outcomes before the randomization at the beginning of period t is drawn. I use $\overline{\mathcal{H}}_t$ to denote the set of all t -period ex ante histories, with $\overline{\mathcal{H}}_0 = \{\emptyset\}$ the singleton set comprising the initial null history. Let $\overline{\mathcal{H}} = \bigcup_{t=0}^{\infty} \overline{\mathcal{H}}_t$ be the set of all ex ante histories. Finally, let $\mathcal{H}_t \equiv \overline{\mathcal{H}}_t \times \Omega$ denote the set of t -period ex post histories and let $\mathcal{H} \equiv \overline{\mathcal{H}} \times \Omega$ denote the set of all ex post histories.

A *matching process* proposes a stage-game matching following each ex post history: a matching process μ is a mapping $\mu : \mathcal{H} \rightarrow M$. One can interpret a matching process as proposals from a history-dependent matching protocol. I use $\mu(f|h)$ and $\mu(w|h)$ to denote the match partners of firm f and worker w in the stage-game matching $\mu(h)$, respectively.

Let $\overline{\mathcal{H}}_\infty = (\Omega \times M)^\infty$ be the set of outcomes of the repeated matching market. For an outcome $h \in \overline{\mathcal{H}}_\infty$, let $m_\tau(h)$ denote the stage-game matching in the τ th period of h . Following every t -period (ex ante or ex post) history $\hat{h} \in \overline{\mathcal{H}} \cup \mathcal{H}$, let

$$U_f(\hat{h}|\mu) \equiv (1 - \delta)\mathbb{E}_\mu \left[\sum_{\tau=t}^{\infty} \delta^{\tau-t} u_f(m_\tau(h)) \middle| \hat{h} \right]$$

denote the continuation payoff firm f obtains from μ following $\hat{h}$, where the expectation is taken with respect to the measure over $\overline{\mathcal{H}}_\infty$ induced by μ conditional on $\hat{h}$. I will also use $U_f(\mu) \equiv U_f(\emptyset|\mu)$ to denote the payoff firm f obtains from μ at the beginning of the game.

Deviation plan In repeated matching markets, a firm f can participate in a sequence of deviations by forming coalitions with workers across multiple generations. Each of these

⁷In fact, preferences are assumed to be predetermined so far; in Section 3, I augment the stage-game to account for the random preferences realizations.

coalitions must be immediately profitable for the short-lived workers, but not necessarily for f , since it cares about the utility it collects from the entire sequence.

Motivated by this observation, I define a *deviation plan* for a firm f as a complete contingent plan that specifies, at each ex post history, a set of workers with whom f wishes to form a deviating coalition: a deviation plan for f is a mapping $d_f : \mathcal{H} \rightarrow 2^{\mathcal{W}}$. Together with the original matching process, a deviation plan generates an altered distribution over the outcome of the game $\overline{\mathcal{H}}_\infty$. Given a matching process μ and f 's deviation plan d_f , the *manipulated matching process*, denoted $[\mu, (f, d_f)] : \mathcal{H} \rightarrow M$, is a matching process defined by

$$[\mu, (f, d_f)](h) \equiv [\mu(h), (f, d_f(h))] \quad \text{for every } h \in \mathcal{H}.$$

To understand the expression above, note that at every history h of the manipulated matching process $[\mu, (f, d_f)]$, the coalition $f \cup d_f(h)$ deviates from the stage-game matching $\mu(h)$ prescribed by μ , which results in the stage-game matching $[\mu(h), (f, d_f(h))]$. The deviation plan d_f is said to be feasible if at every ex post history h , the following statements hold:

- (i) We have $|d_f(h)| \leq q_f$, so the stage-game deviation respects f 's capacity constraint.
- (ii) We have $f \succ_w \mu(w|h)$ for $w \in d_f(h) \setminus \mu(f|h)$, so any worker specified by d_f either already works for f or finds herself strictly better off to do so.

A deviation plan d_f is profitable for firm f if there exists an ex post history h where $U_f(h|[\mu, (f, d_f)]) > U_f(h|\mu)$; that is, having reached the ex post history h , the continuation value from carrying out the deviation plan exceeds that from following the matching process μ .

Note that if the deviation plan d_f agrees with the matching process μ at every ex post history, i.e., $d_f(h) = \mu(f|\hat{h})$ for all $h \in \mathcal{H}$, then it is equivalent to f following the prescription of μ . For much of the analysis in this paper, I focus on a special class of deviation plans that disagree with μ at only one ex post history: a deviation plan d_f is a *one-shot deviation* from a matching process μ if there is a unique ex post history $\hat{h}$ where $d_f(\hat{h}) \neq \mu(f|\hat{h})$.

Self-enforcing matching process The notion of self-enforcing matching process captures stability in repeated matching markets.

DEFINITION 1. A matching process $\mu : \mathcal{H} \rightarrow M$ is *self-enforcing* if

- (i) $\mu(w|h) \succeq_w w$ for every w at every ex post history h
- (ii) no firm f has a deviation plan that is both feasible and profitable.

The first requirement in Definition 1 guards against deviations by singleton workers in every generation; the second requirement asks that no firm can profit from chaining together a sequence of deviating coalitions, each of which is immediately profitable for

the deviating workers, but not necessarily for the firm itself. Note that these requirements are imposed at every ex post history, including those that are off-path: this embeds sequential rationality in the same way as subgame perfection does in a repeated non-cooperative game. Finally, observe that Definition 1 coincides with the definition of static stable matchings when patience is 0.

Note that Definition 1 focuses only on coalitional deviations that involve a single firm. In dynamic environments, this is generally not equivalent to considering deviating coalitions with multiple firms. However, allowing infinite-horizon deviations by multiple long-lived players creates a conceptual difficulty in assessing whether those deviations themselves can be self-enforcing. One appealing feature of the matching environment studied in this paper is that here it is relatively standard to study deviations in coalitions that comprise a single firm and potentially multiple workers, which avoids these conceptual issues. In a separate paper (Ali and Liu (2020)), we show that if coalitions cannot commit to long-run behavior and are unable to make anonymous transfers, then modeling coalitions with multiple long-run players does not alter the set of sustainable outcomes when players are patient.

Lemma 1 below establishes a one-shot deviation principle for self-enforcing matching processes: to check whether a matching process is self-enforcing, instead of checking all deviation plans, it suffices to focus on those that only depart from the matching process at a single history.

LEMMA 1 (One-Shot Deviation Principle). *A matching process μ is self-enforcing if and only if*

- (i) $\mu(w|h) \succeq_w w$ for every w at every ex post history h
- (ii) no firm f has a one-shot deviation that is both feasible and profitable.

The proof of Lemma 1 follows arguments similar to the one-shot deviation principle for repeated normal-form games and is relegated to Appendix A.1. An immediate implication of Lemma 1 is that self-enforcing matching processes exist at every patience level.

OBSERVATION 1. There exists a self-enforcing matching process for every $0 \leq \delta < 1$.

Recall that we assumed that the firms' preferences are responsive. Standard results in the static matching literature (see, for example, Theorem 6.5 in Roth and Sotomayor (1992)) then ensure the existence of stable static matchings. By Lemma 1, the infinite repetition of a static stable matching is always a self-enforcing matching process.

Active firm, passive workers In the current paper, the stage game of the repeated matching market is modeled as a cooperative game, which is consistent with the tradition of the static matching literature. However, the stage game can also be modeled as a normal-form game with firms as the only players. In this normal-form game, each firm f 's action space consists of subsets of $\mathcal{W}$ with no more than q_f workers—intuitively, think of each action as the firm proposing to the group of workers. By contrast, workers

		Firm f_2			
		$\emptyset$	w_1	w_2	$\{w_1, w_2\}$
Firm f_1	$\emptyset$	(0, 0)	(0, 1)	(0, 2)	(0, 3)
	w_1	(2, 0)	(0, 1)	(2, 2)	(0, 3)
	w_2	(1, 0)	(1, 1)	(1, 0)	(1, 1)

FIGURE 4. Workers as actions.

are not modeled as players in this representation; instead, given an action profile from the firms, we use the imputed worker assignment to determine the firms’ payoffs. In particular, each worker is assigned to her favorite acceptable proposer or remains unmatched if she receives no acceptable proposals. We will refer to this normal-form game as the *active firms, passive workers* (AFPW) representation of the matching market, in contrast to the cooperative-game representation we have focused on thus far.

As an illustration, consider a matching market where the stage game consists of two firms $\mathcal{F} = \{f_1, f_2\}$ and two representative workers $\mathcal{W} = \{w_1, w_2\}$, with firms’ capacities $q_1 = 1$ and $q_2 = 2$. Suppose that each firm derives a payoff of 2 from the worker sharing the same index as itself, a payoff of 1 from the worker with a different index, and 0 payoff from unfilled positions. Assume also that f_2 has additively separable payoffs from its two hiring slots. Each worker prefers to work for the firm with a distinct index over working for the firm with the same index, but working for either firm is better than unemployment.

Figure 4 shows the AFPW representation of the stage game. The pure-strategy Nash equilibria are (w_2, w_1) and $(w_2, \{w_1, w_2\})$, both of which correspond to the unique stable matching where f_1 is matched with w_2 and f_2 is matched with w_1 .

More broadly, in Appendix A.5, I show that for the kind of matching markets considered in the current paper, analyzing self-enforcing matching processes in the repeated matching market is equivalent to analyzing subgame perfect Nash equilibria in the repeated AFPW game. Intuitively, each action profile in the AFPW stage game corresponds to a stage-game matching that is acceptable to all workers. Moreover, the payoffs a firm can achieve by deviating in the AFPW stage game are identical to the payoffs it can obtain through feasible coalitional deviations in the coalitional matching game. As a result, the enforceable outcomes in the repeated AFPW game and the repeated coalitional matching game are identical.

2.2 Top coalition sequence

In the second example in Section 1, history dependence is powerless against top firms due to their immunity to future punishments. In a general matching environment without assuming common worker preferences, the appropriate notion of top players is captured by top coalitions.

Fix an arbitrary subset of firms and workers $F \cup W \subseteq \mathcal{F} \cup \mathcal{W}$. A firm $\widehat{f} \in F$ and a set of workers $\widehat{W} \subseteq W$, $|\widehat{W}| \leq q_{\widehat{f}}$ form a top coalition in $F \cup W$ in the following circumstances:

- (i) If $\widetilde{u}_{\widehat{f}}(\widehat{W}) \geq u_{\widehat{f}}(W')$ for all $W' \subseteq W$ such that $|W'| \leq q_{\widehat{f}}$: subject to capacity constraint $q_{\widehat{f}}$, $\widehat{W}$ is $\widehat{f}$ ’s favorite group of workers among W .

(ii) If $\hat{f} \succeq_w f'$ for all $w \in \hat{W}$ and $f' \in F \cup \{w\}$: $\hat{f}$ is the favorite firm for every worker in $\hat{W}$.

In other words, $\hat{f}$ and $\hat{W}$ are mutual favorites. The top coalition sequence takes this idea further by iteratively finding and eliminating top coalitions in the remaining players until no new top coalition can be found.

DEFINITION 2. The top coalition sequence is the ordered set $\mathcal{T} = \{(\hat{f}_1, \hat{W}_1), (\hat{f}_2, \hat{W}_2), \dots\}$ produced by the following procedure:⁸

- Initialization: Set $\mathcal{T} = \emptyset$.
- New Phase:
 - (i) If $(\mathcal{F} \cup \mathcal{W}) \setminus \mathcal{T}$ contains no top coalition, stop.
 - (ii) If $(\mathcal{F} \cup \mathcal{W}) \setminus \mathcal{T}$ has a top coalition $(\hat{f}, \hat{W})$, add $(\hat{f}, \hat{W})$ to $\mathcal{T}$ and restart New Phase.

The top coalition sequence is related to, but distinct from, the top coalition property studied in various cooperative game settings. While the top coalition sequence is an object constructed from arbitrary player preferences, the top coalition property, by contrast, is an assumption that requires the top coalition sequence to include all players in the stage game.⁹ We summarize this connection as an observation below.

OBSERVATION 2. If the stage game satisfies the top coalition property, then the top coalition sequence includes all firms and workers.

Below are some observations on how the composition of the top coalition sequence may depend on the preference configurations in the market.

OBSERVATION 3. Suppose all firms are acceptable to all workers ($f \succ_w w$ for all $w \in \mathcal{W}$).

- (i) If workers share a common preference ranking over firms, then all players are in the top coalition sequence.
- (ii) When firms share a common utility function over workers, the top coalition sequence may be empty.

The first point in Observation 3 follows from the iterative elimination of the top firm along the workers' shared preference list, just as in the second example in Section 1. The second point is illustrated through the following example.

⁸Whenever it causes no confusion, I will use $\mathcal{T}$ to denote both the set of (f, W) pairs and the set of players who show up in those pairs.

⁹See, for example, [Eckhout \(2000\)](#), [Banerjee, Konishi, and Sönmez \(2001\)](#), and [Pycia \(2012\)](#) and [Wu \(2015\)](#) for applications of the top coalition property. See [Peralta \(2020\)](#) for an example of the application of the top coalition sequence in static matching environments.

EXAMPLE 1. $\mathcal{F} = \{f_1, f_2\}$ and $\mathcal{W} = \{w_1, w_2, w_3, w_4\}$. The two firms f_1 and f_2 are identical: both have capacity $q_{f_1} = q_{f_2} = 2$ and share a common utility function $u_{f_1}(w_i) = u_{f_2}(w_i) = i$. Suppose $\succ_{w_i} = f_1, f_2, w_i$ if i is odd and $\succ_{w_i} = f_2, f_1, w_i$ if i is even.

The procedure in Definition 2 stops at the first step: both firms point to $\{w_3, w_4\}$ as their favorite workers, but since neither f_1 nor f_2 is the favorite for both $\{w_3, w_4\}$, there is no top coalition and $\mathcal{T} = \emptyset$. $\diamond$

2.3 The limit of self-enforcement

The results in this section explore the extent to which history dependence can be used to alter the matches obtained by firms when they are sufficiently patient. In particular, I say that a firm is *untouchable* in the repeated matching market if, regardless of the patience level, the firm always retains the same set of workers at each history in every self-enforcing matching process. Therefore, dynamic enforcement cannot be used to alter the matches obtained by untouchable firms as compared to static stable matchings.

Impossible to motivate top coalition sequence As Theorem 1 shows, the firms in the top coalition sequence are untouchable.

THEOREM 1. Suppose $(\hat{f}, \hat{W})$ is in the top coalition sequence. Then $\hat{f}$ is matched to $\hat{W}$ in all static stable matchings. Moreover, for every $0 < \delta < 1$, $\hat{f}$ is matched to $\hat{W}$ in every self-enforcing matching process at every ex post history.

Theorem 1 states that the firms and workers in the top coalition sequence are always matched together in both static and repeated matching markets. This holds in a stark sense in repeated matching markets since it applies regardless of firms' patience and after every ex post history, including those that are off-path.

The complete proof of Theorem 1 can be found in Appendix A.2. Here are the key steps. First, $(\hat{f}, \hat{W})$ being mutual favorites implies that they must be matched together in any static stable matching. Second, in a repeated matching market, $(\hat{f}, \hat{W})$ being mutual favorites further implies that $\hat{f}$ is not punishable through continuation value: whenever a matching process recommends $\hat{f}$ to match with $W \neq \hat{W}$, everyone in $\hat{W}$ is willing to deviate with $\hat{f}$, so $\hat{f}$'s continuation value cannot be lower than $u_{\hat{f}}(\hat{W})$; at the same time, $\hat{f}$'s continuation value also cannot be higher than $u_{\hat{f}}(\hat{W})$, so $\hat{f}$'s continuation value must be precisely $u_{\hat{f}}(\hat{W})$, no matter what happens in the current period. Without credible changes in its continuation value, $\hat{f}$ behaves just like a short-lived player, so it must always match with $\hat{W}$ at every ex post history. Finally, an inductive argument extends this logic to the entire top coalition sequence.

The following is an immediate corollary of Theorem 1 and Observation 2.

COROLLARY 1. If the stage-game matching market satisfies the top coalition property, there is a unique self-enforcing matching process where the unique static stable matching is played at each history.

Difference from standard folk theorem The “impossibility” implication from Theorem 1 stands in contrast to what one might expect from standard folk theorems for repeated games, where many outcomes are sustainable at high patience levels: here, the matching outcome for players in $\mathcal{T}$ is unique no matter how high the patience is.

To understand why, note that for firm f to deviate and poach a worker, the worker must strictly prefer f over her current match. So when stage-game matching m is recommended, f can choose from workers in $D_f(m) \equiv m(f) \cup \{w \in \mathcal{W} : f \succ_w m(w)\}$: these workers are either already working for f or can be poached by f if it wishes. On the other hand, to ward off deviations by individual workers, a self-enforcing matching process can only recommend stage-game matchings in $M^\circ \equiv \{m \in M : m \succeq_w w \text{ for all } w \in \mathcal{W}\}$: these are the matchings that are acceptable to all workers. The minmax payoff for firm f is the payoff it obtains from its “best response” to the worst possible recommendation, which is given by

$$\underline{u}_f \equiv \min_{m \in M^\circ} \max_{W \subseteq D_f(m), |W| \leq q_f} u_f(W). \quad (1)$$

A naive adaptation of the standard folk theorem would state that any payoff profile that gives every firm f strictly higher than their respective $\underline{u}_f$ can be sustained through a self-enforcing matching process as $\delta \rightarrow 1$.

