

Many-to-many matching and price discrimination

RENATO GOMES

Toulouse School of Economics, CNRS, University of Toulouse Capitole

ALESSANDRO PAVAN

Department of Economics, Northwestern University

We study centralized many-to-many matching in markets where agents have private information about (vertical) characteristics that determine match values. Our analysis reveals how matching patterns reflect cross-subsidization between sides. Agents are *endogenously* partitioned into *consumers* and *inputs*. At the optimum, the costs of procuring agents-inputs are compensated by the gains from agents-consumers. We show how such cross-subsidization can be achieved through matching rules that have a simple threshold structure, and deliver testable predictions relating the optimal price schedules to the distribution of the agents' characteristics. The analysis sheds light on the practice of large matching intermediaries, such as media and business-to-business platforms, advertising exchanges, and commercial lobbying firms.

KEYWORDS. Vertical matching markets, many-to-many matching, asymmetric information, mechanism design, cross-subsidization.

JEL CLASSIFICATION. D82.

1. INTRODUCTION

Matching intermediaries, whose business is to link (or match) agents from multiple sides of a market, have a long history. In England, for example, marriage and employment agencies exist since at least the beginning of 19th century. In the United States,

Renato Gomes: renato.gomes@tse-fr.eu

Alessandro Pavan: alepavan@northwestern.edu

Previous versions were circulated under the titles “Price Discrimination in Many-to-Many Matching Markets” and “Many-to-Many Matching Design.” For comments and useful suggestions, we thank the co-editor, anonymous referees, seminar participants at UC Berkeley, Bocconi, Bonn, Cambridge, Universitat Carlos III, Duke, École Polytechnique, Getulio Vargas Foundation (EPGE and EESP), Northwestern, NYU, Pompeu Fabra, Stanford, Toulouse, Yale, UBC, UPenn, and ECARE-ULB, and conference participants at Columbia–Duke–Northwestern IO Conference, SED, Csef-Igier Conference, Stony Brook Game Theory Festival, ESEM, EARIE, ESWM, and ESSET. We are especially grateful to Jacques Crémer, Jan Eeckhout, Johannes Hörner, Doh-Shin Jeon, Bruno Jullien, Patrick Legros, Nenad Kos, Volker Nocke, Marco Ottaviani, Francisco Ruiz-Aliseda, Ron Siegel, Bruno Strulovici, Jean Tirole, Glen Weyl, and Mike Whinston for helpful conversations and useful suggestions. André Veiga provided excellent research assistance. Pavan also thanks the Haas School of Business and the Economics Department at UC Berkeley for their hospitality during the 2011–2012 academic year, as well as the National Science Foundation (Grant SES-60031281) and Nokia for financial support. The usual disclaimer applies.

Copyright © 2016 Renato Gomes and Alessandro Pavan. Licensed under the [Creative Commons Attribution-NonCommercial License 3.0](https://creativecommons.org/licenses/by-nc/3.0/). Available at <http://econtheory.org>.

DOI: 10.3982/TE1904

the establishment of commercial lobbying firms, matching interest groups with policy-makers, predates World War II.¹ Over the last two decades, the Internet has permitted the development of matching intermediaries of unprecedented scale. Notable examples include advertising exchanges, matching publishers and advertisers with compatible profiles; business-to-business platforms, linking firms on different levels of the supply chain; and dating websites, connecting potential partners with similar interests.

Matching intermediation is also at the heart of novel approaches to fund web content. Recently, many media platforms started offering browsers the option to pay to reduce their exposure to advertising.² The platform's problem in designing such offers can be seen from two perspectives. The more familiar one is that of designing a menu to offer to browsers, where each option in the menu consists of a degree of exposure to advertising and a price. The mirror image of this problem consists in designing a matching schedule for advertisers, where prices are contingent on the browsers that each advertiser is able to reach. Because matching is reciprocal, the menus offered to the browsers determine the matching schedules faced by the advertisers, and vice versa. As a consequence, when designing its price-discriminating menus on each side, platforms have to internalize the effects on profits that each side induces on the other side.

The presence of such cross-side effects is what distinguishes price discrimination in matching markets from price discrimination in markets for standard products. In this paper, we present a tractable model of price discrimination in many-to-many matching markets, and show how subsidization across sides shapes the platforms' matching and price schedules.

Model ingredients

The main ingredients of our model are the following. Agents on each side of the market (e.g., browsers and advertisers) are heterogeneous in, and privately informed about, vertical characteristics that determine their willingness-to-pay for matching plans. For example, browsers differ in their tolerance for advertising, while advertisers differ in their willingness-to-pay for browsers' eyeballs.

In addition, agents differ in their salience (or prominence), that is, in the utility (or disutility) they generate to their matching partners. Importantly, we consider both the case in which willingness-to-pay and salience are positively related and the case in which they are negatively related. For example, the ads of those advertisers with the highest willingness-to-pay may be the least annoying for the browsers, although we also

¹See Jones (1805, p. 329) and Seymour (1928) for early accounts of, respectively, marriage and employment agencies in England. See Allard (2008) for a historical account of lobbying in the United States.

²For example, Google recently launched a service, Google Contributor, that allows browsers to pay a monthly fee to reduce the amount of advertising on affiliated sites. Some newspapers, such as the *Guardian*, allow users to pay to remove advertising in their smartphone and tablet apps. Other newspapers, such as the *Washington Post*, offer cheaper (tabloid) versions with similar content but more ads. Similar funding strategies have been adopted by many app and video game developers, offering two versions of the same product that differ only in the amount of advertising. Online publications follow a similar trend: Next Web, for example, charges a yearly fee of \$36.30 to reduce the browsers' exposure to advertising.

consider the opposite case. Analogously, browsers' tolerance for ads may be either positively or negatively related to the browsers' purchasing habits (which determine their value to the advertisers). More broadly, the agents' willingness-to-pay captures their "consumer value," while their salience captures their "input value," i.e., the utility or disutility they bring to the opposite side.

Another flexible feature of the model is that it allows for either increasing or decreasing marginal utility for matching. For example, the browsers' nuisance costs may be convex in the amount of advertising they are exposed to, and the advertisers' profits may be concave in the number of eyeballs they reach. In these cases, an agent's marginal (dis)utility for an extra match depends on the entire set of matches the agent receives.

We study the matching assignments that maximize either social welfare or profits (as many matching intermediaries are privately owned). For each side of the market, the platform chooses a *pricing rule* and a *matching rule*. Along with the usual incentive compatibility constraints (imposing that, for example, each browser chooses the ad-avoidance plan that maximizes his utility), we require only that matching mechanisms satisfy a minimal feasibility constraint, which we call *reciprocity*. This condition requires that if browser i from side A is matched to advertiser j from side B , then advertiser j is matched to browser i . The cases of welfare and profit maximization can then be treated similarly, after one replaces valuations by their "virtual" counterparts (which discount for informational rents).

Our analysis provides answers to the following questions: What matching patterns arise when agents are *privately informed* about the vertical characteristics that determine match values? How does the private (profit-maximizing) provision of matching services compare with the public (welfare-maximizing) provision? How are matching allocations affected by shocks that alter the distribution of the agents' characteristics?

Main results

The recurring theme of this paper is how matching patterns reflect optimal cross-subsidization between sides. Our first main result identifies conditions on primitives under which optimal matching rules exhibit a *threshold structure*. Under a threshold structure, each browser with advertising tolerance v_A is matched to all advertisers with willingness-to-pay above a threshold $t_A(v_A)$, and, conversely, each advertiser with willingness-to-pay v_B is matched to all browsers whose advertising tolerance is higher than $t_B(v_B)$.

Importantly, the reciprocity constraint described above implies that thresholds are weakly decreasing in the vertical characteristic v_k , $k = A, B$. As such, the advertisers with the highest willingness-to-pay are matched to all browsers who are exposed to *any* advertising. In turn, advertisers with lower willingness-to-pay are matched to only a subset of all browsers (namely, those with high tolerance to advertising). Threshold rules thus capture matching allocations exhibiting vertical separation without segmentation (in the form of mutually exclusive groups). The matching allocations induced by threshold rules are consistent with the practice followed by many media platforms (e.g., newspapers) of exposing all readers to premium ads (displayed in all versions of

the newspaper), but only those readers with high tolerance to advertising to discount ads (displayed only in the tabloid or printed version). They are also consistent with the practice followed by many commercial lobbying firms that match interest groups with high willingness-to-pay for political access to all policymakers in their network of influence, while matching interest groups with lower willingness-to-pay to only those policymakers who are less sensitive to political exposure (for a more detailed discussion of the practices of commercial lobbying firms, see [Kang and You 2016](#), who test the predictions of our model using data on U.S. commercial lobbying).

Our second main result provides a precise characterization of the thresholds. We use variational techniques to obtain a Euler equation that equalizes (i) the marginal gains from expanding the matching sets on one side to (ii) the marginal losses that, by reciprocity, arise on the opposite side of the market. Intuitively, this equation endogenously partitions agents from each side into two groups. The first group consists of agents playing the role of *consumers* (e.g., advertisers with high willingness-to-pay). These agents contribute positively to the platform's objective by "purchasing" sets of agents from the other side of the market. The second group consists of agents playing the role of *inputs* (e.g., browsers with low tolerance for advertisement). These agents contribute negatively to the platform's objective, but are used to "feed" the matching sets of those agents from the opposite side who play the role of consumers (cross-subsidization).

The above two results define the paper's theoretical contribution. The fact that matches are reciprocal, along with the fact that each agent is both a consumer and an input in the matching production function, render the cost of extra matches *endogenous* and dependent in a nontrivial way on the *entire* matching rule. The endogeneity of costs is what distinguishes price discrimination in matching markets from price discrimination in commodity markets, where the cost function is exogenous (see, among others, [Mussa and Rosen 1978](#), and [Maskin and Riley 1984](#)).³

Our characterization results enable us to compare the matching allocations that result from the public provision of matching services (which we assume is motivated by welfare maximization) to those that result from the private provision of such services (which we assume is motivated by profit maximization). Interestingly, profit maximization in vertical matching markets may result in inefficiently small matching sets for *all* agents, including those "at the top" of the distribution (e.g., the advertisers with the highest willingness-to-pay). The reason is that the costs of cross-subsidizing such agents is higher under profit maximization, due to the informational rents that must be given to the agents-inputs.

Our analysis also delivers testable predictions about the effects of shocks that alter the salience of the agents. In particular, we show that a shock that increases the salience of all agents from a given side (albeit not necessarily uniformly across agents) induces a profit-maximizing platform to offer larger matching sets to those agents with *low* willingness-to-pay and smaller matching sets to those agents with *high* willingness-to-pay. In terms of surplus, these shocks make low-willingness-to-pay agents better off

³This extra degree of complexity requires stochastic-order techniques and variational arguments that, to the best of our knowledge, are novel to the literature.

at the expense of high-willingness-to-pay agents. In terms of pricing, an increase in attractiveness induces an anticlockwise rotation of the optimal price schedules.

Although formulated in a two-sided matching environment, all our results have implications also for *one-sided* vertical matching markets. Indeed, the one-sided environment is mathematically equivalent to a two-sided matching market where both sides have symmetric primitives and matching rules are constrained to be symmetric across sides. As it turns out, in two-sided matching markets with symmetric primitives, the optimal matching rules are naturally symmetric, in which case the latter constraint is non-binding. All our results thus have implications also for such problems in organization and personnel economics that pertain to the optimal design of teams in the presence of peer effects.

The rest of the paper is organized as follows. Below, we close the [Introduction](#) by briefly reviewing the most pertinent literature. [Section 2](#) describes the model. [Section 3](#) contains all results. [Section 4](#) discusses a few extensions, while [Section 5](#) concludes. All proofs appear in the [Appendix](#).

Related literature

The paper is primarily related to the following literatures.

Matching intermediation with transfers [Damiano and Li \(2007\)](#) and [Johnson \(2013\)](#) consider a *one-to-one* matching intermediary that faces asymmetric information about the agents' vertical characteristics that determine match values. These papers derive conditions on primitives for a profit-maximizing intermediary to induce positive assortative matching. In contrast to these papers, we study *many-to-many* matching in a flexible setting where agents may differ in their consumer value (willingness-to-pay) and input value (salience).

