

Learning by matching

YI-CHUN CHEN

Department of Economics and Risk Management Institute, National University of Singapore

GAOJI HU

School of Economics, Shanghai University of Finance and Economics

This paper studies a stability notion and matching processes in the job market with incomplete information on the workers' side. Each worker is associated with a type, and each firm cares about the type of her employee under a match. Moreover, firms' information structure is described by partitions over possible worker type profiles. With this firm-specific information, we propose a stability notion which, in addition to requiring individual rationality and no blocking pairs, captures the idea that the absence of rematching conveys no further information. When an allocation is not stable under the status quo information structure, a new pair of an allocation and an information structure will be derived. We show that starting from an arbitrary allocation and an arbitrary information structure, the process of allowing randomly chosen blocking pairs to rematch, accompanied by information updating, will converge with probability one to an allocation that is stable under the updated information structure. Our results are robust with respect to various alternative learning patterns.

KEYWORDS. Two-sided matching, incomplete information, stability, learning-blocking path, convergence.

JEL CLASSIFICATION. C78, D83.

1. INTRODUCTION

Matching is one of the important functions of markets (Roth (2008)). In particular, stable matchings have been connected to both equity and efficiency in resource allocation, two of the most important objectives in economics.¹ In this paper, we study stability

Yi-Chun Chen: ecsycc@nus.edu.sg

Gaoji Hu: hugaoji@mail.shufe.edu.cn

We are indebted to three anonymous referees for their insightful comments. We are especially grateful to Fuhito Kojima and Qingmin Liu for encouragement and stimulating discussions. We also thank David Ahn, In-Koo Cho, Laura Doval, Jiangtao Li, Xiao Luo, George Mailath, Luciano Pomatto, Andy Postlewaite, Ning Sun, Xiang Sun, Satoru Takahashi, Qianfeng Tang, Yan Yan, Zaifu Yang, and Yongchao Zhang for their comments and suggestions. Part of this paper was written while Hu was visiting Stanford University and he would like to thank the institution for its hospitality and support. This work is supported by Singapore MOE AcRF Tier 1 Grant. All remaining errors are our own.

¹See Balinski and Sönmez (1999) and Abdulkadiroğlu and Sönmez (2003) for how stability leads to the elimination of justified envy, a basic fairness property. See Shapley and Shubik (1971) and Liu et al. (2014) for how stability leads to efficiency.

and matching processes in a one-to-one job-market setting. We depart from the prevailing assumption of two-sided matching theory that information is complete, i.e., that the characteristics of all market participants are common knowledge. In particular, we study incomplete information on the workers' side.² We first describe what firms know and how firms update their possibilistic belief about workers' types, and propose an incomplete-information stability notion that allows for arbitrarily heterogeneous information. We then show that with probability one, a random matching process converges to an allocation that is stable with respect to the updated information structure.

Stability under complete information requires individual rationality (i.e., each agent has a nonnegative payoff) and no blocking pairs (i.e., no worker and firm would both prefer being matched with each other at some wage to staying with their current partners). When a firm has incomplete information, however, she may not know her potential employees' types, which reflect their productivity. As a result, the firm would not know whether she would prefer hiring another worker or keeping her current employee. In this situation, the notions of blocking and stability in the complete-information environment become inadequate.

Following Liu et al. (2014) (LMPS for short), we assume that a firm will evaluate her potential employees according to their worst possible types, and that firms can observe the prevailing allocation (i.e., the prevailing matching and wage profile) as well as the types of their own employees. In LMPS, the heterogeneity of firms' information stems only from their observation of their own employees' types. Unlike LMPS, however, we allow firms to have arbitrarily heterogeneous information about workers' types, and describe the firms' information structure by a profile of partitions over possible type profiles of the workers. Given an information structure, we propose a stability notion that extends the notion, proposed in LMPS, of stable matching with incomplete information. (For a formal comparison, see the literature review and Section 3.4.)

In our setting, a state of the market consists of an allocation and an information structure. A state is stable if (i) the allocation is individually rational, (ii) the allocation admits no blocking pair with respect to the information structure, and (iii) individual rationality and the absence of blocking convey *no* further information to the firms. The last requirement, in particular, embodies a notion of "informational stability," which is specific to the incomplete-information setting.

Equipped with the notion of stability, we study a matching process that mimics the behavior of market participants searching for desirable jobs or employees. Indeed, if a worker and a firm find that they would benefit more from being matched with each other than from maintaining the status quo, they will act to realize the improvement. The new matching may again admit a blocking pair, and thus another rematching opportunity, that results in another new matching, and so on. One important question is whether such a process finally stops at a stable matching.³

²We use the job-market setting (with transferable utility) to facilitate the comparison between our stability notion and that of Liu et al. (2014). Nevertheless, our convergence result (Theorems 2–3) can be established without difficulty in models with nontransferable utility, such as the model studied in Bikhchandani (2017).

³Knuth (1976) provides an example of a blocking path that admits a cycle; in other words, any matching on the path is not stable. This has motivated the study of the convergence of blocking paths. The literature

When information is incomplete, each observation of rematching or lack of rematching along the matching process carries additional information to the firms. Information updating refines the firms' partitions. Consequently, firms may become more optimistic about the worst type of a potential employee, which results in a new prospect of rematching. A matching process is thus associated with an information-updating process in which firms draw inferences along with each observation.

In this information-updating process, firms' partitions are refined for three possible reasons: a rematching is not observed, a rematching of other agents is observed, and a firm directly observes her new employee's type. Firms may update their information differently, depending on which one of the three possibilities occurs. In this sense, studying matching processes necessitates modeling stability with heterogeneous information, which we do from the beginning.

For an arbitrary initial market state, this learning and rematching process consists of a sequence of states. We call it a *learning-blocking path*. Our main result shows that, by suitably selecting the blocking pairs to be rematched, we can construct a finite learning-blocking path which reaches a stable state. This construction implies that when each blocking pair is randomly selected with positive probability to be rematched, the resulting learning-blocking path converges to a stable state with probability one. Our result is also robust with respect to alternative learning patterns.⁴ The general convergence extends the result of [Roth and Vande Vate \(1990\)](#) to markets with transferable utility and one-sided incomplete information. (For a formal comparison, see [Section 4.3](#), which also highlights how our argument differs from that of [Roth and Vande Vate \(1990\)](#).)

The rest of this section reviews the literature. [Section 2](#) introduces the model. [Section 3](#) defines stability with incomplete information. [Section 4](#) defines the notion of the learning-blocking path and presents our convergence results. [Section 5](#) discusses several related issues and [Section 6](#) concludes.

The related literature

The seminal paper of [Gale and Shapley \(1962\)](#) pioneered the literature of two-sided matching. Many classical developments are surveyed in [Roth and Sotomayor \(1990\)](#) and more recently by, e.g., [Roth \(2008\)](#). In this literature, a prevalent assumption is that information is complete.

Recently, LMPS introduced a notion of incomplete-information stability. Our notion of stability is consistent with the notion proposed by LMPS when the only source of firms' heterogeneous information is due to the observation of their current employees' types. More precisely, each stable state in our definition induces a stable outcome in LMPS; conversely, every stable outcome in LMPS can be supported as a stable state with respect to a specific partition profile. The latter partition profile can be seen as the one that is constructed in the following way: each firm starts with the sole piece of information of her employee's type and iteratively refines her partition based upon the fact that

demonstrates that the answer is primarily a positive one, although the argument generally varies across different setups.

⁴For example, agents may ignore or forget the information conveyed in some observations, or draw more sophisticated inferences from the observations. See [Section 5.1](#) for more discussion.

the state is not blocked. (See [Theorem 1](#) and [Corollary 1](#) for details.) Since our notion of stability is defined with respect to an arbitrary partition profile, we have the flexibility to study the matching process in which the information structure must be endogenized.

[Bikhchandani \(2017\)](#) proposes a notion of stability that is similar to that of LMPS but which applies to a Bayesian setting with nontransferable utilities. Unlike LMPS and [Bikhchandani \(2017\)](#), [Pomatto \(2019\)](#) considers a noncooperative matching game and uses forward-induction reasoning to derive the set of stable outcomes that is identified in LMPS.⁵ [Anderson and Smith \(2010\)](#) also study an employment model in which workers have unobserved abilities. However, agents in their paper are matched according to publicly observable reputation (e.g., probability of having high ability), which involves no information asymmetry as in our paper. Moreover, in their paper the agents maximize the discounted sum of wages in choosing matches, whereas the social planner maximizes the average present value of output in choosing matchings. In contrast, the matching process in our paper is driven by blocking pairs which are randomly drawn.

Whether a matching process converges to a stable allocation is known as the problem of finding paths to stability. [Roth and Vande Vate \(1990\)](#) provide the first result on paths to stability in marriage markets with complete information.⁶ Among many follow-up works, [Klaus and Klijn \(2007\)](#) establish the corresponding result in a matching-with-couples setup, [Kojima and Ünver \(2008\)](#) in the context of many-to-many matching, and [Chen et al. \(2010\)](#) in a job-market setting with transferable utilities.⁷ Most of the previous papers involve neither private information nor information updating, while both of these play a crucial role in our paper. There are two notable exceptions. [Bikhchandani \(2017\)](#) discusses the path to stability under a Bayesian notion of stability. In his paper, the final matching outcome of a blocking path is Bayesian stable only “conditional on the history,” i.e., agents may still block the final outcome with some of their erstwhile partners. In contrast, we provide a Bayesian version of the path-to-stability problem in [Section 5.4](#) where the final outcome is fully Bayesian stable. [Lazarova and Dimitrov \(2017\)](#) study paths to stability with incomplete information under a permissive/best-case notion of blocking. As a result, their approach is not applicable when more conservative blocking notions are adopted, such as the notions in LMPS, [Bikhchandani \(2017\)](#), and [Pomatto \(2019\)](#).⁸

⁵Another stream of literature studies stable mechanisms (instead of stable matchings) which also involve incomplete information. See, e.g., [Roth \(1989\)](#), [Chakraborty et al. \(2010\)](#), and [Ehlers and Massó \(2007, 2015\)](#).

Our stability notion is also related to the literature on the core, particularly the core in incomplete-information problems. In our context, a coalition is simply a worker-firm pair. See [Wilson \(1978\)](#), [Dutta and Vohra \(2005\)](#), and the comprehensive discussions in LMPS. See [Yenmez \(2013\)](#) for stability notions that are in line with [Dutta and Vohra \(2005\)](#).

⁶See [Ma \(1996\)](#) for a variant, called random-order mechanisms, of the paths studied in [Roth and Vande Vate \(1990\)](#).

⁷The main result of [Chen et al. \(2010\)](#), incorporated into [Chen et al. \(2016\)](#), is the convergence of blocking paths to *competitive equilibrium*, which is stronger than stability. See also [Fujishige and Yang \(2017\)](#).

⁸A *matching outcome* in [Lazarova and Dimitrov \(2017\)](#) consists of a matching and a belief system that specifies all agents' probabilistic beliefs about the type of each agent on the opposite side of the market. In their setting, a matching outcome is said to be *blocked* by a pair of agents, as long as there exist two types, one for each agent, such that (i) both agents prefer their opponent's type to their current partner's type, and (ii) both agents put positive probability on their opponent's type. In other words, agents in their setting are aggressive/optimistic in blocking, which reduces learning to trial-and-error.

2. THE MODEL

We consider the following setup of matching with incomplete information that is based on LMPS. The setup generalizes the complete-information matching models studied by Shapley and Shubik (1971) and Crawford and Knoer (1981).

There is a finite set I of workers to be matched with a finite set J of firms. Denote a generic worker by i and a generic firm by j . While each agent's index i or j is publicly observed, the agent's productivity is determined by the agent's *type*. Let W be the finite set of worker types and F be the finite set of firm types. A type assignment for firms is a mapping $\mathbf{f} : J \rightarrow F$, and similarly a type assignment for workers is another mapping $\mathbf{w} : I \rightarrow W$. We denote by Ω a set of type assignments for workers, i.e., $\Omega \subset W^I$.