The problem with this approach is that whenever there exists a top coalition, say $(\hat{f}, \hat{W}) \in \mathcal{T}$, then the workers in $\hat{W}$ are always available to $\hat{f}$ when $\hat{f}$ deviates. That is, $\hat{W} \subseteq D_{\hat{f}}(m)$ for all $m \in M^\circ$. According to (1), firm $\hat{f}$ ’s minmax payoff therefore satisfies

$$\underline{u}_{\hat{f}} \equiv \min_{m \in M^\circ} \max_{W \subseteq D_{\hat{f}}(m), |W| \leq q_{\hat{f}}} u_{\hat{f}}(W) \geq u_{\hat{f}}(\hat{W}) = \max_{W \subseteq \mathcal{W}, |W| \leq q_{\hat{f}}} u_{\hat{f}}(W),$$

where the last equality follows since $\hat{W}$ is $\hat{f}$ ’s favorite group of employees. The expression above implies that $\hat{f}$ ’s minmax payoff is identical to its highest feasible payoff, so the set of feasible payoff profiles giving $\hat{f}$ strictly higher than $\underline{u}_{\hat{f}}$ is an empty set. A top coalition firm essentially has an action that guarantees itself the highest possible payoff from the stage game independent of the actions of other firms, and the folk theorem is always vacuous for such payoff structures.

Moreover, (1) also leads to incorrect minmax payoffs for firms that are not in $\mathcal{T}$. In light of Theorem 1, any credible recommendation must match $\hat{f}$ with $\hat{W}$, but some stage-game matchings in M° do not. As a result, (1) incorrectly assumes that $\hat{f}$ ’s hiring capacity can be used when punishing other firms, which underestimates the minmax payoffs for firms other than $\hat{f}$.

It should be noted that the subtlety introduced by the top coalition sequence is different from the failure of full dimensionality that is studied in Wen (1994) and Fudenberg, Levine, and Takahashi (2007). In the aforementioned papers, there is a nonempty set of payoff profiles that are both feasible and strictly higher than each player’s minmax; however, this set lacks dimensionality due to long-run players having aligned preferences. The presence of a top coalition, by contrast, leads to an empty set of such payoff profiles. The construction of the top coalition sequence in Definition 2 is an iterative process of finding degenerate payoff dimensions while recalibrating the minmax payoffs for the remaining firms until we arrive at a reduced game without new top coalitions.

Modified folk theorem in the reduced game Motivated by the discussion above, let us introduce a few notations that are useful for analyzing the reduced game after the top coalition sequence has been removed. Let $\mathcal{R} \equiv (\mathcal{F} \cup \mathcal{W}) \setminus \mathcal{T}$ denote the players who are not in the top coalition sequence, and let $M_{\mathcal{R}}$ denote the set of stage-game matchings that ensure the top coalition sequence is matched together:

$$M_{\mathcal{R}} \equiv \{m \in M : m(\hat{f}) = \hat{W} \text{ for all } (\hat{f}, \hat{W}) \in \mathcal{T}\}.$$

Let $M_{\mathcal{R}}^{\circ}$ denote the stage-game matchings in $M_{\mathcal{R}}$ that are acceptable to all workers. Recall that for each stage-game matching m , firm f can form feasible coalitional deviations with workers in $D_f(m) \equiv m(f) \cup \{w \in \mathcal{W} : f \succ_w m(f)\}$. For every firm $f \in \mathcal{F} \cap \mathcal{R}$, its reduced-game minmax is then

$$\underline{u}_f^{\mathcal{R}} \equiv \min_{m \in M_{\mathcal{R}}^{\circ}} \max_{W \subseteq D_f(m), |W| \leq q_f} u_f(W). \quad (2)$$

Notice that for every firm $f \in \mathcal{F} \cap \mathcal{R}$, its reduced-game minmax is higher than the value produced from (1), since the minimization is taken over a more restricted set of stage-game recommendations. Finally, let $\Lambda^* \equiv \{\lambda \in \Delta(M_{\mathcal{R}}^{\circ}) : u_f(\lambda) > \underline{u}_f^{\mathcal{R}} \text{ for all } f \in \mathcal{F} \cap \mathcal{R}\}$ denote the randomizations over $M_{\mathcal{R}}^{\circ}$ that secure each firm strictly higher than its reduced-game minmax.

In contrast to Theorem 1, Theorem 2 shows that history dependence can be used to change the matches obtained by players outside of the top coalition sequence: every random matching in Λ^* can be sustained on-path in a stationary manner in a self-enforcing matching process.

THEOREM 2. *For every $\lambda \in \Lambda^*$, there is a $\underline{\delta}$ such that for every $\delta \in (\underline{\delta}, 1)$, there exists a self-enforcing matching process that randomizes according to λ in every period.*

The first step of the proof is to show that firms' payoffs in the reduced game always satisfy the non-equivalent utilities (NEU) condition: no firm's payoff can be a positive affine transformation of another (Abreu, Dutta, and Smith (1994)). The proof then uses this condition to construct player-specific punishments to deter deviations. Given these punishments, the final step adapts the construction from Fudenberg and Maskin (1986) to show that the payoff profile corresponding to any $\lambda \in \Lambda^*$ can be sustained in a self-enforcing matching process when firms are patient. The complete proof can be found in the Supplementary Appendix, available in a supplementary file on the journal website, <https://econtheory.org/supp/4898/supplement.pdf>.

3. LARGE-MARKET ANALYSIS

We see in Section 2 that in repeated matching markets, firms in the top coalition sequence are untouchable, while the others can be motivated through history dependence. But how large is the aggregate impact of these untouchable firms relative to those that can be motivated dynamically? The goal of this section is to quantify asymptotically the relative size of these two kinds of firms. To do this, I build on the repeated matching

model introduced in Section 2, but augment it with randomly drawn workers and focus on large-market analysis. Section 3.1 introduces this setup; Section 3.2 characterizes the large-market asymptotics in this environment.

3.1 The setup

I consider a sequence of market sizes n , letting n go to infinity. In a market of size n , the stage game consists of n firms $\mathcal{F}_n$ each with a hiring quota q ,¹⁰ and workers $\mathcal{W}_n$ where $|\mathcal{W}_n| = \lceil \beta n q \rceil$ with $\beta > 0$. Let M_n denote the set of stage-game matchings among $\mathcal{F}_n$ and $\mathcal{W}_n$.

I use the finite-tier random preference model to capture positive preference correlations that arise from quality differentiation in the market.¹¹ For every n , firms can be partitioned into K quality classes $\mathcal{F}_n = \{\mathcal{F}_n^1, \mathcal{F}_n^2, \dots, \mathcal{F}_n^K\}$. Every worker prefers a firm from a higher quality class to those from a lower quality class, but each worker's preference ranking over firms within the same quality class $\mathcal{F}_n^k$ is drawn uniformly from all permutations of $\mathcal{F}_n^k$. Let π_n denote a realization of worker preferences that are compatible with this restriction. I assume that the proportion of tier- k firms, $|\mathcal{F}_n^k|/n$, converges to $x_k \geq 0$ for $1 \leq k \leq K$.

Similarly, workers can be partitioned into L quality classes $\mathcal{W}_n = \{\mathcal{W}_n^1, \mathcal{W}_n^2, \dots, \mathcal{W}_n^L\}$. When a firm f matches with a worker $w \in \mathcal{W}_n^l$, the firm receives

$$\tilde{u}_f(w) = V(C_l, \zeta_{f,w}),$$

where C_l is the common value shared by all workers in $\mathcal{W}_n^l$ satisfying $C_l > C_{l'}$ for all $l < l'$, and $\zeta_{f,w}$ is the idiosyncratic match quality between f and w . I assume that the quality component for each tier, C_l , is constant over time, while the idiosyncratic components $\zeta_{f,w}$ are drawn independently for every worker in each cohort from the uniform distribution over $[0, 1]$. The function $V(\cdot, \cdot) : \mathbb{R}_+^2 \rightarrow \mathbb{R}_+$ is continuous and strictly increasing in both arguments, and satisfies $V(C_l, 0) > V(C_{l'}, 1)$ for all $l < l'$ (so there is no overlap between tiers). Firms have additive utilities for each job opening, $\tilde{u}_f(W) = \sum_{w \in W} \tilde{u}_f(w)$, and derive zero utility from unfilled positions. Let $\zeta_n = \{\zeta_{f,w}\}_{f \in \mathcal{F}_n, w \in \mathcal{W}_n}$ denote a realization of the matrix of idiosyncratic match qualities. I assume that the proportion of tier- l workers, $|\mathcal{W}_n^l|/|\mathcal{W}_n|$, converges to $y_l \geq 0$ for $1 \leq l \leq L$.

The timing in each period is as follows. First, a new cohort of workers arrives, and preferences π_n and ζ_n are realized; the public randomization $\omega \in \Omega$ is then realized; based on the realization of (π_n, ζ_n, ω) , a stage-game matching is recommended for $\mathcal{F}_n$ and $\mathcal{W}_n$; players then decide whether to deviate from this recommendation, which determines the outcome of the stage game. I will refer to the realization of $s_n = (\pi_n, \zeta_n, \omega) \in S_n$ together as a *state*. The notions of ex ante and ex post histories, introduced in Section 2, are modified accordingly with s_n replacing ω .

¹⁰The assumption that each firm has an identical quota is only for convenience. The same results continue to hold if each firm has a different quota.

¹¹See, for example, Ashlagi, Kanoria, and Leshno (2017) for an ordinal version of the multiple-tiered preferences, and Che and Tercieux (2019) for a cardinal parameterization of this class of preferences. See also Lee (2016) for a different way to model preference correlation without tiers.

3.2 Which firms are untouchable?

In repeated matching markets, we say a firm is untouchable if it must match with the same set of workers at every history of the market regardless of the patience level. In the current section, since workers are drawn randomly every period, it is impossible for any firm to always match with the same workers. Instead, I say a firm is *untouchable* in large matching markets if at every fixed patience level, the firm's ex ante stage-game payoffs across all histories can be bounded by an arbitrarily small interval as market size grows large.

Indeed, with sufficient patience, the range of payoffs that can be sustained for the vast majority of firms is nondegenerate even as the market size grows to infinity. The only untouchable firms, should they exist, are elite firms that make up at most a vanishingly small fraction of the market.

Specifically, I say that the best class of firms $\mathcal{F}_n^1$ is an elite class if the number of firms it contains is vanishing relative to the best class of workers.

DEFINITION 3. The class of firms $\mathcal{F}_n^1$ is an elite quality class if $|\mathcal{F}_n^1|/|\mathcal{W}_n^1| \rightarrow 0$ as $n \rightarrow \infty$.

A firm in $f \in \mathcal{F}_n^1$ dominates the firms in all other lower-quality classes, so workers can possibly only rank other firms in $\mathcal{F}_n^1$ higher than f . If $\mathcal{F}_n^1$ is also an elite firm class, then the number of firms that can be considered better than f by any worker becomes vanishingly small relative to the workers in $\mathcal{W}_n^1$. As the market size increases, an increasingly large number of workers in $\mathcal{W}_n^1$ will rank f as their top choice, so each elite firm becomes “over-demanded” in large markets. However, note that elite firms may not necessarily exist: if $\lim_{n \rightarrow \infty} \mathcal{F}_n^1/n > 0$, then no firms in the market would qualify as elite firms.

Theorem 3 shows that at every fixed patience level, when the market size is large, the elite firms must obtain almost the maximum possible payoff at every history.

THEOREM 3. Suppose $\mathcal{F}_n^1$ is an elite firm class; then for every discount factor $0 < \delta < 1$ and every $\epsilon > 0$, there exists N such that for all $n \geq N$, the ex ante stage-game payoffs of every self-enforcing matching process μ satisfies

$$\mathbb{E}[u_f(\mu(\bar{h}, s_n))] > qV(C_1, 1) - \epsilon$$

for every $f \in \mathcal{F}_n^1$ and every ex ante history $\bar{h}$, where the expectation is taken with respect to the realization of state s_n .

The proof of Theorem 3 can be found in Appendix A.3. To understand the intuition, note that an elite firm f only faces competition from other elite firms. As market size grows large, elite firms become rare relative to $\mathcal{W}_n^1$, which makes these firms over-demanded. As a result, they are able to hire workers who are close to the top of their preference ranking no matter what, which guarantees a high continuation value. This future payoff guarantee increases with the market size and eventually makes it impossible to motivate f to accept a low (expected) stage-game payoff in the current period.

For concreteness, consider a one-to-one matching market with two firm classes $\mathcal{F}_n^1$ and $\mathcal{F}_n^2$ and only one worker class $\mathcal{W}_n$. Firm classes satisfy $|\mathcal{F}_n^1| = \sqrt{n}$ and $|\mathcal{F}_n^2| = n - \sqrt{n}$,

so $\mathcal{F}_n^1$ is an elite class. Workers $\mathcal{W}_n$ satisfy $|\mathcal{W}_n| = n$, so there is an equal number of firms and workers. Suppose that firms' payoffs from matching with each worker are drawn from n independent and identically distributed (i.i.d.) uniform random variables on $[0, 1]$. Now, in the worst-case scenario, an elite firm f can secure a worker who ranks $\sqrt{n}$ out of the total n workers on its preference list. In other words, in each period f is guaranteed a payoff equal to the $(n - \sqrt{n} + 1)$ th order static among n i.i.d. uniform random variables on $[0, 1]$. This payoff guarantee yields an expected payoff of $(n - \sqrt{n} + 1)/(n + 1)$, which converges to 1 as $n \rightarrow \infty$. Since δ is fixed, this payoff guarantee eliminates the possibility of dynamic enforcement as $n \rightarrow \infty$.

Note that Theorem 3 is not a folk theorem: instead of taking patience δ to 1, I consider an arbitrary fixed δ while taking the market size n to infinity. A constant discount factor creates an important driving force for Theorem 3, as it limits the impact of continuation payoffs compared to stage-game payoffs. On the other hand, for each fixed market size n and $\epsilon > 0$, it is possible to find a sufficiently high δ so that elite firms are willing to accept expected stage-game payoffs that are lower than $qV(C_1, 1) - \epsilon$. What Theorem 3 does highlight is that for elite firms, dynamic enforcement boils down to a race between their patience versus how much they are over-demanded. In other words, for elite firms, the minimum δ required for dynamic enforcement increases with market size n . This contrasts with Theorem 4 below, which will show that for non-elite firms, the possibility of dynamic enforcement is not diminished by increases in market size.

Although Theorem 3 highlights the limits of dynamic enforcement, it is worth remembering that elite firms, by definition, only make up a vanishing fraction of the market, so their impact is negligible in large markets. The next result, Theorem 4, confirms that as long as a firm belongs to a quality class that makes up a nonvanishing fraction of the market, then the range of sustainable payoffs for this firm will be nondegenerate. Importantly, the affirmative message of Theorem 4 is not the result of a race between δ and N : Theorem 4 is valid especially when both δ and N are large, and, perhaps surprisingly, this result also applies to firms in $\mathcal{F}_n^1$.

THEOREM 4. *Suppose $\lim_{n \rightarrow \infty} |\mathcal{F}_n^k|/n > 0$ for some $1 \leq k \leq K$. There exists a discount factor $0 < \underline{\delta} < 1$, market size $N \geq 0$, and payoff interval $[\underline{v}, \widehat{v}]$ with $\underline{v} < \widehat{v}$, such that for all discount factors $\delta > \underline{\delta}$, market sizes $n > N$, and payoff profile $\mathbf{v}^n \in [\underline{v}, \widehat{v}]^{\mathcal{F}_n^k}$, there exists a self-enforcement matching process μ_n^* that satisfies $U_f(\mu_n^*) = \mathbf{v}_f^n$ for all $f \in \mathcal{F}_n^k$.*

The proof of Theorem 4 builds on results from both the large-market matching literature and the repeated games literature. In what follows, I briefly discuss the intuition, leaving the formal arguments to Appendix A.4.

Let us call a firm quality class $\mathcal{F}_n^k$ *occupied* if the hiring capacities at quality classes that are better than $\mathcal{F}_n^k$ do not absorb all the workers in the market; that is,

$$\lim_{n \rightarrow \infty} \frac{\sum_{i=1}^{k-1} |\mathcal{F}_n^i|}{|\mathcal{W}_n|} < 1.$$

Otherwise, we say that the firm quality class $\mathcal{F}_n^k$ is *vacant*. Below, I discuss how the claim is established for occupied firms. The claim for vacant firms follows by constructing self-enforcing matching processes where the occupied firms hire less than what their capacities allow, and the excess workers are allocated to vacant firms to increase their payoffs.