Group design with peer effects [Rayo \(2013\)](#) studies second-degree price discrimination by a monopolist selling a menu of conspicuous goods that serve as signals of consumers' hidden characteristics. Rayo's model can be interpreted as a one-sided matching model where the utility of a matching set is proportional to the average quality of its members. Allowing for more general peer effects, [Board \(2009\)](#) studies the design of groups by a profit-maximizing platform (e.g., a school) that can induce agents to self-select into mutually exclusive groups (e.g., classes).⁴

Price discrimination The availability of transfers and the presence of asymmetric information relates this paper to the literature on second-degree price discrimination (e.g., [Mussa and Rosen 1978](#), [Maskin and Riley 1984](#), [Wilson 1993](#)). Our study of price discrimination in many-to-many matching markets introduces two novel features relative to the standard monopolistic screening problem. First, the platform's feasibility constraint

⁴See also [Arnott and Rowse \(1987\)](#), [Epple and Romano \(1998\)](#), [Helsley and Strange \(2000\)](#), and [Lazear \(2001\)](#) for models of group design under complete information.

(namely, the reciprocity of the matching rule) has no equivalent in markets for commodities.⁵ Second, each agent serves as both a consumer and an input in the matching production function. This feature implies that the cost of procuring an input is endogenous and depends in a nontrivial way on the entire matching rule.

Two-sided markets Markets where agents purchase access to other agents are the focus of the literature on two-sided markets (see [Rysman 2009](#) for a survey, and [Weyl 2010](#), [Bedre-Defolie and Calvano 2013](#), and [Lee 2013](#) for recent developments). This literature, however, restricts attention to a single network or to mutually exclusive networks. Our contribution is in allowing for general matching rules, in distinguishing the agents' willingness-to-pay from their salience, and in accommodating for nonlinear preferences over matching sets.

Decentralized matching Since [Becker \(1973\)](#), matching models have been used to study a variety of markets, including marriage, labor, and education, in which agents are heterogeneous in some *vertical* characteristics that determine the value of the matches (e.g., attractiveness or skill). A robust insight from this literature is that when matches are *one-to-one*, the matching pattern is positive assortative provided that the match value function satisfies appropriate supermodularity conditions, which depend on the presence and nature of frictions, and on the possibility of transfers. See, for example, [Legros and Newman \(2002, 2007\)](#) for a setting where frictions take the form of transaction costs or moral hazard, and [Shimer and Smith \(2000\)](#) and [Eeckhout and Kircher \(2010\)](#) for search/matching frictions. Relative to this literature, we study mediated matching, abstract from search frictions or market imperfections, and consider many-to-many matching rules.

2. MODEL AND PRELIMINARIES

Environment

A platform matches agents from two sides of a market. Each side $k \in \{A, B\}$ is populated by a unit-mass continuum of agents. Each agent from each side $k \in \{A, B\}$ has a type $v_k \in V_k \equiv [\underline{v}_k, \bar{v}_k] \subseteq \mathbb{R}$ that parametrizes both the agent's valuation for matching intensity, that is, the value that the agent assigns to interacting with agents from the opposite side, and the agent's "salience," which we denote by $\sigma_k(v_k) \in \mathbb{R}_+$. Importantly, it is only for simplicity that we assume that salience is a deterministic function of valuations: All our results extend to settings in which salience varies stochastically with valuations, and agents have private information about both their valuations and their salience, in which case an agent's type is given by (v_k, σ_k) (see [Appendix A](#) for details).

Each v_k is drawn from an absolutely continuous distribution F_k (with density f_k), independently across agents. As is standard in the mechanism design literature, we assume that F_k is *regular* in the sense of [Myerson \(1981\)](#), meaning that the virtual valuations for matching $v_k - [1 - F_k(v_k)]/f_k(v_k)$ are continuous and nondecreasing.

⁵A related, but simpler, feasibility constraint is also present in the one-to-one matching models of [Damiano and Li \(2007\)](#) and [Johnson \(2013\)](#).

Given any (Borel measurable) set $\mathbf{s}$ of types from side $l \neq k$, the payoff that an agent from side $k \in \{A, B\}$ with type v_k obtains from being matched, at a price p , to the set $\mathbf{s}$ is given by

$$\pi_k(\mathbf{s}, p; v_k) \equiv v_k \cdot g_k(|\mathbf{s}|_l) - p, \quad (1)$$

where $g_k(\cdot)$ is a positive, strictly increasing, and continuously differentiable function satisfying $g_k(0) = 0$, and where

$$|\mathbf{s}|_l \equiv \int_{v_l \in \mathbf{s}} \sigma_l(v_l) dF_l(v_l)$$

is the *matching intensity* of the set $\mathbf{s}$.

The case where an agent from side k dislikes interacting with agents from the other side is thus captured by a negative valuation $v_k < 0$. To avoid the uninteresting case where no agent from either side benefits from interacting with agents from the opposite side, we assume that $\bar{v}_k > 0$ for some $k \in \{A, B\}$. The functions $g_k(\cdot)$, $k = A, B$, in turn capture increasing (or, alternatively, decreasing) marginal utility (or, alternatively, disutility) for matching intensity.

The payoff formulation in (1) is fairly flexible and accommodates the following examples as special cases.

EXAMPLE 1 (Advertising avoidance). The platform is an online intermediary matching browsers from side A to advertisers from side B . Browsers dislike advertising and their tolerance is indexed by the parameter $v_A \in V_A \subset \mathbb{R}_-$. The nuisance generated by an advertiser with willingness-to-pay $v_B \in V_B \subset \mathbb{R}_+$ to a browser with tolerance v_A is given by

$$v_A \cdot \sigma_B(v_B) \cdot (|\mathbf{s}|_B)^\beta,$$

where $\mathbf{s}$ is the set of ads displayed to the browser, and where $\beta \geq 0$ is the nuisance parameter.⁶ The browser's total payoff is then given by

$$\pi_A(\mathbf{s}, p; v_A) = \int_{v_B \in \mathbf{s}} v_A \cdot \sigma_B(v_B) \cdot (|\mathbf{s}|)^\beta dF(v_B) - p = v_A \cdot g_A(|\mathbf{s}|_B) - p,$$

where p is the price the browser pays to the intermediary (to avoid further advertising), and where the function $g_A(x) = x^{1+\beta}$ (which is strictly convex for $\beta > 0$) captures the increasing “marginal” nuisance of advertising. An increasing salience function $\sigma_B(\cdot)$ then captures the idea that advertisers with a higher willingness-to-pay display, on average, ads that are more annoying to browsers, whereas a decreasing $\sigma_B(\cdot)$ captures the opposite case. For simplicity, advertisers are assumed to have preferences that are linear in the number of browsers reached by their ads,

$$\pi_B(\mathbf{s}, \hat{p}; v_B) = v_B \cdot \int_{v_A \in \mathbf{s}} dF(v_A) - \hat{p},$$

where $\hat{p}$ is the price paid to the platform. ◇

⁶See Kaiser and Wright (2006) and Kaiser and Song (2009) for an empirical assessment of the preferences of browsers vis-à-vis advertisers.

The next example considers a market in which the matching values are supermodular, as in the literature on positive assortative one-to-one matching (e.g., [Damiano and Li 2007](#)).

EXAMPLE 2 (Business-to-business platform). A business-to-business platform matches firms on two levels of the supply chain (identified with sides A and B).⁷ The match between a firm of productivity v_A in level A and a firm of productivity v_B in level B yields a total surplus $v_A v_B$, which is split according to a generalized Nash bargaining protocol. In this specification, the salience of each firm coincides with her productivity (i.e., $\sigma_k(v_k) = v_k$ for all $v_k \in V_k$, with $V_k \subset \mathbb{R}_+$ and $g_k(x) = x$, $k = A, B$). The payoff of a level- A firm is then equal to

$$\pi_A(\mathbf{s}, p; v_A) = \alpha \cdot v_A \cdot \int_{v_B \in \mathbf{s}} v_B dF_B(v_B) - p,$$

whereas the payoff of a level- B firm is

$$\pi_B(\mathbf{s}, p; v_B) = (1 - \alpha) \cdot v_B \cdot \int_{v_A \in \mathbf{s}} v_A dF_A(v_A) - p,$$

where the prices here denote the commissions paid to the platform, and where α is the bargaining weight of level- A firms. $\diamond$

Another special case of our model (where the functions g_k are linear and the functions σ_k are weakly increasing) is developed in [Kang and You \(2016\)](#). This paper brings to data the empirical implications of our model in the context of commercial lobbying firms (matching interest groups and policymakers).

Matching mechanisms

A matching mechanism $M \equiv \{\mathbf{s}_k(\cdot), p_k(\cdot)\}_{k=A,B}$ consists of two pairs (indexed by side) of matching and payment rules. For each type $v_k \in V_k$, the rule $p_k(\cdot)$ specifies the payment asked to an agent from side $k \in \{A, B\}$ with type v_k , while the rule $\mathbf{s}_k(\cdot) \subseteq V_l$ specifies the set of types from side $l \neq k$ included in type v_k 's matching set. Note that $p_k(\cdot)$ maps V_k into $\mathbb{R}$ (both positive and negative payments are allowed), while $\mathbf{s}_k(\cdot)$ maps V_k into the Borel sigma algebra over V_l . With some abuse of notation, hereafter we will denote by $|\mathbf{s}_k(v_k)|_l$ the matching intensity of type v_k 's matching set.⁸

A *matching rule* $\{\mathbf{s}_k(\cdot)\}_{k=A,B}$ is *feasible* if and only if it satisfies the *reciprocity condition*

$$v_l \in \mathbf{s}_k(v_k) \quad \Rightarrow \quad v_k \in \mathbf{s}_l(v_l), \quad (2)$$

⁷[Lucking-Reiley and Spulber \(2001\)](#) and [Jullien \(2012\)](#) survey the literature on business-to-business platforms.

⁸Restricting attention to deterministic mechanisms is without loss of optimality under the assumptions in the model (the proof is based on arguments similar to those in [Strausz 2006](#)). It is easy to see that restricting attention to anonymous mechanisms is also without loss of optimality given that there is no aggregate uncertainty and that individual identities are irrelevant for payoffs.

which requires that if an agent from side l with type v_l is included in the matching set of an agent from side k with type v_k , then any agent from side k with type v_k is included in the matching set of any agent from side l with type v_l .

Next, denote by $\hat{\Pi}_k(v_k, \hat{v}_k; M) \equiv v_k \cdot g_k(|\mathbf{s}_k(\hat{v}_k)|_l) - p_k(\hat{v}_k)$ the payoff that type v_k obtains when reporting type $\hat{v}_k$, and denote by $\Pi_k(v_k; M) \equiv \hat{\Pi}_k(v_k, v_k; M)$ the payoff that type v_k obtains by reporting truthfully. A mechanism M is *individually rational* (IR) if $\Pi_k(v_k; M) \geq 0$ for all $v_k \in V_k$, $k = A, B$, and is *incentive compatible* (IC) if $\Pi_k(v_k; M) \geq \hat{\Pi}_k(v_k, \hat{v}_k; M)$ for all $v_k, \hat{v}_k \in V_k$, $k = A, B$.

A matching rule is *implementable* if there exists a payment rule $\{p_k(\cdot)\}_{k=A,B}$ such that the mechanism $M = \{\mathbf{s}_k(\cdot), p_k(\cdot)\}_{k=A,B}$ is individually rational and incentive compatible.⁹

Welfare and profit maximization

The welfare generated by the mechanism M is given by

$$\Omega^W(M) = \sum_{k=A,B} \int_{V_k} v_k \cdot g_k(|\mathbf{s}_k(v_k)|_l) dF_k(v_k),$$

whereas the profits generated by the mechanism M are given by

$$\Omega^P(M) = \sum_{k=A,B} \int_{V_k} p_k(v_k) dF_k(v_k).$$

A mechanism is efficient (alternatively, profit-maximizing) if it maximizes $\Omega^W(M)$ (alternatively, $\Omega^P(M)$) among all mechanisms that are individually rational, incentive-compatible, and satisfy the reciprocity condition (2). Note that the reciprocity condition implies that the matching rule $\{\mathbf{s}_k(\cdot)\}_{k=A,B}$ can be fully described by its side- k correspondence $\mathbf{s}_k(\cdot)$.

It is standard to verify that a mechanism M is individually rational and incentive-compatible *if and only if* the following conditions jointly hold for each side $k = A, B$:

- (i) The matching intensity of the set $\mathbf{s}_k(v_k)$ is nondecreasing in the valuation v_k .
- (ii) The equilibrium payoffs $\Pi_k(\underline{v}_k; M)$ of the agents with the lowest valuation are nonnegative.
- (iii) The pricing rule satisfies the envelope formula

$$p_k(v_k) = v_k \cdot g_k(|\mathbf{s}_k(v_k)|_l) - \int_{\underline{v}_k}^{v_k} g_k(|\mathbf{s}_k(x)|_l) dx - \Pi_k(\underline{v}_k; M). \quad (3)$$

⁹Implicit in the aforementioned specification is the assumption that the platform must charge the agents before they observe their payoffs. This seems a reasonable assumption in most applications of interest. Without such an assumption, the platform could extract the entire surplus by using payments similar to those in Crémer and McLean (1988); see also Mezzetti (2007).