A match between a worker of type $w \in W$ and a firm of type $f \in F$ gives rise to the *worker premuneration value* $\nu_{wf} \in \mathbb{R}$ and the *firm premuneration value* $\phi_{wf} \in \mathbb{R}$.⁹ The sum of ν_{wf} and ϕ_{wf} is called the *surplus of the match*. Denote these values by $\nu_{\mathbf{w}(i), \mathbf{f}(\emptyset)}$ for unmatched worker i and $\phi_{\mathbf{w}(\emptyset), \mathbf{f}(j)}$ for unmatched firm j , both of which are set to be zero. The functions $\nu : W \times F \rightarrow \mathbb{R}$ and $\phi : W \times F \rightarrow \mathbb{R}$ are common knowledge among the agents. Given a match between worker i (of type $\mathbf{w}(i)$) and firm j (of type $\mathbf{f}(j)$) under some wage $p \in \mathbb{R}$, the worker's payoff and the firm's payoff are, respectively, $\nu_{\mathbf{w}(i), \mathbf{f}(j)} + p$ and $\phi_{\mathbf{w}(i), \mathbf{f}(j)} - p$.¹⁰

A *matching* is a function $\mu : I \rightarrow J \cup \{\emptyset\}$, one-to-one on $\mu^{-1}(J)$, that assigns worker i to firm $\mu(i)$. In case $\mu(i) = \emptyset$, this means that worker i is unemployed; similarly, $\mu^{-1}(j) = \emptyset$ means that firm j does not hire anyone. A *payment scheme* $\mathbf{p}$ associated with a matching μ is a vector that specifies a payment $\mathbf{p}_{i, \mu(i)} \in \mathbb{R}$ for each worker $i \in I$ and a payment $\mathbf{p}_{\mu^{-1}(j), j} \in \mathbb{R}$ for each firm $j \in J$. To avoid nuisance cases, we associate zero payments with unmatched agents, by setting $\mathbf{p}_{\emptyset j} = \mathbf{p}_{i\emptyset} = 0$. Finally, an *allocation* $(\mu, \mathbf{p})$ consists of a matching μ and an associated payment scheme $\mathbf{p}$. We assume that the entire allocation is publicly observable.

As in LMPS, we assume that the type assignment for firms (i.e., $\mathbf{f}$) is common knowledge.¹¹ There is, however, incomplete information about the worker's types. In particular, the only facts that are common knowledge are these: (i) that the workers' type assignment belongs to Ω , (ii) that each worker knows his own type, and (iii) that each firm knows her current employee's type. Beyond the public information, each firm may also have her own private information about the workers' type assignment. Specifically, for every firm j , we describe her information by a partition Π_j over Ω . For any type assignment $\mathbf{w}$, write $\Pi_j(\mathbf{w})$ as the element of partition Π_j that contains $\mathbf{w}$. When the true type assignment is $\mathbf{w}$, firm j regards each type assignment $\mathbf{w}'$ in $\Pi_j(\mathbf{w})$ as possible. Denote the profile of partitions by Π , i.e., $\Pi := (\Pi_1, \dots, \Pi_{|J|})$, which is assumed to be common knowledge.

⁹See Mailath et al. (2013, 2017) for discussions on premuneration values.

¹⁰If we adopt the practice that salaries must be rounded to the nearest dollar or penny, the analysis in this section and the next will go through without any extra difficulty. This more practical restriction will be imposed in Section 4, where we study a matching process.

¹¹It is certainly important to also study matching markets with two-sided incomplete information, which involves the subtle formulation of the agents' higher-order reasoning. See Chen and Hu (2017) for details.

Say partition profile Π' is (weakly) finer than partition profile Π if, for each firm j , we have $\Pi'_j(\mathbf{w}) \subset \Pi_j(\mathbf{w})$ for every type assignment $\mathbf{w} \in \Omega$. Let Π^μ denote the partition profile that is generated by a matching μ , i.e., for every j and every $\mathbf{w}, \mathbf{w}' \in \Pi_j^\mu(\mathbf{w})$ if and only if $\mathbf{w}'(\mu^{-1}(j)) = \mathbf{w}(\mu^{-1}(j))$. Indeed, since each firm can observe the type of her current employee, the partition profile Π^μ captures the basic information of firms. We say a partition profile Π is *consistent* with a matching μ if Π is weakly finer than Π^μ . A *state* of the matching market, $(\mu, \mathbf{p}, \mathbf{w}, \Pi)$, specifies an allocation $(\mu, \mathbf{p})$, a type assignment $\mathbf{w}$, and a partition profile Π which is consistent with μ .

3. STABILITY WITH INCOMPLETE INFORMATION

3.1 Individual rationality

A state is said to be individually rational if each agent receives at least the payoff from remaining unmatched, which is assumed to be zero.

DEFINITION 1. A state $(\mu, \mathbf{p}, \mathbf{w}, \Pi)$ is said to be *individually rational* if

$$\begin{aligned} \nu_{\mathbf{w}(i), \mathbf{f}(\mu(i))} + \mathbf{p}_{i, \mu(i)} &\geq 0 \quad \text{for all } i \in I \quad \text{and} \\ \phi_{\mathbf{w}(\mu^{-1}(j)), \mathbf{f}(j)} - \mathbf{p}_{\mu^{-1}(j), j} &\geq 0 \quad \text{for all } j \in J. \end{aligned}$$

3.2 Blocking

The notion of incomplete-information “blocking” naturally extends its complete-information counterpart. In particular, a matching is blocked if some worker-firm pair (i, j) , where i and j are not matched with each other, can mutually benefit from being matched with each other. In order to accommodate the firms’ arbitrary private information, we propose the following definition of “blocking” which extends LMPS’s notion of blocking. We assume, as in LMPS, that firms care about the worst-case payoff when evaluating a potential worker.

DEFINITION 2. A state $(\mu, \mathbf{p}, \mathbf{w}, \Pi)$ is said to be *blocked* if there exists a worker-firm pair (i, j) and a payment $p \in \mathbb{R}$ such that worker i would switch to firm j at wage p , and that firm j would switch to worker i for any possible type assignments under which the worker would switch at the wage, i.e.,

$$\nu_{\mathbf{w}(i), \mathbf{f}(j)} + p > \nu_{\mathbf{w}(i), \mathbf{f}(\mu(i))} + \mathbf{p}_{i, \mu(i)} \quad \text{and} \quad (1)$$

$$\phi_{\mathbf{w}'(i), \mathbf{f}(j)} - p > \phi_{\mathbf{w}'(\mu^{-1}(j)), \mathbf{f}(j)} - \mathbf{p}_{\mu^{-1}(j), j} \quad (2)$$

for all $\mathbf{w}' \in \Pi_j(\mathbf{w})$ that satisfy

$$\nu_{\mathbf{w}'(i), \mathbf{f}(j)} + p > \nu_{\mathbf{w}'(i), \mathbf{f}(\mu(i))} + \mathbf{p}_{i, \mu(i)}. \quad (3)$$

We call the pair (i, j) a *blocking pair*, and the tuple $(i, j; p)$ a *blocking combination*, for the state $(\mu, \mathbf{p}, \mathbf{w}, \Pi)$ when conditions (1)–(3) are satisfied. For a firm j to participate

in a potential blocking pair at state $(\mu, \mathbf{p}, \mathbf{w}, \Pi)$, she must guarantee an improvement for every relevant type assignment. More precisely, when a firm j considers forming a blocking pair with worker i at some wage p , a type assignment $\mathbf{w}'$ is relevant for firm j when $\mathbf{w}' \in \Pi_j(\mathbf{w})$ and (3) holds. Any type assignment that violates (3) is irrelevant due to the worker's objection.

The following fact says that a blocking opportunity is more likely to exist if firms have more precise information about the workers' types.

FACT 1. *Suppose that Π' is a finer partition profile than Π . If state $(\mu, \mathbf{p}, \mathbf{w}, \Pi)$ is blocked, then state $(\mu, \mathbf{p}, \mathbf{w}, \Pi')$ is also blocked.*

PROOF. If $(\mu, \mathbf{p}, \mathbf{w}, \Pi)$ is blocked, then (1) and (2) hold for all $\mathbf{w}' \in \Pi_j(\mathbf{w})$ that satisfy (3). Since $\Pi'_j(\mathbf{w}) \subset \Pi_j(\mathbf{w})$, it follows that (1) and (2) hold for all $\mathbf{w}' \in \Pi'_j(\mathbf{w})$ that satisfy (3). $\square$

3.3 Stability

When information is complete, stable matching embodies the intuition that when “the agents have a very good idea of one another's preferences and have easy access to each other... , we might expect that stable matching will be especially likely to occur” (Roth (1989, p. 22)).¹² In this case, a stable state is simply a state that is individually rational and not blocked.

In contrast, with incomplete information, we argue that individual rationality and the absence of blocking pair are no longer sufficient to describe a “stable state.” To be precise, the partition Π_j represents firm j 's imprecise idea about the workers' information. As a result, the absence of blocking pair may still provide further information to firms. Once the firms' information partitions become finer, the worst case improves, and hence new blocking pairs may emerge. This is illustrated in Example 1 below.

EXAMPLE 1. Consider a job market in which we have two workers and two firms. In particular, $I = \{\alpha, \beta\}$ and $J = \{a, b\}$. The firms' types are given by $f_a = 4$ and $f_b = 3$. A type assignment for workers in this market is a two-dimensional vector, where the first component is the type for α and the second is the type for β . There are three possible type assignments, i.e., $\Omega = \{\mathbf{w}^{34}, \mathbf{w}^{32}, \mathbf{w}^{12}\}$, where $\mathbf{w}^{34} = (3, 4)$, $\mathbf{w}^{32} = (3, 2)$, and $\mathbf{w}^{12} = (1, 2)$. The premuneration value functions are given by $\phi_{wf} = wf$ and $\nu_{wf} = wf + 4 \cdot \mathbb{I}_{\{w=f\}}$, where $\mathbb{I}_{\{\cdot\}}$ is the indicator function. For the sake of simplicity, we do not allow for payments throughout this example.¹³

Obviously, firms prefer a worker of a higher type. A worker may prefer a firm of a lower type but only if the lower type is the same as his own type. Suppose firm a hires worker α and firm b hires worker β . In other words, a matching μ is given by $\mu(\alpha) = a$ and $\mu(\beta) = b$. It is straightforward to verify that if information is complete, then (i) under $\mathbf{w}^{34}$, (β, a) is the unique blocking pair; (ii) under $\mathbf{w}^{32}$, (α, b) is the unique blocking pair; and (iii) under $\mathbf{w}^{12}$, (β, a) is the unique blocking pair.

¹²Information is *complete* if every agent knows the true type assignment, whatever it is. In our notation, this is to say that $\Pi_j(\mathbf{w}') = \{\mathbf{w}'\}$ for all $\mathbf{w}' \in \Omega$ and all $j \in J$.

¹³A similar yet slightly more complicated example can be constructed when payments are allowed.

Suppose that $\mathbf{w}^{34}$ is the true type assignment. Assume that firms only know their own employee's type in μ , i.e., $\Pi = \Pi^\mu$ where

$$\begin{aligned}\Pi_a &= \{\{\mathbf{w}^{34}, \mathbf{w}^{32}\}, \{\mathbf{w}^{12}\}\} \quad \text{and} \\ \Pi_b &= \{\{\mathbf{w}^{34}\}, \{\mathbf{w}^{32}, \mathbf{w}^{12}\}\}.\end{aligned}$$

Then the state $(\mu, \mathbf{0}, \mathbf{w}^{34}, \Pi)$ is not blocked. We proceed to argue that it should not be stable. First, firms learn from the absence of blocking that the true type assignment must be either $\mathbf{w}^{34}$ or $\mathbf{w}^{32}$ by (iii). Next, with the following updated partition profile,

$$\begin{aligned}\Pi'_a &= \{\{\mathbf{w}^{34}, \mathbf{w}^{32}\}, \{\mathbf{w}^{12}\}\} \quad \text{and} \\ \Pi'_b &= \{\{\mathbf{w}^{34}\}, \{\mathbf{w}^{32}\}, \{\mathbf{w}^{12}\}\},\end{aligned}$$

the state $(\mu, \mathbf{0}, \mathbf{w}^{34}, \Pi')$ is again not blocked. Then firm a learns from the absence of blocking that the true type assignment must be $\mathbf{w}^{34}$ by (ii). Finally, under the updated partition profile, firms have complete information, i.e.,

$$\begin{aligned}\Pi''_a &= \{\{\mathbf{w}^{34}\}, \{\mathbf{w}^{32}\}, \{\mathbf{w}^{12}\}\} \quad \text{and} \\ \Pi''_b &= \{\{\mathbf{w}^{34}\}, \{\mathbf{w}^{32}\}, \{\mathbf{w}^{12}\}\}.\end{aligned}$$

Hence, (β, a) will form a blocking pair by (i). Therefore, individual rationality and no blocking pair are insufficient to capture “stability.”

The “stability” notion which we are about to propose not only requires individual rationality and no blocking pair but also necessitates that satisfying these two requirements provides no further information to agents. This latter requirement embodies a notion of information stability that is specific to the incomplete-information environment.