In large matching markets, triggering with static stable matchings is not enough to deter firm deviations: It is well known that as the market size grows large, firms may obtain efficient payoffs from even the worst stable matching (Pittel (1989), Lee (2016), Ashlagi, Kanoria, and Leshno (2017)), so they cannot serve as effective deterrents. Therefore, the first step in proving Theorem 4 is to find a “punishment matching” that (i) reduces the deviating firm’s expected payoff even when market size is large, and (ii) makes sure that no worker is willing to block this punishment with the deviating firm. Such matchings essentially play the role of minmax action profiles in standard repeated games. I show that the stage-game matchings produced by certain variants of the worker-proposing serial dictatorship satisfy both requirements.

While the worker-proposing serial dictatorship establishes the payoff lower bound $\underline{v}$, the upper bound $\hat{v}$ is achieved by the firm-proposing random serial dictatorship. I build on existing results from Che and Tercieux (2018) to show that the gap between $\hat{v}$ and $\underline{v}$ is nonvanishing as the market size grows to infinity. Finally, I use the classical ideas from Fudenberg and Maskin (1986) to construct self-enforcing matching processes that sustain payoffs between $\underline{v}$ and $\hat{v}$.

Section 3.3 below provides a detailed illustration of these serial dictatorship matching algorithms in the example of a one-to-one matching market; it also illustrates how randomizations over stage-game matchings can be used on-path to generate the correct target payoff for self-enforcing matching processes.

3.3 An example for one-to-one matching

Theorem 3 and Theorem 4 show that elite firms, while untouchable, have no real impact on aggregate allocations. The goal of this example is to provide insights into the scope of the sustainable payoffs for non-elite firms. To this end, we study a simple setting with only one class of firms $\mathcal{F}_n$ and one class of workers $\mathcal{W}_n$. Note that in this case $\mathcal{F}_n$ is not an elite class since $|\mathcal{F}_n|/n = 1$. Matching is one-to-one, and there are an equal number of firms and workers (so $q = \beta = 1$). For tractability, we will also assume that the value of a worker to a firm is uniformly distributed on $[0, 1]$ and vice versa.

For each discount factor δ and market size n , let $\mathcal{E}_{\delta,n}$ denote the set of (firms’) payoff profiles that can be sustained by self-enforcing matching processes. In what follows, I first study the set $\mathcal{E}_{\delta,n}$ under a fixed market size n while letting $\delta \rightarrow 1$. I then provide bounds on the firms’ minmax payoff, and use these bounds to characterize $\mathcal{E}_{\delta,n}$ when both $n \rightarrow \infty$ and $\delta \rightarrow 1$.

The market of fixed size n For each fixed market size n , I first construct a set of payoff profiles that can be sustained by self-enforcing matching processes with patient firms. As we shall see later, this set will be crucial to characterizing the behavior of $\mathcal{E}_{\delta,n}$ in large markets.

Let $\widehat{m}_n$ be the matching produced by the firm-proposing random serial dictatorship. By symmetry, all firms obtain the same payoff from $\widehat{m}_n$, so we will denote $\widehat{u}_n \equiv u_f(\widehat{m}_n)$ for all $f \in \mathcal{F}_n$. For each $F \subseteq \mathcal{F}_n$, let $\widehat{m}_n^{-F}$ denote the matching that is identical to $\widehat{m}_n$, except that all the firms in F and their workers are unmatched from each other (so $\widehat{m}_n^{-F}(f) = \emptyset$ for all $f \in F$). Firms' payoff profiles from these matchings satisfy $\{u(\widehat{m}_n^{-F}) : F \subseteq \mathcal{F}_n\} = \{0, \widehat{u}_n\}^n$, so the payoff profiles generated by $\{\widehat{m}_n^{-F} : F \subseteq \mathcal{F}_n\}$ span all the vertices of an n -dimensional cube. It follows that the randomizations over $\{\widehat{m}_n^{-F} : F \subseteq \mathcal{F}_n\}$ span the entire n -dimensional cube; in other words,

$$\{u(\lambda) : \lambda \in \Delta\{\widehat{m}_n^{-F} : F \subseteq \mathcal{F}_n\}\} = [0, \widehat{u}_n]^n.$$

So for each market size n , the set of feasible stage-game payoff profiles is a superset of $[0, \widehat{u}_n]^n$.

Recall that for each firm f , its stage-game minmax payoff is

$$\underline{u}_n = \mathbb{E} \left[\min_{m \in M_n^\circ} \max_{w \in D_f(m)} u_f(w) \right],$$

where M_n° is the set of stage-game recommendations that are acceptable to all workers; $D_f(m) \equiv m(f) \cup \{w \in \mathcal{W}_n : f \succ_w m(w)\}$ are the workers who are either already working for f or can be poached by f ; last, the expectation is taken over all preference realizations.

The following result can be proved using standard constructions in repeated games with perfect monitoring (see, for example, [Fudenberg and Maskin \(1986\)](#) and [Fudenberg, Kreps, and Maskin \(1990\)](#)).

PROPOSITION 1. *For each fixed n , $(\underline{u}_n, \widehat{u}_n)^n \subseteq \liminf_{\delta \rightarrow 1} \mathcal{E}_{\delta, n} \subseteq \limsup_{\delta \rightarrow 1} \mathcal{E}_{\delta, n} \subseteq (\underline{u}_n, 1]^n$.*

In light of Proposition 1, the set of sustainable payoffs in large markets depends crucially on the behavior of $\underline{u}_n$ and $\widehat{u}_n$ when $n \rightarrow \infty$. Theorem 1 in [Che and Tercieux \(2018\)](#) implies that

$$\widehat{u}_n \rightarrow 1 \quad \text{as } n \rightarrow \infty.$$

Although the precise value of the minmax payoff $\underline{u}_n$ is difficult to compute, I provide bounds on its value by calculating firms' payoffs from two stage-game matchings.

To establish an upper bound for $\underline{u}_n$, note that the “max” operation in the definition of $\underline{u}_n$ implies that to reduce f 's payoff, it is without loss to focus on stage-game matchings that leave f with no feasible and profitable deviations; the “min” operation then selects the worst recommendation for f among these matchings. The upper bound is obtained by constructing a stage-game matching that leaves f with no feasible and profitable deviations, even though this particular matching may not be the worst possible for f .

As for the lower bound for $\underline{u}_n$, note that when f contemplates a deviation in the max operation, the only way to credibly remove a worker from f 's choice set $D_f(\cdot)$ is to assign her to a firm that she prefers over f . However, this is constrained by the capacities available at other firms. We can, therefore, obtain a lower bound for $\underline{u}_n$ by considering a fictitious matching where all other firms face no capacity constraints, so the recommendation in the min operation can remove all workers from f 's choice set, except for those who rank f as their first choice.

An upper bound for minmax payoff To obtain an upper bound for the minmax value $\underline{u}_n$, let us consider the following algorithm designed to punish an arbitrary firm f ; we will use $\underline{m}_n^f$ to denote the stage-game matching produced by this algorithm.

The algorithm is a variant of the worker-proposing serial dictatorship. It first assigns priorities to workers according to f 's preference ranking and then runs the worker-proposing serial dictatorship based on these priorities. Note that as workers propose and exit the market with their matched firms, the algorithm moves down f 's preference list. How high f ranks its matched worker therefore boils down to how early it is picked by a worker.

From the workers' perspective, a new firm is sampled each round without replacement. Since workers' preferences for firms are uniformly random, the number of draws it takes for f to be sampled is uniformly distributed from 1 to n . Note that if f is sampled at the i th draw, then f obtains the $(n + 1 - i)$ th order statistic among n i.i.d. uniform random variables over $[0, 1]$; in particular, i is itself uniformly distributed over $\{1, \dots, n\}$. It is well known that the expected value of this $(n + 1 - i)$ th order statistic is $\frac{n+1-i}{n+1}$, so the expected payoff of firm f is

$$\mathbb{E}[u_f(\underline{m}_n^f)] = \frac{1}{n(n+1)}(1 + 2 + \dots + n) = \frac{1}{2}.$$

Moreover, f can find no feasible and profitable deviations from $\underline{m}_n^f$. To see why, note that there are two possibilities for a worker who did not end up matching with f in $\underline{m}_n^f$: (i) she had higher priority than $\underline{m}_n^f(f)$ and is, therefore, more desirable to f than $\underline{m}_n^f(f)$, but left the market with a firm that is more desirable to her than f , or (ii) she had lower priority than $\underline{m}_n^f(f)$, so she is less desirable to f than $\underline{m}_n^f(f)$. In either case, f is unable to find a better worker who is willing to replace $\underline{m}_n^f(f)$. This implies that $u_f(\underline{m}_n^f) \geq u_f(w)$ for all $w \in D_f(\underline{m}_n^f)$, and as a result we have

$$\underline{u}_n = \mathbb{E}\left[\min_{m \in M_n^\circ} \max_{w \in D_f(m)} u_f(w)\right] \leq \mathbb{E}\left[\max_{w \in D_f(\underline{m}_n^f)} u_f(w)\right] = \mathbb{E}[u_f(\underline{m}_n^f)] = \frac{1}{2}.$$

In summary, we have the following result that provides an upper bound for $\underline{u}_n$.

PROPOSITION 2. *The minmax payoff $\underline{u}_n$ satisfies $\underline{u}_n \leq \mathbb{E}[u_f(\underline{m}_n^f)] = \frac{1}{2}$ for all n .*

A lower bound for minmax payoff To obtain a lower bound on the minmax value $\underline{u}_n$, take an arbitrary firm f , and note that when f contemplates a deviation from any stage-game matching $m \in M_n^\circ$, it can always attract the workers who rank f as their first choice. Let $\tilde{W}_n^f \equiv \{w \in \mathcal{W}_n : f \succeq_w f' \text{ for all } f' \in \mathcal{F}_n\}$ denote these workers; then we have $\tilde{W}_n^f \subseteq D_f(m)$ for all $m \in M_n^\circ$. It follows that

$$\underline{u}_n = \mathbb{E}\left[\min_{m \in M_n^\circ} \max_{w \in D_f(m)} u_f(w)\right] \geq \mathbb{E}\left[\max_{w \in \tilde{W}_n^f} u_f(w)\right].$$

The value $\mathbb{E}[\max_{w \in \tilde{W}_n^f} u_f(w)]$ can be computed through a fictitious stage-game matching¹² $\underline{\psi}_n^f$ produced by the following variant of the worker-proposing serial dictatorship. The algorithm is identical to the one that produced $\underline{m}_n^f$, except that now all firms other than f have their capacities enlarged to $q = \infty$. In particular, the algorithm assigns priorities to workers according to f 's preference ranking and then runs the worker-proposing serial dictatorship based on these priorities. Just as in the algorithm that produced $\underline{m}_n^f$, how high f ranks its matched worker depends on how early it is picked by a worker. The difference from $\underline{m}_n^f$ is that from the workers' perspective, all firms other than f are now being sampled *with* replacement due to their infinite capacities, so f will only be sampled by a worker who ranks it as her first choice. Therefore, we have

$$\underline{u}_n \geq \mathbb{E}\left[\max_{w \in \tilde{W}_n^f} u_f(w)\right] = \mathbb{E}[u_f(\underline{\psi}_n^f)]. \quad (3)$$

Since the workers' preferences for firms are uniformly random, the number of draws it takes for f to be sampled, which I denote by i , follows a truncated geometric distribution with success rate $1/n$. In particular, i is a random variable with support $\{1, \dots, n\} \cup \{\infty\}$, where ∞ represents the event that f is never sampled. Its probability, mass function is given by

$$P(i = l) = \begin{cases} \frac{1}{n} \left(1 - \frac{1}{n}\right)^{l-1} & \text{for } l = 1, \dots, n \\ \left(1 - \frac{1}{n}\right)^n & \text{for } l = \infty. \end{cases}$$

In addition, if f is sampled at the i th draw, then it obtains the $(n+1-i)$ th order statistic among n i.i.d. uniform random variables over $[0, 1]$, so firm f 's expected payoff from $\underline{\psi}_n^f$ satisfies

$$\mathbb{E}[u_f(\underline{\psi}_n^f)] = \sum_{i=1}^n \frac{1}{n} \left(1 - \frac{1}{n}\right)^{i-1} \left(\frac{n+1-i}{n+1}\right) + 0 \left(1 - \frac{1}{n}\right)^n.$$

It is well known that as $n \rightarrow \infty$, the distribution of $\frac{i}{n+1}$ converges to a truncated exponential distribution with arrival rate 1, so the expression above can be approximated by

$$\mathbb{E}[u_f(\underline{\psi}_n^f)] \rightarrow \int_0^1 e^{-x}(1-x) dx = \frac{1}{e} \quad \text{as } n \rightarrow \infty. \quad (4)$$

Combining (3) and (4), we have the following result that provides a lower bound on $\underline{u}_n$.

PROPOSITION 3. *The minmax payoff $\underline{u}_n$ satisfies $\underline{u}_n \geq \mathbb{E}[u_f(\underline{\psi}_n^f)]$ for all n . Furthermore, $\mathbb{E}[u_f(\underline{\psi}_n^f)] \rightarrow \frac{1}{e}$ as $n \rightarrow \infty$.*

¹²We emphasize the word “fictitious” because, strictly speaking, $\underline{\psi}_n^f$ does not satisfy the definition of stage-game matching as it violates capacity constraints.

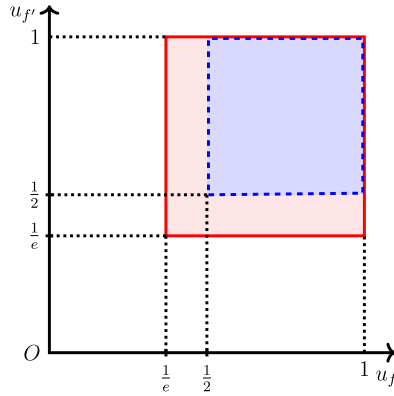


FIGURE 5. An illustration of Proposition 4 when $F = \{f, f'\}$. The limit payoff sets are supersets of the blue square (without borders), and subsets of the red square (with borders).

Limit in large markets The following result is a direct consequence of Propositions 1, 2 and 3. In particular, for every fixed and finite subset of firms F , Proposition 4 characterizes the joint payoffs among firms in F as firms become patient and the market size grows to infinity.

PROPOSITION 4. *Let $F \subseteq \mathcal{F}_n$ be a finite set of firms, and let $\mathcal{E}_{\delta,n}^F$ denote the projection of $\mathcal{E}_{\delta,n}$ onto the payoff space of F . Then we have¹³*

$$\left(\frac{1}{2}, 1\right)^F \subseteq \liminf_{n \rightarrow \infty, \delta \rightarrow 1} \mathcal{E}_{\delta,n}^F \subseteq \limsup_{n \rightarrow \infty, \delta \rightarrow 1} \mathcal{E}_{\delta,n}^F \subseteq \left[\frac{1}{e}, 1\right]^F.$$

See Figure 5 for an illustration of the inner and outer bounds of the limit payoff sets. Note that the interval $[\frac{1}{e}, 1]^F$ is closed on the left since it is possible that $\underline{u}_n < \frac{1}{e}$ for all market size n , which means $\{\frac{1}{e}\}^F$ is a sustainable payoff profile for all n as long as δ is sufficiently high.

The next result, which is a special case of Theorem 4, speaks directly to the payoff space of all firms in large markets.

PROPOSITION 5. *For every payoff interval $[\underline{v}, \widehat{v}] \subseteq (\frac{1}{2}, 1)$, there exist a discount factor $0 < \underline{\delta} < 1$ and market size $N \geq 0$ such that for all $\delta > \underline{\delta}$, $n > N$, and every payoff profile $\mathbf{v}^n \in [\underline{v}, \widehat{v}]^n$, there exists a self-enforcing matching process μ_n^* that satisfies $U_f(\mu_n^*) = \mathbf{v}_f$ for all $f \in \mathcal{F}_n$.*

Note that Proposition 5 does not imply the existence of $\underline{\delta}$ and N such that all payoff profiles in $(\frac{1}{2}, 1)^n$ are sustainable for $\delta > \underline{\delta}$ and $n > N$. To see why, suppose that the

¹³The limits in Proposition 4 are defined as

$$\liminf_{n \rightarrow \infty, \delta \rightarrow 1} \mathcal{E}_{\delta,n}^F \equiv \bigcup_{N \geq 1, \underline{\delta} > 0} \bigcap_{n \geq N, \delta \geq \underline{\delta}} \mathcal{E}_{\delta,n}^F \quad \text{and} \quad \limsup_{n \rightarrow \infty, \delta \rightarrow 1} \mathcal{E}_{\delta,n}^F \equiv \bigcap_{N \geq 1, \underline{\delta} > 0} \bigcup_{n \geq N, \delta \geq \underline{\delta}} \mathcal{E}_{\delta,n}^F.$$

minmax payoff is $\underline{u}_n = \frac{1}{2}$ (so the upper bound derived in Proposition 2 is tight), and consider the payoff profile $\tilde{\mathbf{v}}^n = (\frac{3}{4}, \frac{5}{8}, \dots, \frac{1}{2} + \frac{1}{4n}) \in (\frac{1}{2}, 1)^n$ for all n . There may not be a pair of $\underline{\delta}$ and N such that $\tilde{\mathbf{v}}^n$ is sustainable for all $\delta > \underline{\delta}$ and $n > N$. The problem is that the n th firm's payoff $\frac{1}{2} + \frac{1}{4n} \rightarrow \frac{1}{2}$ as $n \rightarrow \infty$, and the δ required to sustain a payoff may increase as it approaches the minmax, in which case the lower bound $\underline{\delta}$ cannot be chosen independently of the market size n .

