

Institute for Employment
Research

The Research Institute of the
Federal Employment Agency



IAB-Discussion Paper

15/2012

Articles on labour market issues

Assortative matching through signals

Friedrich Poeschel

Assortative matching through signals

Friedrich Poeschel (IAB)

Mit der Reihe „IAB-Discussion Paper“ will das Forschungsinstitut der Bundesagentur für Arbeit den Dialog mit der externen Wissenschaft intensivieren. Durch die rasche Verbreitung von Forschungsergebnissen über das Internet soll noch vor Drucklegung Kritik angeregt und Qualität gesichert werden.

The “IAB Discussion Paper” is published by the research institute of the German Federal Employment Agency in order to intensify the dialogue with the scientific community. The prompt publication of the latest research results via the internet intends to stimulate criticism and to ensure research quality at an early stage before printing.

Contents

Abstract	4
Zusammenfassung	4
1 Introduction	5
2 Related literature	7
3 Model	9
4 Equilibrium	13
4.1 Definition of equilibrium	13
4.2 Putative equilibrium	14
5 Existence proof for the putative equilibrium	15
5.1 Bargaining	15
5.2 Participation and steady state	17
5.3 Signals and beliefs	18
5.4 Marketplace choice and creation	24
6 Equilibrium properties	25
6.1 Efficiency	25
6.2 Stability	26
6.3 Uniqueness	27
6.4 Discussion	28
7 Conclusions	29
A Omitted proofs	31
References	36

Abstract

The matching of likes is a frequently observed phenomenon. However, for such assortative matching to arise in a search model, often implausibly strong conditions are required. This paper shows that, once signals are introduced, a search model can generate even perfect assortative matching under weak conditions: supermodularity of the match production function is a necessary and sufficient condition. It simultaneously drives sorting and functions as a single-crossing property ensuring that agents choose truthful signals. The information thereby transmitted allows agents to avoid all unnecessary costs of random search, which creates in effect an almost frictionless environment. Hence the unique separating equilibrium in the model achieves nearly unconstrained efficiency despite frictions.

Zusammenfassung

Verbindungen zwischen ähnlichen Subjekten sind ein häufig beobachtetes Phänomen. Damit solche Sortierungen in einem Suchmodell auftreten, müssen oft jedoch überraschend starke Bedingungen erfüllt sein. Diese Studie zeigt, dass ein um Signale erweitertes Suchmodell sogar vollkommene Sortierungen unter schwachen Bedingungen generieren kann: Supermodularität der Produktionsfunktion für die Verbindung ist notwendige und hinreichende Bedingung. Sie wirkt zugleich als Ursprung der Sortierung und als "single-crossing"-Eigenschaft, die die Subjekte zutreffende Signale wählen lässt. Die dadurch verbreiteten Informationen erlauben es den Subjekten, alle unnötigen Kosten einer zufallsgeleiteten Suche zu vermeiden, so dass effektiv eine Umgebung ohne Friktionen entsteht. Daher zeichnet sich das einzige Separationsgleichgewicht des Modells durch nahezu uneingeschränkte Effizienz trotz Friktionen aus.

JEL classification: J64, D83, C78

Keywords: Assortative matching, sorting, search, signals

Acknowledgements: This paper is based on chapter 3 of my thesis at the University of Oxford. I would like to thank especially Godfrey Keller for many helpful discussions. I am grateful to Hiroyuki Adachi, Siddhartha Bandyopadhyay, Melvyn Coles, Bruno Decreuse, Jan Eeckhout, Pieter Gautier, Elke Jahn, Leo Kaas, Philipp Kircher, Barbara Mendolicchio, Franco Peracchi, Christopher Pissarides, Lones Smith, Margaret Stevens and to seminar participants at the Royal Economic Society (RES) Easter School 2008, the Econometric Society World Congress 2010, the XXth Aix-Marseille Doctoral Spring School, the 2012 ADRES doctoral conference, the RES Annual Conference 2012 and at the universities of Hannover, Oxford, Rome "Tor Vergata", Konstanz, Hamburg, Munich and at the Ecole Polytechnique for very useful comments and suggestions. The paper was partly written at the Ecole normale supérieure, and I thank this institution for its hospitality. Financial support from the ESRC (PTA-031-2004-00250) and from the EU's 6th Framework Programme are gratefully acknowledged. All errors are my own.

1 Introduction

A number of important markets, notably the labour market, bring agents together in pairs that form durable matches. The dominant approach in the analysis of decentralised interaction on such matching markets is the search and matching model (see e.g. the survey by Rogerson/Shimer/Wright (2005)): agents engage in sequential search for match partners, and frictions in the search process make this costly. In recent years, it has been intensely debated whether the model can replicate important empirical regularities. Alongside volatility over the business cycle and price or wage dispersion, the literature has sought to explain the pattern of sorting among heterogeneous agents. Indeed, that likes tend to match with likes is a pervasive empirical phenomenon, known as *positive assortative matching* (PAM). For example, more productive workers tend to be hired by more productive firms and more educated women tend to marry more educated men.¹ Can the search and matching model generate this phenomenon under plausible conditions?

For a decade, the answer appeared to be No. In their influential contribution, Shimer/Smith (2000) identified three conditions that have to hold simultaneously for PAM to arise in their search and matching model. One of these conditions requires that the match production function, which specifies how matched agents' inputs translate into output, be *supermodular*. It holds when one input's marginal effect on output is increasing in the other input. This mild and intuitive condition suffices in Becker's (1973) seminal model without frictions where it even generates perfect PAM: only equal types match.

The presence of frictions in Shimer/Smith (2000), however, seems to necessitate two additional (and less intuitive) conditions: in addition to the match production function, the logarithm of its first derivative and the logarithm of its cross-partial derivative also have to be supermodular for PAM to arise. In a comparable search and matching model analysed by Smith (2006), the logarithm of the match production function has to be supermodular, so that one input's marginal effect on output is also as a proportion increasing in the other input. Eeckhout/Kircher (2010) show that the conditions in Shimer/Smith (2000) are jointly at least as strong as in Smith (2006) (provided match production is non-decreasing in inputs). At the same time, both these strong conditions generate PAM only in the sense that the lowest and highest types an agent would match with are non-decreasing in her own type.

The combination of three conditions in Shimer/Smith (2000) is criticised by Atakan (2006) as "quite restrictive"; Goldmanis/Ray/Stuart (2009) find it "quite troubling" that the mild condition from Becker (1973) does not suffice to generate PAM in Smith (2006).² Goldmanis/Ray/Stuart (2009) also point out that situations in which there is hardly any or no sorting at all satisfy the formal criterion for PAM that Shimer/Smith (2000) and also Smith (2006) employ. In short, there is a paradox here: agents in many real-world markets sort into PAM, but theoretical models of these markets require improbably strong conditions to generate even a weak form of PAM.

¹ As an exemplary reference for these stylised facts, see Mare (1991).

² See Atakan (2006), p. 667 and Goldmanis/Ray/Stuart (2009), p. 8 and p. 12.

This paper offers a solution to the paradox by appealing to another pervasive phenomenon: the use of signals to transmit information. For example, job advertisements and applications might serve as signals on the labour market. Provided they do transmit information, such signals reduce the effect of frictions: search with better information takes less time and can avoid costly but unsuccessful meetings. What we have in mind is a website where it takes but a click to view the next signal, and meetings happen only in case of mutual interest. Whenever signals make search essentially costless, the situation approaches the frictionless model in which mild conditions suffice for PAM to arise.

In this spirit, the paper introduces signals into a search and matching model close to that in Shimer/Smith (2000). The necessary and sufficient condition for PAM in our model is supermodularity of the match production function. Under this condition, the model generates even perfect PAM. Hence we obtain exactly the same mild condition as in Becker's (1973) frictionless model and the same extreme form of sorting, despite the presence of frictions in our model. This paper therefore reconciles the search and matching model with potentially very many decentralised markets that exhibit PAM in practice but do not meet the strong theoretical conditions required in Shimer/Smith (2000).

For the signals in our analysis, we cannot assume a single-crossing property as in Spence (1973), in the sense that some types can send a certain signal at lower cost than other types. After all, when agents use applications or advertisements, writing a forged CV is as costly as writing a truthful CV, and painting an advertised job in unduly bright colours is as costly as honestly laying out its dull nature. Hence, the costs of signals in this paper are normalised to zero. As shown by Menzio (2007) for a directed search model with strategic bargaining, signals can still be informative in such an environment of *cheap talk*. The model in this paper also features strategic bargaining, but in contrast to Menzio (2007) it even achieves full information transmission. Supermodularity is central to this result because it introduces a single-crossing property into agents' marginal productivity, rather than into the cost of signals as in Spence (1973). Hence the conditions for PAM and for truthful signals exactly coincide in our model.

To obtain truthful signals, one has to show in particular that low types do not want to imitate high types. If a low type in our model signals like a high type, meets a high type, and then reneges, bargaining will fail. The high type then prefers meeting another agent to a second round of bargaining with the low type. When reneging is therefore not an option, the low type has to conceal the difference between expected and actual match production by reducing her share accordingly. If the match production function is supermodular, this reduction will outweigh the gain from higher match production with a high type. Hence no-one deviates from truthful signals and perfect PAM is to be expected: with fully informative signals, agents can replace (almost all) costly search via meetings by costless search via signals and can therefore behave like in a frictionless setting.

The separating equilibrium we can thus identify is the unique separating equilibrium in the model, and it has a number of desirable efficiency properties. Above all, agents match at the very first opportunity, so that no time (or money) is wasted on unsuccessful search. In labour market terms, this would mean that frictional unemployment is reduced to a mini-

mum. The signals allow agents to first locate their best feasible match partner at no cost, so that the only costs are created by the meeting with this agent. The source of mismatch in random search models, the incentive to accept less than the best feasible match so as to avoid further search costs, is therefore absent here. As there is no mismatch, the equilibrium matching is stable and maximises aggregate output. Overall, unconstrained efficiency is nearly achieved here because the frictions are in effect overcome. Only the costs associated with the single meeting that precedes each match distinguish the separating equilibrium from the first-best outcome that could be centrally imposed by a benevolent social planner.

The paper proceeds as follows. After further related literature has been discussed in section 2, section 3 specifies a frictional matching market and the procedures of search. Section 4 defines equilibrium in the model and proposes a separating equilibrium in which supermodularity suffices for perfect PAM. Its existence is proven step by step through a series of lemmas in section 5. The separating equilibrium is found to be unique as well as efficient in section 6 before section 7 concludes.

2 Related literature

There is first evidence that models with more information in the search process only require weaker conditions for PAM. In a recent contribution, Eeckhout/Kircher (2010) investigate PAM in a model of *directed search*: sellers post offers and commit to them; having observed the offers, buyers then simultaneously choose which seller to visit. For common meeting technologies, PAM will arise in this context if the square root of the match production function is supermodular. This condition is weaker than in Shimer/Smith (2000) and Smith (2006), but still stronger than in Becker (1973) and this paper. In a related analysis, Shimer (2005) discusses PAM in a directed search model of the labour market. He finds, at least for the case of only two worker types, that there will be some stochastic form of PAM as long as workers of low type do not have a comparative advantage when working for employers of high type.

For a number of reasons, however, a directed search model is not easily comparable to the search and matching model in Shimer/Smith (2000). Above all, the frictions differ. The only way frictions are introduced in Eeckhout/Kircher (2010) and Shimer (2005) is through congestion: buyers cannot coordinate, so that queues result and only some buyers can buy. In Shimer/Smith (2000) and also in Smith (2006), the frictions are instead due to agents' discounting. Directed search also differs from Shimer/Smith (2000) in how agents split output. The literature refers to *non-transferable utility* (NTU) whenever agents divide output according to a pre-imposed split. In cases of *transferable utility* (TU), there is no pre-imposed split and agents have to bargain. Shimer/Smith (2000) and Becker (1973) feature TU. Smith (2006) features NTU; he also shows that, with NTU, a frictionless model à la Becker (1973) leads to PAM even without supermodularity. Overall, to generate PAM, models with TU appear to require conditions at least as strong as those in models with NTU. On this background, it might be important that directed search features NTU: by committing to a price (wage), the sellers (employers) impose a split ex ante.

Because of such differences, it is not clear whether directed search can be regarded as a solution to the paradox we have outlined. Similar conclusions apply to Morgan (1998) and Atakan (2006): while supermodularity as such gives rise to PAM in both models,³ the only frictions in these models are explicit costs that agents pay out of pocket for each meeting (as opposed to implicit costs from discounting). Yet limiting oneself to explicit costs cannot resolve the paradox, since real-world agents do discount. By contrast, the model in this paper remains very close to Shimer/Smith (2000): it notably features TU and discounting, while allowing in addition for explicit costs. This set-up corresponds in particular to many labour market contexts.