To formulate information stability, we define a set of type assignments as follows:

$$N_{\mu, \mathbf{p}, \Pi} := \{\mathbf{w} \in \Omega : (\mu, \mathbf{p}, \mathbf{w}, \Pi) \text{ is individually rational and not blocked}\}.$$

Intuitively, by the public information $(\mu, \mathbf{p}, \Pi)$ and the absence of blocking, firms know that the true type assignment lies in $N_{\mu, \mathbf{p}, \Pi}$. Let K_Π denote the meet (i.e., finest common coarsening) of the partition profile Π . Then, given a state $(\mu, \mathbf{p}, \mathbf{w}, \Pi)$, the set $K_\Pi(\mathbf{w})$ is the cell of the common knowledge partition that contains the true type assignment $\mathbf{w}$. An implication of [Example 1](#) is that upon observing the absence of blocking, each firm j should refine their partitions within $K_\Pi(\mathbf{w})$ instead of only $\Pi_j(\mathbf{w})$.¹⁴ Moreover, we would also like to make the notion of information stability a local property that depends only on partition within $K_\Pi(\mathbf{w})$; hence, we do not refine the partition outside $K_\Pi(\mathbf{w})$. For

¹⁴In [Example 1](#), given the initial state that is not blocked, firms first update their partitions from Π to Π' , in which neither of the two firm's partition cell at the true type assignment $\mathbf{w}^{34}$ is refined. Then as the new state with partition profile Π' is still not blocked, firms will further update the partition to Π'' , in which firm a 's partition cell at the true type assignment $\mathbf{w}^{34}$ is refined. This leads to a blocking of the state $(\mu, \mathbf{0}, \mathbf{w}^{34}, \Pi'')$.

notational convenience, we denote by $\mathcal{N}_{\mu, \mathbf{p}, \Pi}$ the binary partition that is induced by $N_{\mu, \mathbf{p}, \Pi}$, i.e., $\mathcal{N}_{\mu, \mathbf{p}, \Pi} := \{N_{\mu, \mathbf{p}, \Pi}, \Omega \setminus N_{\mu, \mathbf{p}, \Pi}\}$.

We now formally define an operator $\hat{H}_{\mu, \mathbf{p}}(\cdot)$ to represent the information refinement:

$$[\hat{H}_{\mu, \mathbf{p}}(\Pi)]_j(\mathbf{w}') := \begin{cases} \Pi_j(\mathbf{w}') \cap \mathcal{N}_{\mu, \mathbf{p}, \Pi}(\mathbf{w}'), & \text{if } \mathbf{w}' \in K_{\Pi}(\mathbf{w}); \\ \Pi_j(\mathbf{w}'), & \text{otherwise.} \end{cases} \quad (4)$$

If $\hat{H}_{\mu, \mathbf{p}}(\Pi) = \Pi$, then the fact of individual rationality and no blocking pair provides no further information to firms (in addition to their common knowledge $K_{\Pi}(\mathbf{w})$).¹⁵

A state is said to be stable if it is individually rational and not blocked, and if no further information can be inferred from the fact of individual rationality and no blocking.

DEFINITION 3. A state $(\mu, \mathbf{p}, \mathbf{w}, \Pi)$ is said to be *stable* if it satisfies the following three requirements:

- (i) $(\mu, \mathbf{p}, \mathbf{w}, \Pi)$ is individually rational.
- (ii) $(\mu, \mathbf{p}, \mathbf{w}, \Pi)$ is not blocked.
- (iii) $\hat{H}_{\mu, \mathbf{p}}(\Pi) = \Pi$.

When information is complete, i.e., if $\Pi_j(\mathbf{w}) = \{\mathbf{w}\}$ for every $j \in J$ and every $\mathbf{w} \in \Omega$, then Π is a fixed point of $\hat{H}_{\mu, \mathbf{p}}(\cdot)$ regardless of $(\mu, \mathbf{p})$. Hence, Definition 3 reduces to the standard definition of stable matching.¹⁶ In this case, a stable state exists (see Theorem 2 of Crawford and Knoer (1981)).

Up until now, we analyze a static setting and ask whether or not a state $(\mu, \mathbf{p}, \mathbf{w}, \Pi)$ is stable under an exogenously given information structure Π .¹⁷ In Section 4, we will study matching processes, which are dynamic, that consist of blocking and information updating in each step. Then the information structure Π will serve as both an input and an output variable.

3.4 Equivalence of two stability notions

In this subsection, we compare Definition 3 with the notion of stability defined in LMPS. The stability notion of LMPS is *ex ante* in that it is defined independently of the true type assignment and firms' heterogeneous belief. One can imagine an outside analyst who knows the model except for $\mathbf{w}$ and Π , and who wants to identify possible stable outcomes for the market. As usual, stable outcomes are individually rational and immune

¹⁵We thank an anonymous referee for suggesting the condition $\hat{H}_{\mu, \mathbf{p}}(\Pi) = \Pi$, which is equivalent to requiring that $K_{\Pi}(\mathbf{w}) \subset N_{\mu, \mathbf{p}, \Pi}$.

¹⁶Suppose that $(\mu, \mathbf{p}, \mathbf{w})$ is a complete-information stable outcome. Given an arbitrary Π , the state $(\mu, \mathbf{p}, \mathbf{w}, \Pi)$ is not necessarily stable in our sense. This is different from LMPS, where a complete-information stable outcome is always incomplete-information stable. However, the only reason for $(\mu, \mathbf{p}, \mathbf{w}, \Pi)$ being unstable is that it is not informational stable, i.e., it is individually rational and not blocked whatever the partition profile Π is.

¹⁷The “seemingly dynamic” Example 1 and the information updating in (4) are only used to motivate and to facilitate the introduction of information stability, which itself is a static fixed-point condition.

to blocking pairs. Formally, a *matching outcome* $(\mu, \mathbf{p}, \mathbf{w})$ specifies an allocation and a type assignment. The individual rationality of a matching outcome is defined as in [Definition 1](#), i.e., each agent has a nonnegative payoff. The blocking notion of LMPS is designed to exclude only outcomes that the analyst can be certain are “blocked.”

DEFINITION 4 (LMPS). Let Σ be a nonempty subset of individually rational matching outcomes. A matching outcome $(\mu, \mathbf{p}, \mathbf{w}) \in \Sigma$ is Σ -*blocked* if there exists a worker-firm pair (i, j) and a payment $p \in \mathbb{R}$ that satisfy

$$\nu_{\mathbf{w}(i), \mathbf{f}(j)} + p > \nu_{\mathbf{w}(i), \mathbf{f}(\mu(i))} + \mathbf{p}_{i, \mu(i)} \quad \text{and} \quad (5)$$

$$\phi_{\mathbf{w}'(i), \mathbf{f}(j)} - p > \phi_{\mathbf{w}'(\mu^{-1}(j)), \mathbf{f}(j)} - \mathbf{p}_{\mu^{-1}(j), j} \quad (6)$$

for all $\mathbf{w}' \in \Omega$ satisfying

$$(\mu, \mathbf{p}, \mathbf{w}') \in \Sigma \quad (7)$$

$$\mathbf{w}'(\mu^{-1}(j)) = \mathbf{w}(\mu^{-1}(j)) \quad (8)$$

$$\nu_{\mathbf{w}'(i), \mathbf{f}(j)} + p > \nu_{\mathbf{w}'(i), \mathbf{f}(\mu(i))} + \mathbf{p}_{i, \mu(i)}. \quad (9)$$

A matching outcome $(\mu, \mathbf{p}, \mathbf{w}) \in \Sigma$ is Σ -*stable* if it is not Σ -blocked.

Condition (5) says that worker i prefers firm j at wage p to his current match. Conditions (7)–(9) mean that firm j considers only “reasonable” type assignments, which are consistent with (i) the outcome set Σ , (ii) her “observation” $\mathbf{w}(\mu^{-1}(j))$, and (iii) worker i ’s willingness to block the outcome $(\mu, \mathbf{p}, \mathbf{w})$ with j at p . Condition (6) says that under any “reasonable” type assignments, firm j prefers worker i at p to her current match. Intuitively, the blocking conditions in [Definition 4](#) say that if $\mathbf{w}$ were the true type assignment, then $(\mu, \mathbf{p}, \mathbf{w})$ would be blocked based on the information of Σ , i.e., only outcomes in Σ are possible.

The set of outcomes that are immune to the blocking described in [Definition 4](#) is given by the iteration below. Let Σ^0 be the set of all individually rational outcomes. For $k \geq 1$, define

$$\Sigma^k := \{(\mu, \mathbf{p}, \mathbf{w}) \in \Sigma^{k-1} : (\mu, \mathbf{p}, \mathbf{w}) \text{ is } \Sigma^{k-1}\text{-stable}\}.$$

The set of *incomplete-information stable outcomes* in LMPS is given by $\Sigma^\infty := \bigcap_{k=1}^\infty \Sigma^k$.

The following theorem establishes the equivalence between our stability notion and that of LMPS. On the one hand, as long as $(\mu, \mathbf{p}, \mathbf{w}) \in \Sigma^\infty$, we can find at least one partition profile Π such that $(\mu, \mathbf{p}, \mathbf{w}, \Pi)$ is a stable state. That is, each stable outcome in LMPS can be supported as a part of some stable market state. On the other hand, as long as we can find one partition profile Π to support $(\mu, \mathbf{p}, \mathbf{w})$, the outcome $(\mu, \mathbf{p}, \mathbf{w})$ must be stable in the sense of LMPS. See [Appendix A](#) for the proof.

THEOREM 1. $(\mu, \mathbf{p}, \mathbf{w}) \in \Sigma^\infty$ if and only if there exists a partition profile Π such that $(\mu, \mathbf{p}, \mathbf{w}, \Pi)$ is stable.

Fix an outcome $(\mu, \mathbf{p}, \mathbf{w})$. Define $\Pi^{\mu, \mathbf{p}, 0} := \Pi^\mu$ and $\Pi^{\mu, \mathbf{p}, k} := \hat{H}_{\mu, \mathbf{p}}(\Pi^{\mu, \mathbf{p}, k-1})$ for every $k \geq 1$. Let $\Pi^{\mu, \mathbf{p}, \infty}$ be the limit of the increasingly finer partitions $\Pi^{\mu, \mathbf{p}, k}$. Then $\Pi^{\mu, \mathbf{p}, \infty}$ is the specific partition profile reflecting that firms just know (i) their own worker's type, and (ii) no blocking pair. The following corollary says that an outcome $(\mu, \mathbf{p}, \mathbf{w})$ is stable in the sense of LMPS if and only if the outcome, together with $\Pi^{\mu, \mathbf{p}, \infty}$, constitute a stable state. See [Appendix A](#) for the proof.

COROLLARY 1. $(\mu, \mathbf{p}, \mathbf{w}) \in \Sigma^\infty$ if and only if $(\mu, \mathbf{p}, \mathbf{w}, \Pi^{\mu, \mathbf{p}, \infty})$ is stable.

4. MATCHING PROCESSES WITH INCOMPLETE INFORMATION

In this section, we study matching processes and whether a matching process must lead to a stable state. Specifically, we consider a job market in which any worker and any firm can freely choose to be matched to each other, and any agent can freely opt to be unmatched. Suppose also that the agents are myopic, i.e., once an agent or a worker-firm pair finds an opportunity to improve their status quo, they will do so by either switching to be unmatched or finding a new partner. These individual and/or pairwise rematchings lead to a sequence of market states, which is referred to as a matching process. Note that at each rematching, a state/information structure is both input (i.e., the status quo) and output (i.e., the new state obtained from rematching).

We show that with probability one an arbitrary, random matching process (i.e., a matching process in which each blocking combination is randomly selected with positive probability to be rematched) converges to an incomplete-information stable state after finitely many rematchings. Throughout this section, we will fix a realized type assignment $\mathbf{w}^*$, which will be omitted for notational simplicity, i.e., we will write $(\mu, \mathbf{p}, \Pi)$ for the state $(\mu, \mathbf{p}, \mathbf{w}^*, \Pi)$.

4.1 Learning-blocking paths

With incomplete information, a matching process is necessarily associated with a *learning process*. In the current setup, a learning process corresponds to a sequence of partial information structures. Each information partition is updated from the previous one according to a new observation. More precisely, given a state $(\mu, \mathbf{p}, \Pi)$, firms may observe one of the following two situations:

- (i) there is no rematching; or
- (ii) there is a rematching that satisfies a blocking combination $(i, j; p)$.