4. CONCLUSION

This paper provides a framework and solution concept for studying stability in repeated matching markets that combines elements from repeated noncooperative games with the cooperative stability notion for two-sided matching markets. While history-dependent play can be used to alter the matches obtained by some firms, some other firms are immune to such influences due to the unique payoff structure in matching environments. To understand the impact of such firms, I consider large matching markets with correlated preferences. I find that in large matching markets, these elite firms make up at most a negligible fraction of the market.

It is useful to compare the results in this paper with the rural hospital theorem (Roth (1986)), which states that any firm that does not fill its quota at some static stable matching is assigned precisely the same set of workers at every stable matching. In my setting, when patience is 0, both the top-coalition firms and “rural firms” (i.e., firms that do not fill their quotas at some static stable matching) retain the same matched workers across all static stable matchings. However, this happens due to different reasons for rural firms and top-coalition firms. The rural firms have this feature because they are under-demanded, and they (by definition) have vacancies; the top-coalition firms have this feature because they are over-demanded and typically have no vacancies in stable matchings unless they find fewer workers acceptable than their hiring quotas. When patience goes to 1, only the top-coalition firms’ workers must remain unchanged regardless of the patience level; by contrast, we are able to change all other firms’ workers no matter whether they are rural firms in the static setting or not.

An interesting open question for future research is in understanding whether matching processes like simple exclusion can be approximately self-enforcing sense in large matching markets.

APPENDIX

A.1 Preliminary lemmas

PROOF OF LEMMA 1. At each history h , firm f faces a decision regarding whether and how to deviate with workers who are currently in the market. Suppose firm f has a deviation plan d_f from matching process μ that is both feasible and profitable. Since stage-game payoffs are bounded for firm f and there is discounting, the standard one-shot deviation principle for individual decision making (Blackwell (1965)) implies that there exists a history $\hat{h} \in \mathcal{H}$ such that

$$(1 - \delta)\tilde{u}_f(d_f(\hat{h})) + \delta U_f(\hat{h}, [\mu(\hat{h}), (f, d_f(\hat{h}))]|\mu) > U_f(\hat{h}|\mu).$$

Consider the deviation plan d_f^o that satisfies $d_f^o(h) = d_f(h)$ if $h = \hat{h}$ and $d_f^o(h) = \mu(f|h)$ otherwise. The plan d_f^o is a profitable one-shot deviation plan for firm f . $\square$

Lemma 2 shows that if a firm deviates from the recommended matching with a group of workers, then this deviating firm can be uniquely identified by comparing the resulting matching against the recommended matching.

LEMMA 2. *Let m be a static matching. If $[m, (f, W)] = [m, (f', W')] \neq m$, then $f = f'$.*

PROOF. Let $\bar{m} \equiv [m, (f, W)]$ and $\hat{m} \equiv [m, (f', W')]$. Suppose by contradiction that $f \neq f'$, but $\hat{m} = \bar{m} \neq m$. There are three cases to consider:

(i) If $W \subseteq m(f)$, then $\bar{m}(f') = m(f') \neq \hat{m}(f')$, so $\bar{m} \neq \hat{m}$, a contradiction.

(ii) If $W \not\subseteq m(f)$, then $\bar{m}(f) \not\subseteq m(f)$, but $\hat{m}(f) \subseteq m(f)$, so $\bar{m} \neq \hat{m}$, a contradiction.

Therefore, $[m, (f, W)] = [m, (f', W')] \neq m$ implies $f = f'$. $\square$

A.2 Proof of Theorem 1

Suppose $\mathcal{T} = \{(\hat{f}_k, \hat{W}_k)\}_{k=1}^K$. In both one-shot and repeated matching environments, the proof proceeds by induction.

Static stable matching In every static stable matching, $\hat{f}_1$ and $\hat{W}_1$ must be matched together since they are mutual favorites. Suppose $\hat{f}_i$ and $\hat{W}_i$ must be matched together in all stable matchings for $1 \leq i \leq k-1$, but suppose by contradiction that there is a stable matching m where $\hat{f}_k$ and $\hat{W}_k$ are not matched together. By the induction hypothesis, $m(\hat{f}_k) \subseteq \mathcal{W} \setminus \bigcup_{i=1}^{k-1} \hat{W}_i$, so $\tilde{u}_{\hat{f}_k}(\hat{W}_k) > \tilde{u}_{\hat{f}_k}(m(\hat{f}_k))$, and $m(w) \in \mathcal{F} \setminus \{\hat{f}_i : 1 \leq i \leq k-1\}$, so $\hat{f}_k \succ_w m(w)$ for all $w \in \hat{W}_k \setminus m(\hat{f}_k)$. This is a contradiction to m being a stable matching, since $\hat{f}_k$ and $\hat{W}_k$ find it profitable to jointly deviate. So $\hat{f}_k$ and $\hat{W}_k$ must be matched together in all stable matchings. This completes the induction step.

Self-enforcing matching process The proof again proceeds by induction. First, I prove that in every self-enforcing matching process μ , $\mu(\hat{f}_1|h) = \hat{W}_1$ for all $h \in \mathcal{H}^F$. Suppose by contradiction that $\mu(\hat{f}_1|\tilde{h}) \neq \hat{W}_1$ at some $\tilde{h} \in \mathcal{H}$. Consider the deviation plan d_1 defined by $d_1(h) = \hat{W}_1$ for all $h \in \mathcal{H}$. Plan d_1 is clearly feasible for $\hat{f}_1$ since $\hat{f}_1$ is every worker's favorite firm. In addition,

$$\begin{aligned} U_{\hat{f}_1}(\tilde{h}|\mu) &= (1 - \delta)\tilde{u}_{\hat{f}_1}(\mu(\hat{f}_1|\tilde{h})) + \delta U_{\hat{f}_1}(\tilde{h}, \mu(\tilde{h})|\mu) \\ &< (1 - \delta)\tilde{u}_{\hat{f}_1}(\hat{W}_1) + \delta u_{\hat{f}_1}(\hat{W}_1) = U_{\hat{f}_1}(\tilde{h}[\mu, (\hat{f}_1, d_1)]), \end{aligned}$$

so d_1 is also profitable for $\hat{f}_1$. This is a contradiction to μ being self-enforcing. So $\mu(\hat{f}_1|h) = \hat{W}_1$ for all $h \in \mathcal{H}$.

Suppose it has been shown that in every self-enforcing matching process μ , $\mu(\hat{f}_i|h) = \hat{W}_i$ for $i = 1, \dots, k-1$ at every ex post history $h \in \mathcal{H}$, and suppose by contradiction that there is a self-enforcing matching process μ such that $\mu(\hat{f}_k|\tilde{h}) \neq \hat{W}_k$ for some $\tilde{h} \in \mathcal{H}$. By

the inductive hypothesis, $\mu(\widehat{f}_k|h) \subseteq \mathcal{W} \setminus \bigcup_{i=1}^{k-1} \widehat{W}_i$ at all $h \in \mathcal{H}$, so $\widetilde{u}_{\widehat{f}_k}(\mu(\widehat{f}_k|\widetilde{h})) < \widetilde{u}_{\widehat{f}_k}(\widehat{W}_k)$, and $U_{\widehat{f}_k}(\widetilde{h}|\mu) \leq \widetilde{u}_{\widehat{f}_k}(\widehat{W}_k)$.

Consider the deviation plan d_k defined by $d_k(h) = \widehat{W}_k$ for all $h \in \mathcal{H}$. Plan d_k is feasible for $\widehat{f}_k$ since for every worker in $\mathcal{W} \setminus \bigcup_{i=1}^{k-1} \widehat{W}_i$, $\widehat{f}_k$ is the best firm among $\mathcal{F} \setminus \{f_i\}_{i=1}^{k-1}$. In addition

$$\begin{aligned} U_{\widehat{f}_k}(\widetilde{h}|\mu) &= (1 - \delta)\widetilde{u}_{\widehat{f}_k}(\mu(\widehat{f}_k|\widetilde{h})) + \delta U_{\widehat{f}_k}(\widetilde{h}, \mu(\widetilde{h})|\mu) \\ &< (1 - \delta)\widetilde{u}_{\widehat{f}_k}(\widehat{W}_k) + \delta \widetilde{u}_{\widehat{f}_k}(\widehat{W}_k) = U_{\widehat{f}_k}(\widetilde{h}|\mu, (\widehat{f}_k, d)), \end{aligned}$$

so d_k is both feasible and profitable, contradicting the assumption that μ is self-enforcing. So in every self-enforcing matching process μ , $\mu(\widehat{f}_k|h) = \widehat{W}_k$ at every ex post history $h \in \mathcal{H}$. This completes the induction.

A.3 Proof of Theorem 3

A.3.1 Preliminaries I first establish a few preliminary results so as to prove Theorem 3. Lemma 3 proves that when market size is sufficiently large, an elite firm f is very likely able to fill its positions with workers who (i) rank f as their favorite firm and (ii) give f close to the highest possible stage-game utility.

For each firm $f \in \mathcal{F}_n^1$ and $r > 0$, let $\widehat{W}_n^1(f, \epsilon) \equiv \{w \in \mathcal{W}_n^1 : f \succeq_w f' \text{ for all } f' \in \mathcal{F}_n^1 \text{ and } u_f(w) > V(C_1, 1) - r\}$.

LEMMA 3. *Suppose $|\mathcal{F}_n^1|/|\mathcal{W}_n^1| \rightarrow 0$. In the stage game, for every $r > 0$, there exists N such that $P(|\widehat{W}(f, r, k)| > q) > 1 - r$ for all $n > N$ and every $f \in \mathcal{F}_n^1$.*

PROOF. Let $\underline{\zeta} \in [0, 1)$ be a number such that $V(C_1, \underline{\zeta}) > V(C_1, 1) - r$. Define $\widetilde{W}_n^1(f, r) \equiv \{w \in \mathcal{W}_n^1 : f \succeq_w f' \text{ for all } f' \in \mathcal{F}_n^1 \text{ and } \zeta_{f,w} > \underline{\zeta}\}$, so $\widetilde{W}_n^1(f, r) \subseteq \widehat{W}_n^1(f, r)$.

From the perspective of firm $f \in \mathcal{F}_n^1$, every worker $w \in \mathcal{W}_n^1$ satisfies $\zeta_{f,w} > \underline{\zeta}$ with probability $1 - \underline{\zeta} > 0$. Furthermore, a worker is equally likely to rank any firm within $\mathcal{F}_n^1$ as her top choice within $\mathcal{F}_n^1$. Let $Y_w(f, r)$ be the Bernoulli random variable that takes value 1 if $w \in \widetilde{W}_n^1(f, r)$ and 0 otherwise. Let $\phi \equiv (1 - \underline{\zeta})/|\mathcal{F}_n^1|$. Note that the random variables $\{Y_w(f, r) : w \in \mathcal{W}_n^1\}$ are independently and identically distributed with rate ϕ , and as a result, $|\widetilde{W}_n^1(f, r)| = \sum_{w \in \mathcal{W}_n^1} Y_w(f, r)$ follows binomial distribution $B(|\mathcal{W}_n^1|, \phi)$. By the Chernoff bound,

$$P(|\widetilde{W}_n^1(f, r)| \leq q) \leq \min_{t>0} e^{tq} \prod_{w \in \mathcal{W}_n^1} \mathbb{E}[e^{-tY_w(f, r)}].$$

Since $\mathbb{E}[e^{-tY_w(f, r)}] = 1 + \phi(e^{-t} - 1) \leq e^{\phi(e^{-t}-1)}$ for all $w \in \mathcal{W}_n^1$, we have

$$P(|\widetilde{W}_n^1(f, r)| \leq q) \leq \min_{t>0} e^{tq} \cdot e^{|\mathcal{W}_n^1|\phi(e^{-t}-1)} \leq e^q \cdot e^{|\mathcal{W}_n^1|\phi(e^{-1}-1)},$$

where the second inequality above follows from setting $t = 1$. Since $|\mathcal{W}_n^1|\phi = (1 - \underline{\zeta})|\mathcal{W}_n^1|/|\mathcal{F}_n^1| \rightarrow \infty$ and $e^{-1} - 1 < 0$, we have $P(|\widetilde{W}_n^1(f, r)| \leq q) \rightarrow 0$ as $n \rightarrow \infty$. Finally, since $\widetilde{W}_n^1(f, r) \subseteq \widehat{W}_n^1(f, r)$, it follows that $P(|\widehat{W}_n^1(f, r)| \leq q) \rightarrow 0$ as $n \rightarrow \infty$. $\square$

Lemma 4 shows that an elite firm can secure a high continuation value in large markets.

LEMMA 4. *Suppose $|\mathcal{F}_n^1|/|\mathcal{W}_n^1| \rightarrow 0$. For every $r > 0$ and discount factor $0 < \delta < 1$, there exists N such that for all $n > N$, the continuation value of every $f \in \mathcal{F}_n^1$ in every self-enforcing matching process μ satisfies $U_f(\bar{h}|\mu) > qV(C_1, 1) - r$ at every ex ante history $\bar{h}$.*

PROOF. Suppose that the stage game satisfies $P(|\hat{\mathcal{W}}_n^1(f, \tilde{r})| > q) > 1 - \tilde{r}$ for some $\tilde{r}$, and consider the following deviation plan d_f by an arbitrary firm $f \in \mathcal{F}_n^1$: at every ex post history, given the realized preferences in the current stage game,

$$d_f(h) = \arg \max_{W \subseteq \hat{\mathcal{W}}_n^1(f, \tilde{r}), |W| \leq q} \sum_{w \in W} u_f(w).$$

I first show that with sufficiently small $\tilde{r}$, firm f can guarantee itself a strictly higher payoff than $qV(C_1, 1) - r$ by using d_f to deviate from any matching process. The claim of the lemma then follows by invoking Lemma 3.

Fix an arbitrary matching process μ . Note first that since every worker in $\hat{\mathcal{W}}_n^1(f, \tilde{r})$ ranks f as their favorite firm, d_f is a feasible deviation plan for f by construction.

To see that d_f guarantees $qV(C_1, 1) - r$ for sufficiently small $\tilde{r}$, let T be large enough so that $\delta^T V(C_1, 1)q < r/2$. At every ex ante history $\bar{h}$, firm f 's continuation payoff from the manipulated matching process $[\mu, (f, d_f)]$ satisfies

$$\begin{aligned} U_f(\bar{h}[\mu, (f, d_f)]) &\geq (1 - \delta^T) \{ (1 - \tilde{r})^T q(V(C_1, 1) - \tilde{r}) + [1 - (1 - \tilde{r})^T] 0 \} + \delta^T \cdot 0 \\ &= (1 - \delta^T)(1 - \tilde{r})^T q(V(C_1, 1) - \tilde{r}). \end{aligned} \quad (5)$$

Inequality (5) above follows from decomposing f 's continuation payoff between those accrued within the first T and those after period T ; for the first T periods, the expected payoff is further decomposed by whether $|\hat{\mathcal{W}}_n^1(f, \tilde{r})| \geq q$ is satisfied all the way through this phase or not.

Since $(1 - \delta^T)(1 - \tilde{r})^T q(V(C_1, 1) - \tilde{r}) \rightarrow (1 - \delta^T)qV(C_1, 1)$ as $\tilde{r} \rightarrow 0$, we can find $\underline{r}$ such that $(1 - \delta^T)(1 - \underline{r})^T q(V(C_1, 1) - \underline{r}) > (1 - \delta^T)qV(C_1, 1) - r/2$. By Lemma 3, there exists N such that $P(|\hat{\mathcal{W}}_n^1(f, \underline{r})| > q) > 1 - \underline{r}$ for all $n > N$. So for all $n > N$, we have

$$\begin{aligned} U_f(\bar{h}[\mu, (f, d_f)]) &> (1 - \delta^T)qV(C_1, 1) - r/2 \\ &= (1 - \delta^T)qV(C_1, 1) + \delta^T qV(C_1, 1) - r/2 - \delta^T qV(C_1, 1) \\ &> (1 - \delta^T)qV(C_1, 1) + \delta^T qV(C_1, 1) - r/2 - r/2 \\ &= qV(C_1, 1) - r. \end{aligned}$$

Finally, if μ is a self-enforcing matching process, then it must satisfy

$$U_f(\bar{h}|\mu) \geq U_f(\bar{h}[\mu, (f, d_f)]) > qV(C_1, 1) - r$$

at every ex ante history $\bar{h}$. □

A.3.2 Proof of Theorem 3 By setting $r = (1 - \delta)\epsilon$ in Lemma 4, we know there exists N_1 such that if $n > N_1$,

$$U_f(\bar{h}|\mu) > qV(C_1, 1) - (1 - \delta)\epsilon \quad (6)$$

for every self-enforcing matching process μ at every ex ante history $\bar{h}$.