The key difference between Shimer/Smith (2000) and this paper is that we allow for ex ante information transmitted through signals, so that search becomes non-random. The focus in our analysis is on links between complementarities in match production and agents' incentives to signal truthfully. Eeckhout/Kircher (2010) instead focus on links between these complementarities and agents' individual matching rates. By assuming commitment to posted offers (which is crucial for their analysis), they abstract from the issue of truthful signals; in turn, we abstract in effect from differences in matching rates by allowing for any number of marketplaces with constant returns to scale.

Some more papers consider sorting in the context of a matching market with signals. Hoppe/Moldovanu/Sela (2009) and Hopkins (2012) build two similar models of a *matching tournament* with signalling: match partners are essentially prizes for ex-ante choices of costly signals. In both models, agents first select a costly signal of their privately observed type like in Spence (1973) and then match roughly like in Becker (1973). Hopkins (2012) assumes a single-crossing property and Hoppe/Moldovanu/Sela (2009) assume a specific multiplicative match production function that satisfies log-supermodularity. In the symmetric equilibrium, agents' signals are then strictly increasing in their types. This leads to perfect PAM at the matching stage - just as one would have expected, given Becker's (1973) findings. However, since search frictions do not exist in matching tournaments, neither of the two papers helps us resolve the paradox in Shimer/Smith (2000).

A search model built by Chade (2006) features discounting and noisy signals uncontrolled by the agents. Yet these signals are not observed before agents meet. Rather, when agents do meet, they do not observe each others' true types but only the noisy signal. Hence search is still random in this model, and the noisy signals in fact add information frictions to search frictions. Assuming that the noisy signals exogenously carry some information, matching is shown to exhibit PAM in a very weak sense: the distribution of types that a high type might match with first-order stochastically dominates this distribution for a low type. This paper primarily differs from Chade (2006) in that signals are observed before meetings and allow agents to avoid search costs, thereby tending to reduce the effect of frictions. Moreover, signals in our model are not informative by assumption but are deliberately and strategically chosen by agents. We would argue that real-world agents will exert as much control over the signals as possible, given how important they can be for their payoffs.

³ In Atakan (2006), this result crucially depends on search costs being identical for all agents.

A contribution by Lentz (2010) does not feature any signals but allows for search on the job (more generally, search while matched) with endogenous search intensity in a model with TU and discounting. While search is still random, higher types gain more than others from search on the job if the match production function is supermodular. Higher types thus search more intensively, which leads to PAM in terms of stochastic dominance as in Chade (2006). A related model by Goldmanis/Ray/Stuart (2009) features NTU, discounting, and search on the job. If only one agent in each match can switch to another match, PAM will result if the match output function is log-supermodular. If either agent can switch, a condition has to be met so that agents gain from switching to matches with higher types. Then the situation would evolve into perfect PAM over time, were it not always set back as matches randomly break up and agents begin climbing up the ladder anew. Therefore, perfect PAM is only achieved in the limit as the rate of random break-ups tends to zero. In any case, agents in Lentz (2010) and Goldmanis/Ray/Stuart (2009) sort only over time. By contrast, the fundamentally different sorting mechanism in our model can explain PAM already among graduates in their first job, without invoking stronger conditions.

Our model is finally related to Jacquet/Tan (2007). They consider an environment with discounting, NTU, and a particular log-supermodular match production function. For such an environment, Burdett/Coles (1997) found that types segregate into classes and match exclusively within these classes. Building on this, Jacquet/Tan (2007) let agents establish any number of marketplaces and show that each marketplace is populated by only one class in equilibrium (while perfect PAM cannot be achieved). By going to the appropriate marketplace, each agent can thus avoid meetings that do not lead to a match and can instead match after the first meeting. Agents can do the same in our model if signals are informative: they can use the signals to create any number of marketplaces, which might be websites. While it is left open in Jacquet/Tan (2007) how marketplaces are established, in this paper they are established simply by requiring certain signals. Apart from this difference, our environment features TU and a general match production function.

3 Model

The market in our model consists of heterogeneous agents who match among themselves. Agents are indexed by a discrete productivity type $x \in \Theta$, where $\Theta = \{\underline{x}, \dots, \bar{x}\}$ with $\underline{x} > 0$. For each discrete type, there is a continuum of agents and the overall mass of agents is normalised to 1. The measure of agents with types weakly below $x \in \Theta$ is denoted $L(x)$, where $L(\cdot)$ is a cumulative distribution function with probability mass function $l(\cdot)$. The mass of agents of type x is thus given by $l(x)$, and we require $l(x) > 0, \forall x$.

Time is continuous with an infinite horizon. Each agent is always in one of four states: matched, searching (that is, unmatched but participating), waiting (for continued bargaining, as explained below), and not participating. We denote the mass of waiting agents of type x by $\kappa(x) \leq l(x)$ and that of non-participating agents by $\nu(x) \leq l(x)$. Searching agents can create marketplaces to meet on. We index the N marketplaces agents use by n , and N may be countably infinite. Agents cannot be on several marketplaces simultaneously (i.e. their search activity is indivisible), but they can always switch between

marketplaces without incurring any cost. When they match they immediately leave the marketplace. Let $u^n(x) \leq l(x)$ represent the mass of searching agents of type x on marketplace n , while the mass of matched agents is thus $l(x) - \sum_{n=1}^N u^n(x) - \kappa(x) - \nu(x)$. All these quantities are determined endogenously. When indifferent whether to engage in search, whether to accept a match, and whether to stay in a marketplace or switch, an agent respectively searches, accepts the match, and stays.

Types are exogenously given, but only privately observable. Every searching agent chooses a costless signal $\tilde{x} \in \Theta$ to convey information about her type to other agents on the same marketplace. The signal $\tilde{x}$ may or may not be equal to her true type x , and it can always be instantly and costlessly changed. By contrast, agents who are matched, waiting, or non-participating do not send any signal and are not on any marketplace.

Since searching agents can condition meetings on signals, meetings are non-random. Agents can influence whom they meet through their choice of marketplace: each marketplace n is characterised by a set R^n of required signals, such that each agent who chooses this marketplace and sends a signal $\tilde{x} \in R^n$ can meet all other agents on the marketplace who also send a required signal. We let R^n be public information, as agents can in any case very quickly infer R^n from the signals they observe on marketplace n .

Inside a marketplace, meetings are random and are described by a meeting function $m(\cdot)$. With a mass of agents

$$\lambda^n = \sum_{x \in \Theta} u^n(x)$$

the flow of meetings in marketplace n equals $m(\lambda^n) \leq \lambda^n$, and $m(0) = 0$. The meeting rate on the marketplace is

$$\eta^n = \begin{cases} \frac{m(\lambda^n)}{\lambda^n} & \text{if } \tilde{x} \in R^n \text{ and } \lambda^n \neq 0 \\ 0 & \text{if } \tilde{x} \notin R^n \text{ or } \lambda^n = 0 \end{cases} \quad (1)$$

We assume constant returns to scale in meeting, so that agent x faces the same meeting rate $\eta^n = \eta$ across all N marketplaces, provided she always chooses a required signal. Then x must choose her marketplace by the agents she wants to meet, as she would meet all agents equally quickly. When she is indifferent, she randomises over her most preferred marketplaces. Finally, any marketplace can be created at no cost but must attract agents in order to last. The agent(s) creating marketplace n choose R^n , which cannot be changed thereafter.

Before two agents can match, a meeting between them will have to occur. To distinguish between the agents, we will denote one's type by x and the other's by y . Normalising the flow output generated by an unmatched agent to zero, a match between types x and y generates a constant flow output $f(x, y)$.

Assumption 1 (Regularity conditions). *The match production function $f(\cdot, \cdot)$ is symmetric ($f(x, y) \equiv f(y, x)$), strictly increasing in both arguments, and, for any existing types, takes only positive values ($f : \Theta^2 \mapsto \mathbb{R}_{++}$).*

Because productivity types are scalars, there can only be gains from specialisation in production if one role in production rewards productivity more than the other role does. This specialisation will remain possible despite the assumption of symmetry if the more productive agent always assumes the role that rewards productivity more, so that the output of the match is maximised.

By observing a signal $\tilde{y}$ agent x can only form a belief about the true type y behind the signal. Agents never directly observe each other's actual types. Let h^x be the history of the interaction with some agent as observed by x , i.e. a set of actions such as the observed signal. We represent a belief as a probability distribution $\Psi(\cdot)$. Concretely, for each h^x , the belief held by agent x of the other agent's true type y is the probability distribution $\Psi(\cdot|h^x)$ over Θ . Then x , having observed h^x , believes that the other's type is y with probability mass $\psi(y|h^x)$. All agents use Bayes' rule to form and update their beliefs.

Match output must also be unobservable when types are unobservable: knowing $f(\cdot, \cdot)$, x could otherwise infer y from the observed output $f(x, y)$. To keep the notation simple, let $g(x|h^x)$ denote the match output that x expects based on her belief after observing h^x :

$$g(x|h^x) = \sum_{y \in \Theta} f(x, y) \psi(y|h^x)$$

Agents in a meeting bargain over the division of the match output that they would produce between them. We model this using a strategic bargaining procedure with alternating offers. It is useful to imagine that $f(x, y)$ is contained in a pot that agents can only take from but cannot look into.⁴ When agents first meet, either of them is randomly selected with probability $\frac{1}{2}$ to move first. An agent x who moves first takes a share $\pi(x|y)$ for herself from the pot, which is not observed by y . Then y takes the remainder $f(x, y) - \pi(x|y)$ from the pot (which may be negative), unobserved by x . Now y has three options: she can accept the share left for her in the pot, reject this share but stay, or reject it and walk away to immediately continue searching.⁵ If y rejects but stays, x can walk away; otherwise, the same two agents meet again for the next round of bargaining, in which one agent is again selected with probability $\frac{1}{2}$ to move first, and so on. Shares offered in previous rounds cannot be accepted ex post, and if players never agree nor walk away, both will obtain 0. If y accepts, agents match immediately and obtain their respective share as a flow utility for the duration of the match. As each agent can assure herself flow utility 0 by not participating, negative shares will always be rejected and thus never arise as a flow utility. Finally, matches dissolve exogenously at constant rate δ .

In order to be more precise, let us formalise this bargaining game. The players are 'nature' Q and the agents x and y who meet. The history h^x records the actions that x has observed thus far, h^y records what y has observed, and we simply index histories in chronological

⁴ The function of this pot is to ensure that agents do not agree on shares that sum to more than $f(x, y)$. Alternatively, one can let agents make any agreement and note that it will break down when $f(x, y)$ cannot satisfy the demands agreed. As such a break-down would occur immediately, it would effectively be the same as bargaining failure. The results with this approach would be the same.

⁵ The fact that agents have met implies that these agents prefer engaging in search to not participating. It is thus without loss of generality that non-participation is not a further outside option here.

order. When x and y first meet, they already know both signals, so that $h_1^x = h_1^y = \{\tilde{x}, \tilde{y}\}$. A player function $P(\cdot, \cdot)$ assigns to each history pair (h^x, h^y) (except any terminal history pairs) a player who moves at this history pair. $P(h_1^x, h_1^y) = Q$ selects x and y each with probability $\frac{1}{2}$ to move first. If x is selected, then $h_2^x = h_2^y = h_1^x \cup \{x\}$ and $P(h_2^x, h_2^y) = x$. Agent x now chooses some $\pi(x|y)$ according to her bargaining strategy $B(x)$ that assigns an action to every possible history pair for which $P(h^x, h^y) = x$. As y observes only the remainder, $h_3^y = h_2^y \cup \{f(x, y) - \pi(x|y)\}$ while $h_3^x = h_2^x \cup \{\pi(x|y)\}$, and y responds according to $B(y)$ by choosing an action from the set {"accept", "reject but stay", "reject and walk away"}. If she chooses "accept" or "reject and walk away", then (h_4^x, h_4^y) will be a terminal history pair. If she chooses "reject but stay", then $h_4^x = h_3^x \cup \{\text{"reject but stay"}\}$ and $P(h_4^x, h_4^y) = x$ who chooses from {"continue", "walk away"}. If x does not walk away, bargaining will continue in the next meeting with $P(h_5^x, h_5^y) = Q$, and so on.

Assumption 2 (Common expected delay). *A further meeting with the same agent arrives at the same rate as a new meeting with another agent.*

That is, a time $1/\eta$ elapses in expectation before another bargaining round, so that both a meeting to continue bargaining and a meeting with a different agent arrive at rate η . While there is nothing in our model that would cause these meeting rates to differ systematically, assumption 2 is obviously a simplification. Recall that waiting for another bargaining round is a separate state that an agent can be in, so that waiting agents neither send a signal nor meet anyone else.

All agents are risk-neutral, discount future utility at discount rate r (with $0 < r < \infty$), and seek to maximise the present discounted value (pdv) of their expected utility. Throughout the paper, 'payoff' refers to the pdv, not to the flow utility. Because of discounting, the time that elapses before a meeting makes meetings costly. In addition, we include a second kind of search friction by allowing for explicit cost $c \geq 0$ that an agent incurs each time she attends a meeting.