It is also easy to see that in any mechanism that maximizes the platform's profits, the IR constraints of those agents with the lowest valuations bind, i.e., $\Pi_k(\underline{v}_k; M^P) = 0$, $k = A, B$. Using the expression for payments (3), it is then standard practice to rewrite the platform's profit maximization problem in a manner analogous to the welfare maximization problem. One simply needs to replace the true valuations with their virtual analogs (i.e., with the valuations discounted for informational rents). Formally, for any $k = A, B$ and any $v_k \in V_k$, let $\varphi_k^W(v_k) \equiv v_k$ and $\varphi_k^P(v_k) \equiv v_k - [1 - F_k(v_k)]/f_k(v_k)$. Using the superscript $h = W$ (or, alternatively, $h = P$) to denote welfare (or, alternatively, profits), the platform's problem then consists in finding a matching rule $\{\mathbf{s}_k(\cdot)\}_{k=A,B}$ that maximizes

$$\Omega^h(M) = \sum_{k=A,B} \int_{V_k} \varphi_k^h(v_k) \cdot g_k(|\mathbf{s}_k(v_k)|_I) dF_k(v_k) \quad (4)$$

among all rules that satisfy the monotonicity constraint (i) and the reciprocity condition (2). Bearing these observations in mind, hereafter, we will say that a matching rule $\{\mathbf{s}_k^h(\cdot)\}_{k=A,B}$ is h -optimal if it solves the above h problem. For future reference, for both $h = W, P$, we also define the *reservation value* $r_k^h \equiv \inf\{v_k \in V_k : \varphi_k^h(v_k) \geq 0\}$ when $\{v_k \in V_k : \varphi_k^h(v_k) \geq 0\} \neq \emptyset$.

3. OPTIMAL MATCHING RULES

We start by introducing an important class of matching rules and by identifying natural conditions under which restricting attention to such rules entails no loss of optimality. We then proceed by studying properties of optimal rules and conclude with comparative statics.

3.1 Threshold rules

Consider the following class of matching rules.

DEFINITION 1 (Threshold rules). A matching rule is a threshold rule if, for any $v_k \in V_k$, $k = A, B$,

$$\mathbf{s}_k(v_k) = \begin{cases} [t_k(v_k), \bar{v}_I] & \text{if } v_k \geq \omega_k \\ \emptyset & \text{otherwise,} \end{cases}$$

where the exclusion type $\omega_k \in V_k$ is the valuation below which types are excluded. In this case, we say that the matching rule exhibits the threshold structure $\{t_k(\cdot), \omega_k\}_{k=A,B}$.

Matching rules with a threshold structure are remarkably simple. Any type below ω_k is excluded, while a type $v_k > \omega_k$ is matched to any agent from the other side whose type is above the threshold $t_k(v_k)$. To satisfy the reciprocity condition (2), the threshold functions $\{t_k(\cdot)\}_{k=A,B}$ have to satisfy the constraints identified in the next lemma.

LEMMA 1 (Feasible threshold rules). *Consider a matching rule exhibiting the threshold structure $\{t_k(\cdot), \omega_k\}_{k=A,B}$. This rule is feasible if and only if the following conditions jointly hold:*

(i) For all $v_k \in [\omega_k, \bar{v}_k]$, $k, l = A, B$, $l \neq k$,

$$t_k(v_k) = \min\{v_l : t_l(v_l) \leq v_k\}. \quad (5)$$

(ii) For all $k = A, B$, $t_k(\cdot)$ is a weakly decreasing function.

Condition (i) requires that each threshold function $t_k(\cdot)$, $k = A, B$, coincide with the generalized inverse of the threshold function on the other side of the market. In turn, condition (ii) requires that threshold functions be weakly decreasing in the agents' valuation for matching intensity. The formal proof that these two conditions are jointly equivalent to feasibility is given in the [Appendix](#). The sufficiency claim is proved by directly verifying that any threshold rule satisfying the above two conditions is reciprocal in the sense of (2). The necessity claim is proved by contradiction.

Lemma 1 also implies that matching rules with a threshold structure exhibit a form of negative assortativeness at the margin: Those agents with low valuations are matched only to those agents from the opposite side whose valuation is sufficiently high. Furthermore, the matching sets are ordered across types, in the weak (inclusion) set-order sense, i.e., if $v_k < \hat{v}_k$, then $\mathbf{s}_k(v_k) \subseteq \mathbf{s}_k(\hat{v}_k)$.

REMARK 1 (Implementability). **Lemma 1** implies that feasible threshold rules are always implementable (they generate matching sets whose matching intensity is nondecreasing in v_k). Yet, many implementable matching rules do not exhibit a threshold structure. Incentive compatibility simply requires the matching intensity to be nondecreasing in valuations, but imposes no restrictions on the *composition* of the matching sets. To see this, suppose, for example, that v_k is drawn uniformly from $V_k = [0, 1]$ and that $\sigma_k(v_k) = 1$ for all $v_k \in V_k$, $k = A, B$. Then partitional rules of the type

$$\mathbf{s}_k(v_k) = \begin{cases} [\frac{1}{2}, 1] & \text{if } v_k \in [\frac{1}{2}, 1] \\ [0, \frac{1}{2}] & \text{if } v_k \in [0, \frac{1}{2}], \end{cases}$$

are clearly implementable but do not exhibit a threshold structure. In fact, the matching sets induced by incentive-compatible rules need not be nested or connected.¹⁰

We proceed by identifying weak conditions on payoffs that, along with incentive compatibility, make threshold rules optimal.

CONDITION TP (Threshold primitives). *One of the following two sets of conditions holds:*

¹⁰For example, continue to assume that valuations are drawn uniformly from $V_k = [0, 1]$ but now let $\sigma_k(v_k) = 1 - v_k$ for all $v_k \in V_k$, $k = A, B$. The following matching rule, described by its side- k correspondence, is implementable:

$$\mathbf{s}_k(v_k) = \begin{cases} [0, \frac{1}{3}] \cup [\frac{2}{3}, 1] & \text{if } v_k \in [1 - \frac{\sqrt{2}}{2}, 1] \\ [\frac{1}{3}, \frac{2}{3}] & \text{if } v_k \in [0, 1 - \frac{\sqrt{2}}{2}]. \end{cases}$$

- (a) *The functions $g_k(\cdot)$ are weakly concave and the functions $\sigma_k(\cdot)$ are weakly increasing for both $k = A$ and $k = B$.*
- (b) *The functions $g_k(\cdot)$ are weakly convex and the functions $\sigma_k(\cdot)$ are weakly decreasing for both $k = A$ and $k = B$.*

Condition TP covers two alternative scenarios. The first scenario is one where, on both sides, agents have (weakly) concave preferences for matching intensity. In this case, **Condition TP** also requires that, on both sides, salience increases (weakly) with valuations. The second scenario covers a symmetrically opposite situation, where agents have (weakly) convex preferences for matching intensity and salience decreases (weakly) with valuations. To illustrate, note that the preferences in **Example 1** (advertising avoidance) satisfy **Condition TP** as long as salience on side B is nonincreasing in valuations, meaning that the ads of those advertisers with the highest willingness-to-pay are seen, on average, as being the least annoying. The preferences in **Example 2** (business-to-business platform) also satisfy **Condition TP**, for in this case preferences are linear and salience is increasing in valuations on both sides.

PROPOSITION 1 (Optimality of threshold rules). *Assume **Condition TP** holds. Then both the profit-maximizing and the welfare-maximizing rules have a threshold structure.*

Below, we illustrate heuristically the logic behind the arguments that lead to the result in **Proposition 1** (the formal proof in the **Appendix** is significantly more complex and uses results from the theory of stochastic orders to verify the heuristics described below).

SKETCH OF THE PROOF OF PROPOSITION 1. Consider an agent for whom $\varphi_k^h(v_k) \geq 0$. In case of welfare maximization ($h = W$), this is an agent who values positively interacting with agents from the other side. In case of profit maximization ($h = P$), this is an agent who contributes positively to profits, even when accounting for informational rents. Ignoring for a moment the monotonicity constraints, it is easy to see that it is always optimal to assign to this agent a matching set $\mathbf{s}_k(v_k) \supset \{v_l : \varphi_l^h(v_l) \geq 0\}$ that includes all agents from the other side whose φ_l^h value is nonnegative. This is because (i) these latter agents contribute positively to type v_k 's payoff and (ii) these latter agents have nonnegative φ_l^h values, which implies that adding type v_k to these latter agents' matching sets (as required by reciprocity) never reduces the platform's payoff.

Next, consider an agent for whom $\varphi_k^h(v_k) < 0$. It is also easy to see that it is never optimal to assign to this agent a matching set that contains agents from the opposite side whose φ_l^h values are also negative. The reason is that matching two agents with negative valuations (or, alternatively, virtual valuations) can only decrease the platform's payoff.

These general observations do not hinge on **Condition TP**. Moreover, they say nothing about how to optimally match agents with a positive (virtual) valuation to agents from the opposite side with a negative (virtual) valuation (*cross-subsidization*). This is

where [Condition TP](#), along with the fact that valuations are private information, plays a role.

Consider first the scenario of [Condition TP\(a\)](#), where $g_k(\cdot)$ is weakly concave and $\sigma_k(\cdot)$ is weakly increasing. Pick an agent from side k with $\varphi_k^h(v_k) > 0$ and suppose that the platform wants to assign to this agent a matching set whose intensity

$$q = |\mathbf{s}|_l > \int_{[r_l^h, \bar{v}_l]} \sigma_l(v_l) dF_l(v_l)$$

exceeds the matching intensity of those agents from side l with nonnegative φ_l^h values (i.e., for whom $v_l \geq r_l^h$). The combination of the assumptions that (i) salience is weakly increasing in valuations, (ii) g_l are weakly concave, and (iii) valuations are private information implies that the *least costly* way to deliver intensity q to such an agent is to match him to all agents from side l whose $\varphi_l^h(v_l)$ is the least negative. This is because (a) these latter agents are the most attractive and (b) by virtue of g_l being concave, using the same agents from side l with a negative φ_l^h valuation *intensively* is less costly than using different agents with negative φ_l^h valuations. This, in turn, means that type v_k 's matching set takes the form $[t_k(v_k), \bar{v}_l]$, where the threshold $t_k(v_k)$ is computed so that

$$\int_{[t_k(v_k), \bar{v}_l]} \sigma_l(v_l) dF_l(v_l) = q.$$

Building on the above ideas, the formal proof in the [Appendix](#) uses results from the monotone concave order to verify that when [Condition TP\(a\)](#) holds, starting from any incentive-compatible matching rule, one can construct a threshold rule that weakly improves upon the original. The idea is that threshold rules *minimize the costs of cross-subsidization* by delivering to those agents who play the role of consumers (i.e., whose φ_k^h valuation is nonnegative) matching sets of high quality in the most economical way. Note that the threshold rule constructed above is implementable provided that the original matching rule is implementable. In particular, under the new rule, among those agents with negative φ_l^h valuations, those with higher valuations may receive larger matching sets.

Next, consider the scenario of [Condition TP\(b\)](#), where $g_k(\cdot)$ is weakly convex on both sides and where $\sigma_k(\cdot)$ is weakly decreasing. Then pick a type v_k from side k with $\varphi_k^h(v_k) < 0$. Recall that using such an agent is costly for the platform. Now imagine that the platform wanted to assign to this type a matching set of strictly positive intensity, $|\mathbf{s}_k(v_k)|_l > 0$. The combination of the assumptions that (i) salience decreases with valuations, (ii) $g_l(\cdot)$ are weakly convex, and (iii) types are private information then implies that the most profitable way to use type v_k as an input is to match him to those agents from side l with the highest positive φ_l^h valuations. This is because (a) these latter types are the ones that benefit the most from interacting with type v_k (indeed, as required by incentive compatibility, these types have the matching sets with the highest intensity and, hence, by the convexity of $g_l(\cdot)$, the highest marginal utility for meeting additional agents) and (b) these latter types are the least salient ones and, hence, exert the lowest

negative externalities on type v_k (recall that $\varphi_k^h(v_k) < 0$). In the scenario covered by [Condition TP\(b\)](#), the reason why a threshold structure is thus optimal is that it *maximizes the welfare benefits (or profits) of cross-subsidization*. $\square$

Discussion

Before moving to the characterization of optimal threshold rules, we discuss the role of [Condition TP](#) and of private information for the result in [Proposition 1](#). Considered in isolation, neither [Condition TP](#) nor incentive compatibility is itself sufficient for the optimality of threshold rules. It is the *combination* of the cross-subsidization logic outlined in the proof sketch of [Proposition 1](#) with the monotonicity requirements of incentive compatibility that leads to the optimality of threshold rules. To illustrate this point, we exhibit two examples. The first one shows that threshold rules may not be optimal when information is incomplete but [Condition TP](#) fails. The second one shows that threshold rules may not be optimal when [Condition TP](#) holds but information is complete. The logic behind these examples clarifies what can “go wrong” once we dispense with either one of these conditions.