Let μ be a self-enforcing matching process, and suppose by contradiction that

$$\mathbb{E}[u_f(\mu(\hat{h}, s_n))] \leq qV(C_1, 1) - \epsilon$$

at some ex ante history $\hat{h}$. It follows then that at $\hat{h}$, firm f 's continuation payoff satisfies

$$\begin{aligned} U_f(\hat{h}|\mu) &\leq (1 - \delta)\mathbb{E}[u_f(\mu(\hat{h}, s_n))] + \delta qV(C_1, 1) \\ &\leq (1 - \delta)[qV(C_1, 1) - \epsilon] + \delta qV(C_1, 1) \\ &\leq qV(C_1, 1) - (1 - \delta)\epsilon, \end{aligned}$$

which is a contradiction to (6).

A.4 Proof of Theorem 4

The proof of Theorem 4 is divided into three parts. The first part, Appendix A.4.1, focuses on the “submarket” faced by an occupied firm quality class $\mathcal{F}_n^k$: this is a matching market that comprises only $\mathcal{F}_n^k$ and the workers who these firms would obtain in a static stable matching. Within this submarket, I show how one can construct stage-game matchings that reward firms for cooperation, as well as stage-game matchings that can be used to punish deviating firms. The second part, Appendix A.4.2, returns to the matching market at large, and uses the insights from submarkets to prove Theorem 4 for occupied firm quality classes. Finally, Appendix A.4.3, proves Theorem 4 for vacant quality classes.

A.4.1 Submarkets For each $1 \leq l \leq L$, let $Q_n^{\mathcal{W}}(l) \equiv \sum_{l' \leq l} |\mathcal{W}_n^{l'}|$ denote the number of workers who are in a quality class no worse than $\mathcal{W}_n^l$; similarly, for each $1 \leq k \leq K$, let $Q_n^{\mathcal{F}}(k) \equiv q \sum_{k' \leq k} |\mathcal{F}_n^{k'}|$ denote the number of firm seats that are in a quality class no worse than $\mathcal{F}_n^k$.

I say that a worker quality class $\mathcal{W}_n^l$ is achievable by firms in quality class $\mathcal{F}_n^k$ if

$$\lim_{n \rightarrow \infty} \frac{Q_n^{\mathcal{W}}(l)}{Q_n^{\mathcal{F}}(k-1)} \geq 1 \quad \text{and} \quad \lim_{n \rightarrow \infty} \frac{Q_n^{\mathcal{W}}(l-1)}{Q_n^{\mathcal{F}}(k)} \leq 1.$$

The first inequality above ensures that not all workers in $\mathcal{W}_n^l$ can be absorbed by firms in higher quality classes than $\mathcal{F}_n^k$; the second inequality ensures that not all seats in $\mathcal{F}_n^k$ can be filled by workers in higher quality classes than $\mathcal{W}_n^l$. Let $\mathcal{A}(k) \subseteq \{1, \dots, L\}$ denote the set of worker quality classes achievable by firm quality class k .

In the analysis in this section, when market size is sufficiently large, the only relevant workers are those who are in an achievable quality class. I therefore focus on the submarket that consists of only firms in $\mathcal{F}_n^k$ and workers in its achievable quality classes.

Let the numbers $\alpha_1, \alpha_2, \dots, \alpha_I \geq 0$ be numbers that satisfy $\sum_{i=1}^{I-1} \alpha_i < 1$ and $\sum_{i=1}^I \alpha_i \geq 1$. Consider a sequence of submarkets: for every n , the firm side is made up of $\mathcal{F}_n^k$, while the worker side consists of $\bigcup_{i=1}^I \mathcal{V}_n^i$, where for each i ,

$$\mathcal{V}_n^i \subseteq \mathcal{W}_n^i \quad \text{and} \quad \frac{|\mathcal{V}_n^i|}{|\mathcal{F}_n^k|q} \rightarrow \alpha_i \quad \text{as } n \rightarrow \infty.$$

I treat each seat on the firm side of the market as an individual player who inherits the preference of its firm. For each seat s , let $\tilde{u}_s(\cdot)$ denote the utility function of the seat, which is identical to $\tilde{u}_f(\cdot)$ of the firm that it belongs to.

Reward matching Let $\hat{\phi}_n$ denote the matching resulting from the seat-proposing random serial dictatorship in the submarket. Matching $\hat{\phi}_n$ is played as a reward for the firms when they comply with capacity reduction. Lemma 6 characterizes the payoff from this reward. Noting that the matching $\hat{\phi}_n$ is Pareto efficient, I will make use of the following result from Che and Tercieux (2018).

LEMMA 5 (Che and Tercieux (2018)). As $n \rightarrow \infty$,

$$\frac{\sum_s u_s(\hat{\phi}_n(s))}{|\mathcal{F}_n^k|q} \xrightarrow{p} \sum_{i=1}^{I-1} \alpha_i V(C_i, 1) + \left(1 - \sum_{i=1}^{I-1} \alpha_i\right) V(C_I, 1).$$

Lemma 6 shows that in $\hat{\phi}_n$, every firm obtains a randomly assigned common value from its matched workers, but obtains a close to maximum idiosyncratic component from each worker.

LEMMA 6. For every $\epsilon > 0$, there exists N such that

$$\mathbb{E}[u_f(\hat{\phi}_n)] > q \left[\sum_{i=1}^{I-1} \alpha_i V(C_i, 1) + \left(1 - \sum_{i=1}^{I-1} \alpha_i\right) V(C_I, 1) \right] - \epsilon$$

for all market sizes $n > N$ and every $f \in \mathcal{F}_n^k$.

PROOF. To prove the claim, it suffices to prove that $\mathbb{E}[u_s(\hat{\phi}_n)] > \sum_{i=1}^{I-1} \alpha_i V(C_i, 1) + (1 - \sum_{i=1}^{I-1} \alpha_i) V(C_I, 1) - \epsilon$ for all $n > N$ and every individual seat s . This is what I will prove below.

Let

$$\bar{u}_n \equiv \frac{\sum_s u_s(\hat{\phi}_n(s))}{|\mathcal{F}_n^k|q}$$

denote the average realized utilities for individual seats, and let $\bar{F}_n$ denote the probability distribution of this random variable $\bar{u}_n$. Let $A_n^-(\epsilon) \equiv \{\bar{u}_n \leq \sum_{i=1}^{I-1} \alpha_i V(C_i, 1) + (1 - \sum_{i=1}^{I-1} \alpha_i) V(C_I, 1) - \frac{1}{2}\epsilon\}$ be the event that $\bar{u}$ is less than the payoff upper bound, and let $A_n^+(\epsilon) \equiv \{\bar{u}_n > \sum_{i=1}^{I-1} \alpha_i V(C_i, 1) + (1 - \sum_{i=1}^{I-1} \alpha_i) V(C_I, 1) - \frac{1}{2}\epsilon\}$ be its complement.

Note that by the law of total expectation, for every seat s ,

$$\mathbb{E}[u_s(\hat{\phi}_n)] = \int \mathbb{E}[u_s(\hat{\phi}_n)|\bar{u}_n] d\bar{F}_n = \int_{A_n^-(\epsilon)} \mathbb{E}[u_s(\hat{\phi}_n)|\bar{u}_n] d\bar{F}_n + \int_{A_n^+(\epsilon)} \mathbb{E}[u_s(\hat{\phi}_n)|\bar{u}_n] d\bar{F}_n.$$

Since seats are treated symmetrically in a random serial dictatorship, it follows that $\mathbb{E}[u_s(\hat{\phi}_n)|\bar{u}_n] = \bar{u}_n$. So for every seat s ,

$$\begin{aligned} \mathbb{E}[u_s(\hat{\phi}_n)] &= \int_{A_n^-(\epsilon)} \bar{u}_n d\bar{F}_n + \int_{A_n^+(\epsilon)} \bar{u}_n d\bar{F}_n \\ &\geq P(A_n^-(\epsilon)) \cdot 0 + P(A_n^+(\epsilon)) \left[\sum_{i=1}^{I-1} \alpha_i V(C_i, 1) + \left(1 - \sum_{i=1}^{I-1} \alpha_i\right) V(C_I, 1) - \frac{1}{2}\epsilon \right]. \end{aligned}$$

By Lemma 5, $P(A_n^+(\epsilon)) \rightarrow 1$ as $n \rightarrow \infty$. As a result, there exists N such that $\mathbb{E}[u_s(\hat{\phi}_n)] \geq \sum_{i=1}^{I-1} \alpha_i V(C_i, 1) + (1 - \sum_{i=1}^{I-1} \alpha_i) V(C_I, 1) - \epsilon$ for all $n > N$ and $s \in \mathcal{F}_n^k$. $\square$

Minmax matching

DEFINITION 4. Given a set of firms $\mathcal{F}'$, workers $\mathcal{W}'$, and a firm $f \in \mathcal{F}'$, the punitive matching for f , ϕ_n^f , is the matching produced by the following procedure:

Step 1. Set $S_0 = \emptyset$ and $G_0 = \emptyset$.

Step 2. If either $\mathcal{F}' \setminus S_{k-1} = \emptyset$ or $\mathcal{W}' \setminus G_{k-1} = \emptyset$, stop; otherwise, let $\hat{w}$ be f 's favorite worker in $\mathcal{W}' \setminus G_{k-1}$ and let $\hat{f}$ be $\hat{w}$'s favorite firm among $\mathcal{F}' \setminus S_{k-1}$. Match $\hat{f}$ and $\hat{w}$, and set $S_k = S_{k-1} \cup \{\hat{f}\}$ and $G_k = G_{k-1} \cup \{\hat{w}\}$. Go to Step $k+1$.

For a seat s belonging to a firm f , let $\underline{R}_n^s$ denote f 's ranking of the worker matched to seat s in the matching ϕ_n^f : that is, $\underline{R}_n^s = j$ if $\phi_n^f(s)$ is the j th favorite worker according to the preference of f . The next lemma shows that $\underline{R}_n^s$ is uniformly distributed.

LEMMA 7. For every seat s belonging to f and $1 \leq r \leq |\mathcal{F}_n^k|q$, $P(\underline{R}_n^s = r) = \frac{1}{|\mathcal{F}_n^k|q}$.

PROOF. In the procedure in Definition 4, starting from f 's favorite worker, each worker takes a turn to pick a seat from the remaining seats. As we move down f 's preference list, the first worker to pick seat s will determine $\underline{R}_n^s$.

We can equivalently think of $\underline{R}_n^s$ as being determined through sampling *without* replacement from an urn that contains a red ball (representing seat s) and $(|\mathcal{F}_n^k|q - 1)$ black balls (representing all other seats). Specifically, $\underline{R}_n^s$ is the number of draws it takes for the red ball to be sampled. Since the red ball is equally likely to be in any position in the sequence of balls to be drawn, $P(\underline{R}_n^s = r) = \frac{1}{|\mathcal{F}_n^k|q}$ for $1 \leq r \leq |\mathcal{F}_n^k|q$. $\square$

Lemma 8 shows that if a firm is excluded from the very top section of its preference list, then with high probability it obtains a low utility.

LEMMA 8. Fix any $0 < \epsilon < 1$ and $0 < \gamma < \epsilon$. For every $f \in \mathcal{F}_n^k$, let X_n^f denote the event that, according to f 's preference, the worst $\gamma|\mathcal{V}_n^I|$ workers in $\mathcal{V}_n^I$ all satisfy $\zeta_{f,w} < \epsilon$. Then $P(X_n^f) \rightarrow 1$ as $n \rightarrow \infty$.

PROOF. For each $f \in \mathcal{F}_n^k$, let K_n^f be the number of workers in $\mathcal{V}_n^I$ such that $\zeta_{f,w} < \epsilon$. I first prove that $P(K_n^f < \gamma|\mathcal{V}_n^I|) \rightarrow 0$ as $n \rightarrow \infty$.

Since $\zeta_{f,w}$ is uniformly distributed over $[0, 1]$, K_n^f follows binomial distribution $B(|\mathcal{V}_n^I|, \epsilon)$. By Hoeffding's inequality, for every $f \in \mathcal{F}_n^k$,

$$P(K_n^f < \gamma|\mathcal{V}_n^I|) \leq \frac{1}{2} \exp \left\{ -2 \frac{(\epsilon|\mathcal{V}_n^I| - \gamma|\mathcal{V}_n^I|)^2}{|\mathcal{V}_n^I|} \right\} \rightarrow 0$$

as $n \rightarrow \infty$. The claim of the lemma follows since $\{K_n^f \geq \gamma|\mathcal{V}_n^I|\} \subseteq X_n^f$. $\square$

Lemma 9 proves that as market gets large, the payoff f obtains from the punishment algorithm in Definition 4 is bounded away from what it obtains from the reward matching $\hat{\phi}_n$.

LEMMA 9. There exists $g > 0$ and N such that

$$\mathbb{E}[u_f(\underline{\phi}_n^f)] < q \left[\sum_{i=1}^{I-1} \alpha_i V(C_i, 1) + \left(1 - \sum_{i=1}^{I-1} \alpha_i \right) V(C_I, 1) \right] - g$$

for all $f \in \mathcal{F}_n^k$ and $n > N$.

PROOF. To prove the claim, it suffices to prove that there exists $g > 0$ and N such that if $n > N$, then

$$\mathbb{E}[u_s(\underline{\phi}_n^f)] < \sum_{i=1}^{I-1} \alpha_i V(C_i, 1) + \left(1 - \sum_{i=1}^{I-1} \alpha_i \right) V(C_I, 1) - g$$

for every firm $f \in \mathcal{F}_n^k$ and seat s belonging to f . This is what I will prove below.

For every $f \in \mathcal{F}_n^k$, every seat s belonging to f , and $1 \leq i \leq I$, let

$$A_n^s(i) \equiv \left\{ \sum_{j \leq i-1} |\mathcal{V}_n^j| < R_n^s \leq \sum_{j \leq i} |\mathcal{V}_n^j| \right\}$$

denote the event that s is matched to a worker in $\mathcal{V}_n^i$ in the matching $\underline{\phi}_n^f$. By Lemma 7, R_n^s is uniformly distributed for every s and f , so $P(A_n^s(i))$ is independent of f or s , and I will write $p_n^i \equiv P(A_n^s(i))$ for every $1 \leq i \leq I$.

Fix $(\sum_{i=1}^I \alpha_i - 1)/\alpha_I < \epsilon < 1$ and $(\sum_{i=1}^I \alpha_i - 1)/\alpha_I < \gamma < \epsilon$.¹⁴ Since the function $V(C_I, \cdot)$ is strictly increasing, there exists $g_I > 0$ such that $V(C_I, \zeta) < V(C_I, 1) - g_I$ for

¹⁴Note that $\alpha_I > 0$ and $0 < (\sum_{i=1}^I \alpha_i - 1)/\alpha_I = 1 - (1 - \sum_{i=1}^{I-1} \alpha_i)/\alpha_I < 1$. So such ϵ always exists.

all $\zeta < \epsilon$. For every seat s belonging to any $f \in \mathcal{F}_n^k$, let

$$\Gamma_n^s \equiv \left\{ \sum_{i=1}^{I-1} |\mathcal{V}_n^i| + (1 - \gamma) |\mathcal{V}_n^I| \leq \underline{R}_n \leq \sum_{i=1}^I |\mathcal{V}_n^i| \right\}$$

denote the event that s is matched to the γ -tail section of $\mathcal{V}_n^I$ in the matching ϕ_n^f . Again, by Lemma 7, $\underline{R}_n^s$ is uniformly distributed, so $P(\Gamma_n^s)$ is independent of f or s . I will write $p_n^\Gamma \equiv P(\Gamma_n^s)$.

For each firm $f \in \mathcal{F}_n^k$, let X_n^f denote the event that, according to any firm f 's preference, workers in the γ -tail section of $\mathcal{V}_n^I$ all satisfy $V(C_I, \zeta_{f,w}) < V(C_I, 1) - g_I$. Since all firms in $\mathcal{F}_n^k$ are symmetric, the probability $P(X_n^f)$ does not depend on f , and I will simply write $p_n^X \equiv P(X_n^f)$.

By the definition of the events Γ_n^s and X_n^f , we have, for every $f \in \mathcal{F}_n^k$ and seat s belonging to f ,

$$\begin{aligned} P(u_s(\phi_n^f) < V(C_I, 1) - g_I | A_n^s(I)) &\geq P(\Gamma_n^s \cap X_n^f | A_n^s(I)) \\ &= P(\Gamma_n^s \cap X_n^f \cap A_n^s(I)) / P(A_n^s(I)) \\ &= P(\Gamma_n^s \cap X_n^f) / P(A_n^s(I)) \\ &= [p_n^\Gamma - P(\Gamma_n^s \cap (X_n^f)^c)] / p_n^I \\ &\geq [p_n^\Gamma - (1 - p_n^X)] / p_n^I. \end{aligned} \quad (7)$$

Note that the second expression above follows since $\Gamma_n^s \subseteq A_n^s(I)$.