In case (i), it is known among the firms that $(\mu, \mathbf{p}, \Pi)$ is not blocked, an event that can be distinguished from the event of the state being blocked. Then firms update their information by aggregating two pieces of information Π and $\mathcal{N}_{\mu, \mathbf{p}, \Pi}$. This aggregated information is represented by the join of the two partitions,¹⁸ i.e.,

$$H_{\mu, \mathbf{p}}(\Pi) := \Pi \vee \mathcal{N}_{\mu, \mathbf{p}, \Pi}. \quad (10)$$

¹⁸The *join* of two partitions is the coarsest common refinement of them. See [Aumann \(1976\)](#). We denote the join operator by $\vee$. The join of a partition profile Π and another partition $\mathcal{N}_{\mu, \mathbf{p}, \Pi}$ is a new partition profile such that $[\Pi \vee \mathcal{N}_{\mu, \mathbf{p}, \Pi}]_j = \Pi_j \vee \mathcal{N}_{\mu, \mathbf{p}, \Pi}$ for all $j \in J$.

That is, at each (hypothetically) true type assignment $\mathbf{w}$, firm j knows that the true type assignment lies in the set $[H_{\mu, \mathbf{p}}(\Pi)]_j(\mathbf{w}) = \Pi_j(\mathbf{w}) \cap \mathcal{N}_{\mu, \mathbf{p}, \Pi}(\mathbf{w})$. In the absence of re-matching, we only require that the firms be “level-1 sophisticated” in updating their belief once. In other words, we only build in a naive behavioral rule in describing firms’ inferences from the lack of re-matching.

In case (ii), a re-matching of the blocking combination $(i, j; p)$ is observed. In this case, we may consider two different situations, depending on whether or not a firm is firm j . First, firm j will observe worker i ’s type after they are matched. Second, all the other firms may update their information about worker i ’s type to exclude type assignments under which worker i would not have found it profitable to block the status quo with firm j at wage p . In fact, our result does not depend on the precise specification of how firms other than j update their belief. To allow for flexible belief updating, we assume in condition (iv) below only that the firms update their information partition profile from Π to another profile Π' that is (weakly) finer than $\Pi \vee \Pi^{\mu'}$.

We denote by $(\mu, \mathbf{p}, \Pi) \uparrow_{(i, j; p)}$ the state that is derived from state $(\mu, \mathbf{p}, \Pi)$ by satisfying a blocking combination $(i, j; p)$ for $(\mu, \mathbf{p}, \Pi)$. Formally, we define the state $(\mu', \mathbf{p}', \Pi') = (\mu, \mathbf{p}, \Pi) \uparrow_{(i, j; p)}$ such that:

- (i) worker i and firm j are rematched at salary p , i.e., $\mu'(i) = j$ and $\mathbf{p}'_{i, j} = p$;
- (ii) the previous partners of i and j , if any, become unmatched, i.e., $\mu'(\mu^{-1}(j)) = \emptyset$ if $\mu^{-1}(j) \neq \emptyset$, and $(\mu')^{-1}(\mu(i)) = \emptyset$ if $\mu(i) \neq \emptyset$;
- (iii) other parts of the allocation remain the same as $(\mu, \mathbf{p})$, i.e.,

$$\mu'(i') = \mu(i') \quad \text{and} \quad \mathbf{p}'_{i', \mu'(i')} = \mathbf{p}_{i', \mu(i')} \quad \text{for any } i' \in I \setminus \{i, \mu^{-1}(j)\};$$

- (iv) each firm updates her information according to her observation of the re-matching, i.e., Π' is finer than $\Pi \vee \Pi^{\mu'}$.¹⁹

In defining $(\mu, \mathbf{p}, \Pi) \uparrow_{(i, j; p)}$, we allow either agent i or agent j to be $\emptyset$, in which case $p = 0$. In particular, $i = \emptyset$ means that firm j dismisses her employee $\mu^{-1}(j)$, whereas $j = \emptyset$ means that worker i resigns from his firm $\mu(i)$. Thus, the operation $\uparrow_{(i, j; p)}$ and the term “re-matching” apply to both pairs and individuals. For notational convenience, we also set $(\mu, \mathbf{p}, \Pi) \uparrow_{(\emptyset, \emptyset; 0)} := (\mu, \mathbf{p}, \Pi)$.

DEFINITION 5. A *learning-blocking path* is a sequence of states $\{(\mu^l, \mathbf{p}^l, \Pi^l)\}_{l=0}^L$ such that for any $l \geq 0$, the following hold:

- (i) if $(\mu^l, \mathbf{p}^l, \Pi^l)$ is not blocked, then $(\mu^{l+1}, \mathbf{p}^{l+1}) = (\mu^l, \mathbf{p}^l)$ and $\Pi^{l+1} = H_{\mu^l, \mathbf{p}^l}(\Pi^l)$;

¹⁹For instance, we may set $\Pi'_j = \Pi_j \vee \Pi^{\mu'}_j$. For firm $j' \neq j$, define

$$B := \{\mathbf{w} \in \Omega : (\mu, \mathbf{p}, \mathbf{w}, \Pi) \text{ is blocked by } (i, j; p)\}.$$

That is, B consists of all type assignments of the workers that are consistent with the observation that $(\mu, \mathbf{p}, \Pi)$ is blocked by $(i, j; p)$. Then we may set $\Pi'_{j'}$ as the join of $\Pi_{j'}$ and $\{B, \Omega \setminus B\}$.

- (ii) moreover, if $(\mu^l, \mathbf{p}^l, \Pi^l)$ is blocked, then $(\mu^{l+1}, \mathbf{p}^{l+1}, \Pi^{l+1}) = (\mu^l, \mathbf{p}^l, \Pi^l) \uparrow_{(i,j;p)}$, where $(i, j; p)$ is a blocking combination for $(\mu^l, \mathbf{p}^l, \Pi^l)$.

We say that a learning-blocking path is *finite* if its length L is finite. A learning-blocking path $\{(\mu^l, \mathbf{p}^l, \Pi^l)\}_{l=0}^L$ is said to *converge within finitely many steps* if there exists a finite $T < L$ such that $(\mu^l, \mathbf{p}^l, \Pi^l) = (\mu^T, \mathbf{p}^T, \Pi^T)$ for every $l \geq T$. Indeed, a complete-information blocking path is a special case of a learning-blocking path. In this case, the partition profile is already the finest (i.e., $\Pi_j(\mathbf{w}) = \{\mathbf{w}\}$), which implies that no extra information can be obtained from any observation. Thus, a learning-blocking path is simply a blocking path in the literature, i.e., a sequence of allocations where each allocation is derived from its preceding allocation by satisfying one of the preceding allocation's blocking combinations (see, e.g., [Roth and Vande Vate \(1990\)](#) and [Chen et al. \(2010\)](#)).

We close this subsection by recording the following simple lemma which highlights a difference between our notion of learning-blocking path and that of [Roth and Vande Vate \(1990\)](#). The lemma also demonstrates how far agents can go when they are only “level-1 sophisticated”: applying $H_{\mu, \mathbf{p}}(\cdot)$ step by step leads to either a fixed point or a blocking opportunity. In this sense, our analysis shares a similar spirit as the literature on learning in game theory.²⁰

LEMMA 1. *Suppose that a state $(\mu, \mathbf{p}, \Pi)$ admits no blocking pair. Then there exists a finite learning-blocking path to either (i) a stable state or (ii) a state which admits a blocking pair and has a partition profile that is strictly finer than Π .*

PROOF. Let $\Pi^0 := \Pi$ and $\Pi^k := H_{\mu, \mathbf{p}}(\Pi^{k-1})$ for every $k \geq 1$. Observe that Π^k is increasingly (weakly) finer in k . Since Ω is finite, there is some finite k^* such that $\Pi^k = \Pi^{k^*}$ for every $k \geq k^*$. If $(\mu, \mathbf{p}, \Pi^{k^*})$ admits no blocking pair, it must be a stable state since $H_{\mu, \mathbf{p}}(\Pi^{k^*}) = \Pi^{k^*}$ implies $\hat{H}_{\mu, \mathbf{p}}(\Pi^{k^*}) = \Pi^{k^*}$. Since none of the intermediate state $(\mu, \mathbf{p}, \Pi^k)$, $0 \leq k < k^*$, is blocked due to [Fact 1](#), we know that $\{(\mu, \mathbf{p}, \Pi^k)\}_{k=0}^{k^*}$ is a learning-blocking path.

If $(\mu, \mathbf{p}, \Pi^{k^*})$ admits a blocking pair, then we let k^{**} be the smallest $k \leq k^*$ such that $(\mu, \mathbf{p}, \Pi^k)$ admits a blocking pair. Obviously, $\{(\mu, \mathbf{p}, \Pi^k)\}_{k=0}^{k^{**}}$ is a learning-blocking path. $\square$

4.2 Convergence of learning-blocking paths

When a rematching happens on a learning-blocking path, the rematched worker and firm both become better off while their previous partners become unmatched. It is then easier for these unmatched agents to find blocking opportunities. New blocking opportunities may, in turn, drag down the payoffs of the agents who previously became better off. As a result, there may be cycles along a learning-blocking path, i.e., a learning-blocking path may not converge. This is illustrated in the example provided by [Knuth \(1976\)](#) in an ordinal-preference setting, and one can also construct an example with cycles in our transferable-utility setting.

²⁰See, e.g., [Fudenberg and Levine \(1998\)](#) for a comprehensive survey.

However, learning-blocking paths are not completely chaotic. Given an initial state, we first show a deterministic convergence result. That is, given an arbitrary initial state, we can suitably choose the blocking combination to be satisfied (when there are many blocking combinations) such that the resulting learning-blocking path must converge to a stable state within finitely many steps (see [Theorem 2](#) below). Then we show a random convergence result. That is, given an arbitrary initial state, the learning-blocking path that is resulted from randomly satisfying blocking combinations (when there are many blocking combinations) must almost surely converge to a stable state within finitely many steps (see [Theorem 3](#) below). The deterministic convergence has to do with whether there exists *one* learning-blocking path that leads to a stable state, while the random convergence has to do with the probability of reaching a stable state when blocking combinations are randomly satisfied.

To obtain our results, we need to impose the following assumption that reflects the fact that payments in practice are measured in monetary units, and hence integers.²¹

ASSUMPTION 1. *Payments permitted in the job market are integers.*²²

Given an arbitrary initial state, we show that by carefully choosing blocking pairs at each state, we can construct one finite learning-blocking path that ends with a stable state. The following Theorems 2 and 3 extend the result in ([Roth and Vande Vate \(1990\)](#)), henceforth, RV) to accommodate incomplete information. The proof of [Theorem 2](#) is in [Section 4.4](#).

THEOREM 2. *Suppose that [Assumption 1](#) holds. Then starting from an arbitrary initial state, there exists a finite learning-blocking path that leads to a stable state.*

Now, following RV, we consider a random process which starts with an arbitrary state. The process proceeds to generate a random learning-blocking path, i.e., whenever an intermediate state is blocked by many combinations, the process randomly satisfies one of them. In particular, the blocking combination to be satisfied is drawn from a distribution which has full support on the set of blocking combinations, i.e., each blocking combination is satisfied with strictly positive probability. Moreover, the distribution depends only on the state but not on the history. This random process mimics the practical situations in the labor market: agents meet and negotiate randomly until they expect no more improvement. The theorem below follows from [Theorem 2](#).

²¹Under [Assumption 1](#), we can analogously define the notion of stable states, and the existence of stable states is still guaranteed (see Theorem 1 of [Crawford and Knoer \(1981\)](#)). In the rest of this section, we refer to notions of blocking and stability as those defined under [Assumption 1](#).

²²As we mentioned in [Footnote 10](#), salaries must be rounded to the nearest dollar or penny. This is a technical assumption to ensure *finite* bargaining choices when a worker-firm pair negotiates, as well as (more importantly) a realistic situation in decentralized market practice which our matching process mimics. See [Crawford and Knoer \(1981\)](#), [Kelso and Crawford \(1982\)](#), and [Chen et al. \(2016\)](#) for similar integral assumptions when *finite* matching processes are studied. In marriage models in which our results hold and no payment is involved, this assumption is of course not necessary any more.

Moreover, one can easily construct an example in which (i) two firms compete for one worker, (ii) the salary increment converges to zero, and (iii) the limit salary still permits a blocking. Therefore, without [Assumption 1](#), a finite path cannot be guaranteed even in the complete-information environment.

THEOREM 3. *Suppose that [Assumption 1](#) holds. Then starting from an arbitrary initial state, the random learning-blocking path converges with probability one to a stable state.*

We briefly explain why [Theorem 2](#) implies [Theorem 3](#). Indeed, thanks to [Assumption 1](#), we may restrict attention to finitely many states.²³ Since each blocking combination will be selected according to a distribution that depends on the state but not the history, we may regard the random learning-blocking path as a finite-state Markov chain. Obviously, any stable state is absorbing, i.e., it is impossible to leave it. Then the existence of stable states implies that the Markov chain has at least one absorbing state. Moreover, by [Theorem 2](#), from every state it is possible to go to an absorbing state (not necessarily in one step). Therefore, the Markov chain is absorbing, and our [Theorem 3](#) follows immediately from Theorem 11.3 of [Grinstead and Snell \(1997\)](#).