EXAMPLE 3 (Sub-optimality of threshold rules I). Agents from sides A and B have their valuations drawn uniformly from $V_A = [0, 1]$ and $V_B = [-1, 0]$, respectively. The salience of the side- B agents is constant and normalized to 1, i.e., $\sigma_B(v_B) \equiv 1$ for all $v_B \in V_B$, while the salience of the side- A agents is given by

$$\sigma_A(v_A) = \begin{cases} 10 & \text{if } v_A \in [\frac{9}{10}, 1] \\ \frac{10}{9} & \text{if } v_A \in [0, \frac{9}{10}]. \end{cases}$$

Preferences for matching intensity are linear on side A (that is, g_A is the identity function), whereas preferences on side B are given by the convex function¹¹

$$g_B(x) = \begin{cases} x & \text{if } x \leq 1 \\ +\infty & \text{if } x > 1. \end{cases}$$

In this environment, the welfare-maximizing threshold rule is described by the threshold function $t_A(v) = -v/10$, with exclusion types $\omega_A = \frac{9}{10}$ and $\omega_B = -\frac{1}{10}$, as can be easily verified from [Proposition 2](#) below. Total welfare under such a rule is $\frac{4}{10^3}$. Now consider the following non-threshold rule, which we describe by its side- A correspondence:

$$\mathbf{s}_A(v_A) = \begin{cases} [-\frac{1}{10}, 0] & \text{if } v_A \in [\frac{9}{10}, 1] \\ [-\frac{2}{10}, -\frac{1}{10}] & \text{if } v_A \in [0, \frac{9}{10}]. \end{cases}$$

It is easy to check that this matching rule is implementable. Total welfare under this rule equals $\frac{3}{10^2} > \frac{4}{10^3}$. $\diamond$

¹¹That the function g_B jumps at infinity at $x = 1$ simplifies the exposition but is not important for the result; the sub-optimality of threshold rules clearly extends to an environment identical to that in the example but where the function g_B is replaced by a sufficiently close smooth convex approximation.

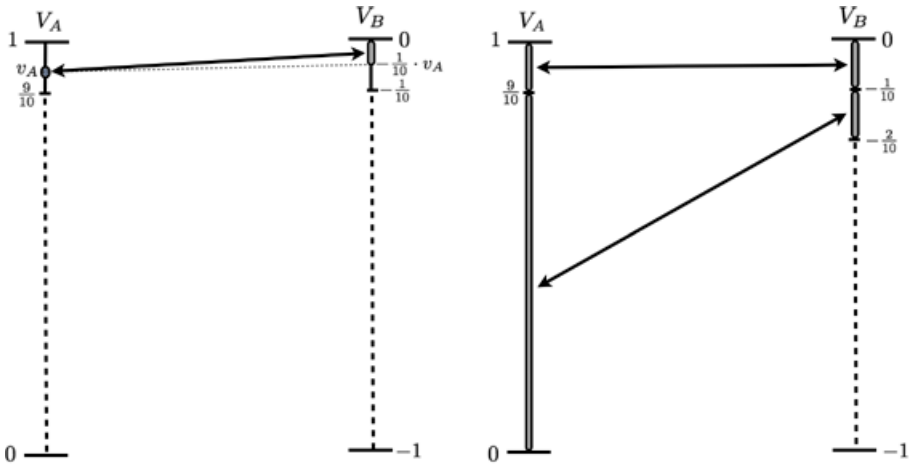


FIGURE 1. The welfare-maximizing rule among those with a threshold structure (left) and the welfare-improving nonthreshold rule (right) from Example 3.

The matching rules in this example are illustrated in Figure 1. To understand the logic of the example, let agents from side A be advertisers and agents from side B be browsers. The advertisers with the highest willingness-to-pay for eyeballs, $v_A \in [\frac{9}{10}, 1]$, are the most salient (i.e., their ads are perceived as the most annoying by the browsers). Browsers have convex nuisance costs, as described by the function g_B (in particular, the disutility from advertising becomes arbitrarily large once the salience-adjusted mass of advertising exceeds 1). In this example, the optimal threshold rule matches advertisers with a high willingness-to-pay with those browsers whose tolerance for advertising is sufficiently high, and assigns empty matching sets to all other advertisers and browsers. Because of the convexity of nuisance costs, few advertisers are matched to browsers under such a rule. The alternative rule proposed in the example better distributes advertisers to browsers. Under the proposed rule, advertisers with willingness-to-pay $v_A \in [0, \frac{9}{10}]$ (whose ads are not particularly annoying) are matched to browsers with moderate tolerance for advertising (i.e., $v_B \in [-\frac{2}{10}, -\frac{1}{10}]$), while advertisers with a high willingness-to-pay (whose ads are the most annoying) are matched with browsers whose tolerance for advertising is the highest (i.e., the matching allocation exhibits segmentation). Welfare under the proposed rule is almost 10 times higher than under the optimal threshold rule.

The example above violates Condition TP by exhibiting a salience function that is nondecreasing in valuations and preferences that are strictly convex in matching intensity. Similarly, one can show that threshold rules may fail to be optimal when salience is nonincreasing but preferences are strictly concave (see Appendix B for an example with this structure). The next example illustrates the role of private information for the result in Proposition 1.

EXAMPLE 4 (Sub-optimality of threshold rules II). Agents from sides A and B have valuations drawn uniformly from $V_A = [0, 1]$ and $V_B = [-2, 0]$, respectively. Preferences are

linear on both sides, that is, $g_k(x) = x$, $k = A, B$. The salience function on side A is constant, $\sigma_A(v_A) = 1$ for all $v_A \in V_A$, whereas the salience function on side B is given by

$$\sigma_B(v_B) = \begin{cases} 1 & \text{if } v_B \in [-1, 0] \\ 8 & \text{if } v_B \in [-2, -1]. \end{cases}$$

These preferences clearly satisfy [Condition TP\(b\)](#). Now suppose that valuations are publicly observable on both sides and that the platform maximizes welfare. The optimal matching rule is then given by

$$t_A(v_A) = \begin{cases} [-2, -1] \cup [-v_A, 0] & \text{if } v_A \geq \frac{1}{4} \\ [-8v_A, -1] \cup [-v_A, 0] & \text{if } \frac{1}{8} \leq v_A < \frac{1}{4} \\ [-v_A, 0] & \text{if } 0 \leq v_A < \frac{1}{8}. \end{cases}$$

Furthermore, no threshold rule yields the same welfare as the above rule. ◇

The key ingredient of the above example is that salience is decreasing in valuations on side B (which is the “input” side, as $v_B \leq 0$). As such, some of the most “expensive” agents from side B are the most attractive ones to the side- A agents. The welfare-maximizing rule (under complete information) then proceeds by evaluating separately each possible match between agents from the two sides (note that this follows from the fact that, in this example, g is linear on both sides, which implies that preferences are separable in the matches). It is then welfare-enhancing to include in the matching sets of side- A agents (whose valuation is positive) a disjoint collection of types from side B . The matching rule in the example, however, fails the monotonicity condition required by incentive compatibility (that is, the total salience of the matching sets is nonmonotone in valuations). As such, it is not implementable when types are private information.

Interestingly, threshold rules are more likely to be optimal when information is incomplete. This is discussed in the next remark.

REMARK 2 (Threshold rules: complete and incomplete information). Consider a welfare-maximizing platform. By virtue of [Lemma 1](#), whenever a threshold rule is optimal under complete information, it is also optimal under incomplete information; the reason is that feasible threshold rules are always implementable (see [Remark 1](#)). The converse is, however, false, as demonstrated by [Example 4](#).

Further assume that preferences are linear on both sides, that is, $g_k(x) = x$, $k = A, B$. In this case, the welfare-maximizing rule is obtained by evaluating separately each possible match between agents from the two sides. That is, agents with valuations v_A and v_B are matched if and only if

$$v_A \sigma_B(v_B) + v_B \sigma_A(v_A) \geq 0 \iff \frac{v_A}{\sigma_A(v_A)} + \frac{v_B}{\sigma_B(v_B)} \geq 0.$$

Together with [Lemma 1](#), the last inequality implies that a threshold rule is welfare-maximizing under complete information if and only if $v_k/\sigma_k(v_k)$ is weakly increasing, $k = A, B$. As discussed below, a similar condition determines whether welfare-maximizing matching rules under incomplete information exhibit bunching (see [Remark 6](#) below).

REMARK 3 (One-to-one matching). The optimality of threshold rules hinges on the assumption that the attractiveness of any set of agents is determined by the intensity of the set. When, instead, the attractiveness of a set is determined either by the average, or by the maximal salience, of its members, optimal rules are typically one-to-one and exhibit positive assortativeness (i.e., $\mathbf{s}_k(v_k) = F_l^{-1}(F_k(v_k))$ for all $v_k \in V_k$, $k = A, B$) when salience is increasing in values on both sides and negative assortativeness (i.e., $\mathbf{s}_k(v_k) = F_l^{-1}(1 - F_k(v_k))$ for all $v_k \in V_k$, $k = A, B$) when salience is decreasing in values on both sides.¹² Interestingly, when preferences are linear on both sides, that is, when $g_k(x) = x$, $k = A, B$, one can also reinterpret our results as describing the optimal matching rule between the types of a given pair of agents. As in [Myerson and Satterthwaite \(1983\)](#), in this case, the matching rule specifies whether matching between any pair of types of the two agents should occur.¹³

3.2 Properties of optimal threshold rules

Assuming throughout the rest of the paper that [Condition TP](#) holds, we then proceed by further investigating the properties of optimal threshold rules. To conveniently describe the agents' payoffs, we introduce the function $\hat{g}_k : V_l \rightarrow \mathbb{R}_+$ defined by

$$\hat{g}_k(v_l) \equiv g_k\left(\int_{v_l}^{\bar{v}_l} \sigma_l(x) dF_l(x)\right),$$

$k, l = A, B$, $l \neq k$. The utility that an agent with type v_k obtains from a matching set $[t_k(v_k), \bar{v}_l]$ can then be written concisely as $v_k \cdot \hat{g}_k(t_k(v_k))$. Note that $\hat{g}_k(t_k(v_k))$ is decreasing in $t_k(v_k)$, as increasing the threshold $t_k(v_k)$ reduces the intensity of the matching set.

Equipped with this notation, we can then recast the platform's problem as choosing a pair of threshold functions $(t_k^h(\cdot))_{k \in \{A, B\}}$ along with two scalars (ω_A, ω_B) so as to maximize the platform's objective subject to the conditions of [Lemma 1](#). Note that the reciprocity constraint (5) renders the platform's problem a nonstandard control problem (as each of the two controls $t_k(\cdot)$, $k \in \{A, B\}$, is required to coincide with the generalized inverse of the other).

The next definition extends to our two-sided matching setting the notion of separating schedules, as it appears, for example, in [Maskin and Riley \(1984\)](#).

DEFINITION 2 (Separation). The h -optimal matching rule entails the following qualities:

- (i) *It exhibits separation* if there exists a (positive measure) set $\hat{V}_k \subset V_k$ such that, for any $v_k, v'_k \in \hat{V}_k$, $t_k^h(v_k) \neq t_k^h(v'_k)$.
- (ii) *It exhibits exclusion at the bottom* on side k if $\omega_k^h > \underline{v}_k$.

¹²The result follows from arguments similar to those in [Damiano and Li \(2007\)](#).

¹³The same is true when g is nonlinear, but in this case our preference representation is no longer consistent with expected utility.

(iii) *It exhibits bunching at the top on side k if $t_l^h(\omega_l^h) < \bar{v}_k$.*

The rule is *maximally separating* if $t_k^h(\cdot)$ is strictly decreasing over the interval $[\omega_k^h, t_l^h(\omega_l^h)]$, which, hereafter, we refer to as the *separating range*.