As $n \rightarrow \infty$, $p_n^I \rightarrow 1 - \sum_{i=1}^{I-1} \alpha_i > 0$ and $p_n^\Gamma \rightarrow 1 - \sum_{i=1}^{I-1} \alpha_i - (1 - \gamma)\alpha_I > 0$. Importantly, these limits are all strictly positive numbers.¹⁵ In addition, by Lemma 8, $p_n^X \rightarrow 1$ as $n \rightarrow \infty$. So inequality (7) implies that there exists N_I such that if $n \geq N_I$,

$$P(u_s(\phi_n^f) < V(C_I, 1) - g_I | A_n^s(I)) \geq \frac{p_n^\Gamma}{2p_n^I}$$

for all $f \in \mathcal{F}_n^k$ and seat s belonging to f , and, therefore,

$$\mathbb{E}(u_s(\phi_n^f) | A_n^s(I)) < V(C_I, 1) - g_I \frac{p_n^\Gamma}{2p_n^I}. \quad (8)$$

Meanwhile, for $1 \leq i \leq I - 1$, $p_n^i \rightarrow \alpha_i > 0$ and

$$\mathbb{E}(u_s(\phi_n^f) | A_n^s(i)) \leq V(C_i, 1). \quad (9)$$

¹⁵To see why, note that by our choice of γ , $p_n^\Gamma = 1 - \sum_{i=1}^{I-1} \alpha_i - (1 - \gamma)\alpha_I > 1 - \sum_{i=1}^{I-1} \alpha_i - [1 - (\sum_{i=1}^I \alpha_i - 1)/\alpha_I]\alpha_I = 0$.

Define $g \equiv (1 - \sum_{i=1}^{I-1} \alpha_i) g_I \frac{p_n^\Gamma}{2p_n^I} > 0$. Combining inequalities (8) and (9), there must exist N such that for all $n > N$,

$$\begin{aligned} \mathbb{E}(u_s(\phi_n^f)) &= \sum_{1 \leq i \leq m} p_n^i \mathbb{E}(u_s(\phi_n^f) | A_n^s(i)) \\ &< \sum_{i=1}^{I-1} \alpha_i V(C_i, 1) + \left(1 - \sum_{i=1}^{I-1} \alpha_i\right) V(C_I, 1) - g \end{aligned}$$

for all $f \in \mathcal{F}_n^k$ and seat s belonging to f . This completes the proof. $\square$

Finally, Lemma 10 proves that when a firm f is being punished, it cannot find any profitable deviations with workers.

LEMMA 10. *For every $f \in \mathcal{F}_n^k$, no coalition in the form of (f, W) can be a profitable deviation from the matching ϕ_n^f .*

PROOF. Suppose there is a profitable coalition (f, W) for some W . Let $w \in W \setminus \phi_n^f(f)$ be any new worker in W who is not originally matched to f in ϕ_n^f . I will show that $u_f(w') > u_f(w)$ for all $w' \in \phi_n^f(f)$. This will be a contradiction to (f, W) being a profitable coalition, because in this case any new worker in W is worse than the worker in $\phi_n^f(f)$ she replaced. It is then impossible for f to prefer W over $\phi_n^f(f)$ since firms' preferences are responsive.

To this end, first observe that if there exists $w' \in \phi_n^f(f)$ such that $u_f(w') < u_f(w)$, then it means that in Step A.4.2 in the punishment algorithm in Definition 4, w chose $\phi_n^f(w)$ over f when f still had vacancy (which would later be filled by w') so $\phi_n^f(w) \succ_w f$, and w does not find this deviation profitable. This is a contradiction. So $u_f(w') > u_f(w)$ for all $w' \in \phi_n^f(f)$. This completes the proof. $\square$

A.4.2 Proof of Theorem 4 for occupied quality classes Recall that $Q_n^{\mathcal{W}}(l) \equiv \sum_{l' \leq l} |\mathcal{W}_n^{l'}|$ denotes the number of workers who are in a quality class no worse than $\mathcal{W}_n^l$, while $Q_n^{\mathcal{F}}(k) \equiv q \sum_{k' \leq k} |\mathcal{F}_n^{k'}|$ denotes the number of firm seats that are in a quality class no worse than $\mathcal{F}_n^k$. Let $\mathcal{A}(k) = \{j+1, \dots, j+I\}$ be the set of worker quality classes achievable by firms in quality class k . By definition, there exists N_0 such that for all $n > N_0$ and every $l \notin \mathcal{A}(k)$,

$$\text{either } Q_n^{\mathcal{W}}(l) < Q_n^{\mathcal{F}}(k-1) \quad \text{or} \quad Q_n^{\mathcal{W}}(l-1) > Q_n^{\mathcal{F}}(k). \quad (10)$$

I focus on market sizes greater than this N_0 . By the inequalities in (10), in all the matchings defined below, workers in any unachievable quality class by $\mathcal{F}_n^k$ are either already matched to firms in a higher quality class or ranked too low to be relevant for $\mathcal{F}_n^k$. Therefore, to analyze the payoff for a firm $f \in \mathcal{F}_n^k$, it suffices to isolate $\mathcal{F}_n^k$ and $\bigcup_{l \in \mathcal{A}(k)} \mathcal{W}_n^l$ as if they are the only firms and workers in the market. For each $l \in \mathcal{A}(k)$, let

$$\alpha_l \equiv \lim_{n \rightarrow \infty} \left\{ \frac{Q_n^{\mathcal{W}}(l) - \max\{Q_n^{\mathcal{W}}(l-1), Q_n^{\mathcal{F}}(k-1)\}}{q |\mathcal{F}_n^k|} \right\}$$

denote the asymptotic share of seats in $\mathcal{F}_n^k$ filled by workers in $\mathcal{W}_n^l$. Since $|\mathcal{W}_n| = \lceil \beta n q \rceil$, we know all α_{lj} are finite numbers. Note that they also satisfy $\sum_{i=1}^{I-1} \alpha_{j+i} < 1$ and $\sum_{i=1}^I \alpha_{j+i} \geq 1$.

Minmax matching For each $f \in \mathcal{F}_n^k$, consider the matching $\underline{m}_n^f$ defined by the following procedure:

- Step 1. Let $\phi_n^<$ be a stable matching between $\bigcup_{j < k} \mathcal{F}_n^j$ and $\mathcal{W}_n$. Set $\underline{m}_n^f(f') = \phi_n^<(f')$ for every $f' \in \bigcup_{j < k} \mathcal{F}_n^j$.
- Step 2. Let ϕ_n^f be the punitive matching for f among $\mathcal{F}_n^k$ and $\mathcal{W}_n \setminus \phi_n^<(\bigcup_{j < k} \mathcal{F}_n^j)$. Set $\underline{m}_n^f(f') = \phi_n^f(f')$ for every $f' \in \mathcal{F}_n^k$.
- Step 3. Let $\phi_n^>$ be a stable matching between $\bigcup_{j > k} \mathcal{F}_n^j$ and $\mathcal{W}_n \setminus [\phi_n^<(\bigcup_{j < k} \mathcal{F}_n^j) \cup \phi_n^f(\mathcal{F}_n^k)]$. Set $\underline{m}_n^f(f') = \phi_n^>(f')$ for every $f' \in \bigcup_{j > k} \mathcal{F}_n^j$.

By Lemma 9, there exists $g > 0$ and market size N_1 such that

$$0 \leq \mathbb{E}[u_f(\underline{m}_n^f)] < \sum_{i=1}^{I-1} \alpha_{j+i} V(C_{j+i}, 1) + \left(1 - \sum_{i=1}^{I-1} \alpha_{j+i}\right) V(C_{j+I}, 1) - 2g \quad (11)$$

for all $f \in \mathcal{F}_n^k$ and $n > N_1$.

Reward matching Consider the matching $\widehat{m}_n$ defined by the following procedure:

- Step 1. Let $\phi_n^<$ be a stable matching between $\bigcup_{j < k} \mathcal{F}_n^j$ and $\mathcal{W}_n$. Set $\widehat{m}_n(f) = \phi_n^<(f)$ for every $f \in \bigcup_{j < k} \mathcal{F}_n^j$.
- Step 2. Let $\widehat{\phi}_n$ be the matching resulting from seat-proposing random serial dictatorship between $\mathcal{F}_n^k$ and $\mathcal{W}_n \setminus \phi_n^<(\bigcup_{j < k} \mathcal{F}_n^j)$. Set $\widehat{m}_n(f) = \widehat{\phi}_n(f)$ for all $f \in \mathcal{F}_n^k$.
- Step 3. Let $\phi_n^>$ be a stable matching between $\bigcup_{j > k} \mathcal{F}_n^j$ and $\mathcal{W}_n \setminus [\phi_n^<(\bigcup_{j < k} \mathcal{F}_n^j) \cup \widehat{\phi}_n(\mathcal{F}_n^k)]$. Set $\widehat{m}_n(f) = \phi_n^>(f)$ for every $f \in \bigcup_{j > k} \mathcal{F}_n^j$.

By Lemma 6, there exists market size N_2 such that

$$\mathbb{E}[u_f(\widehat{m}_n)] > \sum_{i=1}^{I-1} \alpha_{j+i} V(C_{j+i}, 1) + \left(1 - \sum_{i=1}^{I-1} \alpha_{j+i}\right) V(C_{j+I}, 1) - g \quad (12)$$

for all $f \in \mathcal{F}_n^k$ and $n > N_2$.

I focus on market sizes greater than $N \equiv \max\{N_0, N_1, N_2\}$. In particular, combining (11) and (12) yields

$$\mathbb{E}[u_f(\widehat{m}_n)] > \mathbb{E}[u_f(\underline{m}_n^f)] + g \quad (13)$$

for all $n > N$ and all $f \in \mathcal{F}_n^k$.

Payoff interval In this section we construct the payoff interval $(\underline{v}, \widehat{v})$ mentioned in the statement of Theorem 4 as well as the randomizations over stage-game matchings that achieve the payoffs therein.

For each subset of firms $F \subseteq \mathcal{F}_n^k$, we construct a matching $\widehat{m}_n^F$ through the procedure outlined below. In particular, $\widehat{m}_n^F$ is constructed using the same procedure as the reward matching $\widehat{m}_n$ except that no workers are matched to firms in F .

- Step 1. Let $\phi_n^<$ be a stable matching between $\bigcup_{j < k} \mathcal{F}_n^j$ and $\mathcal{W}_n$. Set $\widehat{m}_n^F(f) = \phi_n^<(f)$ for every $f \in \bigcup_{j < k} \mathcal{F}_n^j$.
- Step 2. Let $\widehat{\phi}_n$ be the matching resulting from seat-proposing random serial dictatorship between $\mathcal{F}_n^k$ and $\mathcal{W}_n \setminus \phi_n^<(\bigcup_{j < k} \mathcal{F}_n^j)$. Let $\widehat{\phi}_n^F(f) = \widehat{\phi}_n(f)$ for all $f \notin F$, but $\widehat{\phi}_n^F(f) = \emptyset$ for all $f \in F$. Set $\widehat{m}_n^F(f) = \widehat{\phi}_n^F(f)$ for all $f \in \mathcal{F}_n^k$.
- Step 3. Let $\phi_n^>$ be a stable matching between $\bigcup_{j > k} \mathcal{F}_n^j$ and $\widehat{\phi}_n(F) \cup \mathcal{W}_n \setminus [\phi_n^<(\bigcup_{j < k} \mathcal{F}_n^j) \cup \widehat{\phi}_n(\mathcal{F}_n^k)]$. Set $\widehat{m}_n(f) = \phi_n^>(f)$ for every $f \in \bigcup_{j > k} \mathcal{F}_n^j$.

Note that for each firm $f \in \mathcal{F}_n^k$, we have

$$\mathbb{E}[u_f(\widehat{m}_n^F)] = \begin{cases} 0 & \text{for all } f \in F \\ \mathbb{E}[u_f(\widehat{m}_n)] & \text{for all } f \notin F, \end{cases}$$

so together, the deterministic matchings $\{\widehat{m}_n^F : F \subseteq \mathcal{F}_n^k\}$ span all payoff vectors in $[0, \widehat{u}]^{\mathcal{F}_n^k}$, while randomizations over $\{\widehat{m}_n^F : F \subseteq \mathcal{F}_n^k\}$ span $[0, \widehat{u}]^{\mathcal{F}_n^k}$. Let us set $\widehat{v} \equiv \widehat{u}$ and $\underline{v} = \underline{u} + \frac{2}{3}g$. We will show that all payoff vectors in $(\underline{v}, \widehat{v})^{\mathcal{F}_n^k} \subseteq [0, \widehat{u}]^{\mathcal{F}_n^k}$ can be sustained when patience is sufficiently high.

Firm-specific punishments Recall that randomizations over $\{\widehat{m}_n^F : F \subseteq \mathcal{F}_n^k\}$ span $[0, \widehat{u}]^{\mathcal{F}_n^k}$. For each firm $f \in \mathcal{F}_n^k$, let $\lambda_n^f \in \Delta(\{\widehat{m}_n^F : F \subseteq \mathcal{F}_n^k\})$ be the random matching that satisfies

$$\mathbb{E}[u_{f'}(\lambda_n^f)] = \begin{cases} \underline{u} + \frac{1}{3}g & \text{if } f' = f \\ \underline{u} + \frac{2}{3}g & \text{if } f' \neq f. \end{cases}$$

Fix any payoff vector $v \in (\underline{v}, \widehat{v})^{\mathcal{F}_n^k}$, and let $\lambda^0 \in \Delta(\{\widehat{m}_n^F : F \subseteq \mathcal{F}_n^k\})$ be the random matching that satisfies $\mathbb{E}[u_f(\lambda_n^0)] = v_f$ for all $f \in \mathcal{F}_n^k$. The random matchings $\{\lambda_n^f : f \in \mathcal{F}_n^k\}$ form a system of player-specific punishments for v . In particular,

$$\mathbb{E}[u_f(\lambda_n^f)] \leq v_f - \frac{1}{3}g \quad (14)$$

$$\mathbb{E}[u_f(\lambda_n^f)] \leq \mathbb{E}[u_f(\lambda_n^{f'})] - \frac{1}{3}g \quad \text{for all } f \neq f'. \quad (15)$$

Also note that by construction,

$$\mathbb{E}[u_f(\lambda_n^f)] \geq \underline{u} + \frac{1}{3}g \quad \text{for all } f. \quad (16)$$

Note that, in particular, the payoff gap is guaranteed to be at least $1/3g$ regardless of the choice of target payoff vector $v \in (\underline{v}, \widehat{v})^{\mathcal{F}_n^k}$.

Self-enforcing matching process In what follows, I will focus on verifying the incentives firms in $\mathcal{F}_n^k$, since the firms in other tiers will have no incentive to deviate by construction. Consider the matching process represented by the automaton $(\Theta, \gamma^0, O, \kappa)$, where the following statements hold:

- (i) The equality $\Theta_n = \{\theta(e, m) : e \in \mathcal{F}_n^k \cup \{0\}; m \in M_n\} \cup \{\underline{\theta}(f, t) : f \in \mathcal{F}_n^k, 0 \leq t < L\}$ is the set of all possible states.
- (ii) The variable γ^0 is the initial distribution over states, which satisfies $\gamma^0(\theta(0, m)) = \lambda^0(m)$ for all $m \in M_n$.
- (iii) The function $O : \Theta \rightarrow M_n$ is the output function, where $O(\theta(f, m)) = m$ and $O(\underline{\theta}(f, t)) = \underline{m}_n^f$.
- (iv) The function $\kappa : \Theta \times M_n \rightarrow \Delta(\Theta)$ is the transition function defined as follows.
For states $\{\underline{\theta}(f, t) | 0 \leq t < L - 1\}$, κ is defined as

$$\begin{aligned} & \kappa(\underline{\theta}(f, t), m') \\ &= \begin{cases} \underline{\theta}(f', 0) & \text{if } m' \neq \underline{m}_n^f; m' = [\underline{m}_n^f, (f', W)] \text{ for some } f' \in \mathcal{F}_n^k \text{ and } W \subseteq \mathcal{W}_n \\ \underline{\theta}(f, t+1) & \text{otherwise.} \end{cases} \end{aligned}$$

For states $\underline{\theta}(f, L - 1)$, the transition is defined as

$$\begin{aligned} & \kappa(\underline{\theta}(f, L - 1), m') \\ &= \begin{cases} \underline{\theta}(f', 0) & \text{if } m' \neq \underline{m}_n^f; m' = [\underline{m}_n^f, (f', W)] \text{ for some } f' \in \mathcal{F}_n^k \text{ and } W \subseteq \mathcal{W}_n \\ \gamma^f & \text{otherwise,} \end{cases} \end{aligned}$$

where γ^f is the distribution over states that satisfies $\gamma^f(\theta(f, m)) = \lambda^f(m)$ for all $f \in \mathcal{F}_n^k$ and $m \in M_n$.

For states $\theta(e, m)$, the transition is

$$\begin{aligned} & \kappa(\theta(e, m), m') \\ &= \begin{cases} \underline{\theta}(f', 0) & \text{if } m' \neq m; m' = [m, (f', W)] \text{ for some } f' \in \mathcal{F}_n^k \text{ and } W \subseteq \mathcal{W}_n \\ \gamma^e & \text{otherwise,} \end{cases} \end{aligned}$$

where γ^f and γ^0 are defined as above.

The matching process represented by the above automaton randomizes over M_n according to λ^0 in every period. It remains to verify that no firm has profitable one-shot deviations in any automaton state.

For states of the form $\theta(e, m)$ Consider a one-shot deviation (f', W) by firm f' . There are two cases to consider.