Assumption 3 (Gains from trade). *The output produced in a match between two agents of the lowest type, discounted at effective discount rate $r + \delta$, can reimburse both agents' explicit costs of one meeting:*

$$2c \leq \frac{f(\underline{x}, \underline{y})}{r + \delta}$$

While explicit costs always have to be limited relative to the available payoffs to ensure agents' participation, note that assumption 3 is particularly mild. For example, we do not assume that each agent is in fact reimbursed in the event of a match, nor that match output is sufficient to reimburse the costs of the expected number of meetings before a match. Finally, agents know everything except the true type of any other agent and, by consequence, the actual match output $f(x, y)$.

4 Equilibrium

4.1 Definition of equilibrium

We begin by defining three expected present values: $U^n(x)$ as the value to x of searching in marketplace n , $V(x|y)$ as the value to x of waiting for another bargaining round with y , and $W(x|y)$ as the value to x from being matched with y . Let the set $A(h^x, h^y)$ comprise of all combinations of bargaining strategies $(B(x), B(y))$ that lead to a subgame-perfect equilibrium (SPE) of the bargaining game given history pair (h^x, h^y) , so that an agreement is reached immediately and agents match. Let $\alpha(\cdot, \cdot)$ be an indicator function such that $\alpha(B(x), B(y)) = 1$ if $(B(x), B(y)) \in A(h^x, h^y)$ and 0 otherwise. In exact analogy, also define $\Omega(h^x, h^y)$ as the set of bargaining strategies that lead to another round of bargaining given (h^x, h^y) , and $\omega(\cdot, \cdot)$ as an indicator function such that $\omega(B(x), B(y)) = 1$ if $(B(x), B(y)) \in \Omega(h^x, h^y)$. Then the following *asset equation* expresses, for one marketplace, the expected return on searching as the expected gain from a meeting net of search cost c :

$$rU^n(x) = \eta^n \left(-c + \sum_{y \in \Theta} \alpha(B(x), B(y)) [W(x|y) - U^n(x)] \psi(y|h^x = \{\tilde{y} \in R^n\}) + \sum_{y \in \Theta} \omega(B(x), B(y)) [V(x|y) - U^n(x)] \psi(y|h^x = \{\tilde{y} \in R^n\}) \right) \quad (2)$$

where $\psi(y|h^x = \{\tilde{y} \in R^n\})$ is the probability mass of y that x believes conditional on meeting y in marketplace n (and thus after observing a required signal $\tilde{y} \in R^n$). Whenever x does not send a required signal herself, $rU^n(x) = 0$.

Let us define $U(x)$ as the value of $U^n(x)$ that x obtains in equilibrium. As is natural when signals are involved, we look for a *perfect Bayesian equilibrium* (PBE) of our model. We will focus our attention on separating equilibria that survive the Intuitive Criterion.⁶ Because signals are costless all PBE will necessarily be cheap-talk equilibria. A steady-state PBE of our model, separating or not, requires that the flows into and out of matches balance for every type (a pointwise steady state), that all agents choose all their strategies optimally, and that agents' beliefs are consistent with all agents' actual equilibrium behaviour.

Definition 1 (Search equilibrium with signals). *In a steady-state PBE of the model, each agent $x \in \Theta$*

- (i) *engages in search if and only if $U(x) \geq 0$*
- (ii) *optimally chooses a marketplace such that $\forall n, U(x) \geq U^n(x)$ given $B(x), B(y)$ for all $y \in \Theta$, and $(R^n)_{n=1}^N$, where $U^n(x)$ is determined by equation (2)*
- (iii) *chooses her signal optimally as $\arg \max_{\tilde{x}} rU^n(x)$ given $B(x), B(y)$, and R^n , noting that η^n depends on $\tilde{x}$ as specified by equation (1)*

⁶ Kübler/Müller/Normann (2008) report experimental evidence suggesting that pooling equilibria never arise when some types can benefit from the effective use of signals.

- (iv) chooses a stationary subgame-perfect bargaining strategy as $\arg \max_{B(x)} rU^n(x)$ given all $B(y)$ and R^n , noting that $W(x|y)$ depends on the share obtained in bargaining
- (v) holds beliefs that are formed using Bayes' rule where possible and that are consistent with equilibrium play: given an equilibrium history h^x , $\psi(y|h^x) = u^n(y|h^x)$ where $u^n(y|h^x)$ is the true probability mass of y in marketplace n conditional on h^x

and the matching market is in a pointwise steady state, so that the flows into and out of $\sum_{n=1}^N u^n(x) + \kappa(x)$ balance for each $x \in \Theta$. Marketplaces are created until there is no new marketplace n^0 such that $U^{n^0}(x) > U(x)$ holds for any $x \in \Theta$.

A PBE only requires agents' beliefs to be consistent with equilibrium play, not with actions out of equilibrium. As is well known, a PBE can therefore depend on unreasonable off-equilibrium beliefs because these beliefs are never tested in equilibrium. Since unreasonable beliefs are not needed for any of our results, we rule out beliefs that are unreasonable in the sense of the Intuitive Criterion. To do this formally, let us call the choices of n , $\tilde{x}$, and $B(x)$ the 'grand strategy' of agent x , denoted $GS(x) = (n, \tilde{x}, B(x))$. Also define $BR(x|h^x)$ as the set of continuation strategies $GS(x|h^x)$ that are best responses for x . To apply the Intuitive Criterion as an equilibrium refinement, we have to define the notion of equilibrium domination in our model:

Definition 2 (Equilibrium domination). Given a PBE of the model, the continuation strategy $GS(x|h^x)$ is equilibrium-dominated at history pair (h^x, h^y) if

$$U(x) > \max_{GS(y|h^y) \in BR(y|h^y)} U(x|GS(x|h^x))$$

where $U(x|GS(x))$ is the present value to x of searching with strategy $GS(x|h^x)$.

The Intuitive Criterion then demands that the beliefs of y place probability 0 on any type x who would have to pursue equilibrium-dominated strategies to reach the respective history: $\psi(x|h^y) = 0$ if, at a history up to h^y , x would have had to play an equilibrium-dominated strategy $GS(x|h^x)$.

4.2 Putative equilibrium

We next propose that a particular separating equilibrium exists under a simple condition on the match production function $f(\cdot, \cdot)$. All we need is a weak and intuitive form of complementarity known as strict supermodularity (or increasing differences): the marginal product of one agent in a match is strictly increasing in the type of the other agent.⁷

Definition 3 (Supermodularity). The match production function $f(\cdot, \cdot)$ is strictly supermodular if, for all $x_H > x_L$ and $y_H > y_L$,

$$f(x_H, y_H) - f(x_L, y_H) > f(x_H, y_L) - f(x_L, y_L)$$

⁷ A stronger form of complementarity is strict log-supermodularity, which is defined using $\ln f(\cdot, \cdot)$ instead of $f(\cdot, \cdot)$ in definition 3, so that the proportional marginal product of one agent is increasing in the other's type.

Further, we refer to the sorting with $x = y$ in all matches as *perfect positive assortative matching* (PPAM). We can now propose existence of the following PBE in our model:

Proposition 1 (Existence). *Let agents' beliefs place probability 0 on the occurrence of equilibrium-dominated actions. Then for any type distribution $L(x)$, strict supermodularity of $f(\cdot, \cdot)$ is necessary and sufficient for the existence of a separating PBE in which each agent $x \in \Theta$*

- (i) *engages in search: $U(x) \geq 0$*
- (ii) *chooses a marketplace where she meets exclusively agents of her own type*
- (iii) *signals truthfully: $\tilde{x} = x$*
- (iv) *reaches a bargaining agreement in the first meeting and thus matches:
 $\alpha(B(x), B(y)) = 1$ and $\omega(B(x), B(y)) = 0$ for $x = y$*
- (v) *correctly believes all signals to be truthful:
 $\psi(y|h^x = \{\tilde{y}\}) = u^n(y|h^x = \{\tilde{y}\}) = 1$ for all $y = \tilde{y}$.*

The market is in pointwise steady state and is perfectly segmented, so that there is only one type $x \in \Theta$ on each marketplace. The equilibrium matching is PPAM.

A proof of proposition 1 will thus establish that not only PAM, but even PPAM arises in our model under the same weak condition as in a frictionless model, although our model allows for two kinds of frictions. In a model with frictions only from discounting, Shimer/Smith (2000) establish PAM, albeit not PPAM, under the condition that the match production function $f(\cdot, \cdot)$, the logarithm of its first derivative, and the logarithm of its cross-partial derivative are all supermodular. These conditions are directly comparable to our condition and are unambiguously more restrictive: proposition 1 claims that our model achieves PPAM with just the first of these conditions, which is also the most intuitive.

The next section proves proposition 1 through a series of lemmas. Each time, we separately consider a component of proposition 1, taking as given that all other components are indeed as specified in proposition 1. We verify for the component in question, as applicable, that it is optimal for agents to behave as specified, that a steady state results, and that beliefs are consistent with equilibrium play.

5 Existence proof for the putative equilibrium

5.1 Bargaining

We first determine the expected present values in the putative equilibrium situation. Given that beliefs are consistent with equilibrium play, we have

$$\psi(y|h^x = \{\tilde{y} \in R^n\}) = u^n(y|h^x = \{\tilde{y} \in R^n\})$$

If x only meets agents of her own type, then

$$u^n(y|h^x = \{\tilde{y} \in R^n\}) = 0 \quad \forall y \neq x \quad (3)$$

Since every meeting in the putative equilibrium leads to match,

$$\alpha(B(x), B(y)) = 1 \quad \text{for } y = x \quad \text{and} \quad \omega(B(x), B(y)) = 0 \quad \text{for } y \neq x \quad (4)$$

For the marketplace chosen in the putative equilibrium, equation (2) thus simplifies to

$$rU(x) = \eta[W(x|y) - c - U(x)] \quad (5)$$

with $y = x$. Hence the rate of matches equals the rate of meetings, and an agent effectively incurs costs c each time she matches. Next, the expected return on being matched with y is the expected flow utility while matched and the loss from match dissolution at rate δ :

$$rW(x|y) = \sigma(x|y) - \delta[W(x|y) - U(x)] \quad (6)$$

where $\sigma(x|y)$ denotes the expected share that x obtains when bargaining with y over the flow of match output $f(x, y)$, which is in effect known from truthful signals:

$$\sigma(x|y) = \frac{1}{2}\pi(x|y) + \frac{1}{2}[f(x, y) - \pi(y|x)] \quad (7)$$

One can solve equation (5) for $U(x)$ and equation (6) for $W(x|y)$, then use the latter to substitute for $W(x|y)$ in the former to obtain

$$rU(x) = \beta[\sigma(x|y) - (r + \delta)c] \quad (8)$$

where $\beta = \eta/(r + \delta + \eta)$. Now suppose y has been randomly selected to move first in the bargaining game. In response to the share left for her, x can reject it and continue searching, which carries the value $U(x)$, or she can reject this share and wait for another round of bargaining, which carries a value $V(x|y)$. Note that the first mover y cannot hope to attain a better position than she currently has: at best, she will find herself as first mover again in a later meeting, be it with the same agent x or another agent of the same type. As delay is costly, y seeks to seize the opportunity and to ensure that x accepts her offer. In turn, x will accept any implicitly offered payoff $W^O(x|y)$ that satisfies

$$W^O(x|y) \geq \max[V(x|y), U(x)] \quad (9)$$

as she would otherwise reject the offer. When x moves first, y requires

$$W^O(y|x) \geq \max[V(y|x), U(y)] \quad (10)$$

In case of a second meeting, the same logic as before implies that the first mover seeks to ensure agreement, so that the second meeting can be expected to result in a match. By assumption 2, the second meeting happens at rate η , so that

$$rV(x|y) = \eta[W(x|y) - c - V(x|y)] \quad (11)$$

in the putative equilibrium. Solving equation (11) for $V(x|y)$ and equation (5) for $U(x)$ establishes that $V(x|y) = U(x)$, since x meets a type $y = x$ after an expected delay of $1/\eta$ in any case. Hence the outside option $U(x)$ is not binding. As we also require bargaining strategies to be stationary, the game reduces to a variant of Rubinstein's (1982) set-up, and we have the following result:

Lemma 1 (Bargaining equilibrium). *Given truthful signals and given marketplace choices as in the putative equilibrium situation, the following stationary strategies form the unique SPE of the bargaining game:*

(i) *for herself, agent x always proposes*

$$\pi^*(x|y) = \left(1 - \frac{\beta}{2}\right) f(x, y) + \beta(r + \delta)c \quad (12)$$

When y proposes $\pi(y|x)$, x always accepts if and only if $\pi(y|x) \leq \pi^(y|x)$.*

(ii) *for herself, y always proposes $\pi^*(y|x) = \pi^*(x|y)$. When x proposes $\pi(x|y)$, y always accepts if and only if $\pi(x|y) \leq \pi^*(x|y)$*

Agreement is reached in the first round of bargaining.