Accordingly, separation occurs when *some* agents on the same side receive different matching sets. Exclusion at the bottom occurs when all agents in a neighborhood of $\underline{v}_k$ are assigned empty matching sets. Bunching at the top occurs when all agents in a neighborhood of $\bar{v}_k$ receive identical matching sets. In turn, maximal separation requires that, as valuations increase, matching sets strictly expand whenever they are “interior” (in the sense that $t_k^h(v_k) \in (\omega_l^h, t_l^h(\omega_l^h))$).

The following regularity condition guarantees that the optimal rules are maximally separating.

CONDITION MR (Match regularity). *The functions $\psi_k^h : V_k \rightarrow \mathbb{R}$ defined by*

$$\psi_k^h(v_k) \equiv \frac{f_k(v_k) \cdot \varphi_k^h(v_k)}{-\hat{g}_l'(v_k)} = \frac{\varphi_k^h(v_k)}{g_l'([v_k, \bar{v}_k]_k) \cdot \sigma_k(v_k)}$$

are strictly increasing, $k = A, B$, $h = W, P$.

As will be clear shortly, the optimal matching rules entail maximal separation if and only if **Condition MR** holds for every valuation in the separating range. Accordingly, this condition is the analog of Myerson’s standard regularity condition in two-sided matching problems.

To understand the condition, take the case of profit maximization, $h = P$. The numerator in $\psi_k^h(v_k)$ accounts for the effect on the platform’s revenue of an agent from side k with valuation v_k as a *consumer* (as his virtual valuation $\varphi_k^h(v_k)$ is proportional to the marginal revenue produced by the agent). In turn, the denominator accounts for the effect on the platform’s revenue of this agent as an *input* (as $-\hat{g}_l'(v_k)$ is proportional to the marginal utility brought by this agent to every agent from side l who is already matched to any other agent from side k with valuation above v_k). Therefore, the above regularity condition requires that, under a threshold rule, the contribution of an agent as a consumer (as captured by his virtual valuation) increases faster than his contribution as an input.

REMARK 4 (Conditions MR and TP). Note that, except in the uninteresting case in which $\varphi_k^h(v_k) < 0$ (alternatively, $\varphi_k^h(v_k) \geq 0$) for all $v_k \in V_k$, $k = A, B$ and $h = W, P$, **Condition MR** is not implied by, nor implies, **Condition TP**. In particular, **Condition MR** requires that $\varphi_k^h(v_k)$ increases faster than $g_l'([v_k, \bar{v}_k]_k)\sigma_k(v_k)$ over $[r_k^h, \bar{v}_k]$ (that is, over the subset of V_k in which $\varphi_k^h(v_k) > 0$) and that $|\varphi_k^h(v_k)|$ decreases slower than $g_l'([v_k, \bar{v}_k]_k)\sigma_k(v_k)$ over $[\underline{v}_k, r_k^h]$ (that is, over the subset of V_k in which $\varphi_k^h(v_k) < 0$) for $k = A, B$ and $h = W, P$.

To better appreciate the platform's trade-offs at the optimum, it is convenient to introduce the *marginal surplus function* $\Delta_k^h : V_k \times V_l \rightarrow \mathbb{R}$ defined by

$$\Delta_k^h(v_k, v_l) \equiv -\hat{g}'_k(v_l) \cdot \varphi_k^h(v_k) \cdot f_k(v_k) - \hat{g}'_l(v_k) \cdot \varphi_l^h(v_l) \cdot f_l(v_l)$$

for $k, l \in \{A, B\}$, $l \neq k$. Note that $\Delta_A^h(v_A, v_B) = \Delta_B^h(v_B, v_A)$ represents the marginal effect on the platform's objective of decreasing the threshold $t_A^h(v_A)$ below v_B , while, by reciprocity, also reducing the threshold $t_B^h(v_B)$ below v_A .

PROPOSITION 2 (Optimal rules). *Assume Conditions TP and MR hold. Then, for both $h = W$ and $h = P$, the h -optimal matching rules are such that $\mathbf{s}_k^h(v_k) = V_l$ for all $v_k \in V_k$, $k = A, B$, if $\Delta_k^h(\underline{v}_k, \underline{v}_l) \geq 0$.¹⁴*

When, instead, $\Delta_k^h(\underline{v}_k, \underline{v}_l) < 0$, the h -optimal matching rule is maximally separating and entails the following components:

- (i) *It has bunching at the top on side k and no exclusion at the bottom on side l if $\Delta_k^h(\bar{v}_k, \underline{v}_l) > 0$.*
- (ii) *It has exclusion at the bottom on side l and no bunching at the top on side k if $\Delta_k^h(\bar{v}_k, \underline{v}_l) < 0$.¹⁵*

Finally, the threshold function $t_k^h(\cdot)$ is implicitly defined by the Euler equation

$$\Delta_k^h(v_k, t_k^h(v_k)) = 0 \quad (6)$$

for any v_k in the separating range $[\omega_k^h, t_l^h(\omega_l^h)]$.

The optimal matching rule thus entails separation whenever the marginal surplus function evaluated at the lowest valuations on both sides of the market is negative: $\Delta_k^h(\underline{v}_k, \underline{v}_l) < 0$. When, instead, this condition fails, each agent from each side is matched to any other agent from the opposite side: $\mathbf{s}_k^h(v_k) = V_l$ for all $v_k \in V_k$, $k = A, B$.

When separation occurs, **Proposition 2** sheds light on the optimal cross-subsidization strategy employed by the platform. To illustrate, consider the case of profit maximization (the arguments for welfare maximization are analogous), and let $\underline{v}_k < 0$ for $k = A, B$. An important feature of the profit-maximizing rule is that $t_k^P(v_k) \leq r_l^P$ if and only if $v_k \geq r_k^P$, where the reservation type r_k^P is the lowest type for whom $\varphi_k^P(v_k) \geq 0$. This implies that agents from each side of the market are endogenously partitioned into two groups. Those agents with positive virtual valuations (equivalently, with valuations $v_k \geq r_k^P$) play the role of *consumers*, “purchasing” sets of agents from the other side of the market (these agents contribute positively to the platform's profits). In turn, those agents with negative virtual valuations (equivalently, with valuation $v_k < r_k^P$) play the role of *inputs* in the matching process, providing utility to those agents from the opposite side

¹⁴The statement above holds true even when **Condition MR** is violated, provided that either $\varphi_k^h(\underline{v}_k) < 0$ for $k = A, B$ or $\varphi_k^h(\underline{v}_k) > 0$ for $k = A, B$. **Condition MR** is only needed in the case where $\varphi_k^h(\underline{v}_k) > 0 > \varphi_l^h(\underline{v}_l)$ for $k, l \in \{A, B\}$.

¹⁵In the knife-edge case where $\Delta_k^h(\bar{v}_k, \underline{v}_l) = 0$, the h -optimal rule entails neither bunching at the top on side k nor exclusion at the bottom on side l .

they are matched to (these agents contribute negatively to the platform's profits). At the optimum, the platform recovers the "costs" of procuring agents-inputs from the gains obtained by agents-consumers.

The Euler equation (6) in the proposition then describes the optimal level of cross-subsidization for each type. In particular this equation can be rewritten as

$$\underbrace{-\hat{g}'_k(t_k^P(v_k)) \cdot \varphi_k^P(v_k) \cdot f_k(v_k)}_{\text{marginal gains}} = \underbrace{\hat{g}'_l(v_k) \cdot \varphi_l^P(t_k^P(v_k)) \cdot f_l(t_k^P(v_k))}_{\text{marginal losses}}. \quad (7)$$

At the optimum, the platform equalizes the marginal gains and the marginal losses of expanding the matching set of each agent in the separating range. When v_k corresponds to an agent-consumer (i.e., when $\varphi_k^P(v_k) > 0$), the left-hand side of (7) is the marginal revenue of expanding the agent's matching set, starting from a situation in which the agent is matched already to all agents from the other side whose valuation is above $t_k^P(v_k)$. In turn, the right-hand side of (7) is the marginal cost associated with procuring extra agents-inputs from the opposite side; under a threshold rule, this cost is the loss that the platform incurs by expanding the matching set of an agent from side l whose valuation is $v_l = t_k^P(v_k)$, starting from a situation where such an agent is already matched to all agents from side k whose valuation exceeds $v_k = t_l^P(v_l)$, as required by reciprocity (recall that $t_l^P(t_k^P(v_k)) = v_k$). The terms $\hat{g}'_k(t_k^P(v_k))$ and $\hat{g}'_l(v_k)$ adjust the marginal utilities to account for the effect of the new matches on the supramarginal matches (i.e., those matches above the profit-maximizing thresholds).

Note that optimality also implies that there is bunching at the top on side k if and only if there is no exclusion at the bottom on side l . In other words, bunching can only occur at the top due to binding capacity constraints, that is, when the "stock" of agents from side $l \neq k$ has been exhausted.

REMARK 5. **Condition MR** is necessary and sufficient for the marginal surplus function $\Delta_k^h(v_k, v_l)$ to satisfy the following single-crossing property: whenever $\Delta_k^h(v_k, v_l) \geq 0$, then $\Delta_k^h(v_k, \hat{v}_l) > 0$ for all $\hat{v}_l > v_l$ and $\Delta_k^h(\hat{v}_k, v_l) > 0$ for all $\hat{v}_k > v_k$. As can be seen from the Euler equation (6), this single-crossing property is equivalent to the threshold function $t_k^h(\cdot)$ being strictly decreasing over the separating range. Therefore, **Condition MR** is the "weakest" regularity condition that rules out nonmonotonicities (or bunching) in the matching rule.

REMARK 6. Consider a welfare-maximizing platform and assume that preferences are linear on both sides, that is, $g_k(x) = x$, $k = A, B$. Then the following statements are equivalent: (i) **Condition MR** holds; (ii) the optimal matching rule under complete information exhibits a threshold structure; (iii) the optimal matching rule under incomplete information is maximally separating.

Consider the environment described in **Example 4**. It follows from the above remark that, under incomplete information, the welfare-maximizing matching rule is not maximally separating (therefore exhibiting some interior bunching interval). Accordingly, threshold rules are welfare-maximizing under incomplete information, *but not*

under complete information, exactly when the monotonicity constraint associated with incentive compatibility is binding at the optimum. The presence of this constraint explains why incomplete information is more conducive to threshold rules than complete information.

The next example illustrates the characterization of [Proposition 2](#).

EXAMPLE 5 (Advertising avoidance). Consider an online intermediary matching advertisers to browsers with convex nuisance costs, as in [Example 1](#). Assume that browsers and advertisers have valuations drawn from a uniform distribution over $V_A = (-1, 0)$ and $V_B = (0, 1)$, respectively. The advertisers' salience function is $\sigma_B(v_B) = 1/v_B$. It is easy to check that Conditions [TP\(b\)](#) and [MR](#) are satisfied. From [Proposition 2](#), the welfare-maximizing rule is described by the threshold function

$$t_B^W(v_B) = -\frac{v_B^2}{(1 + \beta)[- \log v_B]^\beta}$$

for any v_B in the separating range. Moreover, there is bunching at the top of side B , i.e., all advertisers with high enough valuation are matched to all browsers. As the nuisance cost β increases, all advertisers obtain weakly smaller matching sets (strictly so for advertisers in the separating range). $\diamond$

3.3 Welfare-maximizing versus profit-maximizing rules

We now turn to the distortions brought in by profit maximization relative to the welfare-maximizing matching rule. Consider the following example.

EXAMPLE 6 (Business-to-business platform). Let the environment be as in [Example 2](#) and assume that all v_k are drawn from a uniform distribution over $[\underline{v}, \bar{v}]$, with $\underline{v} > 0$ and $2\underline{v} < \bar{v}$, $k = A, B$. Because linking any two firms generates positive surplus, the welfare-maximizing rule matches all firms on each level of the supply chain. Next consider the profit-maximizing rule. It is easy to check that Conditions [TP\(a\)](#) and [MR](#) are satisfied. Because $\Delta_k^P(\underline{v}, \underline{v}) = \underline{v}(2\underline{v} - \bar{v}) < 0$, it follows from [Proposition 2](#) that the profit-maximizing rule entails separation and is described by the threshold function

$$t_A^P(v_A) = \frac{\bar{v}(1 - \alpha)v_A}{2v_A - \alpha\bar{v}}$$

defined over $(\omega_k^P, \bar{v}) = (\alpha\bar{v}/(1 + \alpha), \bar{v})$. Under profit maximization, there is exclusion at the bottom on both levels and each firm that is not excluded from the platform is matched to a strict subset of its efficient matching set. $\diamond$

The matching rules in this example are illustrated in [Figure 2](#). As indicated in the next proposition, the distortions in this example are general properties of profit-maximizing rules (the proof follows directly from [Proposition 2](#)).