4.3 A comparison of [Theorem 2](#) and RV's theorem

[Theorem 2](#) is parallel to RV's theorem. RV constructs a sequence of subsets of agents, $\{A(l)\}_{l=1}^{\bar{l}}$, and correspondingly, a sequence of matchings $\{\mu^l\}_{l=1}^{\bar{l}+1}$ such that at each step $l \geq 1$, there is no blocking pair for μ^{l+1} that is contained in $A(l)$. Thus, a blocking pair for matching μ^{l+1} , if any, must involve at least one agent outside $A(l)$. In the subsequent step $l + 1$, the set $A(l + 1)$ is obtained by taking the union of $A(l)$ and at least one of those outside agents. As a result, the set $A(l)$ expands in the sense of set inclusion as the number of steps grows. Since there are only finitely many agents, the $\bar{l}$ can be chosen to be large enough such that the set $A(\bar{l})$ includes everyone in the market. By construction of the sequences, there is no blocking pair for $\mu^{\bar{l}+1}$ that is contained in $A(\bar{l})$, i.e., $\mu^{\bar{l}+1}$ is stable.

At each step of their construction, Roth and Vande Vate start by matching an outside agent α with her favorite partner in $A(l)$ among those who are willing to form a blocking pair with her. Then the agent who was abandoned by the favorite partner chosen by α , if any, will be denoted as α' . RV then match agent α' with her favorite partner in $A(l)$ among those who are willing to form a blocking pair with her. Again, there may be an agent α'' who was abandoned by the favorite partner chosen by α' , and for whom we repeat the argument, and so on. This chain will exhaust all blocking opportunities within $A(l + 1)$ and will produce the desired matching μ^{l+2} .

When a firm's information is incomplete, it is no longer clear who is her favorite worker. More precisely, without knowing the workers' types, firms do not know which worker to favor among those who are willing to form a blocking pair with them. As a result, a firm may even be unwilling to form a blocking pair with its *de facto* favorite worker due to worry about his worst possible type. Moreover, with incomplete information, observing either a rematching or the absence of blocking pairs leads to information updating. Hence, even if we manage to construct μ^{l+1} and an $A(l)$ which contains no blocking pair for μ^{l+1} , this no-blocking property need not be preserved upon satisfying

²³More precisely, since there are finitely many type assignments, there is a bounded set of integers $\bar{P}$ such that a state is blocked by $(i, j; p)$ only if it is blocked by $(i, j; p)$ with $p \in \bar{P}$. Clearly, the set of states with wages either in $\bar{P}$ or as in the initial state is finite.

a new blocking pair (to form μ^{l+2}). This is because the new blocking pair, or its absence, carries a new piece of information which may refine the firms' partitions, improve the worst case in their mind, and thereby open up new blocking prospects.

Due to these two issues, we cannot adapt RV's proof to the situation with incomplete information, although their argument still plays a role here. In proving [Theorem 2](#), as is done in RV, we construct a sequence of subsets of agents that contain no blocking pairs. We show that at any initial state there is a path that leads to either (a) a larger subset of agents that contains no blocking pairs (which is what RV shows) or (b) some firm's partition being strictly refined with the new observations along the path ([Lemma 2](#)). Since there are only finitely many type assignments, (b) cannot recur infinitely often. Hence, the resulting path leads to a stable state by the repetition of (a). We provide a formal sketch and proof in the next subsection.

4.4 Proof of [Theorem 2](#)

To provide a constructive proof of [Theorem 2](#), we consider two cases: the initial state admits a blocking pair or otherwise. If the initial state $(\mu^0, \mathbf{p}^0, \Pi^0)$ admits no blocking pair, then by [Lemma 1](#), there is a finite learning-blocking path to either a stable state or a state with a blocking pair and a partition profile strictly finer than Π^0 . Hence, we only need to focus on the case where a state admits a blocking pair.

DEFINITION 6. A set of agents $A \subset I \cup J$ is *internally stable* under state $(\mu, \mathbf{p}, \Pi)$ if the following hold:

- (i) Agents in A are only matched with agents in A .
- (ii) The set A does not contain two agents who form a blocking pair for $(\mu, \mathbf{p}, \Pi)$.
- (iii) Moreover, every matched agent in A has a strictly positive payoff.

Pick an internally stable set A at state $(\mu^0, \mathbf{p}^0, \Pi^0)$. If $(\mu^0, \mathbf{p}^0, \Pi^0)$ admits a blocking pair $(\bar{i}, \bar{j})$, then either one or both of agents $\bar{i}$ and $\bar{j}$ are outside A . We deal with the case where exactly one agent is in A in [Lemma 2](#) below and then the other case in the proof of [Theorem 2](#). To be precise, we show that we can construct a finite learning-blocking path that leads to a state $(\bar{\mu}, \bar{\mathbf{p}}, \bar{\Pi})$ under which either (a) we obtain a strict superset of A which is still internally stable; or (b) there is a firm j such that $\bar{\Pi}_j$ is strictly finer than Π_j . The former case resembles RV, and the latter is new.

LEMMA 2. Suppose that [Assumption 1](#) holds. Let A^0 be a set of agents which is internally stable under a state $(\mu^0, \mathbf{p}^0, \Pi^0)$. Suppose that worker $i^0 \notin A^0$ (resp., firm $j^0 \notin A^0$) forms a blocking pair for $(\mu^0, \mathbf{p}^0, \Pi^0)$ with a firm (resp., a worker) in A^0 . Then, starting from state $(\mu^0, \mathbf{p}^0, \Pi^0)$, there exists a finite learning-blocking path to a state $(\bar{\mu}, \bar{\mathbf{p}}, \bar{\Pi})$ under which either (a) the internally stable set is expanded, i.e., $A^0 \cup \{i^0\}$ (resp., $A^0 \cup \{j^0\}$) is internally stable; or (b) there exists a firm j whose information partition is strictly refined, i.e., $\bar{\Pi}_j$ is strictly finer than Π_j^0 .

We document the following simple lemma that will be used in the proof of [Lemma 2](#) and [Theorem 2](#).

LEMMA 3. *Suppose that there is a finite learning-blocking path starting from state $(\mu, \mathbf{p}, \Pi)$ to state $(\mu', \mathbf{p}', \Pi')$; moreover, $(i, j; p)$ is a blocking combination for state $(\mu', \mathbf{p}', \Pi')$ but not for state $(\mu, \mathbf{p}, \Pi)$. If both worker i and firm j get no less payoff in state $(\mu', \mathbf{p}', \Pi')$ than in state $(\mu, \mathbf{p}, \Pi)$, then Π_j'' is strictly finer than Π_j , where $(\mu'', \mathbf{p}'', \Pi'') = (\mu', \mathbf{p}', \Pi') \uparrow_{(i,j;p)}$.*

PROOF. Suppose to the contrary that $\Pi_j'' = \Pi_j$. Since i and j are matched under $(\mu'', \mathbf{p}'', \Pi'')$, firm j knows the type of worker i , which implies that firm j knows the type of worker i under $(\mu, \mathbf{p}, \Pi)$, and thus also under $(\mu', \mathbf{p}', \Pi')$. Since worker i and firm j gets no less payoff in state $(\mu', \mathbf{p}', \Pi')$ than in state $(\mu, \mathbf{p}, \Pi)$, $(i, j; p)$ being a blocking combination of $(\mu', \mathbf{p}', \Pi')$ implies that it is also a blocking combination for state $(\mu, \mathbf{p}, \Pi)$. This is a contradiction. $\square$

To prove [Lemma 2](#), we will explicitly construct a learning-blocking path which either (a) outputs a state such that the internally stable set is expanded in the sense of set inclusion; or (b) identifies a combination $(i, j; p)$ that satisfies the conditions in [Lemma 3](#), satisfying which leads to a strictly finer partition for firm j . The lemma is proved using what we call the **WORKER-ADDING (FIRM-ADDING) ALGORITHM**. The algorithm resembles the argument of RV by trying to expand the internally stable set but differs in an essential manner, i.e., the internally stable set may shrink to an empty set due to information updating.

PROOF OF LEMMA 2. Consider the case of a worker $i^0 \notin A^0$ who forms a blocking pair for state $(\mu^0, \mathbf{p}^0, \Pi^0)$ with firm in A^0 . To prove the claim, we input state $(\mu, \mathbf{p}, \Pi) = (\mu^0, \mathbf{p}^0, \Pi^0)$, $A = A^0$, worker $\alpha = i^0$, and $(i, j; p) = (\emptyset, \emptyset; 0)$ into the following **WORKER-ADDING ALGORITHM**. The case of a firm $j^0 \notin A^0$ who forms a blocking pair with a worker in A^0 can be similarly proved by switching the roles of worker and firm. In describing the algorithm, for each input value x of a variable, we denote by x' the updated value of x .

THE WORKER-ADDING ALGORITHM

START. Input state $(\mu, \mathbf{p}, \Pi)$, a (possibly empty) subset A of $I \cup J$, worker α , and $(i, j; p)$ where $i \neq \alpha$. Consider four mutually exclusive cases:

Case 1. There exists a blocking combination for $(\mu, \mathbf{p}, \Pi)$ which includes worker α and some firm in A .

Pick an arbitrary blocking combination $(\alpha, \bar{j}; \bar{p})$ with firm $\bar{j} \in A$ for $(\mu, \mathbf{p}, \Pi)$. Set $(\mu', \mathbf{p}', \Pi') = ((\mu, \mathbf{p}, \Pi) \uparrow_{(\alpha, \bar{j}; \bar{p})}) \uparrow_{(i, j; p)}$, $A' = A \cup \{\alpha\}$, worker $\alpha' = \alpha$, and $(i', j'; p') = (\mu^{-1}(\bar{j}), \bar{j}; \mathbf{p}_{\mu^{-1}(\bar{j}), \bar{j}})$. Go to **START**.

Case 2. There exists no blocking combination for $(\mu, \mathbf{p}, \Pi)$ which includes worker α and a firm in A but there exists a blocking combination for $(\mu, \mathbf{p}, \Pi)$ which includes worker $i \in I$ and some firm in A .

Pick an arbitrary blocking combination $(i, \bar{j}; \bar{p})$ with firm $\bar{j} \in A$ for $(\mu, \mathbf{p}, \Pi)$. Set $(\mu', \mathbf{p}', \Pi') = (\mu, \mathbf{p}, \Pi) \uparrow_{(i, \bar{j}; \bar{p})}$, $A' = A$, worker $\alpha' = i$, and $(i', j'; p') = (\mu^{-1}(\bar{j}), \bar{j}, \mathbf{p}_{\mu^{-1}(\bar{j})})$. Go to START.

Case 3. There exists no blocking combination for $(\mu, \mathbf{p}, \Pi)$ which includes either worker α or worker i and a firm in A . However, there exists a blocking combination $(\bar{i}, \bar{j}; \bar{p})$ for $(\mu, \mathbf{p}, \Pi)$ with both the firm and the worker in A .

Set $(\mu', \mathbf{p}', \Pi') = (\mu, \mathbf{p}, \Pi) \uparrow_{(\bar{i}, \bar{j}; \bar{p})}$ and $A' = \emptyset$. Go to END.

Case 4. There exists no blocking combination for $(\mu, \mathbf{p}, \Pi)$ which includes a pair of agents in $A \cup \{\alpha\}$.

Set $(\mu', \mathbf{p}', \Pi') = (\mu, \mathbf{p}, \Pi)$ and $A' = A$. Go to END.

END. Output $\bar{A} := A'$ and $(\bar{\mu}, \bar{\mathbf{p}}, \bar{\Pi}) := (\mu', \mathbf{p}', \Pi')$.

The algorithm keeps track of the following variables to be updated in each step: a state $(\mu, \mathbf{p}, \Pi)$, a set A of agents, a worker α , and a “potential blocking combination $(i, j; p)$.” In each step, exactly one of the four cases will be triggered. Case 1 says that worker α has the first priority to block the state: as long as he still wants to do so, we update the state by satisfying one such blocking combination, say matching worker α and firm $\bar{j}$; at the same time, we satisfy $(i, j; p)$, update A to $A \cup \{\alpha\}$, and update the “potential blocking combination” to $(i', j'; p')$ where j' becomes firm $\bar{j}$, and i' and p' become the employee of firm $\bar{j}$ and the wage which firm $\bar{j}$ paid to him before she left for worker α , respectively. Hence, when Case 1 is triggered at a state in which firm j and worker α are matched, the “potential blocking combination” $(i, j; p)$ is actually firm j ’s match immediately before she is matched with worker α .