Proof. See appendix. $\square$

The essence of the bargaining SPE is that each agent makes offers that leave the other indifferent, and each agent accepts offers that make her indifferent or better off: the first-mover takes a share $\pi^*(x|y)$ such that the second-mover share

$$f(x, y) - \pi^*(x|y) = \frac{\beta}{2} f(x, y) - \beta(r + \delta)c$$

is just enough to prevent the second mover from rejecting. The second-mover share will still be weakly positive if

$$\frac{\beta}{2} f(x, y) \geq \beta(r + \delta)c \quad \Leftrightarrow \quad 2c \leq \frac{f(x, y)}{r + \delta}$$

which by assumption 3 even holds for $f(x, y) = f(\underline{x}, \underline{y})$. The two indifference conditions in equations (9) and (10), depending on who moves first, then together pin down a unique SPE. Finally, expected shares in the SPE are

$$\sigma(x|y) = \sigma(y|x) = \frac{1}{2}\pi^*(x|y) + \frac{1}{2}[f(x, y) - \pi^*(x|y)] = \frac{1}{2}f(x, y) \quad (13)$$

as one would expect when everything is symmetric.

5.2 Participation and steady state

To ensure that all agents engage in search, c must not be so high that $U(x)$ becomes negative for some x , since each agent can obtain a payoff 0 by not participating.

Lemma 2 (Participation). *Assumption 3 is necessary and sufficient for all agents to prefer engaging in search to non-participation.*

Proof. As match output is the only source of utility in the model, agents who do not engage in search obtain payoff 0. Then agent x will only engage in search if $U(x) \geq 0$. By equation (8), this requires

$$c \leq \frac{\sigma(x|y)}{r + \delta} \Leftrightarrow 2c \leq \frac{f(x, y)}{r + \delta}$$

using equation (13). If this holds for $f(\underline{x}, \underline{y})$, as stated in assumption 3, then it will also hold for the output generated in any other match because $f(x, y)$ is strictly increasing in x and y by assumption 1. $\square$

As agents prefer search to non-participation, $\nu(x) = 0, \forall x \in \Theta$. Moreover, recall that agents in the putative equilibrium reach an agreement in the first bargaining round, so that $\kappa(x) = 0, \forall x \in \Theta$. Hence agents only flow from searching to being matched (at rate η) and back (at rate δ). Equating these flows, we obtain the pointwise steady state in the putative equilibrium:

$$\delta \left[l(x) - \sum_{\mathcal{N}(x)} u^n(x) \right] = \eta \sum_{\mathcal{N}(x)} u^n(x) \quad \forall x \in \Theta \quad (14)$$

where $\mathcal{N}(x) \equiv \{n | R^n = \{x\}\}$ is the set of all marketplaces on which x meets exclusively her own type when signals are truthful.

5.3 Signals and beliefs

In this section, we examine whether any one agent has an incentive to unilaterally deviate from the putative equilibrium by choosing a different signal. Therefore, we take as given that all other agents signal truthfully, that all believe signals to be truthful, as well as the other components of the putative equilibrium. We proceed by identifying first the conditions under which every agent prefers her match in the putative equilibrium (henceforth the *equilibrium match*) to any other match that is available to her (i.e. mutually acceptable). From this, we infer under which conditions there will be no incentive to deviate from the truthful signal.

There are two reasons why we need to worry about false signals. First, because true types are only privately observable, agents can perfectly imitate agents of other types by sending their signal and bargaining as these types would. Second, agents might just imitate another type's signal and then renege on it in the meeting. Since search frictions make switching to another meeting costly, the other agent in the meeting might still accept the match. For example, consider a rather high type y_H who matches with x_H in the putative equilibrium. If y_H finds she has been lured into a meeting with a type $x_L < x_H$ by a false signal, she will nevertheless grudgingly accept whenever her share of $f(x_L, y_H)$ is not so far below her expected share of $f(x_H, y_H)$ that the costs of another meeting would be justified. Therefore, there can in principle be an incentive to send false signals.

We first compare the equilibrium match to matches with lower types. Without loss of generality, let us take the perspective of some agent with a type $x_H > \underline{x}$, so that lower types necessarily exist. We thus want to compare being matched with $y_H = x_H$ to being matched with $y_L < x_H$. The expected present discounted values of these matches are $W(x_H|y_H)$ and $W(x_H|y_L)$, respectively. In the spirit of the one-deviation principle, x reverts to the putative equilibrium strategies after the deviation. Hence, the asset equations for both $rW(x_H|y_H)$ and $rW(x_H|y_L)$ in analogy to equation (6) depend on the same $U(x_H)$ and thus differ only in the expected shares. Solving these two asset equations respectively for $W(x_H|y_H)$ and $W(x_H|y_L)$, we therefore find that

$$W(x_H|y_H) > W(x_H|y_L) \Leftrightarrow \sigma(x_H|y_H) > \sigma(x_H|y_L)$$

where $\sigma(x_H|y_H)$ and $\sigma(x_H|y_L)$ denote the expected share obtained by x_H in a match with y_H and y_L , respectively.

Thus suppose a type $x_H > \underline{x}$ signals to be of type x_L in order to meet a type y_L . Further suppose that agent x_H continues to behave like a type x_L so as to conform to the beliefs of y_L , given that all other agents signal truthfully. Recall from section 5.1 that neither agent's signal implies a binding outside option. Hence the bargaining equilibrium described by lemma 1 will be reached in the first round of bargaining. Then the expected flow utility for x_H in the match with y_L is

$$\begin{aligned} \sigma(x_H|y_L) &= \frac{1}{2} \left[f(x_H, y_L) - \frac{\beta}{2} f(x_L, y_L) + \beta(r + \delta)c \right] \\ &\quad + \frac{1}{2} \left[f(x_H, y_L) - \left(1 - \frac{\beta}{2}\right) f(x_L, y_L) - \beta(r + \delta)c \right] \\ &= f(x_H, y_L) - \frac{1}{2} f(x_L, y_L) \end{aligned} \tag{15}$$

If x_H moves first (with probability $\frac{1}{2}$), she leaves a second-mover share to y_L as if output was $f(x_L, y_L)$ and keeps the rest of the actual output $f(x_H, y_L)$. If y_L moves first, y_L takes the first-mover share of $f(x_L, y_L)$ for herself and x_H obtains the actual remainder. In an equilibrium match, by contrast, x_H would obtain

$$\begin{aligned} \sigma(x_H|y_H) &= \frac{1}{2} \left[\left(1 - \frac{\beta}{2}\right) f(x_H, y_H) + \beta(r + \delta)c \right] + \frac{1}{2} \left[\frac{\beta}{2} f(x_H, y_H) - \beta(r + \delta)c \right] \\ &= \frac{1}{2} f(x_H, y_H) \end{aligned} \tag{16}$$

Comparing $\sigma(x_H|y_L)$ and $\sigma(x_H|y_H)$, we find the following:

Lemma 3 (Matches with lower types). *In the putative equilibrium, strict supermodularity of $f(\cdot, \cdot)$ is necessary and sufficient for any agent $x \in \Theta$ to strictly prefer the equilibrium match to a match with a lower type in which she perfectly imitates the lower type.*

Proof. Any agent $x_H > \underline{x}$ will strictly prefer the equilibrium match to a match with a lower type y_L if $W(x_H|y_H) > W(x_H|y_L)$. As argued above, this is equivalent to

$$\sigma(x_H|y_H) > \sigma(x_H|y_L)$$

$$\Rightarrow f(x_H, y_H) - f(x_H, y_L) > f(x_H, y_L) - f(x_L, y_L) \quad (17)$$

using equations (15) and (16). Next, note that we can write

$$f(x_H, y_L) = f(y_H, x_L) = f(x_L, y_H) \quad (18)$$

where the first equality holds because $x_H = y_H$ and $y_L = x_L$, while the second equality holds by symmetry of $f(\cdot, \cdot)$ (see assumption 1). Therefore substituting $f(x_L, y_H)$ for $f(x_H, y_L)$ on the left-hand side of equation (17) only, we obtain the equation in definition 3. By this definition, strict supermodularity of $f(\cdot, \cdot)$ is necessary and sufficient for the equation to hold. Finally, the type $\underline{x}$ matches with $\underline{y}$ in the putative equilibrium, so that a lower type than in the equilibrium match does not exist in this case. $\square$

Next suppose that x_H has signalled to be of type x_L , has thus met a type y_L , but now wants to renege on the signal. We will find below that x_H has to let at least one round of bargaining fail to actually convince y_L of her true type. Here we ask whether reneging could possibly make the deviation to a match with a lower type worthwhile. By considering the hypothetical extreme case that y_L instantly observes the true type x_H , we obtain an envelope result and thereby a negative answer:

Lemma 4 (Reneging in matches with lower types). *Suppose types were instantly observable in meetings. Consider a type x_H who deviates from the putative equilibrium situation and meets a type $y_L < x_H$.*

a) If neither agent's outside option is binding, the following stationary strategies will form the unique SPE of the bargaining game and lead to agreement in the first round:

(i) *for herself, agent x_H always proposes*

$$\pi^*(x_H|y_L) = \frac{2r + \eta}{2(r + \eta)} \left[f(x_H, y_L) + \frac{\beta\delta}{2r} \left[f(x_L, y_L) - \frac{\eta}{2r + \eta} f(x_H, y_H) \right] \right] + \beta(r + \delta)c$$

When y_L proposes $\pi(y_L|x_H)$, x_H always accepts if and only if $\pi(y_L|x_H) \leq \pi^(y_L|x_H)$.*

(ii) *for herself, y_L always proposes*

$$\pi^*(y_L|x_H) = \frac{2r + \eta}{2(r + \eta)} \left[f(x_H, y_L) + \frac{\beta\delta}{2r} \left[f(x_H, y_H) - \frac{\eta}{2r + \eta} f(x_L, y_L) \right] \right] + \beta(r + \delta)c$$

When x_H proposes $\pi(x_H|y_L)$, y_L always accepts if and only if $\pi(x_H|y_L) \leq \pi^(x_H|y_L)$.*

b) If only the outside option of y_L binds, the shares in the unique SPE become

$$\begin{aligned} \pi^*(x_H|y_L) &= f(x_H, y_L) - \frac{\beta}{2} f(x_L, y_L) + \beta(r + \delta)c \\ \pi^*(y_L|x_H) &= \frac{1}{2r + \eta} \left[2r f(x_H, y_L) + \delta\beta f(x_H, y_H) + \frac{\eta\beta}{2} f(x_L, y_L) \right] + \beta(r + \delta)c \end{aligned}$$

c) Strict supermodularity of $f(\cdot, \cdot)$ is sufficient for any agent $x \in \Theta$ to strictly prefer the equilibrium match to this deviation.

Proof. See appendix. $\square$

Part a) of lemma 4 presents an expression for $\pi^*(x_H|y_L)$ that is increasing in $f(x_L, y_L)$ and decreasing in $f(x_H, y_H)$. This has nothing to do with outside options, as they were assumed non-binding. Rather, it reflects that x_H is less patient than y_L , since bargaining delay is more costly when one expects a share $\frac{1}{2}f(x_H, y_H)$ eventually after a match break-up than when one expects only $\frac{1}{2}f(x_L, y_L)$. This tends to reduce $\pi^*(x_H|y_L)$ and raise $\pi^*(y_L|x_H)$ (but does not drive any of our results). For the same reason, $\sigma(x_H|y_L)$ turns out to be slightly less than $\frac{1}{2}f(x_H, y_L)$ (see the proof of part c)). The expression for $\pi^*(y_L|x_H)$ is interpreted analogously. The shares in part b) exhibit the pattern one would expect, given that $f(x_L, y_L)$ drives the binding outside option of y_L . If, however, x_H, y_H, x_L , and y_L are all set equal in lemma 4, both expressions for $\pi^*(x_H|y_L)$ will collapse into that for $\pi^*(x|y)$ in lemma 1, and likewise for $\pi^*(y_L|x_H)$. Hence parts a) and b) of lemma 4 may be regarded as a generalisation of lemma 1 to an asymmetric case.