PROPOSITION 3 (Distortions). *Assume Conditions [TP](#) and [MR](#) hold. Relative to the welfare-maximizing rule, the profit-maximizing rule has the following qualities:*

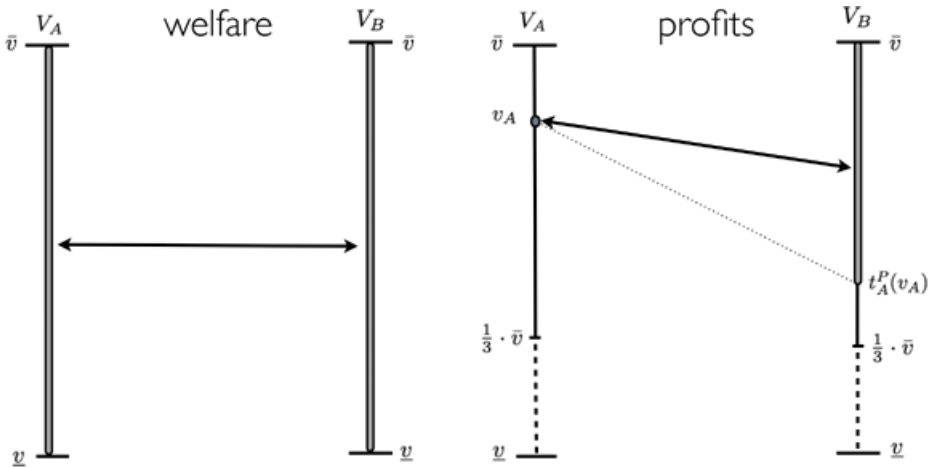


FIGURE 2. The welfare-maximizing matching rule (left) and the profit-maximizing matching rule (right) from Example 6 when $\alpha = \frac{1}{2}$.

- (i) It completely excludes a larger group of agents (exclusion effect), i.e., $\omega_k^P \geq \omega_k^W$, $k = A, B$.
- (ii) It matches each agent who is not excluded to a subset of his efficient matching set (isolation effect), i.e., $\mathbf{s}_k^P(v_k) \subseteq \mathbf{s}_k^W(v_k)$ for all $v_k \geq \omega_k^P$, $k = A, B$.

The intuition for both effects can be seen by comparing $\Delta_k^P(v_k, v_l)$ with $\Delta_k^W(v_k, v_l)$: under profit maximization, the platform only internalizes the cross-effects on marginal revenues (which are proportional to the virtual valuations), rather than the cross-effects on welfare (which are proportional to the true valuations). Contrary to other mechanism design problems, inefficiencies do not necessarily vanish as agents' types approach the "top" of the distribution (i.e., the highest valuation for matching intensity). The reason is that although virtual valuations converge to the true valuations as agents' types approach the top of the distribution, the cost of cross-subsidizing these types remains strictly higher under profit maximization than under welfare maximization, due to the inframarginal losses implied by reciprocity on the opposite side.

3.4 Comparative statics: The detrimental effects of becoming more attractive

Shocks that alter the cross-side effects of matches are common in vertical matching markets. Changes in productivity, for example, affect the pricing strategies of business-to-business platforms, for they affect the attractiveness of business connections for the same population of firms.

The next definition formalizes the notion of a change in attractiveness. We restrict the attention here to a platform maximizing profits in a market where all agents from each side value positively interacting with agents from the opposite side (i.e., $v_k \geq 0$ for $k = A, B$). For simplicity, we also restrict attention to markets in which preferences for matching intensity are linear (i.e., $g_A(x) = g_B(x) = x$ for all $x \in \mathbb{R}_+$).

DEFINITION 3 (Higher attractiveness). Consider a market in which all agents value positively interacting with agents from the opposite side, i.e., $\underline{v}_k \geq 0$ for $k = A, B$. Side k is more attractive under $\hat{\sigma}_k(\cdot)$ than under $\sigma_k(\cdot)$ if $\hat{\sigma}_k(v_k) \geq \sigma_k(v_k)$ for all $v_k \in V_k$, with the inequality strict for a positive-measure subset of V_k .

The definition above does not impose that the attractiveness of side- k agents increases uniformly across agents. For instance, it allows that only the attractiveness of agents with high valuations increases, while that of other agents remains the same. The next proposition describes how the profit-maximizing matching rule changes as side k becomes more attractive.

PROPOSITION 4 (Increase in attractiveness). *Consider a market in which (a) Conditions **TP** and **MR** hold, (b) all agents value positively interacting with agents from the opposite side (i.e., $\underline{v}_k \geq 0$ for $k \in \{A, B\}$), and (c) preferences for matching intensity are linear (i.e., $g_A(x) = g_B(x) = x$ for all $x \in \mathbb{R}_+$). Suppose side k becomes more attractive. Then a profit-maximizing platform switches from a matching rule $\mathbf{s}_k^P(\cdot)$ to a matching rule $\hat{\mathbf{s}}_k^P(\cdot)$ such that the following statements hold:*

- (i) *The matching sets on side k increase for those agents with a low valuation and decrease for those agents with a high valuation, i.e., $\hat{\mathbf{s}}_k^P(v_k) \supseteq \mathbf{s}_k^P(v_k)$ if and only if $v_k \leq r_k^P$.*
- (ii) *Low-valuation agents from side k are better off, whereas the opposite is true for high-valuation agents, i.e., there exists $\hat{v}_k \in (r_k^P, \bar{v}_k]$ such that $\Pi_k(v_k; \hat{M}^P) \geq \Pi_k(v_k; M^P)$ if and only if $v_k \leq \hat{v}_k$.*

Intuitively, an increase in the attractiveness of side- k agents alters the costs of cross-subsidization between the two sides. Recall that agents with $v_k \geq r_k^P$ are valued by the platform mainly as consumers. As these agents become more attractive, the costs of cross-subsidizing their “consumption” using agents from side l with negative virtual valuation increases, whereas the revenue gains on side k are unaltered. As a consequence, the matching sets of these agents shrink. The opposite is true for those agents with valuation $v_k \leq r_k^P$. These agents are valued by the platform mainly as inputs; as they become better inputs, their matching sets expand.

Perhaps surprisingly, an agent from side k can suffer from a positive shock to his own attractiveness (holding constant the attractiveness of all other agents). To understand why, consider the case where salience $\sigma_k(v_k)$ is strictly increasing and take an agent from side k with the highest possible valuation, i.e., for whom $v_k = \bar{v}_k$. Assume that $\sigma_k(\bar{v}_k)$ increases by $\delta > 0$ while $\sigma_k(v_k)$ for all $v_k < \bar{v}_k$ remains constant. Because the revenues collected from any agent with valuation $\bar{v}_k$ are unaltered, at the optimum, the matching set of any agent with valuation $\bar{v}_k$ must shrink. At the same time, because the size of the matching sets must be monotone in types, the matching sets of all agents whose valuation is close to $\bar{v}_k$ must also shrink. As a result, an agent with valuation $\bar{v}_k$ is negatively affected by an increase in his own attractiveness, even if all other agents’ attractiveness does not change.

This result has interesting implications in terms of payoffs. For all $v_k \leq r_k^P$, we can evaluate payoffs according to

$$\Pi_k(v_k; M^P) = \int_{\underline{v}_k}^{v_k} |\mathbf{s}_k(x)|_l dx \leq \int_{\underline{v}_k}^{v_k} |\hat{\mathbf{s}}_k(x)|_l dx = \Pi_k(v_k; \hat{M}^P),$$

meaning that all agents from side k with valuation $v_k \leq r_k^h$ are necessarily better off. Alternatively, since $|\hat{\mathbf{s}}_k(v_k)|_l \leq |\mathbf{s}_k(v_k)|_l$ for all $v_k \geq r_k^h$, then either payoffs increase for all agents from side k or there exists a threshold $\hat{v}_k > r_k^h$ such that the payoff of each agent from side k is higher under the new rule than under the original one if and only if $v_k \leq \hat{v}_k$.

In many applications, the agents' payoffs and matching sets are not observable, whereas the prices charged by the platform are publicly available (e.g., business-to-business platforms do not offer precise descriptions of how the matching sets assigned to firms are determined, yet, the prices charged are clear). To derive a testable implication of [Proposition 4](#), the next corollary studies the impact of the agents' attractiveness on the platform's pricing policy.

For any matching intensity q_k , let $\rho_k^P(q_k)$ denote the total price that each agent from side k has to pay for any matching set of intensity q_k under the profit-maximizing mechanism M^P . By optimality, the tariff $\rho_k^P(\cdot)$ has to satisfy

$$\rho_k^P(q_k) = p_k^P(v_k) \quad \text{for all } v_k \text{ such that } |\mathbf{s}_k^P(v_k)|_l = q_k.$$

We then have the following result.

COROLLARY 1 (Effect of an increase in attractiveness on prices). *Under the assumptions of [Proposition 4](#), if the attractiveness of side k increases (in the sense of [Definition 3](#)), the platform's price schedule rotates anticlockwise. That is, the platform switches from a price schedule $\rho_k^P(\cdot)$ to a price schedule $\hat{\rho}_k^P(\cdot)$ such that $\hat{\rho}_k^P(q_k) \leq \rho_k^P(q_k)$ for any matching set of intensity $q_k \leq \hat{q}_k$, where $\hat{q}_k > |\mathbf{s}_k^P(r_k^P)|_l = |\hat{\mathbf{s}}_k^P(r_k^P)|_l$.*

An increase in the attractiveness of side k thus triggers a reduction in the price that the platform charges on side k for matching sets of low intensity, possibly along with an increase in the price it charges for matching sets of high intensity.

4. EXTENSIONS

The analysis developed above can accommodate a few simple enrichments, which we discuss hereafter.

Imperfect correlation between salience and valuation

To simplify the exposition, the baseline model assumes that salience is a deterministic function of the valuations. As mentioned above, all our results extend to environments in which the two dimensions are imperfectly correlated and in which agents have private

information about both dimensions. We formally establish this result in the [Appendix](#) by first relaxing [Condition TP](#) to require that salience and valuation be positively (or, alternatively, negatively) affiliated. This is the natural generalization of the assumption that σ_k be increasing (or, alternatively, decreasing) in v_k , as required by [Condition TP](#). Under this condition, we then show that the optimal matching rules have a threshold structure, with the thresholds depending on valuations but not on salience. Note that the result is not a mere consequence of the fact that individual preferences are invariant in the agents' own salience. Combined with incentive compatibility, the latter property only implies that the matching intensity is invariant in the agent's own salience, thus permitting the *composition* of the matching sets to depend on salience. Once this result is established, it is then immediate that all other results in the paper extend to this richer environment.

The group design problem

Consider now the problem of how to assign agents to different “teams” in the presence of peer effects, which is central to the theory of organizations and to personnel economics. As anticipated in the [Introduction](#), such a one-sided matching problem is a special case of the two-sided matching problems studied in this paper. To see this, note that the problem of designing nonexclusive groups in a one-sided matching setting is mathematically equivalent to the problem of designing an optimal matching rule in a two-sided matching setting where (i) the preferences and type distributions of the two sides coincide, and (ii) the matching rule is required to be symmetric across sides, i.e., $\mathbf{s}_A(v) = \mathbf{s}_B(v)$ for all $v \in V_A = V_B$.

Under the new constraint that matching rules be symmetric across the two sides, maximizing (4) is equivalent to maximizing twice the objective function associated with the one-sided matching problem. As it turns out, the symmetry constraint is never binding in a two-sided matching market in which the two sides are symmetric (in which case $\psi_l^h(\cdot) = \psi_k^h(\cdot)$). Indeed, the characterization from [Proposition 2](#) reveals that, at any point where the threshold rule $t_k^h(\cdot)$ is strictly decreasing, $t_k^h(v) = (\psi_l^h)^{-1}(-\psi_k^h(v)) = (\psi_k^h)^{-1}(-\psi_l^h(v)) = t_l^h(v)$. It is also easy to see that the symmetry condition is satisfied when the optimal rule entails bunching at the top.

Coarse matching

In reality, platforms typically offer menus with finitely many alternatives. As pointed out by [McAfee \(2002\)](#) and [Hoppe et al. \(2011\)](#), the reason for such *coarse matching* is that platforms may face costs for adding more alternatives to their menus.¹⁶ It is easy to see that the analysis developed above extends to a setting where the platform can include no more than N plans in the menus offered to each side. Furthermore, as the number of plans increases (e.g., because menu costs decrease), the solution to the platform's problem uniformly converges to the h -optimal rule identified in the paper.¹⁷ In other

¹⁶See also [Wilson \(1989\)](#).