We keep triggering Case 1 until there is no more blocking combination which involves worker α . Then we turn to trigger Case 2 if there is a blocking combination which involves worker i . In this case, we update the state by satisfying one such blocking combination, say matching worker i and firm $\bar{j}$; moreover, we let worker i become the new worker α and update the “potential blocking combination” to $(i', j'; p')$ where j' becomes firm $\bar{j}$, and i' and p' become the employee of firm $\bar{j}$ and the wage which firm $\bar{j}$ paid to him before she left for worker α , respectively. The algorithm will stop, once there is no blocking combination which involves either worker α or worker i . Then the algorithm outputs the updated state and $\emptyset$ if there is still a blocking combination with two agents both in A ; otherwise, it outputs the state and the set A of the final step.

First of all, we claim that the algorithm produces a learning-blocking path. It suffices to clarify that in Case 1, $(i, j; p)$ is a blocking combination for $(\mu, \mathbf{p}, \Pi) \uparrow_{(\alpha, \bar{j}; \bar{p})}$. For the initial input $(i, j; p) = (\emptyset, \emptyset; 0)$, this is a dummy condition. If $(i, j; p)$ is updated after either Case 1 or Case 2 is triggered, worker i is unmatched; moreover, after we satisfy a new blocking combination $(\alpha, \bar{j}; \bar{p})$ in Case 1, firm j also becomes unmatched. Since every matched agent in A has a strictly positive payoff (and remains so along the path), we know that i and j both prefer being rematched with each other at p to standing alone, i.e., $(i, j; p)$ is indeed a blocking combination for $(\mu, \mathbf{p}, \Pi) \uparrow_{(\alpha, \bar{j}; \bar{p})}$.

Second, we claim that the constructed path is finite. Indeed, by [Assumption 1](#), some firm’s payoff is strictly increased when Case 1 or Case 2 is triggered; moreover, a firm’s payoff never decrease in the algorithm unless she is firm j in Case 1, in which case her

payoff drops after a temporary increase and, at the same time, some other firm's payoff must strictly increase. Since we have only finitely many firms in $\mathcal{A}$, the constructed path can be infinite only if some firm's payoff is improved indefinitely. However, this contradicts individual rationality of workers and that the surpluses are uniformly bounded across all matches.

Third, if the path terminates by triggering Case 4, then $\bar{\mathcal{A}} = \mathcal{A}^0 \cup \{i^0\}$ is internally stable. Hence, it suffices to argue that when the path terminates by triggering Case 3, the updated partition of the blocking firm $\bar{j}$ in Case 3 must be strictly finer than her partition under the initial state. To see this, we denote by $\{(\alpha^k, i^k)\}_{k=1}^K$ the sequence of workers who serve in order the role of worker α and the role of worker i in Case 2 up until Case 3 is triggered. Let $k^* \leq K - 1$ be the maximal k such that $\bar{i} = \alpha^k$ ($k^* := 0$ if $\bar{i}$ has never served the role of worker α). Hence, worker $\bar{i}$ is neither α^k nor i^k for every $k \geq k^* + 1$.

Since only worker α or worker i changes his partner or wage in the algorithm before it terminates, worker $\bar{i}$ remains matched with the firm and wage which he settled as worker α^{k^*} (or at the initial state $(\mu^0, \mathbf{p}^0, \Pi^0)$ if $k^* = 0$). Moreover, no firm's payoff goes down in the algorithm. Finally, Case 2 must be triggered when worker α^{k^*+1} succeeds worker $\bar{i}$ to become a new worker α and when there was no more blocking combination which involves worker $\bar{i}$. In particular, $(\bar{i}, \bar{j}; \bar{p})$ was not a blocking combination for the state in the underlying Case 2 but becomes one when Case 3 is triggered. It follows from [Lemma 3](#) that firm $\bar{j}$ has an updated partition $\bar{\Pi}_{\bar{j}}$ that is strictly finer than $\Pi_{\bar{j}}^0$. $\square$

PROOF OF THEOREM 2. Consider an initial state. Without loss of generality, we assume that the initial state is individually rational (otherwise, break up all pairs with an agent who obtains a negative payoff). Furthermore, it is straightforward that any learning-blocking path preserves individual rationality. Thus, as long as we construct a learning-blocking path, the terminal state is individually rational.

We take an initial set of agents $\mathcal{A}^0 = \emptyset$, where the variable $\mathcal{A}$ will be updated during the construction below. For each state $(\mu, \mathbf{p}, \Pi)$, we distinguish two cases: (i) $(\mu, \mathbf{p}, \Pi)$ admits no blocking pair; (ii) $(\mu, \mathbf{p}, \Pi)$ admits a blocking pair. First, by [Lemma 1](#), for each state $(\mu, \mathbf{p}, \Pi)$ in Case (i), there is a finite learning-blocking path to either a stable state or a state in Case (ii) with a partition profile that is strictly finer than Π , where we update $\mathcal{A}$ to $\mathcal{A}' = \emptyset$ and proceed. Second, for each state $(\mu, \mathbf{p}, \Pi)$ in Case (ii) with a set of agents $\mathcal{A}$ which is internally stable under $(\mu, \mathbf{p}, \Pi)$, consider two subcases: First, if some blocking combination involves an agent in $\mathcal{A}$ and an agent outside $\mathcal{A}$, then by [Lemma 2](#), starting from $(\mu, \mathbf{p}, \Pi)$, there exists a finite learning-blocking path to a state $(\mu', \mathbf{p}', \Pi')$ under which either (a) a set $\mathcal{A}' \supsetneq \mathcal{A}$ is internally stable; or (b) Π' is strictly finer than Π , where we set $\mathcal{A}' = \emptyset$ and proceed. Second, if every blocking combination involves only two agents outside $\mathcal{A}$, then we satisfy a blocking combination $(\bar{i}, \bar{j}; \bar{p})$ to obtain $(\mu, \mathbf{p}, \Pi) \uparrow_{(\bar{i}, \bar{j}; \bar{p})}$. If $\mathcal{A}' = \mathcal{A} \cup \{\bar{i}, \bar{j}\}$ is internally stable under $(\mu, \mathbf{p}, \Pi) \uparrow_{(\bar{i}, \bar{j}; \bar{p})}$, then (a) holds. If $\mathcal{A}' = \mathcal{A} \cup \{\bar{i}, \bar{j}\}$ contains a blocking pair (i, j) , with wage p , for $(\mu, \mathbf{p}, \Pi) \uparrow_{(\bar{i}, \bar{j}; \bar{p})}$, then by [Lemma 3](#), (b) holds for $(\mu', \mathbf{p}', \Pi') := ((\mu, \mathbf{p}, \Pi) \uparrow_{(\bar{i}, \bar{j}; \bar{p})}) \uparrow_{(i, j; p)}$, where we set $\mathcal{A}' = \emptyset$ and proceed.

To sum up, for each state in Case (ii) with a set $\mathcal{A}$ which is internally stable, we can construct a finite learning-blocking path to a state under which either (a) a set $\mathcal{A}' \supsetneq \mathcal{A}$ is internally stable; or (b) the partition profile is strictly finer. Moreover, for each state in

Case (i), we can also construct a finite learning-blocking path to either (a') a stable state or (b') a state in Case (ii) with a strictly finer partition profile. Since Ω is finite, the partition profile cannot be refined indefinitely. Hence, along the path which we construct by applying [Lemma 1](#) and [Lemma 2](#), eventually neither (b) nor (b') happens. Therefore, (a) keeps enlarging the internally stable set until (a') happens. That is, there is a finite learning-blocking path which leads to a stable state. $\square$

5. DISCUSSIONS

5.1 Robustness of convergence

The learning-blocking path described in [Section 4.1](#) includes three kinds of inferences that can be drawn from different observations. We use $H_{\mu, \mathbf{p}}(\Pi)$, defined in (10), to describe agents' updated information structure when they observe no rematching. Note that we only require that the firms be "level-1 sophisticated" in applying the operator $H_{\mu, \mathbf{p}}(\cdot)$ once. When agents observe a rematching, we allow for the flexibility of the updated information structure. More precisely, we put no additional restrictions on Π' in condition (iv) beyond requiring that it is finer than $\Pi \vee \Pi^{\mu'}$. To sum up, we only build in naive behavioral rules in defining a learning-blocking path, yet our result shows that a stable state can be reached with probability one. In this sense, our result share a similar spirit as the literature on learning in game theory.

In fact, it is clear from its proof that [Theorem 2](#) remains valid even if equation (10), as an information-updating rule, and condition (iv) are violated for finitely many times along each learning-blocking path. Such violations include situations in which the firms do not have perfect recall (i.e., Π' is not necessarily finer than $\Pi \vee \Pi^{\mu'}$ as required in condition (iv)) or in which it takes some time for firms to learn how to draw the inference embodied in (10).

5.2 Initial states and limit states

In general, the set of stable states that can arise from a learning-blocking path depends on its initial state. For instance, when the initial state is stable, there is only a trivial learning-blocking path which leads to itself. RV shows that in a complete-information setting, starting from an initial matching where all agents stand alone imposes no restriction on the limit matching beyond stability. That is, every stable matching can be achieved by a blocking path that starts from the no-match status. As a result, a random blocking path starting at a no-match initial status will achieve every stable matching with strictly positive probability. However, this is no longer the case if information is incomplete.

EXAMPLE 2. Consider a job market with only one worker α and one firm a . The firm's type is given by $f_a = 1$. The worker's type is either $w_\alpha^1 = 1$ or $w_\alpha^{-1} = -1$. The remuneration value functions are given by $\phi_{wf} = wf$ and $\nu_{wf} = |wf|$.

Suppose that w_α^1 is the true type for the worker. In this market, any state with a match $\mu(\alpha) = a$, wage payment p in $[-1, 1]$, and the firm's partition given by $\Pi_a = \{\{w_\alpha^1\}, \{w_\alpha^{-1}\}\}$,

is stable. Moreover, the no-match state with $\mu(\alpha) = \emptyset$, $p = 0$, and $\Pi_a = \{\{w_\alpha^1, w_\alpha^{-1}\}\}$, is also stable. Obviously, starting from the no-match state, the stable states with a match cannot be achieved by any learning-blocking path.

Given an arbitrary initial state, we know that the limiting stable state necessarily has a partition profile that is finer than the initial partition profile. That is, only a state which is stable with a finer partition profile can be a limit. However, [Example 2](#) shows that the converse is not true, i.e., not every such stable state can be achieved from the given initial state via learning-blocking paths.

There are also situations where all initial states lead to a unique limit matching (up to the relabeling of agents who have the same type). For example, if agents have one-dimensional types and the value functions satisfy monotonicity and supermodularity, then every stable matching outcome is efficient (see Proposition 3 in LMPS). It then follows from [Theorem 1](#) that every stable state is efficient, where the matching is positive assortative and thereby unique up to the relabeling of agents with the same type.²⁴

5.3 Allocation change along learning-blocking paths

As RV observed, the blocking path to stability which they construct is closely related to the deferred acceptance (DA) algorithm proposed by [Gale and Shapley \(1962\)](#). To be precise, consider the blocking path in a marriage matching market with an initial matching in which all agents are single. Under the initial matching, the set of men is internally stable. Then a woman is added to the set and a blocking path is triggered. Once the blocking path stops (i.e., once the expanded set becomes internally stable), another woman is added to the expanded set. Inductively, women are added one by one until a stable matching is reached. The sequence of matchings which RV construct in proving their convergence result is precisely the sequence of matchings which occur in the women-proposing DA algorithm. As in DA, along the path every man gets better and better partners while every woman gets worse and worse partners.

With transfers and incomplete information, the learning-blocking path which we construct in the proof of [Theorem 2](#) exhibits a similar kind of monotonicity. Consider the situation in [Lemma 2](#) with $i^0 \notin A$. Suppose that the WORKER-ADDING ALGORITHM outputs $A \cup \{i^0\}$, i.e., that the internally stable set is expanded. We observe that the wage which a firm pays to the same worker must be monotonically decreasing. This is because the payoff of every firm in A monotonically increases as the “competition” of workers intensifies. Similarly, in the FIRM-ADDING ALGORITHM, i.e., $j^0 \notin A$ in the statement of [Lemma 2](#), if the algorithm outputs $A \cup \{j^0\}$, then the wage which a firm pays to the same worker must be monotonically increasing. Hence, along the learning-blocking path that we constructed, the wage which a firm pays to the same worker must be piecewise monotonic, except for some construction steps where we simultaneously add a worker-firm pair to the internally stable set A .