Crucially, part c) finds that even if x_H could immediately convince y_L of her true type, x_H would strictly prefer the equilibrium match, as she does when she would have to imitate some lower type. Based on lemmas 3 and 4, we show below that higher types never have an incentive to deviate from the putative equilibrium to matches with lower types if $f(\cdot, \cdot)$ is supermodular, for any beliefs that lower types might hold about deviants. In turn, whenever a deviant causes bargaining to fail, the other agent thus knows that she faces a strictly lower type: for a weakly higher type, a deviation would be equilibrium-dominated. Next recall from assumption 2 that in expectation the same time $1/\eta$ elapses before another round of bargaining as before a meeting with a different agent. When bargaining fails due to a deviant, the other agent (whose type was observable from a truthful signal) now prefers by lemma 4 to meet a different agent: she simply chooses her equilibrium match rather than a match with some strictly lower type after the same expected delay.⁸ As this insight is central to our argument, we formally prove it:

Lemma 5 (Equilibrium-dominated strategies). *Let agents' beliefs place probability 0 on the occurrence of equilibrium-dominated actions, let $f(\cdot, \cdot)$ be strictly supermodular, and consider a meeting between some x and y in the putative equilibrium. If x deviates such that bargaining fails, y will correctly believe to face a lower type and will choose to walk away.*

Proof. We first establish that any agent $x_H > \underline{x}$ always prefers, for any beliefs of $y_L \leq x_H$, her equilibrium match to a deviation such that she meets a weakly lower type with whom bargaining fails. For $x_H = y_L$, lemma 1 implies that x_H would have preferred reaching a bargaining agreement with y_L . For $x_H > y_L$, we have to consider all possible beliefs held by y_L about the potential match output $f(x, y)$ when bargaining fails:

- (i) $g(y_L|h^y) = f(x_H, y_L)$ so that y_L believes to face the true type x_H . By lemma 4, x_H strictly prefers her equilibrium match.

⁸ This logic will also apply if a deviation is only detected after the start of the match: it can only be detected when agents' initial bargaining agreement breaks down, so that there is no basis for further production while agents wait for the new round of bargaining required for renegotiation.

- (ii) $g(y_L|h^y) > f(x_H, y_L)$ so that y_L overestimates potential match output and thus believes to face a type even higher than x_H . By the same argument as in the proof of part c) of lemma 4, y_L does not believe the outside option of x_H to bind: if it did, x would have had to pursue an equilibrium-dominated strategy. Observe that both $\pi^*(y_L|x_H)$ and $f(x_H, y_L) - \pi^*(x_H|y_L)$ in lemma 4 are non-decreasing in x_H , whether or not the outside option of y_L binds. Hence y_L demands weakly higher shares than under (i). Because x_H strictly prefers her equilibrium match under (i), she still prefers her equilibrium match when y_L is more demanding.
- (iii) $f(x_H, y_L) > g(y_L|h^y) > f(x_L, y_L)$ so that y_L underestimates potential match output but still believes to face a higher type. Note that $f(x_L, y_L)$ is then a lower bound for $g(y_L|h^y)$. By lemma 3, x_H strictly prefers her equilibrium match if y_L believes to face x_L (and x_H imitates x_L to avoid bargaining failure). By the same arguments as under (ii), if y_L believes to face a higher type $x_H > x_L$, she will not believe the outside option of x_H to bind and will demand weakly higher shares. Then x_H still prefers her equilibrium match.
- (iv) $g(y_L|h^y) = f(x_L, y_L)$ so that y_L believes to face the same type as her own type. By lemma 3, x_H strictly prefers her equilibrium match.
- (v) $g(y_L|h^y) < f(x_L, y_L)$ so that y_L believes to face a lower type. As we consider a unilateral deviation from the putative equilibrium by x_H , y_L has sent a truthful signal, so that her type has been disclosed to x_H . By lemma 4, y_L then strictly prefers her equilibrium match to a match with x_H who is perceived as a lower type. Hence y_L walks away to meet another agent, as the expected delay is the same.

Hence the deviation in question is equilibrium-dominated for weakly higher types than y_L . Now requiring that agents' beliefs place probability 0 on the occurrence of equilibrium-dominated actions, y_L must believe to face a lower type when bargaining fails, $g(y_L|h^y) < f(x_L, y_L)$. By the argument under (v), y_L thus walks away when bargaining fails. As we supposed that $y_L \leq x_H$, the entire reasoning applies to any $y \in \Theta$. $\square$

Let us turn to the incentive for lower types to deviate to a match with a higher type. Without loss of generality, consider some agent with a type $x_L < \bar{x}$, so that higher types necessarily exist. Now we want to compare being matched with an exactly corresponding type $y_L = x_L$, as in the equilibrium match, to being matched with a higher type $y_H > x_L$. The lower type x_L has two possibilities: she can either perfectly imitate x_H , or she can signal to have type x_H in order to meet y_H but then renege on the signal. We have just shown that, if x_L reneges such that bargaining fails, y_H will walk away and x_L does not gain from the deviation. Next note that, once signals have been observed, x_L cannot further influence the beliefs y_H holds before bargaining begins. In particular, should x_L simply claim to have a still higher type, she would thereby claim to have taken an equilibrium-dominated action (see lemma 5). Provided the beliefs of y_H rule out equilibrium-dominated actions, this claim will not be taken seriously. Should x_L claim to have a lower type, y_H still will not have any reason to believe this:

Lemma 6 (Irrelevant communication). *Any claims by agents to have a lower type than signalled will not be credible if $f(\cdot, \cdot)$ is strictly supermodular.*

Proof. See appendix. $\square$

Hence, unless x_L herself walks away (without gain from the deviation), she will have to bargain under two constraints: y_H believes to face a type x_H and bargaining must not fail. Now recall that these are exactly the constraints under which the bargaining strategy of x_H in the putative equilibrium is optimal (see lemma 1), so that x_L cannot do better than perfectly imitate x_H : if she is more demanding than x_H , bargaining will fail, and if she is less demanding, she will not be optimising. When x_L therefore perfectly imitates x_H , the expected flow utility for x_L is

$$\begin{aligned}\sigma(x_L|y_H) &= \frac{1}{2} \left[f(x_L, y_H) - \frac{\beta}{2} f(x_H, y_H) + \beta(r + \delta)c \right] \\ &\quad + \frac{1}{2} \left[f(x_L, y_H) - \left(1 - \frac{\beta}{2}\right) f(x_H, y_H) - \beta(r + \delta)c \right] \\ &= f(x_L, y_H) - \frac{1}{2} f(x_H, y_H)\end{aligned}\tag{19}$$

If x_L moves first, she has to leave y_H the second-mover share of $f(x_H, y_H)$ to avoid being found out and can thus take whatever is left of the actual output $f(x_L, y_H)$. If y_H moves first, y_H takes the first-mover share of $f(x_H, y_H)$ for herself and x_L obtains the actual remainder. By contrast, the expected flow utility for x_L from her equilibrium match would be

$$\sigma(x_L|y_L) = \frac{1}{2} f(x_L, y_L)\tag{20}$$

A comparison of $\sigma(x_L|y_H)$ and $\sigma(x_L|y_L)$ yields the following result:

Lemma 7 (Matches with higher types). *In the putative equilibrium, strict supermodularity of $f(\cdot, \cdot)$ is necessary and sufficient for any agent $x \in \Theta$ to strictly prefer the equilibrium match to a match with a higher type in which she perfectly imitates the higher type.*

Proof. Any agent $x_L < \bar{x}$ will strictly prefer the equilibrium match to a match with a higher type y_H if $W(x_L|y_L) > W(x_L|y_H)$, which is equivalent to

$$\begin{aligned}\sigma(x_L|y_L) &> \sigma(x_L|y_H) \\ \Rightarrow f(x_H, y_H) - f(x_L, y_H) &> f(x_L, y_H) - f(x_L, y_L)\end{aligned}$$

using equations (19) and (20). By equation (18), we can replace $f(x_L, y_H)$ on the right-hand side by $f(x_H, y_L)$. Hence strict supermodularity is necessary and sufficient for this equation to hold. Finally, for the type $\bar{x}$, a higher type than in the equilibrium match does not exist. $\square$

We have thus identified conditions under which each agent $x \in \Theta$ strictly prefers her equilibrium match to a deviation to any other match. Also recall that matching rates are the same across marketplaces, so that matching rates do not reverse this preference. Corollary 1 collects the implications of this section for agents' signals and beliefs:

Corollary 1 (Truthful signals). *Let agents' beliefs place probability 0 on the occurrence of equilibrium-dominated actions. Then strict supermodularity of $f(\cdot, \cdot)$ is necessary and sufficient for each agent $x \in \Theta$ in the putative equilibrium to strictly prefer a truthful signal $\tilde{x} = x$. Given $h^x = \{\tilde{y}\}$, the only beliefs consistent with truthful signals are $\psi(y|h^x = \{\tilde{y}\}) = u^n(y|h^x = \{\tilde{y}\}) = 1$ for all $y = \tilde{y}$.*

Proof. Choose and fix some arbitrary unmatched agent with a type $x \in \Theta$ and call this exemplary type x_E . Given the choices of marketplace and bargaining strategy in the putative equilibrium, an agent of type x_E matches with an agent of type $y_E = x_E$ at rate η unless there is a deviation. Further, types x_E and y_E also meet unless there is a deviation. Hence, if x_E does not deviate, but sends a truthful signal $\tilde{x}_E = x_E$, it must be that $\tilde{x}_E \in R^n$ for the chosen marketplace n . Given that agents meet exclusively their own type in the putative equilibrium, $|R^n| = 1$. Then we have $\tilde{x}'_E \notin R^n$ for any non-truthful signal $\tilde{x}'_E \neq x_E$. Hence x_E has to signal truthfully to obtain her equilibrium match. By lemmas 3 through 7, she will strictly prefer this match to any other mutually acceptable match if $f(\cdot, \cdot)$ is supermodular and agents' beliefs rule out equilibrium-dominated actions. Because type x_E was arbitrarily chosen, the reasoning extends to any type $x \in \Theta$. Finally, if all signals are truthful, then $u^n(y|h^x = \{\tilde{y}\}) = 1$ for all $y = \tilde{y}$, and agents' beliefs can only be consistent if $\psi(y|h^x = \{\tilde{y}\}) = 1$ for all $y = \tilde{y}$. $\square$

In short, each agent $x \in \Theta$ finds it optimal to signal truthfully because this is the only way to obtain her equilibrium match, which she prefers to a deviation. As all agents therefore indeed signal truthfully, only beliefs that signals are truthful can be consistent with equilibrium play.

In conclusion, this section has presented an extensive but essentially simple reasoning. We found that higher types will never deviate from the putative equilibrium to match with lower types if $f(\cdot, \cdot)$ is supermodular. An agent who detects a deviation should therefore believe to face a lower type; when she can choose between continued bargaining with a lower type and her equilibrium match, she prefers the latter. Lower types can thus only match with higher types by imitating them, but they will not gain from such a deviation if $f(\cdot, \cdot)$ is supermodular. Then each agent prefers not to deviate and consequently finds it optimal to signal truthfully.

5.4 Marketplace choice and creation

In this section, we take it as given that signals are truthful and concentrate on the creation of marketplaces and on agents' optimal choice among them. Consider three types x_L, x_M , and x_H , with $\underline{x} \leq x_L < x_M < x_H \leq \bar{x}$ (while the argument generalises to all types). Suppose these types search in the same marketplace, so that each of them can meet with y_L, y_M , or y_H . We know from lemma 4 that each x_H would prefer a match with y_H to a match with y_M or y_L . Using the truthful signals, the agents of type x_H can profitably set up a new marketplace where $R^n = \{x_H\}$ so that agents of type x_H exclusively meet each other. In the initial marketplace, they would also meet other types although matches with these types would be less desirable, which is not offset by any advantage in meeting rates.

By setting up an exclusive marketplace, the congestion externality imposed by these other types is avoided (see Jacquet/Tan (2007) for details of this logic).

Given that signals are truthful, the remaining types x_M and x_L can no longer meet with y_H , as they would have to send a false signal to join the marketplace where y_H can be met. Among the possible matches, x_M prefers by lemma 4 the match with y_M , so that all agents of type x_M now set up an exclusive marketplace with $R^n = \{x_M\}$, leaving the initial marketplace to the agents of type x_L . This logic applies to any marketplace with different types, so that all types have their own exclusive marketplaces in equilibrium. (We will generalise this logic in section 6.3 to show that it does not only apply in the putative equilibrium, but always when signals are truthful.) There may be several exclusive marketplaces for the same type in equilibrium ($|\mathcal{N}(x)| \geq 1$), as none of our conclusions is affected by their exact number due to constant returns to scale in meeting. Formally, each agent x thus optimally chooses a marketplace $n \in \mathcal{N}(x)$ and thereby obtains the present value $U(x) \geq U^n(x), \forall n$. Given the optimal bargaining strategies in lemma 1, every meeting in the putative equilibrium then leads to a match, as one would expect when truthful signals allow agents to know everything in advance.

By way of summary, this subsection and the preceding have each shown a component of the putative equilibrium situation to hold, given the other components. We thus found the pointwise steady state in the PBE. Given a supermodular match production function and beliefs that rule out equilibrium-dominated actions, agents seek to meet only exactly corresponding types. All agents then signal their types truthfully and correctly believe that all other agents signal truthfully. They match only with exactly corresponding types, so that the resulting equilibrium matching of agents is PPAM. Our model thus leads to PPAM under the same weak condition as in Becker's (1973) frictionless model, despite two kinds of search frictions. The next section discusses key properties of the separating equilibrium.