Case 1: $f' \neq e$. Choose a number $Z > qV(C_1, 1)$, so no firm can derive payoff higher than Z from any deviation. Without deviation, f' has value $(1 - \delta)u_{f'}(m) + \delta\mathbb{E}[u_{f'}(\lambda_n^e)]$. After deviation, f' yields less than $(1 - \delta)Z + \delta(1 - \delta^L)\mathbb{E}[u_{f'}(\underline{m}_n^{f'})] + \delta^{L+1}\mathbb{E}[u_{f'}(\lambda_n^{f'})]$. There is no profitable one-shot deviation for f' if

$$(1 - \delta)u_{f'}(m) + \delta\mathbb{E}[u_{f'}(\lambda_n^e)] \geq (1 - \delta)Z + \delta(1 - \delta^L)\mathbb{E}[u_{f'}(\underline{m}_n^{f'})] + \delta^{L+1}\mathbb{E}[u_{f'}(\lambda_n^{f'})].$$

A sufficient condition for the inequality above is

$$(1 - \delta)0 + \delta\mathbb{E}[u_{f'}(\lambda_n^e)] \geq (1 - \delta)Z + \delta(1 - \delta^L)\mathbb{E}[u_{f'}(\lambda_n^{f'})] + \delta^{L+1}\mathbb{E}[u_{f'}(\lambda_n^{f'})],$$

which is equivalent to

$$\delta[\mathbb{E}u_{f'}(\lambda_n^e) - \mathbb{E}u_{f'}(\lambda_n^{f'})] \geq (1 - \delta)Z.$$

By (14) and (15), $\mathbb{E}[u_{f'}(\lambda_n^e)] - \mathbb{E}[u_{f'}(\lambda_n^{f'})] \geq \frac{1}{3}g$ for all $n > N$ and $f' \in \mathcal{F}_n^k$. Let $\underline{\delta}_1$ be sufficiently high such that the left-hand side (LHS) is greater than the right-hand side (RHS). Note that the inequality above does not depend on the choice of target payoff v . It follows that for all target payoffs $v \in (\underline{v}, \widehat{v})^{\mathcal{F}_n^k}$, these deviations are not profitable as long as $\delta > \underline{\delta}_1$ and $n > N$.

Case 2: $f' = e$. Without deviation, f' has value $(1 - \delta)u_{f'}(m) + \delta\mathbb{E}[u_{f'}(\lambda_n^{f'})]$. After deviation, f' yields less than $(1 - \delta)Z + \delta(1 - \delta^L)\mathbb{E}[u_{f'}(\underline{m}_n^{f'})] + \delta^{L+1}\mathbb{E}[u_{f'}(\lambda_n^{f'})]$. There is no profitable one-shot deviation for f' if

$$(1 - \delta)u_{f'}(m) + \delta\mathbb{E}[u_{f'}(\lambda_n^{f'})] \geq (1 - \delta)Z + \delta(1 - \delta^L)\mathbb{E}[u_{f'}(\underline{m}_n^{f'})] + \delta^{L+1}\mathbb{E}[u_{f'}(\lambda_n^{f'})].$$

The inequality is equivalent to

$$Z - u_{f'}(m) \leq \delta(1 + \dots + \delta^{L-1})[\mathbb{E}[u_{f'}(\lambda_n^{f'})] - \mathbb{E}[u_{f'}(\underline{m}_n^{f'})]].$$

Since $u_{f'}(m) \geq 0$ and, by (16), $\mathbb{E}[u_{f'}(\lambda_n^{f'})] - \mathbb{E}[u_{f'}(\underline{m}_n^{f'})] \geq \frac{1}{3}g$ for all $f' \in \mathcal{F}_n^k$ and $n > N$, a sufficient condition for the inequality above is

$$Z \leq \delta(1 + \dots + \delta^{L-1})\frac{1}{3}g.$$

Choose L large enough so that $\frac{1}{3}Lg > Z$. As $\delta \rightarrow 1$, the LHS remains unchanged while the RHS converges to $\frac{1}{3}Lg$, so there exists $\underline{\delta}_2$ such that no deviation is profitable for $\delta > \underline{\delta}_2$ and $n > N$. Note again that the choice of $\underline{\delta}_2$ does not depend on the target payoff v , so no such deviations are profitable regardless of $v \in (\underline{v}, \widehat{v})^{\mathcal{F}_n^k}$.

For states of the form $\theta(f, t)$ There are two cases to consider.

Case 1: $f' \neq f$. Without deviation, f' has payoff $(1 - \delta^{L-t})\mathbb{E}[u_{f'}(\underline{m}_n^f)] + \delta^{L-t} \times \mathbb{E}[u_{f'}(\lambda_n^f)]$. With any deviation, f' has payoff less than $(1 - \delta)Z + \delta(1 - \delta^L)\mathbb{E}[u_{f'}(\underline{m}_n^{f'})] +$

$\delta^{L+1}\mathbb{E}[u_{f'}(\lambda_n^{f'})]$. There is no profitable one-shot deviation for f' if

$$(1 - \delta^{L-t})\mathbb{E}[u_{f'}(\underline{m}_n^f)] + \delta^{L-t}\mathbb{E}[u_{f'}(\lambda_n^f)] \geq (1 - \delta)Z + \delta(1 - \delta^L)\mathbb{E}[u_{f'}(\underline{m}_n^{f'})] \\ + \delta^{L+1}\mathbb{E}[u_{f'}(\lambda_n^{f'})].$$

Note that $\mathbb{E}[u_{f'}(\underline{m}_n^f)] \geq 0$ and $\mathbb{E}[u_{f'}(\underline{m}_n^{f'})] \leq \mathbb{E}[u_{f'}(\lambda_n^{f'})]$; in addition, the inequality above is most demanding when $t = 0$. So a sufficient condition for the inequality above is

$$\delta^L\mathbb{E}[u_{f'}(\lambda_n^f)] \geq (1 - \delta)Z + \delta\mathbb{E}[u_{f'}(\lambda_n^{f'})].$$

As $\delta \rightarrow 1$, the LHS converges to $\mathbb{E}[u_{f'}(\lambda_n^f)]$ for all t such that $0 \leq t \leq L$, while the RHS converges to $\mathbb{E}[u_{f'}(\lambda_n^{f'})]$. By (15), $\mathbb{E}[u_{f'}(\lambda_n^f)] > \mathbb{E}[u_{f'}(\lambda_n^{f'})] + \frac{1}{3}g$ for all $f' \neq f$ and $n > N$. So there exists $\underline{\delta}_3$ such that for all $\delta > \underline{\delta}_3$ and $n > N$, there is no profitable deviation. Once again, note that the choice of $\underline{\delta}_3$ is independent of the target payoff v .

Case 2: $f' = f$. Without deviation, firm f' has payoff

$$(1 - \delta^{L-t})\mathbb{E}[u_{f'}(\underline{m}_n^{f'})] + \delta^{L-t}\mathbb{E}[u_{f'}(\lambda_n^{f'})].$$

When deviating from $\underline{m}_{k'}$, by Lemma 10, f' 's stage-game payoff is at most $\mathbb{E}[u_{f'}(\underline{m}_n^{f'})]$. So f' 's discounted expected payoff from deviation is at most

$$(1 - \delta)\mathbb{E}[u_{f'}(\underline{m}_n^{f'})] + \delta(1 - \delta^L)\mathbb{E}[u_{f'}(\underline{m}_n^{f'})] + \delta^{L+1}\mathbb{E}[u_{f'}(\lambda_n^{f'})] \\ = (1 - \delta^{L+1})\mathbb{E}[u_{f'}(\underline{m}_n^{f'})] + \delta^{L+1}\mathbb{E}[u_{f'}(\lambda_n^{f'})].$$

Firm f' has no profitable deviation if

$$(1 - \delta^{L-t})\mathbb{E}[u_{f'}(\underline{m}_n^{f'})] + \delta^{L-t}\mathbb{E}[u_{f'}(\lambda_n^{f'})] \geq (1 - \delta^{L+1})\mathbb{E}[u_{f'}(\underline{m}_n^{f'})] + \delta^{L+1}\mathbb{E}[u_{f'}(\lambda_n^{f'})]$$

or

$$\mathbb{E}[u_{f'}(\lambda_n^{f'})] > \mathbb{E}[u_{f'}(\underline{m}_n^{f'})],$$

which is true for all $f' \in \mathcal{F}_n^k$, $n > N$, and regardless of δ . So no firm has profitable one-shot deviations of this kind for all $n > N$ and all δ . To sustain the payoff interval

Define $\underline{\delta} \equiv \max\{\delta_1, \delta_2, \delta_3\}$. For every target payoff $v \in (\underline{v}, \widehat{v})$, there is no profitable one-shot deviation in any states of the automaton as long as $\delta > \underline{\delta}$ and $n > N$. This completes the proof.

A.4.3 Proof of Theorem 4 for vacant quality classes

Reward matching Let $\{\mathcal{F}_n^k : k = 1, \dots, J\}$ be the occupied quality classes and let $\{\mathcal{F}_n^k : k = J + 1, \dots, K\}$ be the vacant quality classes.

Consider the matching $\widehat{x}_n$ defined by the following procedure:

Step 1. Let $\widehat{\phi}_n^1$ be the matching resulting from seat-proposing random serial dictatorship between $\mathcal{F}_n^1$ and $\mathcal{W}_n$. Set $\widehat{x}_n(f) = \widehat{\phi}_n^1(f)$ for all $f \in \mathcal{F}_n^1$.

Step 2. Let $\hat{\phi}_n^k$ be the matching resulting from seat-proposing random serial dictatorship between $\mathcal{F}_n^k$ and $\mathcal{W}_n \setminus \phi_n^< (\bigcup_{j < k} \mathcal{F}_n^j)$. Set $\hat{x}_n(f) = \hat{\phi}_n^k(f)$ for all $f \in \mathcal{F}_n^k$.

For each occupied quality class $\mathcal{F}_n^k$, let $[\underline{v}_k, \hat{v}_k]$ be the payoff interval in Theorem 4 corresponding to $\mathcal{F}_n^k$. Note that there exists an N such that $U_f(\hat{x}_n) \geq \hat{v}_k$ for all $k = 1, \dots, J$ and $f \in \mathcal{F}_n^k$. Since $\hat{v}_k > \underline{v}_k$ for each $1 \leq k \leq J$, there exists $p^x \in (0, 1)$ such that $(\underline{v}_k + \hat{v}_k)/2 \leq p^x \hat{v}_k + (1 - p^x)0$ for each $1 \leq k \leq J$.

Let $\underline{\mu}^k$ be the corresponding self-enforcing matching process that satisfies $U_f(\underline{\mu}^k) = \underline{v}_k$ for all $f \in \mathcal{F}_n^k$.

Payoff interval Consider the matching $\hat{y}_n$ defined by the following procedure:

Step 1. Set $\hat{y}_n(f) = \emptyset$ for every $f \in \mathcal{F}_n^1 \cup \dots \cup \mathcal{F}_n^J$.

Step 2. Let $\hat{\phi}_n$ be the matching resulting from seat-proposing random serial dictatorship between $\mathcal{F}_n^{J+1} \cup \dots \cup \mathcal{F}_n^K$ and $\mathcal{W}_n$. Set $\hat{y}_n(f) = \hat{\phi}_n(f)$ for all $f \in \mathcal{F}_n^{J+1} \cup \dots \cup \mathcal{F}_n^K$.

Note that $\mathbb{E}[u_f(\hat{y}_n)] > 0$ for every $f \in \mathcal{F}_n^{J+1} \cup \dots \cup \mathcal{F}_n^K$. For each subset of firms $F \subseteq \mathcal{F}_n^{J+1} \cup \dots \cup \mathcal{F}_n^K$, consider the matching $\hat{y}_n^F$ defined by the following procedure:

Step 1. Set $\hat{y}_n^F(f) = \hat{y}_n(f)$ for every $f \in \mathcal{F}_n \setminus F$.

Step 2. Set $\hat{y}_n^F(f) = \hat{y}_n(f)$ for every $f \in F$.

Using a similar argument as the one in Appendix A.4.2, the payoffs that firms in $\mathcal{F}_n^{J+1} \cup \dots \cup \mathcal{F}_n^K$ obtain from $\Delta(\{\hat{m}_n^F : F \subseteq \mathcal{F}_n^{J+1} \cup \dots \cup \mathcal{F}_n^K\})$ span $[0, \mathbb{E}[u_f(\hat{y}_n)]]^{\mathcal{F}_n^{J+1} \cup \dots \cup \mathcal{F}_n^K}$. Observe also that $\mathbb{E}[u_f(\hat{y}_n^F)] = 0$ for every $f \in \mathcal{F}_n^1 \cup \dots \cup \mathcal{F}_n^J$ and all $F \subseteq \mathcal{F}_n^{J+1} \cup \dots \cup \mathcal{F}_n^K$.

Define $\hat{w} \equiv (1 - p^x)\mathbb{E}[u_f(\hat{y}_n)] > 0$. The payoff interval for vacant firms is $[\frac{1}{2}\hat{w}, \hat{w}]$. Take any $w \in [0, \hat{w}]^{\mathcal{F}_n^{J+1} \cup \dots \cup \mathcal{F}_n^K}$. Then we have $w = (1 - p^x)\tilde{w}$ for some $\tilde{w} \in [0, \mathbb{E}[u_f(\hat{y}_n)]]^{\mathcal{F}_n^{J+1} \cup \dots \cup \mathcal{F}_n^K}$. Let $\lambda_n^y \in \Delta(\{\hat{m}_n^F : F \subseteq \mathcal{F}_n^{J+1} \cup \dots \cup \mathcal{F}_n^K\})$ be a random matching that satisfies $\mathbb{E}[u_f(\lambda_n^y)] = \tilde{w}_f$ for all $f \in \mathcal{F}_n^{J+1} \cup \dots \cup \mathcal{F}_n^K$. Let λ_n^y denote the (compound) lottery that assigns p^x weight on $\hat{x}_n$ and $(1 - p^x)$ weight on λ_n^y .

We have

$$\mathbb{E}[u_f(\lambda_n^*)] \geq p^x \hat{v}_k + (1 - p^x)0 \geq \frac{1}{2}(\hat{v}_k + \underline{v}_k) \quad \text{for all } f \in \mathcal{F}_n^k, 1 \leq k \leq J \quad (17)$$

$$\mathbb{E}[u_f(\lambda_n^*)] = p^x 0 + (1 - p^x)\tilde{w}_f = w_f \quad \text{for all } f \in \mathcal{F}_n^k, J+1 \leq k \leq K. \quad (18)$$

Self-enforcing matching process Consider a triggering matching process μ_n^* that randomizes according to λ_n^* on-path. If a firm in $\mathcal{F}_n^k$, $k = 1, \dots, J$, deviates, the matching process switches to permanently playing $\underline{\mu}_n^k$; if any firm in $\mathcal{F}_n^{J+1} \cup \dots \cup \mathcal{F}_n^K$ deviates, the matching process permanently switches to playing the worker-proposing static stable matching.

A firm belonging to an occupied quality class, $\mathcal{F}_n^k$, obtains at least $\frac{1}{2}(\hat{v}_k + \underline{v}_k)$ on-path and only $\underline{v}_k$ after deviation, so there exists $\underline{\delta}_1$ so that no firm in $\mathcal{F}^1 \cup \dots \cup \mathcal{F}_n^J$ finds

it profitable to deviate when $\delta > \underline{\delta}_1$. Meanwhile, a firm f belonging to a vacant quality class $\mathcal{F}^{J+1} \cup \dots \cup \mathcal{F}_n^K$ obtains $w_f \geq \frac{1}{2}\widehat{w} > 0$ on-path, but 0 if it were to deviate, so there exists $\underline{\delta}_2$ so that no firm in $\mathcal{F}^1 \cup \dots \cup \mathcal{F}_n^J$ finds it profitable to deviate when $\delta > \underline{\delta}_2$. So when $\delta > \underline{\delta} \equiv \max\{\underline{\delta}_1, \underline{\delta}_2\}$, no firm has profitable deviations. In addition, no worker can deviate profitably by construction. So the matching process is self-enforcing and sustains payoff vector w .

A.5 Active firms, passive workers

Consider a stage-game matching market that consists of firms $\mathcal{F}$ and workers $\mathcal{W}$. The active firms, passive workers (AFPW) stage game is a normal-form game, where $\mathcal{F}$ is the set of players and each $f \in \mathcal{F}$ has action set $A_f \equiv \{W \subseteq \mathcal{W} : |W| \leq q_f\}$; that is, each firm's action set consists of (possibly empty) subsets of workers that do not violate its capacity constraint. For each action profile $a = \times_{f \in \mathcal{F}} a_f$ and worker $w \in \mathcal{W}$, define $P^w(a) \equiv \{f \in \mathcal{F} : w \in a_f, f \succ_w w\} \cup \{w\}$ as the set of proposals that are acceptable to w . We associate each action profile a in the AFPW stage game with a stage-game matching m^a defined by

$$m^a(w) = \max_{\succ_w} P^w(a).$$

In the matching m^a , each worker is matched to her favorite acceptable proposer. Note that the maximizer above is unique since worker preferences are strict. Finally, in the AFPW stage game, each firm f 's payoff from action profile a , $v_f(a)$, is the payoff it obtains from the imputed matching m^a : that is, $v_f(a) = u_f(m^a)$ for all $a \in A \equiv \times_{f \in \mathcal{F}} A_f$.