²⁴With discrete transfer ([Assumption 1](#)), we can show that when the monetary unit is sufficiently small, the monotonicity and supermodularity of value functions still imply the efficiency of stable states.

5.4 Bayesian stability

One crucial assumption which we made is that firms evaluate the prospect of a blocking opportunity according to the worst possible scenario. We adopt this assumption from LMPS to obtain a stability notion which is comparable to theirs ([Theorem 1](#)). Here, we demonstrate that we can instead adopt a Bayesian perspective in defining stability and still prove the result on paths to stability.

Consider a Bayesian setting in which we fix the firms' common prior λ , which has full support on Ω . Given an allocation $(\mu, \mathbf{p})$, let $D_{ijp}^{(\mu, \mathbf{p})}$ denote the set of type assignments under which worker i gains after the combination $(i, j; p)$ is satisfied, i.e.,

$$D_{ijp}^{(\mu, \mathbf{p})} := \{\mathbf{w} \in \Omega : \nu_{\mathbf{w}(i), \mathbf{f}(j)} + p > \nu_{\mathbf{w}(i), \mathbf{f}(\mu(i))} + \mathbf{p}_{i, \mu(i)}\}.$$

Drawing on [Bikhchandani \(2017\)](#), we can define a Bayesian blocking notion which accommodates heterogeneous information among firms. Given a state $(\mu, \mathbf{p}, \mathbf{w}, \Pi)$ and a combination $(i, j; p)$, let $D_{ijp}^{(\mu, \mathbf{p}, \mathbf{w}, \Pi)}$ be the set of type assignments which is consistent with firm j 's partition and under which worker i finds the combination $(i, j; p)$ profitable, i.e.,

$$D_{ijp}^{(\mu, \mathbf{p}, \mathbf{w}, \Pi)} := D_{ijp}^{(\mu, \mathbf{p})} \cap \Pi_j(\mathbf{w}).$$

DEFINITION 7. A state $(\mu, \mathbf{p}, \mathbf{w}, \Pi)$ is said to be *Bayesian blocked* if there exists a worker-firm pair (i, j) and a payment $p \in \mathbb{R}$ such that worker i and firm j both prefer to be re-matched with each other at wage p , i.e.,

$$\begin{aligned} & \nu_{\mathbf{w}(i), \mathbf{f}(j)} + p > \nu_{\mathbf{w}(i), \mathbf{f}(\mu(i))} + \mathbf{p}_{i, \mu(i)} \quad \text{and} \\ & \mathbb{E}_\lambda[\phi_{\mathbf{w}'(i), \mathbf{f}(j)} | D_{ijp}^{(\mu, \mathbf{p}, \mathbf{w}, \Pi)}] - p > \phi_{\mathbf{w}(\mu^{-1}(j)), \mathbf{f}(j)} - \mathbf{p}_{\mu^{-1}(j), j}.^{25} \end{aligned}$$

Equipped with [Definition 7](#), we can define $N_{\mu, \mathbf{p}, \Pi}^B$ and $\hat{H}_{\mu, \mathbf{p}}^B(\cdot)$ in a way that is similar to how we define $N_{\mu, \mathbf{p}, \Pi}$ and $\hat{H}_{\mu, \mathbf{p}}(\cdot)$ in [Section 3.3](#). Then we say that a state $(\mu, \mathbf{p}, \mathbf{w}, \Pi)$ is *Bayesian stable* if it is individually rational ([Definition 1](#)) and not Bayesian blocked, and if it satisfies $\hat{H}_{\mu, \mathbf{p}}^B(\Pi) = \Pi$. Clearly, if a state is blocked, then it must be Bayesian blocked. Hence, if a state $(\mu, \mathbf{p}, \mathbf{w}, \Pi)$ is Bayesian stable, then the state $(\mu, \mathbf{p}, \mathbf{w}, \hat{H}_{\mu, \mathbf{p}}^\infty(\Pi))$ is stable, where $\hat{H}_{\mu, \mathbf{p}}^k(\cdot) := \hat{H}_{\mu, \mathbf{p}}(\hat{H}_{\mu, \mathbf{p}}^{k-1}(\cdot))$.²⁶ Thus, a Bayesian stable outcome must be a stable outcome. Example 2 (with minor modifications) demonstrates that the converse is not true.

EXAMPLE 2 (Revisited). Let λ be a full-support prior over w_α^1 and w_α^{-1} . Recall that the no-match state with $\mu(\alpha) = \emptyset$, $p = 0$, and $\Pi_\alpha = \{\{w_\alpha^1, w_\alpha^{-1}\}\}$, is stable. However, according

²⁵Note that the conditional expectation in the second condition is well-defined given the first condition because λ has full support on Ω .

²⁶Since blocking implies Bayesian blocking, $(\mu, \mathbf{p}, \mathbf{w}, \Pi)$ is not blocked and $N_{\mu, \mathbf{p}, \Pi}^B \subset N_{\mu, \mathbf{p}, \Pi}$. An argument similar to that in the proof of [Corollary 1](#) shows that $(\mu, \mathbf{p}, \mathbf{w}, \hat{H}_{\mu, \mathbf{p}}^\infty(\Pi))$ is not blocked. Since $\hat{H}_{\mu, \mathbf{p}}^\infty(\Pi)$ is a fixed point of $\hat{H}_{\mu, \mathbf{p}}$, we know that $(\mu, \mathbf{p}, \mathbf{w}, \hat{H}_{\mu, \mathbf{p}}^\infty(\Pi))$ is stable.

to [Definition 7](#), the no-match state can be Bayesian blocked. To be precise, consider the wage $p = 1/2[-1 + \lambda(w_\alpha^1) - \lambda(w_\alpha^{-1})]$. In the match with p , the worker's payoff is given by

$$1 + p = \frac{1}{2} + \frac{1}{2}[\lambda(w_\alpha^1) - \lambda(w_\alpha^{-1})].$$

Thus, since $\lambda(w_\alpha^1) - \lambda(w_\alpha^{-1}) > -1$, it follows that the worker's payoff is greater than the payoff of being unmatched (i.e., zero). Similarly, the firm's expected payoff in the match with p is

$$[\lambda(w_\alpha^1) \cdot 1 + \lambda(w_\alpha^{-1}) \cdot (-1)] - p = \frac{1}{2} + \frac{1}{2}[\lambda(w_\alpha^1) - \lambda(w_\alpha^{-1})],$$

which, again, is greater than zero. Therefore, the no-match state can be Bayesian blocked.

We can also study the path-to-stability problem under the notion of Bayesian stability. To prove results parallel to [Theorems 2–3](#), we need to make only two changes: (i) replace blocking combinations by Bayesian blocking combinations and the operator $H_{\mu, \mathbf{p}}$ by $H_{\mu, \mathbf{p}}^B$; and (ii) modify the proof of [Lemma 2](#). We provide the details in [Appendix B](#).

As in the rest of our paper, we adopt the interim notion of stability with arbitrary information structure to study the path-to-stability problem. In contrast, [Bikhchandani \(2017\)](#), like LMPS, considers a setting where firms' private information is only the observed types of their own employees. [Liu \(2017\)](#) defines an ex ante notion of Bayesian stable matching which refers to neither the true type assignment nor the information partition. In particular, an allocation $(\mu, \mathbf{p})$ is ex ante Bayesian blocked by $(i, j; p)$ if $\lambda(D_{ijp}^{(\mu, \mathbf{p})}) > 0$ and

$$\mathbb{E}_\lambda[\phi_{\mathbf{w}(i), \mathbf{f}(j)} | D_{ijp}^{(\mu, \mathbf{p})}] - p > \max\{0, \mathbb{E}_\lambda[\phi_{\mathbf{w}(\mu^{-1}(j)), \mathbf{f}(j)} | D_{ijp}^{(\mu, \mathbf{p})}] - \mathbf{p}_{\mu^{-1}(j), j}\}.$$

An allocation $(\mu, \mathbf{p})$ is *ex ante Bayesian stable* if it is (ex ante) individually rational and not ex ante Bayesian blocked by any $(i, j; p)$. In this stability notion, information stability becomes irrelevant because neither the information partition nor the true type assignment is fixed.

6. CONCLUDING REMARKS

In this paper, we propose a notion of stability for matching with one-sided incomplete information, a notion which accommodates arbitrary heterogeneity of firms' information. Moreover, we show the convergence of random learning-blocking paths to stable states; the convergence extends the result due to [Roth and Vande Vate \(1990\)](#) to an incomplete-information environment. For the current paper, it is crucial to describe what firms know and how firms update their possibilistic information. From this perspective, our analysis complements that of [Liu et al. \(2014\)](#) and provides a benchmark for studying the dynamic decentralized foundation of incomplete-information stability.²⁷

²⁷See, e.g., [Lauermann and Nöldeke \(2014\)](#) for a formulation of decentralized matching markets.

APPENDIX A: PROOFS FOR SECTION 3

PROOF OF THEOREM 1. *Necessity.* Suppose that $(\mu, \mathbf{p}, \mathbf{w}) \in \Sigma^\infty$. Define

$$\Omega^0 = \{\mathbf{w}' \in \Omega : (\mu, \mathbf{p}, \mathbf{w}') \in \Sigma^\infty\}.$$

Define Π as the partition induced by μ and Ω^0 , i.e., $\Pi := \Pi^\mu \vee \{\Omega^0, \Omega \setminus \Omega^0\}$. Since $(\mu, \mathbf{p}, \mathbf{w})$ is individually rational, it follows that $(\mu, \mathbf{p}, \mathbf{w}, \Pi)$ is individually rational. Since $(\mu, \mathbf{p}, \mathbf{w})$ is not Σ^∞ -blocked, it follows that $(\mu, \mathbf{p}, \mathbf{w}, \Pi)$ is not blocked. Moreover, for each $\mathbf{w}' \in \Omega^0$ we have $(\mu, \mathbf{p}, \mathbf{w}') \in \Sigma^\infty$, and hence $(\mu, \mathbf{p}, \mathbf{w}', \Pi)$ is not blocked. This implies that

$$\Omega^0 \subset N_{\mu, \mathbf{p}, \Pi}. \quad (11)$$

Since $\Pi_j(\mathbf{w}') \subset \Omega^0$ for every $\mathbf{w}' \in \Omega^0$ and every $j \in J$, it follows from (11) that $\Pi_j(\mathbf{w}') \cap N_{\mu, \mathbf{p}, \Pi} = \Pi_j(\mathbf{w}')$ for every $\mathbf{w}' \in \Omega^0$, i.e., $\hat{H}_{\mu, \mathbf{p}}(\Pi) = \Pi$. Therefore, $(\mu, \mathbf{p}, \mathbf{w}, \Pi)$ is stable.