6 Equilibrium properties

6.1 Efficiency

The separating equilibrium we have identified is efficient in a number of important respects. First and foremost, search costs are minimised, both for each agent individually and overall: in equilibrium, truthful signals allow each agent to ensure that no meeting is wasted, but that every meeting she attends results in a match. Hence, whenever an agent searches, she matches after an expected search time of $1/\eta$. This is the minimum delay because a meeting necessarily precedes a match. In a random search model, each match would typically be preceded by a number of unsuccessful meetings, and only by chance will the first meeting of an agent result in a match. Therefore, search costs in random search models are at least as high from the individual perspective as in our model with truthful signals, and strictly higher in expectation as well as on aggregate. Second, note that all agents match in equilibrium so that there is no unrealised surplus left in the form of agents who never match. On the contrary, Becker (1973) proved the following result:

Corollary 2 (Output efficiency). *If the match production function is strictly supermodular, PPAM will maximise aggregate output.*

Proof. See appendix. $\square$

Random search models, be it with or without supermodularity of the match production function, do in general not maximise aggregate match output, as they lead to a certain degree of mismatch instead of PPAM. Finally, among the mutually acceptable matches, agents in the equilibrium we found always obtain the match they most prefer. This again contrasts starkly with random search models, where the match an agent expects is the expectation over the mutually acceptable matches, not the most preferred one of them.

6.2 Stability

In this section, we examine whether the equilibrium matching we found is a *stable matching*. Because this equilibrium is symmetric, our notation can abstract from the distinction between types and individual agents without loss of generality. Suffice to let $\sigma(x)$ denote the expected flow utility that an agent of type x obtains under a particular matching. Recall that $\sigma(x) = \sigma(x|y)$ if x and y are matched in this matching and $\sigma(x) = 0$ if x remains unmatched. We can then define stability in a symmetric equilibrium as follows:

Definition 4 (Stable matching). *A matching is stable if $\sigma(x) \geq 0$ for all $x \in \Theta$ and there is no match between any agents with types x and y such that $\sigma(x|y) > \sigma(x)$ and $\sigma(y|x) > \sigma(y)$.*

Becker (1973) instead used the concept of the *core* in his seminal work. In words, the core is the set of matchings such that no legal coalition of agents can ensure more flow utility for all its members than obtained in the matching (see e.g. the definition in Telser (1978)). However, if only the sum of utility obtained by the coalition counts, the possibility of side payments within the coalition is implicitly assumed. Side payments are crucial in Becker's (1973) reasoning: an individual agent then always prefers, among all matches available to her, the match generating the highest match output, since her partner in this match will use the extra output to outbid any other potential match partners. Yet where output is divided through bargaining, an agent's share in the match generating the highest output may fall short of her share in another match.

We would thus have to modify the definition of the core to ensure that each agent's $\sigma(x)$ weakly exceeds the utility she obtains in any legal coalition available to her. In our model, legal coalitions always have one or two members. Then a modified definition reduces to two requirements: $\sigma(x)$ has to weakly exceed the utility of being single, and no match is available to x in which she obtains strictly more than $\sigma(x)$. When agents go to perfectly segmented marketplaces, we can identify a match that is available to x with a match where the match partner is better off than in any other available match. Then these requirements coincide with those in definition 4.

A proof that PPAM is a stable matching would thus also prove that it is a matching in the core of our model. We find that supermodularity of the match production function is a sufficient condition here for PPAM to be a stable matching:

Corollary 3 (Stability of PPAM). *Whenever it exists, the separating equilibrium described by the putative equilibrium leads to a stable matching.*

Proof. Recall that the separating equilibrium above exists provided $f(\cdot, \cdot)$ is strictly supermodular and that it leads to PPAM. Now suppose that PPAM is not a stable matching. Then there must be a match between unequal types that is preferred by both types to matches with exactly corresponding types. However, given strict supermodularity of $f(\cdot, \cdot)$, matching with a lower type is an equilibrium-dominated action for the higher type in any match between unequal types, by the proof of lemma 5. Hence the higher type would then always prefer a match with an exactly corresponding type, so that a match between unequal types that is preferred by both does not exist. Finally, lemma 1 implies together with assumption 1 that $\sigma(x) \geq 0 \forall x \in \Theta$ under PPAM. $\square$

A stable matching is a most unusual result in a model with search frictions. In random search models, agents cannot search selectively and might thus be matched with any type from a certain range of types. Of course, many of these types are only accepted because search frictions make continued search undesirable. A stable matching cannot be expected to arise under such circumstances and is very unlikely to arise by chance whenever the number of different types is not trivially small. Stable matchings normally only arise in frictionless models. We attribute the reason that a stable matching is achieved here despite search frictions to the signals: they allow agents to pursue their search almost as if there were no search frictions.

Adachi (2003) shows for a fairly general search model that the set of equilibria will reduce to the set of stable matchings in a model à la Gale/Shapley (1962) if search frictions become negligible. Our result in this section qualifies this finding in so far as search frictions remain in our model because agents do not meet immediately ($\eta < \infty$) and incur costs from meetings ($c \geq 0$), and yet a stable matching results. This suggests that frictions do not prevent a stable matching in a search model as long as they do not keep agents from meeting only specifically chosen types. Intuitively, arbitrarily high frictions do not have any effect when agents find ways to match like in a frictionless environment. If search costs are only incurred at the end of an otherwise costless search process, the matching will be as in the absence of any costs, provided agents still participate.

6.3 Uniqueness

While we have shown that a particular separating equilibrium exists, this section argues that it is unique. By its very nature, a separating equilibrium is characterised by truthful signals.⁹ In section 5.4, truthful signals lead to marketplaces where agents meet exclusively

⁹ We ignore separating equilibria where signals are not truthful yet still informative because they are linked by a one-to-one mapping to agents' true types, and this mapping forms the basis of agents' correct beliefs. Such equilibria would only be variants of equilibria with truthful signals.

their own type. This result generalises:

Lemma 8 (Market segmentation). *Agents will meet only their own type in any separating equilibrium.*

Proof. See appendix. $\square$

Since agents only meet their own type in any separating equilibrium, they can only match with their own type. Therefore, PPAM is the unique matching that may result in any separating equilibrium of our model. We can now conclude more comprehensively:

Proposition 2 (Uniqueness). *Whenever it exists, the separating equilibrium described by the putative equilibrium is unique up to off-equilibrium beliefs.*

No formal proof is needed, as this follows from our earlier results. We know from lemma 8 that any separating equilibrium would have to lead to PPAM, so that other separating equilibria would have to differ in agents' signals, their beliefs, their choice of marketplace, their bargaining strategy, or in the steady state. However, there is only one way for each agent to signal truthfully. When signals are truthful, lemma 8 means that choosing a marketplace $n \in \mathcal{N}(x)$, as in section 5.4, is the uniquely optimal choice rule for x . Section 5.1 identified the unique bargaining SPE in this context. Then only one specification of beliefs about equilibrium actions will be compatible with these choices.

Finally, as the bargaining SPE ensures agreement in the first round of bargaining, $\kappa(x) = 0$, for all $x \in \Theta$ and in any separating equilibrium. Since this agreement to match is reached with an agent of the same type, assumption 3 is sufficient to ensure participation of all types, as shown in section 5.2. Hence also $\nu(x) = 0, \forall x \in \Theta$, so that equation (14) applies to the steady state and determines a unique mass for the matched and for the unmatched agents of each type. Hence, separating equilibria other than the putative equilibrium can only differ in beliefs about off-equilibrium actions.

6.4 Discussion

Let us clarify the role that supermodularity plays for our results. Since types are only privately observable and nothing keeps agents from imitating other types, an agent may match incognito with any type she likes. However, because actual match output then differs from the match output suggested by the signals, the deviant will only remain incognito if she bears the necessary adjustment: she has to give up as much of her own share as is necessary to bridge the gap when actual output is lower (otherwise bargaining fails and the other agent walks out), and she quietly pockets the excess output when actual output is higher. To explain why a lower type x_L would then not match incognito with a higher type $y_H > x_L$, supermodularity is key: $f(x_H, y_H) - f(x_L, y_H)$ is the necessary adjustment when y_H otherwise matches with x_H in equilibrium, while $f(x_L, y_H) - f(x_L, y_L)$ is the extra output produced in comparison to the equilibrium match of x_L . With $f(x_L, y_H) = f(x_H, y_L)$ in the latter, as established by equation (18), the necessary adjustment will exceed the extra

output if $f(\cdot, \cdot)$ is strictly supermodular. From the perspective of a lower type, any possible gains from higher output with a higher type are therefore more than outweighed by the costs from adjustment.

More generally, supermodularity has in effect assumed the role of a single-crossing property in our model. This way, we obtain a fully separating equilibrium even though signals are costless. Separation is therefore not driven by differences in the cost of signals, but by differences in marginal productivity of the same agent over different matches. However, it can be shown that these differences in themselves do not deliver a separating equilibrium for a general type distribution and for unrestricted sets of signals that agents can choose from. Yet under the realistic assumption that true types are always only privately observable, as in this model, supermodularity does deliver full separation for a general type distribution and unrestricted signal sets.

Sorting in the separating PBE is driven by a logic apparently new to the literature. In intuitive terms, agents are effectively bound by their signal, so that a low type can only choose between “being herself” in a match with an equally low type and behaving like a high type in a match with a high type. The most desirable option, “being herself” in a match with a high type, is not available. When behaving like a high type implies disproportionate sacrifices due to supermodularity, low types prefer matches with equally low types.

To put this into a real-world labour market context, suppose a low-skilled worker faces the choice between working at McDonalds and working at McKinsey. While McKinsey would presumably pay a significantly higher wage, the low-skilled worker would have to perform at McKinsey like her high-skilled colleagues. It is easy to imagine that the sheer effort and the extra hours needed to reach this performance outweigh the benefit of a higher salary, so that the low-skilled worker actually prefers working at McDonalds. Whenever this is the case, McKinsey does not even have to check whether applications are truthful, but would still meet only those who claim to be high-skilled. Indeed, this seemingly paradoxical behaviour that our model rationalises appears to be widespread in recruiting.

While lies in applications and job advertisements are certainly more frequent in practice than in the separating PBE, they seem much less frequent than one might think, given how easy it is to lie in applications and job advertisements. This suggests that most real-world agents consciously choose not to lie and thus do behave as in the separating PBE. It also makes sense in practice to dismiss applicants who are known to have lied in their application: firstly, it is very likely that such applicants are underqualified rather than overqualified. Secondly, as our model suggests, it appears easier to find a replacement rather than to disentangle lies from truth for such applicants, thereby determine their actual qualifications, and then adjust the job requirements to fit these qualifications.

7 Conclusions

This paper has introduced costless signals into a search model with transferable utility. We find a unique separating equilibrium characterised by perfect positive assortative matching,

minimised search duration and search costs, and maximised overall match output. These efficiency benchmarks are virtually never met by random search models because frictions lead to lengthy search and to some mismatch. In our model, signals allow agents to avoid this, so that signals largely offset the effect of frictions on efficiency. The role of signals reflects the pervasive use of effective communication in real-world matching markets that facilitates search.

Positive assortative matching in the separating equilibrium only requires supermodularity of the match production function. Supermodularity simultaneously ensures enough complementarity for sorting and replaces a single-crossing condition that is normally needed for truthful signals. Our model thereby proposes a solution to the paradox in Shimer/Smith (2000): supermodularity as such is unambiguously a weaker condition than the conditions they identified. In fact, our particularly mild condition does not merely ensure positive assortative matching, but even perfect positive assortative matching. To the best of our knowledge, ours is the only model that generates perfect sorting despite discounting or explicit search costs.

We conclude that positive assortative matching as an empirical regularity can be replicated by search models under plausible conditions. This is demonstrated by the model in this paper for the most extreme form of sorting; less pronounced sorting can presumably be obtained by adding noise to various elements of the model. Compared to models with random search, a model with more information in the search process thus appears to generate sorting more easily. We have found this for the realistic case that agents control the additional information flows and may manipulate them strategically. Our results would therefore explain why sorting is much more frequent across many different real-world matching markets than one would expect, given the findings in previous models.

Sorting is likely to become more important as technological and societal progress favours specialisation. At the same time, many new means have appeared of effective and rapid communication that might, as in our paper, support sorting. The combination of these two developments offers ample scope for further research.