Below I show that analyzing subgame perfect Nash equilibria in the AFPW game is equivalent to analyzing self-enforcing matching processes in the repeated coalitional matching game. The argument proceeds in two steps. Claim 1 below establishes that the two game forms have the same set of stage-game outcomes. Claim 2 further shows that for each "recommended" stage-game outcome, the payoffs that firm f can achieve by deviating in the AFPW stage game are identical to the payoffs that it can obtain through feasible coalitional deviations in the coalitional matching stage game. Together, Claim 1 and Claim 2 imply that the enforceable outcomes in the repeated AFPW game and the repeated coalitional matching game are identical.

Let M° denote the set of stage-game matchings that are acceptable to all workers. Our first claim shows that each action profile in the AFPW stage game corresponds to a stage-game matching in M° and vice versa.

CLAIM 1. For every action profile a in the AFPW stage game, we have $m^a \in M^\circ$. Conversely, for each matching $m \in M^\circ$, there exists an action profile a in the AFPW stage game such that $m = m^a$.

PROOF. The first half of the claim follows by the construction of m^a . To see the second half, consider any matching $m \in M^\circ$, and let firm f 's action in the AFPW stage game be defined by $a_f = m(f)$. It is straightforward to see that $m = m^a$. $\square$

The next result shows that for each outcome, both game forms offer firms the same opportunities to deviate. In particular, I show that the set of payoffs that firm f can achieve by deviating from an action profile a is identical to the payoffs that it can obtain through feasible coalitional deviations from the matching m^a .

For each stage-game matching m , let

$$D(m, f) \equiv \{[m, (f, W)] : \{f\} \cup W \text{ is a feasible coalitional deviation from } m\}$$

denote the set of stage-game matchings that can result from a feasible coalitional deviation from m involving f . For each action profile a in the AFPW stage game, let

$$\tilde{D}(a, f) \equiv \{a' \in A : a'_{-f} = a_{-f}\}$$

denote the set of action profiles that can result from f deviating from action profile a .

CLAIM 2. Let a be an action profile in the AFPW stage game and let m^a be its corresponding stage-game matching, and let $f \in \mathcal{F}$ be any firm. Then for every action profile $a' \in \tilde{D}(a, f)$, there exists $m' \in D(m^a, f)$ such that $v_f(a') = u_f(m')$. Conversely, for each matching $m' \in D(m^a, f)$, there exists action profile $a' \in \tilde{D}(a, f)$ such that $u_f(m') = v_f(a')$.

PROOF. For the first half of the claim, let $W = m^a(f)$ and $W' = m^{a'}(f)$ be f 's match partners in m^a and $m^{a'}$, respectively. For every worker $w \in W'$, we know from the definition of $m^{a'}$ that

$$f \succ_w f' \quad \text{for every } f' \in P^w(a') \setminus \{f\}. \quad (19)$$

But since $a_{-f} = a'_{-f}$, we know that w received the same set of alternative proposals under a and a' , so $P^w(a) \setminus \{f\} = P^w(a') \setminus \{f\}$. As a result, (19) implies that $f \succ_w f'$ for every $w \in W'$ and $f' \in P^w(a) \setminus \{f\}$. In other words, if m^a is the status quo matching and a worker $w \in W'$ receives a proposal from f , then f would be the most attractive proposal to her. This implies that $\{f\} \cup W'$ is a feasible deviation, and the matching $m' \equiv [m, (f, W')]$ satisfies $u_f(m') = \tilde{u}(W') = v_f(a')$.

For the second half of the claim, suppose without loss of generality that $m' = [m^a, (f, W')]$ for some W' , where $\{f\} \cup W'$ is a feasible coalitional deviation from m^a . By the definition of m^a , we know that

$$f \succeq_w f' \quad \text{for all } w \in m^a(f) \text{ and } f' \in P^w(a) \quad (20)$$

and

$$m^a(w) \succeq_w f' \quad \text{for all } w \in W' \setminus m^a(f) \text{ and } f' \in P^w(a). \quad (21)$$

TABLE 3. Preferences.

$u_f(w)$	w_1	w_2	w_3	$\succ$			
f_1	2	9	−1	w_1	f_1	f_3	f_2
f_2	7	7.5	9	w_2	f_2	f_1	f_3
f_3	−1	9	2	w_3	f_3	f_2	f_1

In addition, since $\{f\} \cup W'$ is a feasible deviation from m^a , it must be true that $f \succ_w m^a(w)$ for every $w \in W' \setminus m^a(f)$, so (21) becomes

$$f \succeq_w f' \quad \text{for all } w \in W' \setminus m^a(f) \text{ and } f' \in P^w(a). \tag{22}$$

Let $a'_f \equiv W'$ and $a' \equiv (a'_f, a_{-f})$. We will show that $u_f(m^{a'}) = \tilde{u}(W') = u_f(m')$, which, since $v_f(a') = u_f(m^{a'})$ by definition, would imply $u_f(m') = v_f(a')$ and establish our claim. To see why, first note that every worker $w \in W' \cap m^a(f)$ receives the same set of proposals under a and a' , i.e., $P^w(a) = P^w(a')$. So (20) becomes

$$f \succeq_w f' \quad \text{for all } w \in W' \cap m^a(f) \text{ and } f' \in P^w(a'),$$

which implies $m^{a'}(w) = f$ for all $w \in W' \cap m^a(f)$. Second, note that $P^w(a') = P^w(a) \cup \{f\}$ for every $w \in W' \setminus m^a(f)$, so (22) becomes

$$f \succeq_w f' \quad \text{for all } w \in W' \setminus m^a(f) \text{ and } f' \in P^w(a'),$$

and, therefore, $m^{a'}(w) = f$ for all $w \in W' \setminus m^a(f)$. As a result, $m^{a'}(f) = W'$ and $u_f(m^{a'}) = \tilde{u}(W') = u_f(m')$, which completes the proof. $\square$

A.6 Examples

A.6.1 One-to-one matching market with unique stable matching The goal of this section is to demonstrate that when $q = 1$ and there is a unique stable matching, repetition can still lead to gains for the firms.

Suppose $\mathcal{F} = \{f_1, f_2, f_3\}$ with $q_1 = q_2 = q_3 = 1$, and $\mathcal{W} = \{w_1, w_2, w_3\}$. Firms' payoffs and workers' preferences are shown in Table 3. Assume also that firms obtain zero payoff from unfilled positions and workers prefer any employer over unemployment. In this matching market there is a unique stable matching m^* , as shown in Figure 6 (the uniqueness can be verified by checking that both the worker- and firm-optimal matchings coincide).

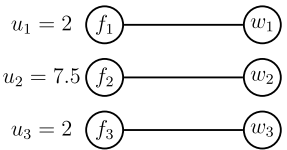


FIGURE 6. Matching m^* : Unique stable matching.

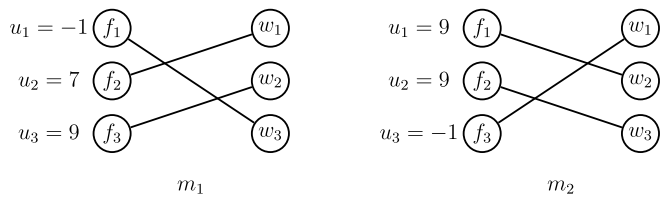


FIGURE 7. Unstable matchings.

Figure 7 shows two unstable matchings, m_1 and m_2 . From the firms’ perspective, the randomization $\hat{\lambda} = \frac{1}{2}m_1 + \frac{1}{2}m_2$ gives the firms a payoff profile of $(4, 8, 4)$, which strictly Pareto dominates their payoff profile $(2, 7.5, 2)$ from m^* . Moreover, $\hat{\lambda}$ can be sustained by the threat of permanently reverting to m^* .

A.6.2 Common ranking over workers The goal of this section is to demonstrate that when firms share a common ranking over workers, repetition can still expand the set of sustainable outcomes.

Suppose $\mathcal{F} = \{f_1, f_2\}$ with $q_1 = q_2 = 2$, and $\mathcal{W} = \{w_1, w_2, w_3, w_4\}$. Firms’ payoffs and workers’ preferences are shown in Table 4. Note that firms share a common ranking over workers. Assume also that firms obtain zero payoff from unfilled positions and that workers prefer any employer over unemployment. In this matching market, there is a unique stable matching m^* , as is shown in the left panel of Figure 8. In the right panel of Figure 8 is an unstable matching $\hat{m}$ that (from the firms’ perspective) strictly Pareto dominates the stable matching m^* . If the matching market is repeated and firms are sufficiently patient, then $\hat{m}$ can be sustained through the threat of permanently reverting back to m^* .

TABLE 4. Preferences.

$u_f(w)$	w_1	w_2	w_3	w_4
f_1	6	5	4	1
f_2	9	3	2	1

$\succ$		
w_1	f_1	f_2
w_2	f_2	f_1
w_3	f_2	f_1
w_4	f_1	f_2

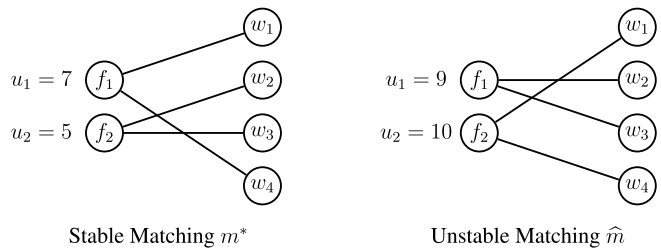


FIGURE 8. Stage-Game Matchings.

REFERENCES

- Abreu, Dilip, Prajit K. Dutta, and Lones Smith (1994), “The folk theorem for repeated games: A NEU condition.” *Econometrica*, 62, 939–948. [1726]
- Acemoglu, Daron and Alexander Wolitzky (2020), “Sustaining cooperation: Community enforcement versus specialized enforcement.” *Journal of the European Economic Association*, 18, 1078–1122. [1716]
- Acemoglu, Daron and Alexander Wolitzky (2021), “A theory of equality before the law.” *The Economic Journal*, 131, 1429–1465. [1716]
- Akbarpour, Mohammad, Shengwu Li, and Shayan Oveis Gharan (2020), “Thickness and information in dynamic matching markets.” *Journal of Political Economy*, 128, 783–815. [1716]
- Ali, S. Nageeb and Ce Liu (2020), “Conventions and coalitions in repeated games.” arXiv preprint arXiv:1906.00280. [1717, 1721]
- Ali, S. Nageeb and David A. Miller (2016), “Ostracism and forgiveness.” *American Economic Review*, 106, 2329–2348. [1716]
- Altınok, Ahmet (2021), “Dynamic many-to-one matching.” Technical Report. [1716]
- Anderson, Ross, Itai Ashlagi, David Gamarnik, and Yash Kanoria (2015), “A dynamic model of barter exchange.” In *Proceedings of the Twenty-Sixth Annual ACM-SIAM Symposium on Discrete Algorithms*, 1925–1933, Society for Industrial and Applied Mathematics. [1716]
- Ashlagi, Itai, Yash Kanoria, and Jacob D. Leshno (2017), “Unbalanced random matching markets: The stark effect of competition.” *Journal of Political Economy*, 125, 69–98. [1716, 1727, 1730]
- Baccara, Mariagiovanna, SangMok Lee, and Leeat Yariv (2019), “Optimal dynamic matching.” *Theoretical Economics*. [1716]
- Banerjee, Suryapratim, Hideo Konishi, and Tayfun Sönmez (2001), “Core in a simple coalition formation game.” *Social Choice and Welfare*, 18, 135–153. [1723]
- Bardhi, Arjada, Yingni Guo, and Bruno Strulovici (2023), “Early-career discrimination: Spiraling or self-correcting?” Technical Report, Working Paper, Duke and Northwestern. [1717]
- Bernheim, B. Douglas and Sita N. Slavov (2009), “A solution concept for majority rule in dynamic settings.” *The Review of Economic Studies*, 76, 33–62. [1717]
- Blackwell, David (1965), “Discounted dynamic programming.” *The Annals of Mathematical Statistics*, 36, 226–235. [1735]
- Che, Yeon-Koo and Olivier Tercieux (2018), “Payoff equivalence of efficient mechanisms in large matching markets.” *Theoretical Economics*, 13, 239–271. [1730, 1731, 1740]

- Che, Yeon-Koo and Olivier Tercieux (2019), "Efficiency and stability in large matching markets." *Journal of Political Economy*, 127, 2301–2342. [1727]
- Corbae, Dean, Ted Temzelides, and Randall Wright (2003), "Directed matching and monetary exchange." *Econometrica*, 71, 731–756. [1716]
- Damiano, Ettore and Ricky Lam (2005), "Stability in dynamic matching markets." *Games and Economic Behavior*, 52, 34–53. [1716]
- Deb, Joyee, Takuo Sugaya, and Alexander Wolitzky (2020), "The folk theorem in repeated games with anonymous random matching." *Econometrica*, 88, 917–964. [1716]
- Doval, Laura (2022), "Dynamically stable matching." *Theoretical Economics*, 17, 687–724. [1716]
- Du, Songzi and Yair Livne (2016), "Rigidity of transfers and unraveling in matching markets." Working Paper. [1716]
- Eeckhout, Jan (2000), "On the uniqueness of stable marriage matchings." *Economics Letters*, 69, 1–8. [1723]
- Ellison, Glenn (1994), "Cooperation in the prisoner's dilemma with anonymous random matching." *The Review of Economic Studies*, 61, 567–588. [1716]
- Fudenberg, Drew, David M. Kreps, and Eric S. Maskin (1990), "Repeated games with long-run and short-run players." *The Review of Economic Studies*, 57, 555–573. [1715, 1731]
- Fudenberg, Drew, David K. Levine, and Satoru Takahashi (2007), "Perfect public equilibrium when players are patient." *Games and Economic Behavior*, 61, 27–49. [1725]
- Fudenberg, Drew and Eric Maskin (1986), "The folk theorem in repeated games with discounting or with incomplete information." *Econometrica*, 54, 533–554. [1726, 1730, 1731]
- Gale, David and Lloyd S. Shapley (1962), "College admissions and the stability of marriage." *The American Mathematical Monthly*, 69, 9–15. [1711]
- Kadam, Sangram V. and Maciej H. Kotowski (2018a), "Multiperiod matching." *International Economic Review*, 59, 1927–1947. [1716]
- Kadam, Sangram V. and Maciej H. Kotowski (2018b), "Time horizons, lattice structures, and welfare in multi-period matching markets." *Games and Economic Behavior*, 112, 1–20. [1716]
- Kandori, Michihiro (1992), "Social norms and community enforcement." *The Review of Economic Studies*, 59, 63–80. [1716]
- Kotowski, Maciej H. (2020), "A perfectly robust approach to multiperiod matching problems." Working Paper. [1716]
- Kurino, Morimitsu (2020), "Credibility, efficiency, and stability: A theory of dynamic matching markets." *The Japanese Economic Review*, 71, 135–165. [1716]

- Lee, SangMok (2016), “Incentive compatibility of large centralized matching markets.” *The Review of Economic Studies*, 84, 444–463. [1716, 1727, 1730]
- Leshno, Jacob (2022), “Dynamic matching in overloaded waiting lists.” *American Economic Review*, 112, 3876–3910. [1716]
- Newton, Jonathan and Ryoji Sawa (2015), “A one-shot deviation principle for stability in matching problems.” *Journal of Economic Theory*, 157, 1–27. [1716]
- Peralta, Esteban (2020), “Participation in matching markets with distributional constraints.” Working Paper. [1723]
- Pittel, Boris (1989), “The average number of stable matchings.” *SIAM Journal on Discrete Mathematics*, 2, 530–549. [1716, 1730]
- Pittel, Boris (1992), “On likely solutions of a stable marriage problem.” *The Annals of Applied Probability*, 2, 358–401. [1716]
- Pycia, Marek (2012), “Stability and preference alignment in matching and coalition formation.” *Econometrica*, 80, 323–362. [1723]
- Roth, Alvin E. (1986), “On the allocation of residents to rural hospitals: A general property of two-sided matching markets.” *Econometrica: Journal of the Econometric Society*, 54, 425–427. [1735]
- Roth, Alvin E. and Marilda Sotomayor (1992), “Two-sided matching.” In *Handbook of Game Theory With Economic Applications*, Vol. 1, 485–541. [1721]
- Ünver, M. Utku (2010), “Dynamic kidney exchange.” *The Review of Economic Studies*, 77, 372–414. [1716]
- Wen, Quan (1994), “The “folk theorem” for repeated games with complete information.” *Econometrica: Journal of the Econometric Society*, 62, 949–954. [1725]
- Wolitzky, Alexander (2013), “Cooperation with network monitoring.” *Review of Economic Studies*, 80, 395–427. [1716]
- Wu, Qinggong (2015), “A finite decentralized marriage market with bilateral search.” *Journal of Economic Theory*, 160, 216–242. [1723]

Co-editor Simon Board handled this manuscript.

Manuscript received 5 June, 2021; final version accepted 24 October, 2022; available online 4 November, 2022.