¹⁷This follows from the fact that any weakly decreasing threshold function $t_k(\cdot)$ can be approximated arbitrarily well by a step function in the sup-norm, i.e., in the norm of uniform convergence.

words, the maximally separating matching rules of [Proposition 2](#) are the limit as N grows large of those rules offered when the number of plans is finite.

Quasi-fixed costs

Permitting an agent to interact with agents from the other side of the market typically involves a quasi-fixed cost. From the perspective of the platform, these costs are quasi-fixed, in the sense that they depend on whether an agent is completely excluded, but not on the composition of the agent's matching set.

The analysis developed above can easily accommodate such costs. Let c_k denote the quasi-fixed cost that the platform must incur for each agent from side k whose matching set is nonempty. The h -optimal mechanism can then be obtained through the following two-step procedure:

Step 1. Ignore quasi-fixed costs and maximize (4) among all weakly decreasing threshold functions $t_k^h(\cdot)$.

Step 2. Given the optimal threshold function $t_k^h(\cdot)$ from Step 1, choose the h -optimal exclusion types ω_A^h, ω_B^h by solving the problem

$$\max_{\omega_A, \omega_B} \sum_{k=A, B} \int_{\omega_k}^{\bar{v}_k} (\hat{g}_k(\max\{t_k^h(v_k), \omega_l\}) \cdot \varphi_k^h(v_k) - c_k) \cdot dF_k(v_k).$$

As the quasi-fixed costs increase, so do the exclusion types $\omega_k^h(c_A, c_B)$, $k = A, B$. For c_k sufficiently high, the exclusion types reach the reservation values r_k^h , in which case the platform switches from offering a menu of matching plans to offering a unique plan. Therefore, another testable prediction that the model delivers is that, *ceteris paribus*, discrimination should be more prevalent in matching markets with low quasi-fixed costs.

Robust implementation

In the direct revelation version of the matching game, each agent from each side is asked to submit a report v_k , which leads to a payment $p_k^h(v_k)$ as defined in (3), and grants access to all agents from the other side of the market who reported a valuation above $t_k^h(v_k)$. This game admits one Bayes–Nash equilibrium implementing the h -optimal matching rule $\mathbf{s}_k^h(\cdot)$, along with other equilibria implementing different rules.¹⁸

As pointed out by [Weyl \(2010\)](#) in the context of a monopolistic platform offering a single plan, equilibrium uniqueness can, however, be guaranteed when network effects depend only on quantities (i.e., when $\sigma_k(\cdot) \equiv 1$ for $k = A, B$).¹⁹ In the context of our

¹⁸In the implementation literature, this problem is referred to as *partial implementation*, whereas in the two-sided market literature, it is referred to as the *chicken and egg* problem (e.g., [Caillaud and Jullien 2001, 2003](#)) or the *failure to launch* problem (e.g., [Evans and Schmalensee 2010](#)). See also [Ellison and Fudenberg \(2003\)](#) and [Ambrus and Argenziano \(2009\)](#).

¹⁹See also [White and Weyl \(2015\)](#).

model, it suffices to replace the payment rule $(p_k^h(\cdot))_{k=A,B}$ given by (3) with the payment rule

$$\begin{aligned} \varrho_k^h(v_k, (v_l^j)^{j \in [0,1]}) &= v_k \cdot g_k(|\{j \in [0, 1] : v_l^j \geq t_k(v_k)\}|_l) \\ &\quad - \int_{\underline{v}_k}^{v_k} g_k(|\{j \in [0, 1] : v_l^j \geq t_k(x)\}|_l) dx, \end{aligned} \quad (8)$$

where $|\{j \in [0, 1] : v_l^j \geq t_k(v_k)\}|_k \equiv \int_{\{j: v_l^j \geq t_k(v_k)\}} d\lambda(j)$ denotes the Lebesgue measure of agents from side $l \neq k$ reporting a valuation above $t_k(v_k)$. Given the above payment rule, it is *weakly dominant* for each agent to report truthfully. This follows from the fact that, *given any* profile of reports $(v_l^j)^{j \in [0,1]}$ by agents from the opposite side, the intensity of the matching set for each agent from side k is increasing in his report, along with the fact that the payment rule $\varrho_k^h(\cdot; (v_l^j)^{j \in [0,1]})$ satisfies the familiar envelope formula with respect to v_k . In the spirit of the Wilson doctrine, this also means that the optimal allocation rule can be robustly fully implemented in weakly undominated strategies.²⁰

5. CONCLUDING REMARKS

The analysis reveals how matching patterns reflect optimal cross-subsidization between sides in centralized markets. We deliver two main results. First, we identify conditions on primitives under which the optimal matching rules have a simple threshold structure, according to which agents with a low valuation for matching are included only in the matching sets of those agents from the opposite side whose valuation is sufficiently high. While these conditions are arguably weak, they cannot be dispensed with. We demonstrate this fact by means of counterexamples highlighting the complementary role that incentive compatibility and the monotonicity requirements on salience and marginal utility play in the optimality of threshold rules.

Second, we show that the optimal matching rules are determined by a simple formula that equalizes the marginal gains in welfare (or, alternatively, in profits) with the cross-subsidization losses that the platform must incur on the opposite side of the market. We show that the optimal rules *endogenously* separate agents into consumers and inputs. At the margin, the “costs” of procuring agents-inputs are recovered from the gains from agents-consumers (cross-subsidization).

The model is flexible enough to permit interesting comparative statics. For example, we show that when the attractiveness of one side increases, a profit-maximizing platform responds by reducing the intensity of the matching sets offered to those agents whose valuation is high, and by increasing the intensity of the matching sets offered to those agents whose valuation is low. This leads to lower (respectively, higher) payoffs to

²⁰With more general preferences, it is still possible to robustly fully implement any monotone matching rule in weakly undominated strategies by replacing the definition of $|\{j \in [0, 1] : v_l^j \geq t_k(v_k)\}|_l$ in (8) with $|\{j \in [0, 1] : v_l^j \geq t_k(v_k)\}|_l \equiv \int_{\{j: v_l^j \geq t_k(v_k)\}} \underline{q}_l d\lambda(j)$, where $\underline{q}_l \equiv \min\{\sigma_l(v_l) : v_l \in V_l\}$. However, these payments generate less revenue than those given in (3), implying that, in general, there is a genuine trade-off between robust full implementation and profit maximization.

those agents at the top (respectively, bottom) of the valuation distribution, and induces an anticlockwise rotation of the price schedule.

The above analysis is worth extending in a few important directions. For example, all the results are established assuming that the utility/profit that each agent derives from any given matching set is independent of who else from the same side has access to the same set. This is a reasonable starting point but is definitely inappropriate for certain markets. In advertising, for example, reaching a certain set of consumers is more profitable when competitors are blocked from reaching the same set. Extending the analysis to accommodate for “congestion effects” and other “same-side externalities” is challenging but worth exploring.

Likewise, the analysis focuses on a market with a single platform. Many matching markets are populated by competing platforms. Understanding to what extent the distortions identified in the present paper are affected by the degree of market competition, and studying policy interventions aimed at inducing platforms “to get more agents on board” (for example, through subsidies, and in some cases the imposition of universal service obligations) are other important venues for future research (see Damiano and Li 2008, Lee 2014, and Jullien and Pavan 2014 for models of platform competition in settings with a limited degree of price discrimination).

APPENDIX A

This appendix collects all proofs omitted in the text.

PROOF OF LEMMA 1. Necessity. We first show that $t_k(\cdot)$ must be weakly decreasing, $k = A, B$. Toward a contradiction, assume that $t_k(\cdot)$ is strictly increasing in an open neighborhood of $v_k \in V_k$. This means that there exists $\varepsilon > 0$ such that $t_k(v_k + \varepsilon) > t_k(v_k)$. Let $\hat{v}_l \equiv \frac{1}{2}t_k(v_k + \varepsilon) + \frac{1}{2}t_k(v_k)$, and note that $\hat{v}_l \in \mathbf{s}_k(v_k)$ and that $t_l(\hat{v}_l) \leq v_k$ (else, reciprocity is violated). Therefore, $v_k + \varepsilon \in \mathbf{s}_l(\hat{v}_l)$ and yet $\hat{v}_l \notin \mathbf{s}_k(v_k + \varepsilon)$, violating reciprocity. Hence, condition (ii) in Lemma 1 must hold.

Next, we show that condition (i) in Lemma 1 must also hold. To this end, let $\tilde{v}_l(v_k) \equiv \inf\{v_l : t_l(v_l) \leq v_k\}$. We first show that $t_k(v_k) = \tilde{v}_l(v_k)$ and then prove that $\tilde{v}_l(v_k) = \min\{v_l : t_l(v_l) \leq v_k\}$ (that is, a minimum exists).

We proceed again by contradiction and assume that there exists some $v_k \in V_k$ such that $t_k(v_k) \neq \tilde{v}_l(v_k)$. If $t_k(v_k) > \tilde{v}_l(v_k)$, there exists $v_l \geq \tilde{v}_l(v_k)$ such that $t_l(v_l) \leq v_k$ and $t_k(v_k) > v_l$. This implies that $v_k \in \mathbf{s}_l(v_l)$ and yet $v_l \notin \mathbf{s}_k(v_k)$, violating reciprocity.

Therefore, it must be that $t_k(v_k) < \tilde{v}_l(v_k)$. Let $\check{v}_l \equiv \frac{1}{2}\tilde{v}_l(v_k) + \frac{1}{2}t_k(v_k) \in (t_k(v_k), \tilde{v}_l(v_k))$ and notice that $t_l(\check{v}_l) > v_k$. But then $\check{v}_l \in \mathbf{s}_k(v_k)$ and yet $v_k \notin \mathbf{s}_l(\check{v}_l)$, again violating reciprocity. We conclude that $t_k(v_k) = \tilde{v}_l(v_k)$.

Finally, suppose that $\min\{v_l : t_l(v_l) \leq v_k\}$ does not exist. Because $t_k(v_k) = \tilde{v}_l(v_k)$, it follows that $t_l(\tilde{v}_l(v_k)) > v_k$, a violation of reciprocity. We conclude that condition (i) in Lemma 1 is also necessary.

Sufficiency. Take any $v_l \in \mathbf{s}_k(v_k)$. By definition of a threshold rule, $v_l \geq t_k(v_k)$. Furthermore, by condition (ii) in the lemma, $t_l(v_l) \leq t_l(t_k(v_k))$. In turn, by condition (i), $t_l(t_k(v_k)) \leq v_k$. This means that $t_l(v_l) \leq v_k$ and hence $v_k \in \mathbf{s}_l(v_l)$. This concludes the proof. $\square$

PROOF OF PROPOSITION 1. Below we prove a stronger result that supports both the claim in the proposition as well as the claim in [Section 4](#) about the optimality of threshold rules in environments where salience is imperfectly correlated with the valuation and where agents have private information about both dimensions.

To this purpose, we enrich the model as follows. For each $v_k \in V_K$, let $\Psi_k(\cdot|v_k)$ denote the conditional distribution of σ_k , given v_k , $k = A, B$, and denote by $\Lambda_k = F_k \cdot \Psi_k$ the measure defined by the product of F_k and Ψ_k . Now assume that agents observe both v_k and σ_k at the time they interact with the platform. Each agent's type is then given by the bi-dimensional vector $\theta_k \equiv (v_k, \sigma_k) \in \Theta_k \equiv V_k \times \Sigma_k$, with $\Sigma_k \subset \mathbb{R}_+$. In this environment, a matching mechanism $M = \{\mathbf{s}_k(\cdot), p_k(\cdot)\}_{k=A,B}$ continues to be described by a pair of matching rules and a pair of payment rules, with the only difference that $p_k(\cdot)$ now maps Θ_k into $\mathbb{R}$, whereas $\mathbf{s}_k(\cdot)$ maps Θ_k into the Borel sigma algebra over Θ_l , $k, l = A, B, l \neq k$. With some abuse of notation, hereafter we will denote by $|\mathbf{s}_k(\theta_k)|_l = \int_{(v_l, \sigma_l) \in \mathbf{s}_k(\theta_k)} \sigma_l d\Lambda_l$ the matching intensity of the set $\mathbf{s}_k(\theta_k)$.

Now consider the following extension of [Condition TP](#) in the main text.

CONDITION TP-EXTENDED. *One of the following two sets of conditions holds for both $k = A$ and $k = B$:*

(i.a) *The function $g_k(\cdot)$ is weakly concave, and (i.b) the random variables $\tilde{\sigma}_k$ and $\tilde{v}_k$ are weakly positively affiliated.*

(ii.a) *The function $g_k(\cdot)$ is weakly convex, and (ii.b) the random variables $\tilde{\sigma}_k$ and $\tilde{v}_k$ are weakly negatively affiliated.*²¹

Below we will prove the following claim.