A Omitted proofs

Proof of lemma 1. Equations (9) and (10) uniquely determine equilibrium, as follows. Whenever y moves first, her optimisation problem is

$$\max \pi(y|x) \quad \text{s.t.} \quad W^O(x|y) \geq \max[V(x|y), U(x)]$$

With $V(x|y) = U(x)$, the constraint becomes $W^O(x|y) - U(x) \geq 0$. When match output is $f(x, y)$ and y takes $\pi(y|x)$ for herself, $f(x, y) - \pi(y|x)$ would be left for x . Therefore,

$$rW^O(x|y) = f(x, y) - \pi(y|x) - \delta[W^O(x|y) - U(x)]$$

Solving this for $W^O(x|y)$, we find that $W^O(x|y) - U(x) \geq 0$ if and only if

$$f(x, y) - \pi(y|x) - rU(x) \geq 0 \quad (21)$$

After substituting for $rU(x)$ and then for $\sigma(x|y)$ from equations (8) and (7), respectively,

$$[f(x, y) - \pi(y|x)]\phi \geq \pi(x|y) - 2(r + \delta)c \quad (22)$$

where $\phi = (1 - \frac{\beta}{2})/\frac{\beta}{2}$. As y raises $\pi(y|x)$ the left-hand side of equation (22) linearly falls, while the right-hand side stays constant. Hence this constraint will hold with equality for the equilibrium value of $\pi(y|x)$. When x moves first, the constraint is analogously found as

$$[f(x, y) - \pi(x|y)]\phi \geq \pi(y|x) - 2(r + \delta)c \quad (23)$$

As binding constraints, equations (22) and (23) are two equations in two unknowns, so that they determine a unique equilibrium. By the symmetry of these equations, we infer that $\pi(x|y) = \pi(y|x)$. When we make this substitution in either equation and solve for $\pi(y|x)$, we obtain

$$\pi(y|x) = \frac{\phi}{1 + \phi} f(x, y) + \frac{2}{1 + \phi} (r + \delta)c = \left(1 - \frac{\beta}{2}\right) f(x, y) + \beta(r + \delta)c$$

which also equals $\pi(x|y)$. Because both first-mover shares have been derived under the constraint that the second mover accepts, agreement is reached in the first round of bargaining. Finally, subgame perfection as in Rubinstein (1982) holds because present values such as $V(x|y)$ and $U(x)$ incorporate optimising behaviour in every later subgame. $\square$

Proof of lemma 4, part a). Agents x_H and y_L would respectively accept if

$$W^O(x_H|y_L) \geq \max[V(x_H|y_L), U(x_H)], \quad W^O(y_L|x_H) \geq \max[V(y_L|x_H), U(y_L)]$$

If outside options are not binding and y_L moves first, she will maximise $\pi(y_L|x_H)$ subject to $W^O(x_H|y_L) \geq V(x_H|y_L)$. As players revert to the putative equilibrium after a match break-up,

$$rW^O(x_H|y_L) = f(x_H, y_L) - \pi(y_L|x_H) - \delta[W^O(x_H|y_L) - U(x_H)] \quad (24)$$

while $W(x_H|y_L)$ and $V(x_H|y_L)$ are determined by

$$rW(x_H|y_L) = \sigma(x_H|y_L) - \delta[W(x_H|y_L) - U(x_H)] \quad (25)$$

$$rV(x_H|y_L) = \eta[W(x_H|y_L) - c - V(x_H|y_L)] \quad (26)$$

Use equation (25) to substitute for $W(x_H|y_L)$ in equation (26) and solve for $V(x_H|y_L)$. After also solving (24) for $W^O(x_H|y_L)$, we can rewrite $W^O(x_H|y_L) \geq V(x_H|y_L)$ as

$$f(x_H, y_L) - \pi(y_L|x_H) + \delta U(x_H) \geq \frac{\eta}{r + \eta} [\sigma(x_H|y_L) + \delta U(x_H) - (r + \delta)c] \quad (27)$$

With $\sigma(x_H|y_L)$ defined in analogy to equation (7), equation (27) becomes

$$(2r + \eta) [f(x_H, y_L) - \pi(y_L|x_H)] \geq \eta\pi(x_H|y_L) - 2[\delta rU(x_H) + \eta(r + \delta)c] \quad (28)$$

after collecting terms. Using the results from lemma 1 in equation (8),

$$rU(x_H) = \beta \left[\frac{1}{2} f(x_H, y_H) - (r + \delta)c \right] \quad (29)$$

Thus substituting for $rU(x_H)$ in equation (28), we obtain

$$(2r + \eta) [f(x_H, y_L) - \pi(y_L|x_H)] \geq \eta\pi(x_H|y_L) - \beta [\delta f(x_H, y_H) + 2(r + \eta)(r + \delta)c] \quad (30)$$

As before, the left-hand side of equation (30) linearly falls as y_L raises $\pi(y_L|x_H)$, while the right-hand side stays constant. This constraint will therefore hold with equality. The same applies to the analogous constraint for the case that x_H moves first:

$$(2r + \eta) [f(x_H, y_L) - \pi(x_H|y_L)] \geq \eta\pi(y_L|x_H) - \beta [\delta f(x_L, y_L) + 2(r + \eta)(r + \delta)c] \quad (31)$$

As a system of two binding constraints in two unknowns, equations (30) and (31) then determine a unique equilibrium. Solving them simultaneously, one obtains the expressions given for $\pi^*(x_H|y_L)$ and $\pi^*(y_L|x_H)$ in lemma 4. The equilibrium is subgame-perfect because the present values incorporate optimising behaviour in following subgames. $\square$

Proof of lemma 4, part b). The proof is very similar to that for part a) and we thus focus on the differences. With a binding outside option, y_L would accept if $W^O(y_L|x_H) \geq U(y_L)$, i.e.

$$f(x_H, y_L) - \pi(x_H|y_L) - rU(y_L) \geq 0 \quad (32)$$

as in equation (21). This will hold with equality by the same logic as before. The other constraint $W^O(x_H|y_L) \geq V(x_H|y_L)$ leads as in part a) to equation (28), which will likewise hold with equality. Solve equation (32) for $\pi(x_H|y_L)$, substitute in equation (28), and solve it for $\pi(y_L|x_H)$ to arrive at

$$\pi(y_L|x_H) = \frac{1}{2r + \eta} [2rf(x_H, y_L) + \eta rU(y_L) + 2[\delta rU(x_H) + \eta(r + \delta)c]]$$

Substituting for $rU(x_H)$ from equation (29) and analogously for $rU(y_L)$, the expression given for $\pi^*(y_L|x_H)$ in part b) of the lemma results. The expression for $\pi^*(x_H|y_L)$ is obtained directly from equation (32) after substituting for $rU(y_L)$. $\square$

Proof of lemma 4, part c). We want to prove that some $x_H > \underline{x}$ will strictly prefer the equilibrium match to a match with a type $y_L < x_H$ when the true type x_H is observed before bargaining begins. First suppose the outside option of x_H binds, $V(x_H|y_L) < U(x_H)$, with $V(x_H|y_L)$ and $U(x_H)$ determined by

$$rV(x_H|y_L) = \eta[W(x_H|y_L) - c - V(x_H|y_L)], \quad rU(x_H) = \eta[W(x_H|y_H) - c - U(x_H)] \quad (33)$$

Solving equation (33) respectively for $V(x_H|y_L)$ and $U(x_H)$, we obtain

$$V(x_H|y_L) = \frac{\eta[W(x_H|y_L) - c]}{r + \eta}, \quad U(x_H) = \frac{\eta[W(x_H|y_H) - c]}{r + \eta} \quad (34)$$

so that $V(x_H|y_L) < U(x_H)$ if and only if $W(x_H|y_L) < W(x_H|y_H)$. Hence x_H strictly prefers her equilibrium match whenever her outside option binds. Therefore suppose instead that neither agent's outside option binds, so that the results from part a) apply. Then

$$\begin{aligned} \sigma(x_H|y_L) &= \frac{1}{2}\pi^*(x_H|y_L) + \frac{1}{2}[f(x_H, y_L) - \pi^*(y_L|x_H)] \\ &= \frac{1}{2}\left[f(x_H, y_L) + \frac{\beta\delta}{2r}[f(x_L, y_L) - f(x_H, y_H)]\right] \end{aligned}$$

Recalling that $\sigma(x_H|y_H) = \frac{1}{2}f(x_H, y_H)$, we will thus have $\sigma(x_H|y_H) > \sigma(x_H|y_L)$ if

$$f(x_H, y_H) > f(x_H, y_L) + \frac{\beta\delta}{2r}[f(x_L, y_L) - f(x_H, y_H)]$$

which holds because $f(x_H, y_H) > f(x_H, y_L)$ and $f(x_L, y_L) - f(x_H, y_H) < 0$. We conclude that x_H strictly prefers her equilibrium match when neither outside option binds. In the final case to consider, only the outside option of y_L binds, so that the results in part b) apply. Then

$$\sigma(x_H|y_L) = \frac{r + \eta}{2r + \eta} \left[f(x_H, y_L) - \frac{\beta}{2} \left[\frac{\delta}{r + \eta} f(x_H, y_H) + f(x_L, y_L) \right] \right]$$

We will thus have $\sigma(x_H|y_H) > \sigma(x_H|y_L)$ if, after collecting terms in $f(x_H, y_H)$,

$$f(x_H, y_H) > \frac{2(r + \eta)}{2r + \eta + \delta\beta} \left[f(x_H, y_L) - \frac{\beta}{2} f(x_L, y_L) \right]$$

Now subtracting $f(x_H, y_L)$ from both sides gives

$$\begin{aligned} f(x_H, y_H) - f(x_H, y_L) &> \frac{\eta - \delta\beta}{2r + \eta + \delta\beta} f(x_H, y_L) - \frac{\beta(r + \eta)}{2r + \eta + \delta\beta} f(x_L, y_L) \\ \Leftrightarrow f(x_H, y_H) - f(x_H, y_L) &> \frac{\beta(r + \eta)}{2r + \eta + \delta\beta} [f(x_H, y_L) - f(x_L, y_L)] \end{aligned}$$

since $\eta - \delta\beta = \beta(r + \eta)$. With the substitution $f(x_H, y_L) = f(x_L, y_H)$ on the left-hand side only, this equation will hold by strict supermodularity of $f(\cdot, \cdot)$ if

$$\beta(r + \eta) \leq 2r + \eta + \delta\beta \quad \Leftrightarrow \quad \eta(r + \eta) \leq \eta(r + \eta) + 2[\eta\delta + r(r + \delta + \eta)]$$

which holds because $\eta\delta + r(r + \delta + \eta) > 0$. Hence, if $f(\cdot, \cdot)$ is strictly supermodular, x_H will strictly prefer her equilibrium match also when only the outside option of y_L binds. $\square$

Proof of lemma 6. We have to show that an agent who does not have a lower type than y_H also has an incentive to make such claims. In particular, consider an agent of type $x_H = y_H$ with whom y_H matches in the putative equilibrium. Suppose y_H believes this agent's claim to be of type $x_L < y_H$, so that bargaining proceeds as in a meeting between x_L and y_H under full information. Recall from lemma 4 c) that $W(y_H|x_L) < W(y_H|x_H)$ if $f(\cdot, \cdot)$ is strictly supermodular. Then an analogy to equation (34) implies $V(y_H|x_L) < U(y_H)$, so that the outside option of y_H is binding. However, while assuming supermodularity of $f(\cdot, \cdot)$ keeps the proof short, this is not a necessary condition. Hence y_H accepts if $W^O(y_H|x_L) \geq U(y_H)$, i.e.

$$f(x_L, y_H) - \pi(x_L|y_H) - rU(y_H) \geq 0 \quad (35)$$

as in equation (21). As before, this will hold with equality, and after substituting for $rU(y_H)$ in analogy to equation (29),

$$\pi^*(x_L|y_H) = f(x_L, y_H) - \frac{\beta}{2}f(x_H, y_H) + \beta(r + \delta)c$$

The outside option of x_L cannot bind, as an analogy to equation (34) implies that x_L then would not have deviated. Hence y_H can expect that x_L accepts if $W^O(x_L|y_H) \geq V(x_L|y_H)$, where

$$rW^O(x_L|y_H) = f(x_L, y_H) - \pi(y_H|x_L) - \delta [W^O(x_L|y_H) - U(x_L)] \quad (36)$$

while $W(x_L|y_H)$ and $V(x_L|y_H)$ are determined by

$$\begin{aligned} rW(x_L|y_H) &= \sigma(x_L|y_H) - \delta [W(x_L|y_H) - U(x_L)] \\ rV(x_L|y_H) &= \eta[W(x_L|y_H) - c - V(x_L|y_H)] \end{aligned}$$

Then follow the same steps as in the proof of lemma 4 a) to obtain an analogy of equation (30):

$$(2r + \eta) [f(x_L, y_H) - \pi(y_H|x_L)] \geq \eta\pi(x_L|y_H) - \beta [\delta f(x_L, y_L) + 2(r + \eta)(r + \delta)c]$$

which will hold with equality as before. Using the result for $\pi^*(x_L|y_H)$,

$$\pi^*(y_H|x_L) = \frac{1}{2r + \eta} \left[2rf(x_L, y_H) + \delta\beta f(x_L, y_L) + \frac{\eta\beta}{2}f(x_H, y_H) \right] + \beta(r + \delta)c$$