Sufficiency. Suppose that there exists a partition profile Π such that Π is consistent with μ and $(\mu, \mathbf{p}, \mathbf{w}, \Pi)$ is a stable state. Then obviously $(\mu, \mathbf{p}, \mathbf{w})$ is individually rational. Define

$$\Sigma^K := \{(\mu, \mathbf{p}, \mathbf{w}') : \mathbf{w}' \in K_\Pi(\mathbf{w})\}.$$

It suffices to show that Σ^K is a self-stabilizing set.²⁸ Then it follows from Proposition 2 of LMPS that $\Sigma^K \subset \Sigma^\infty$; hence, $(\mu, \mathbf{p}, \mathbf{w})$, which belongs to Σ^K , is a stable outcome. To see that Σ^K is self-stabilizing, fix an arbitrary $\mathbf{w}' \in K_\Pi(\mathbf{w})$ and we show that $(\mu, \mathbf{p}, \mathbf{w}')$ is not Σ^K -blocked. First, since $\mathbf{w}' \in K_\Pi(\mathbf{w})$, we have

$$\Pi_j(\mathbf{w}') \subset K_\Pi(\mathbf{w}). \quad (12)$$

Second, since Π is consistent with μ , we also have for each $\mathbf{w}'' \in \Pi_j(\mathbf{w}')$ that

$$\mathbf{w}''(\mu^{-1}(j)) = \mathbf{w}'(\mu^{-1}(j)). \quad (13)$$

Third, since $(\mu, \mathbf{p}, \mathbf{w}, \Pi)$ is a stable state, we have $\hat{H}_{\mu, \mathbf{p}}(\Pi) = \Pi$, which implies $K_\Pi(\mathbf{w}) \subset N_{\mu, \mathbf{p}, \Pi}$, and thus, by (12), $\mathbf{w}' \in N_{\mu, \mathbf{p}, \Pi}$. Finally, since $(\mu, \mathbf{p}, \mathbf{w}, \Pi)$ is not blocked and $\mathbf{w}' \in N_{\mu, \mathbf{p}, \Pi}$, it follows that $(\mu, \mathbf{p}, \mathbf{w}', \Pi)$ is not blocked, either. Thus, by (12), (13), and Fact 1, we conclude that $(\mu, \mathbf{p}, \mathbf{w}')$ is not Σ^K -blocked. $\square$

PROOF OF COROLLARY 1. By Theorem 1, it suffices to prove the necessity part. Consider an outcome $(\mu, \mathbf{p}, \mathbf{w})$ in Σ^∞ . By Theorem 1, there exists a partition profile Π such that $(\mu, \mathbf{p}, \mathbf{w}, \Pi)$ is stable. Hence, $\Pi = \hat{H}_{\mu, \mathbf{p}}(\Pi)$ and $(\mu, \mathbf{p}, \mathbf{w}, \Pi)$ is not blocked. Since Π is consistent with μ , it follows that Π is (weakly) finer than Π^μ . Therefore, by Fact 1, we know that $(\mu, \mathbf{p}, \mathbf{w}, \Pi^\mu)$ is not blocked. Similarly, for every $\mathbf{w}' \in N_{\mu, \mathbf{p}, \Pi}$, $(\mu, \mathbf{p}, \mathbf{w}', \Pi^\mu)$ is not blocked. That is, $N_{\mu, \mathbf{p}, \Pi} \subset N_{\mu, \mathbf{p}, \Pi^\mu}$. Note that $\Pi = \hat{H}_{\mu, \mathbf{p}}(\Pi)$ implies $K_\Pi(\mathbf{w}) \subset N_{\mu, \mathbf{p}, \Pi}$. Thus, we have

$$K_\Pi(\mathbf{w}) \subset N_{\mu, \mathbf{p}, \Pi^\mu}.$$

²⁸A nonempty set of individually rational matching outcomes E is *self-stabilizing* if every $(\mu, \mathbf{p}, \mathbf{w}) \in E$ is E -stable.

As a result, $\Pi_j(\mathbf{w}') \subset \Pi_j^{\mu, \mathbf{p}, 1}(\mathbf{w}')$ for all $\mathbf{w}' \in K_\Pi(\mathbf{w})$ and all j . Inductively, we have

$$K_\Pi(\mathbf{w}) \subset N_{\mu, \mathbf{p}, \Pi^{\mu, \mathbf{p}, k}},$$

which implies that $\Pi_j(\mathbf{w}') \subset \Pi_j^{\mu, \mathbf{p}, k}(\mathbf{w}')$ for all $\mathbf{w}' \in K_\Pi(\mathbf{w})$, all j , and all k . Finally, since $(\mu, \mathbf{p}, \mathbf{w}, \Pi)$ is not blocked, it follows from [Definition 2](#) that $(\mu, \mathbf{p}, \mathbf{w}, \Pi^{\mu, \mathbf{p}, \infty})$ is not blocked. $\square$

APPENDIX B: CONVERGENCE UNDER BAYESIAN STABILITY

Under Bayesian stability, the proof of the convergence of learning-blocking paths is identical to that under (worst-case) stability (i.e., the proof of [Theorem 2](#)), except that we need to modify the proof of [Lemma 2](#). In particular, cases in the WORKER-ADDING ALGORITHM are replaced by the following two cases.

Case 0. $(i, j; p)$ is a Bayesian blocking combination for $(\mu, \mathbf{p}, \Pi)$.

Set $(\mu', \mathbf{p}', \Pi') = (\mu, \mathbf{p}, \Pi) \uparrow_{(i, j; p)}$ and $A' = \emptyset$. Go to END.

Case 1. $(i, j; p)$ is not a Bayesian blocking combination for $(\mu, \mathbf{p}, \Pi)$. Consider four subcases as in the WORKER-ADDING ALGORITHM, which are now referred to as Cases 1.1–1.4.

In contrast to the worst-case notion, a firm may “regret” joining a blocking pair in a Bayesian setting; for example, after she joins a Bayesian blocking pair, she discovers that her payoff under the new employee’s true type ends up being strictly lower than her payoff with the previous employee. The only difference between the modified WORKER-ADDING ALGORITHM and the previous WORKER-ADDING ALGORITHM is that whenever firm j “regrets,” we let firm j hire her previous partner worker i again and terminate the algorithm (Case 0). With this modification, we only need to make the following minor changes in the proof of [Lemma 2](#).

First, the path which we construct is still finite. Indeed, the payoff of any firm in $\mathcal{A}$ never decrease before the algorithm terminates, unless she is the firm j in Case 1.1. In this case, firm j is first abandoned by worker α (and suffer a payoff decrease) and then rematched with the previous worker i . As a result, firm j still gets her payoff before being matched with worker α . Moreover, when Case 1.1 or Case 1.2 is triggered, one of the following two situations must be true: (a) firm $\bar{j}$ ’s payoff is strictly increased; or (b) the information partition of firm $\bar{j}$ gets strictly finer. To see this, suppose firm $\bar{j}$ gets no higher payoff when Case 1.1 is triggered. Since firm $\bar{j}$ is the blocking firm, her expected payoff from being rematched with α at $\bar{p}$ must be strictly higher than her status quo payoff, i.e., the one she gets when she was matched with $\mu^{-1}(\bar{j})$ at $\mathbf{p}_{\mu^{-1}(\bar{j}), \bar{j}}$. Since firm $\bar{j}$ expects strictly higher payoff but ends up getting no higher payoff (from being rematched with α at $\bar{p}$), firm $\bar{j}$ must have imprecise information about worker α ’s type. Therefore, after being rematched, firm $\bar{j}$ ’s partition gets strictly finer because she has learned the true type of α . The argument for Case 1.2 is identical to that for Case 1.1, except that we need to replace α with i . As a result, we only need to focus on situation (a), and thus the situation where the payoff of any firm in $\mathcal{A}$ never decrease (essentially) and that of at least one firm in $\mathcal{A}$ strictly increases. Since we have only finitely many firms

in A , the constructed path can be infinite only if some firm's payoff is improved indefinitely. However, this contradicts individual rationality of workers and that the surpluses are uniformly bounded across all matches.

Second, we argue that if the path terminates by triggering Case 0, then the partition of the blocking firm j in Case 0 must be strictly finer than her partition under the preceding state. Indeed, if firm j “regrets” leaving worker i for worker α , it must be the case that firm j has not learned worker α 's type prior to her match with α . Hence, firm j 's partition becomes strictly finer after observing the type of α .

The rest of the proof of Lemma 2, Lemma 3, and Theorem 2 does not change in the Bayesian setting.

REFERENCES

- Abdulkadiroğlu, Atila and Tayfun Sönmez (2003), “School choice: A mechanism design approach.” *American Economic Review*, 729–747. [29]
- Anderson, Axel and Lones Smith (2010), “Dynamic matching and evolving reputations.” *The Review of Economic Studies*, 77, 3–29. [32]
- Aumann, Robert J. (1976), “Agreeing to disagree.” *The Annals of Statistics*, 1236–1239. [39]
- Balinski, Michel and Tayfun Sönmez (1999), “A tale of two mechanisms: Student placement.” *Journal of Economic Theory*, 84, 73–94. [29]
- Bikhchandani, Sushil (2017), “Stability with one-sided incomplete information.” *Journal of Economic Theory*, 168, 372–399. [30, 32, 50, 51]
- Chakraborty, Archishman, Alessandro Citanna, and Michael Ostrovsky (2010), “Two-sided matching with interdependent values.” *Journal of Economic Theory*, 145, 85–105. [32]
- Chen, Bo, Satoru Fujishige, and Zaifu Yang (2010), “Decentralized market processes to stable job matchings with competitive salaries.” KIER Discussion paper, 749. [32, 41]
- Chen, Bo, Satoru Fujishige, and Zaifu Yang (2016), “Random decentralized market processes for stable job matchings with competitive salaries.” *Journal of Economic Theory*, 165, 25–36. [32, 42]
- Chen, Yi-Chun and Gaoji Hu (2017), “A theory of stability in matching with incomplete information.” Working paper. [33]
- Crawford, Vincent P. and Elsie Marie Knoer (1981), “Job matching with heterogeneous firms and workers.” *Econometrica*, 49, 437–450. [33, 37, 42]
- Dutta, Bhaskar and Rajiv Vohra (2005), “Incomplete information, credibility and the core.” *Mathematical Social Sciences*, 50, 148–165. [32]
- Ehlers, Lars and Jordi Massó (2007), “Incomplete information and singleton cores in matching markets.” *Journal of Economic Theory*, 136, 587–600. [32]

- Ehlers, Lars and Jordi Massó (2015), “Matching markets under (in) complete information.” *Journal of Economic Theory*, 157, 295–314. [32]
- Fudenberg, Drew and David K. Levine (1998), *The Theory of Learning in Games*, volume 2. MIT press. [41]
- Fujishige, Satoru and Zaifu Yang (2017), “On a spontaneous decentralized market process.” *The Journal of Mechanism and Institution Design*, 2, 1–37. [32]
- Gale, David and Lloyd S. Shapley (1962), “College admissions and the stability of marriage.” *The American Mathematical Monthly*, 69, 9–15. [31, 49]
- Grinstead, Charles M. and James Laurie Snell (1997), *Introduction to Probability*. American Mathematical Society. [43]
- Kelso, Alexander S. Jr and Vincent P. Crawford (1982), “Job matching, coalition formation, and Gross substitutes.” *Econometrica: Journal of the Econometric Society*, 50, 1483–1504. [42]
- Klaus, Bettina and Flip Klijn (2007), “Paths to stability for matching markets with couples.” *Games and Economic Behavior*, 58, 154–171. [32]
- Knuth, Donald Ervin (1976), *Mariages Stables et Leurs Relations Avec d'Autres Problèmes Combinatoires*. Presses de l'Université de Montréal. [30, 41]
- Kojima, Fuhito and M. Utku Ünver (2008), “Random paths to pairwise stability in many-to-many matching problems: A study on market equilibration.” *International Journal of Game Theory*, 36, 473–488. [32]
- Lauermann, Stephan and Georg Nöldeke (2014), “Stable marriages and search frictions.” *Journal of Economic Theory*, 151, 163–195. [51]
- Lazarova, Emiliya and Dinko Dimitrov (2017), “Paths to stability in two-sided matching under uncertainty.” *International Journal of Game Theory*, 46, 29–49. [32]
- Liu, Qingmin (2017), “Stable belief and stable matching.” Working paper. [51]
- Liu, Qingmin, George J. Mailath, Andrew Postlewaite, and Larry Samuelson (2014), “Stable matching with incomplete information.” *Econometrica*, 82, 541–587. [29, 30, 51]
- Ma, Jinpeng (1996), “On randomized matching mechanisms.” *Economic Theory*, 8, 377–381. [32]
- Mailath, George J., Andrew Postlewaite, and Larry Samuelson (2013), “Pricing and investments in matching markets.” *Theoretical Economics*, 8, 535–590. [33]
- Mailath, George J., Andrew Postlewaite, and Larry Samuelson (2017), “Premuneration values and investments in matching markets.” *The Economic Journal*. [33]
- Pomatto, Luciano (2019), “Stable matching under forward-induction reasoning.” Working paper. [32]
- Roth, Alvin E. (1989), “Two-sided matching with incomplete information about others' preferences.” *Games and Economic Behavior*, 1, 191–209. [32, 35]

- Roth, Alvin E. (2008), “Deferred acceptance algorithms: History, theory, practice, and open questions.” *International Journal of Game Theory*, 36, 537–569. [29, 31]
- Roth, Alvin E. and Marilda A. Oliveira Sotomayor (1990), *Two-Sided Matching: A Study in Game-Theoretic Modeling and Analysis*. Cambridge University Press. [31]
- Roth, Alvin E. and John H. Vande Vate (1990), “Random paths to stability in two-sided matching.” *Econometrica*, 58, 1475–1480. [31, 32, 41, 42, 51]
- Shapley, Lloyd S. and Martin Shubik (1971), “The assignment game I: The core.” *International Journal of Game Theory*, 1, 111–130. [29, 33]
- Wilson, Robert (1978), “Information, efficiency, and the core of an economy.” *Econometrica: Journal of the Econometric Society*, 46, 807–816. [32]
- Yenmez, M. Bumin (2013), “Incentive-compatible matching mechanisms: Consistency with various stability notions.” *American Economic Journal: Microeconomics*, 5, 120–141. [32]

Co-editor Simon Board handled this manuscript.

Manuscript received 3 December, 2017; final version accepted 20 June, 2019; available online 8 July, 2019.