Now if x_H claims to be x_L and then bargains like x_L , the expected share $\sigma'(x_H|y_H)$ is

$$\begin{aligned} & \frac{1}{2} \left[f(x_H, y_H) - \frac{\beta}{2}f(x_H, y_H) + \beta(r + \delta)c \right] \\ + & \frac{1}{2} \left[f(x_H, y_H) - \frac{1}{2r + \eta} \left[2rf(x_L, y_H) + \delta\beta f(x_L, y_L) + \frac{\eta\beta}{2}f(x_H, y_H) \right] - \beta(r + \delta)c \right] \\ = & f(x_H, y_H) - \frac{1}{2r + \eta} \left[\frac{\beta}{2}(r + \eta)f(x_H, y_H) + rf(x_L, y_H) + \frac{\delta\beta}{2}f(x_L, y_L) \right] \end{aligned}$$

by the same logic that leads to equation (15). Then x_H gains from a false claim to be x_L if $\sigma(x_H|y_H) < \sigma'(x_H|y_H)$, which requires

$$\begin{aligned} \frac{1}{2}f(x_H, y_H) &< f(x_H, y_H) - \frac{1}{2r + \eta} \left[\frac{\beta}{2}(r + \eta)f(x_H, y_H) + rf(x_L, y_H) + \frac{\delta\beta}{2}f(x_L, y_L) \right] \\ \Leftrightarrow & 2rf(x_L, y_H) + \delta\beta f(x_L, y_L) < (2r + \eta - \beta(r + \eta))f(x_H, y_H) \end{aligned}$$

Noting that $2r + \eta - \beta(r + \eta) = 2r + \delta\beta$, this inequality holds because $f(x_H, y_H) > f(x_L, y_H) > f(x_L, y_L)$. Hence agent x_H has an incentive to downplay her type in order to make y_H propose and accept lower shares for herself. $\square$

Proof of corollary 2. The proof given in Becker (1973) applies to our set-up and we essentially repeat it here. Let $f(\cdot, \cdot)$ be strictly supermodular and index the types $x \in \Theta$ by $1, 2, \dots, I$ such that $x_1 < x_2 < \dots < x_I$. If PPAM maximises aggregate output, then

$$\sum_{j=1}^I f(x_j, y_{i_j}) < \sum_{i=1}^I f(x_i, y_i) \quad \text{for all permutations } (i_1, i_2, \dots, i_I) \neq (1, 2, \dots, I)$$

Suppose to the contrary that aggregate output is maximised by some permutation $i_1, i_2, \dots, i_I$ for which $i_1 < i_2 < \dots < i_I$ does not hold. Then the permutation includes at least one j_0 such that

$i_{j_0} > i_{j_0+1}$. By strict supermodularity of $f(\cdot, \cdot)$,

$$f(x_{j_0+1}, y_{i_{j_0}}) - f(x_{j_0}, y_{i_{j_0}}) > f(x_{j_0+1}, y_{i_{j_0+1}}) - f(x_{j_0}, y_{i_{j_0+1}})$$

because $x_{j_0+1} > x_{j_0}$ while $y_{i_{j_0}} > y_{i_{j_0+1}}$. After rewriting this as

$$f(x_{j_0}, y_{i_{j_0+1}}) + f(x_{j_0+1}, y_{i_{j_0}}) > f(x_{j_0}, y_{i_{j_0}}) + f(x_{j_0+1}, y_{i_{j_0+1}})$$

the left-hand side represents the match production under PPAM, while the right-hand side represents the match production under the permutation $i_1, i_2, \dots, i_I$. As the former exceeds the latter, the permutation $i_1, i_2, \dots, i_I$ does not maximise aggregate output. $\square$

Proof of lemma 8. Suppose there is at least one marketplace $\mathcal{M}$ in which, with truthful signals, agents do not only meet their own type, so that two or more types meet. Focus on the lowest type y_L in $\mathcal{M}$. This type must be the most preferred feasible type of some higher type $x_H > y_L$ in $\mathcal{M}$, otherwise the higher types would exclude y_L from $\mathcal{M}$ to reduce congestion.

We will show that such a marketplace $\mathcal{M}$ cannot exist in a separating equilibrium. When x_H and y_L bargain, $V(x_H|y_L) \geq U(x_H)$ because x_H most prefers y_L and continued bargaining guarantees a meeting with y_L at rate η . While $U(y_L)$ is unknown, y_L could choose in any separating equilibrium to meet only agents of her own type on an exclusive marketplace $\mathcal{L}$. As part of a separating equilibrium, the situation in $\mathcal{L}$ would correspond to the putative equilibrium situation, and because of the symmetry when y_L and x_L bargain in $\mathcal{L}$,

$$\pi^*(x_L|y_L) = \pi^*(y_L|x_L) \Rightarrow \sigma(y_L|x_L) = \frac{1}{2}f(x_L, y_L)$$

independently of outside options. As $\mathcal{L}$ is always an option for y_L , the payoff y_L would obtain there constitutes a lower bound for $U(y_L)$, denoted $\underline{U}(y_L)$. With equation (8), it is found as

$$r\underline{U}(y_L) = \beta \left[\frac{1}{2}f(x_L, y_L) - (r + \delta)c \right]$$

Next observe that x_H cannot do better in a match with y_L than to leave y_L only with the payoff $\underline{U}(y_L)$ in expectation, so that the payoff to x_H in this case constitutes an upper bound $\overline{W}(x_H|y_L)$. Now suppose that an agent of type $y_H = x_H$ sets up an exclusive marketplace $\mathcal{H}$ for her type. If this creates a profitable deviation for x_H who currently most prefers y_L , the supposed marketplace $\mathcal{M}$ cannot exist in equilibrium. The symmetry in $\mathcal{H}$ would lead to

$$\pi^*(x_H|y_H) = \pi^*(y_H|x_H) \Rightarrow \sigma(x_H|y_H) = \frac{1}{2}f(x_H, y_H)$$

again as in the putative equilibrium situation. As an envelope case, suppose x_H obtains $\overline{W}(x_H|y_L)$ in a match with y_L in $\mathcal{M}$ and now faces the choice between this match and a match with y_H in $\mathcal{H}$. Part c) of lemma 4 applies to this choice (with $U(y_L) = \underline{U}(y_L)$) and establishes a strict preference for the match with y_H over the match with y_L . As x_H meets y_H at rate η and y_L at most at rate η , this preference also translates into a strict preference for marketplace $\mathcal{H}$. Hence x_H has a profitable deviation from $\mathcal{M}$ to $\mathcal{H}$ even when $\overline{W}(x_H|y_L)$ is obtained in $\mathcal{M}$. By the same reasoning, y_H also gains from setting up $\mathcal{H}$. $\square$

References

- [1] Adachi, Hiroyuki (2003): A Search Model of Two-sided Matching under Nontransferable Utility, in: *Journal of Economic Theory* 113, p. 182-198.
- [2] Atakan, Alp E. (2006): Assortative Matching with Explicit Search Costs, in: *Econometrica* 74, p. 667-680.
- [3] Becker, Gary S. (1973): A Theory of Marriage: Part I, in: *Journal of Political Economy* 81, p. 813-846.
- [4] Burdett, Kenneth; Coles, Melvyn G. (1997): Marriage and Class, in: *Quarterly Journal of Economics* 112, p. 141-168.
- [5] Chade, Hector (2006): Matching with Noise and the Acceptance Curse, in: *Journal of Economic Theory* 129, p. 81-113.
- [6] Eeckhout, Jan; Kircher, Philipp (2010): Sorting and Decentralized Price Competition, in: *Econometrica* 78, p. 539-574.
- [7] Gale, David; Shapley, Lloyd S. (1962): College Admissions and the Stability of Marriage, in: *American Mathematical Monthly* 69, p. 9-15.
- [8] Goldmanis, Maris; Ray, Korok; Stuart, Erik (2009): Until Death Do Us Part? The Marriage Model with Divorce, University of Chicago mimeo.
- [9] Hopkins, Ed (2012): Job Market Signalling of Relative Position, or Becker Married to Spence, in: *Journal of the European Economic Association* 10, p. 290-322.
- [10] Hoppe, Heidrun C.; Moldovanu, Benny; Sela, Aner (2009): The Theory of Assortative Matching Based on Costly Signals, in: *Review of Economic Studies* 76, p. 253-281.
- [11] Jacquet, Nicolas L.; Tan, Serene (2007): On the Segmentation of Markets, in: *Journal of Political Economy* 115, p. 639-664.
- [12] Kübler, Dorothea; Müller, Wieland; Normann, Hans-Theo (2008): Job Market Signaling and Screening: An Experimental Comparison, in: *Games and Economic Behavior* 64, p. 219-236.
- [13] Lentz, Rasmus (2010): Sorting by Search Intensity, in: *Journal of Economic Theory* 145, p. 1436-1452.
- [14] Mare, Robert D. (1991): Five Decades of Educational Assortative Mating, in: *American Sociological Review* 56, p. 15-32.
- [15] Menzio, Guido (2007): A Theory of Partially Directed Search, in: *Journal of Political Economy* 115, p. 748-769.
- [16] Morgan, Peter B. (1998): A Model of Search, Coordination, and Market Segmentation, University at Buffalo mimeo.
- [17] Rogerson, Richard; Shimer, Robert; Wright, Randall (2005): Search-Theoretic Models of the Labor Market: A Survey, in: *Journal of Economic Literature* 43, p. 959-988.

- [18] Rubinstein, Ariel (1982): Perfect Equilibrium in a Bargaining Model, in: *Econometrica* 50, p. 97-110.
- [19] Shimer, Robert (2005): The Assignment of Workers to Jobs in an Economy with Coordination Frictions, in: *Journal of Political Economy* 113, p. 996-1025.
- [20] Shimer, Robert; Smith, Lones (2000): Assortative Matching and Search, in: *Econometrica* 68, p. 343-369.
- [21] Smith, Lones (2006): The Marriage Model with Search Frictions, in: *Journal of Political Economy* 114, p. 1124-1146.
- [22] Spence, Michael (1973): Job Market Signaling, in: *Quarterly Journal of Economics* 87, p. 355-374.
- [23] Telser, Lester G. (1978): *Economic Theory and the Core*. Chicago: University of Chicago Press.

Recently published

No.	Author(s)	Title	Date
1/2012	Stephani, J.	Wage growth and career patterns of German low-wage workers	1/12
2/2012	Grunow, D. Müller, D.	Kulturelle und strukturelle Faktoren bei der Rückkehr in den Beruf	2/12
3/2012	Poeschel, F.	The time trend in the matching function	2/12
4/2012	Schmerer, H.-J.	Foreign direct investment and search unemployment: Theory and evidence	3/12
5/2012	Ellguth, P. Gerner, H.-D. Stegmaier, J.	Wage bargaining in Germany: The role of works councils and opening clauses	3/12
6/2012	Stüber, H.	Are real entry wages rigid over the business cycle? Empirical evidence for Germany from 1977 to 2009	3/12
7/2012	Felbermayr, G. Hauptmann, A. Schmerer, H.-J.	International trade and collective bargaining outcomes: Evidence from German employer-employee data	3/12
8/2012	Engelmann, S.	International trade, technical change and wage inequality in the U.K. Economy	3/12
9/2012	Drasch, K.	Between familial imprinting and institutional regulation: Family related employment interruptions of women in Germany before and after the German reunification	3/12
10/2012	Bernhard, S. Kruppe, Th.	Effectiveness of further vocational training in Germany: Empirical findings for persons receiving means-tested unemployment benefit	4/12
11/2012	Deeke, A. Baas, M.	Berufliche Statusmobilität von Arbeitslosen nach beruflicher Weiterbildung: Ein empirischer Beitrag zur Evaluation der Förderung beruflicher Weiterbildung	4/12
12/2012	Nordmeier, D.	The cyclicity of German worker flows: Inspecting the time aggregation bias	5/12
13/2012	Kröll, A. Farhauer, O.	Examining the roots of homelessness: The impact of regional housing market conditions and the social environment on homelessness in North-Rhine-Westphalia, Germany	5/12
14/2012	Mendolicchio, C. Paolini, D. Pietra, T.	Asymmetric information and overeducation	6/12

As per: 28.06.2012

For a full list, consult the IAB website

<http://www.iab.de/de/publikationen/discussionpaper.aspx>



Imprint

IAB-Discussion Paper 15/2012

Editorial address

Institute for Employment Research
of the Federal Employment Agency
Regensburger Str. 104
D-90478 Nuremberg

Editorial staff

Regina Stoll, Jutta Palm-Nowak

Technical completion

Jutta Sebold

All rights reserved

Reproduction and distribution in any form, also in parts,
requires the permission of IAB Nuremberg

Website

<http://www.iab.de>

Download of this Discussion Paper

<http://doku.iab.de/discussionpapers/2012/dp1512.pdf>

For further inquiries contact the author:

Friedrich Poeschel

Phone +49.911.179 3105

E-mail friedrich.poeschel@iab.de