CLAIM 1. *Assume [Condition TP-extended](#) holds. Then both the profit-maximizing ($h = P$) and the welfare-maximizing ($h = W$) rules discriminate only along the valuation dimension (that is, $s_k^h(v_k, \sigma_k) = s_k^h(v_k, \sigma'_k)$ for any $k = A, B$, $v_k \in V_k$, $\sigma_k, \sigma'_k \in \Sigma_k$, $h = W, P$) and are threshold rules. That is, there exists a scalar $\omega_k^h \in [\underline{v}_k, \bar{v}_k]$ and a nonincreasing function $t_k^h: V_k \rightarrow V_l$ such that, for any $\theta_k = (v_k, \sigma_k) \in \Theta_k$, $k = A, B$,*

$$\mathbf{s}_k^h(v_k, \sigma_k) = \begin{cases} [t_k^h(v_k), \bar{v}_l] \times \Sigma_l & \text{if } v_k \in [\omega_k^h, \bar{v}_k] \\ \emptyset & \text{otherwise.} \end{cases} \quad (9)$$

The case where salience is a deterministic monotone function of the valuation is clearly a special case of affiliation. It is then immediate that the above claim implies the result in [Proposition 1](#).

To establish the claim, we start by observing that if $\varphi_k^h(\underline{v}_k) \geq 0$ for $k = A, B$, then it is immediate from (4) that h -optimality requires that each agent from each side be matched to all agents from the other side, in which case $\mathbf{s}_k^h(\theta_k) = \Theta_l$ for all $\theta_k \in \Theta_k$. This is obviously a threshold rule.

Thus consider the situation where $\varphi_k^h(\underline{v}_k) < 0$ for some $k \in \{A, B\}$. Define $\Theta_k^{h+} \equiv \{\theta_k = (v_k, \sigma_k) : \varphi_k^h(v_k) \geq 0\}$ as the set of types θ_k whose φ_k^h value is nonnegative, and define $\Theta_k^{h-} \equiv \{\theta_k = (v_k, \sigma_k) : \varphi_k^h(v_k) < 0\}$ as the set of types with strictly negative φ_k^h values.

²¹See [Milgrom and Weber \(1982\)](#) for a formal definition of affiliation.

Let $\mathbf{s}'_k(\cdot)$ be any implementable matching rule. We will show that when [Condition TP-extended](#) holds, starting from $\mathbf{s}'_k(\cdot)$, one can construct another implementable matching rule $\hat{\mathbf{s}}_k(\cdot)$ that satisfies the threshold structure described in (9) and that weakly improves upon the original in terms of the platform's objective.

The proof proceeds as follows. First, we establish a couple of lemmas that will be used throughout the rest of the proof. We then consider separately the two sets of primitive conditions covered by [Condition TP-extended](#).

LEMMA 2. *A mechanism M is incentive compatible only if, with the exception of a countable subset of V_k , $|\mathbf{s}_k(v_k, \sigma_k)|_l = |\mathbf{s}_k(v_k, \sigma'_k)|_l$ for all $\sigma_k, \sigma'_k \in \Sigma_k$, $k = A, B$.*

PROOF. To see this, note that incentive compatibility requires that $|\mathbf{s}_k(v_k, \sigma_k)|_l \geq |\mathbf{s}_k(v'_k, \sigma'_k)|_l$ for any (v_k, σ_k) and (v'_k, σ'_k) such that $v_k \geq v'_k$. This in turn implies that $\mathbb{E}[|\mathbf{s}_k(v_k, \tilde{\sigma}_k)|_l]$ must be nondecreasing in v_k , where the expectation is with respect to $\tilde{\sigma}_k$ given v_k . Now at any point $v_k \in V_k$ at which $|\mathbf{s}_k(\sigma_k, v_k)|_l$ depends on σ_k , the expectation $\mathbb{E}[|\mathbf{s}_k(\tilde{\sigma}_k, v_k)|_l]$ is necessarily discontinuous in v_k . Because monotone functions can be discontinuous at most over a countable set of points, this means that the intensity of the matching set may vary with σ_k only over a countable subset of V_k . $\triangleleft$

The next lemma introduces a property for arbitrary random variables that will turn out to be useful to establish the results.

DEFINITION 4 (Monotone concave/convex order). Let F be a probability measure on the interval $[a, b]$ and let $z_1, z_2 : [a, b] \rightarrow \mathbb{R}$ be two random variables defined over $[a, b]$. We say that z_2 is smaller than z_1 in the monotone concave order if $\mathbb{E}[g(z_2(\tilde{\omega}))] \leq \mathbb{E}[g(z_1(\tilde{\omega}))]$ for any weakly increasing and weakly concave function $g : \mathbb{R} \rightarrow \mathbb{R}$. We say that z_2 is smaller than z_1 in the monotone convex order if $\mathbb{E}[g(z_2(\tilde{\omega}))] \leq \mathbb{E}[g(z_1(\tilde{\omega}))]$ for any weakly increasing and weakly convex function $g : \mathbb{R} \rightarrow \mathbb{R}$.

LEMMA 3. (i) *Suppose that $z_1, z_2 : [a, b] \rightarrow \mathbb{R}_+$ are nondecreasing and that z_2 is smaller than z_1 in the monotone concave order. Then for any weakly increasing and weakly concave function $g : \mathbb{R} \rightarrow \mathbb{R}$ and any weakly increasing and weakly negative function $h : [a, b] \rightarrow \mathbb{R}_-$, $\mathbb{E}[h(\tilde{\omega}) \cdot g(z_1(\tilde{\omega}))] \leq \mathbb{E}[h(\tilde{\omega}) \cdot g(z_2(\tilde{\omega}))]$.*

(ii) *Suppose that $z_1, z_2 : [a, b] \rightarrow \mathbb{R}_+$ are nondecreasing and that z_2 is smaller than z_1 in the monotone convex order. Then for any weakly increasing and weakly convex function $g : \mathbb{R} \rightarrow \mathbb{R}$ and any weakly increasing and weakly positive function $h : [a, b] \rightarrow \mathbb{R}_+$, $\mathbb{E}[h(\tilde{\omega}) \cdot g(z_1(\tilde{\omega}))] \geq \mathbb{E}[h(\tilde{\omega}) \cdot g(z_2(\tilde{\omega}))]$.*

PROOF. Consider first the case where z_2 is smaller than z_1 in the monotone concave order, g is weakly increasing and weakly concave, and h is weakly increasing and weakly negative. Let $(h^n)_{n \in \mathbb{N}}$ be the family of weakly increasing and weakly negative step functions $h^n : [a, b] \rightarrow \mathbb{R}$, where n is the number of steps. Because z_2 is smaller than z_1 in the monotone concave order, the inequality in the lemma is obviously true for any one-step negative function h^1 . Induction then implies that it is also true for any n -step negative

function h^n , any $n \in \mathbb{N}$. Because the set of weakly increasing and weakly negative step functions is dense (in the topology of uniform convergence) in the set of weakly increasing and weakly negative functions, the result follows. Similar arguments establish part (ii) in the lemma. $\triangleleft$

The rest of the proof considers separately the two sets of primitive conditions covered by [Condition TP-extended](#).

CASE 1. Consider markets in which the following primitive conditions jointly hold for $k = A, B$: (a) the functions $g_k(\cdot)$ are weakly concave; (b) the random variables $\tilde{\sigma}_k$ and $\tilde{v}_k$ are weakly positively affiliated.

Let $\mathbf{s}'_k(\cdot)$ be the original rule and for any $\theta_k \in \Theta_k^{h+}$, let $\hat{t}_k(v_k)$ be the threshold defined as follows:

(i) If $|\mathbf{s}'_k(\theta_k)|_l \geq |\Theta_l^{h+}|$, then let $\hat{t}_k(v_k)$ be such that

$$|[\hat{t}_k(v_k), \bar{v}_l] \times \Sigma_l|_l = |\mathbf{s}'_k(\theta_k)|_l.$$

(ii) If $|\mathbf{s}'_k(\theta_k)|_l \leq |\Theta_l^{h+}| = |\Theta_l|$, then $\hat{t}_k(v_k) = \underline{v}_l$.

(iii) If $0 < |\mathbf{s}'_k(\theta_k)|_l \leq |\Theta_l^{h+}| < |\Theta_l|$, then let $\hat{t}_k(v_k) = r_l^h$ (note that in this case $r_l^h \in (\underline{v}_l, \bar{v}_l)$).

Now apply the construction above to $k = A, B$ and consider the matching rule $\hat{\mathbf{s}}_k(\cdot)$ such that

$$\hat{\mathbf{s}}_k(\theta_k) = \begin{cases} [\hat{t}_k(v_k), \bar{v}_l] \times \Sigma_l & \Leftrightarrow \theta_k \in \Theta_k^{h+} \\ \{(v_l, \sigma_l) \in \Theta_l^+ : \hat{t}_l(v_l) \leq v_k\} & \Leftrightarrow \theta_k \in \Theta_k^{h-}. \end{cases}$$

By construction, $\hat{\mathbf{s}}_k(\cdot)$ is implementable. Moreover, $g_k(|\hat{\mathbf{s}}_k(\theta_k)|_l) \geq g_k(|\mathbf{s}'_k(\theta_k)|_l)$ for all $\theta_k \in \Theta_k^{h+}$, implying that for $k = A, B$,

$$\int_{\Theta_k^{h+}} \varphi_k^h(v_k) \cdot g_k(|\hat{\mathbf{s}}_k(v_k, \sigma_k)|_l) d\Lambda_k \geq \int_{\Theta_k^{h+}} \varphi_k^h(v_k) \cdot g_k(|\mathbf{s}'_k(\sigma_k, v_k)|_l) d\Lambda_k. \quad (10)$$

Below, we show that the matching rule $\hat{\mathbf{s}}_k(\cdot)$ also reduces the costs of cross-subsidization, relative to the original matching rule $\mathbf{s}'_k(\cdot)$. That is,

$$\int_{\Theta_k^{h-}} \varphi_k^h(v_k) \cdot g_k(|\mathbf{s}'_k(v_k, \sigma_k)|_l) d\Lambda_k \leq \int_{\Theta_k^{h-}} \varphi_k^h(v_k) \cdot g_k(|\hat{\mathbf{s}}_k(v_k, \sigma_k)|_l) d\Lambda_k. \quad (11)$$

We start with the following result.

LEMMA 4. Consider the two random variables $z_1, z_2 : [\underline{v}_k, r_k^h] \rightarrow \mathbb{R}_+$ given by $z_1(v_k) \equiv \mathbb{E}_{\tilde{\sigma}_k} [|\mathbf{s}'_k(v_k, \tilde{\sigma}_k)|_l | v_k]$ and $z_2(v_k) \equiv \mathbb{E}_{\tilde{\sigma}_k} [|\hat{\mathbf{s}}_k(v_k, \tilde{\sigma}_k)|_l | v_k]$, where the distribution over

$[\underline{v}_k, r_k^h]$ is given by $F_k(v_k)/F_k(r_k^h)$. Then z_2 is smaller than z_1 in the monotone concave order.

PROOF. From (i) the construction of $\hat{\mathbf{s}}_k(\cdot)$, (ii) the assumption of positive affiliation between valuations and salience, (iii) the fact that the measure $F_k(v_k)$ is absolute continuous with respect to the Lebesgue measure, and (iv) Lemma 2, we have that for all $x \in [\underline{v}_k, r_k^h]$,

$$\int_{\underline{v}_k}^x \int_{\Sigma_k} |\mathbf{s}'_k(v_k, \sigma_k)|_l d\Lambda_k \geq \int_{\underline{v}_k}^x \int_{\Sigma_k} |\hat{\mathbf{s}}_k(v_k, \sigma_k)|_l d\Lambda_k$$

or, equivalently,

$$\int_{\underline{v}_k}^x z_1(v_k) dF_k(v_k) \geq \int_{\underline{v}_k}^x z_2(v_k) dF_k(v_k).$$

The result in the lemma clearly holds if for all $v_k \in [\underline{v}_k, r_k^h]$, $z_1(v_k) \geq z_2(v_k)$. Thus consider the case where $z_1(v_k) < z_2(v_k)$ for some $v_k \in [\underline{v}_k, r_k^h]$, and denote by $[\dot{v}_k^1, \dot{v}_k^2]$, $[\dot{v}_k^3, \dot{v}_k^4]$, $[\dot{v}_k^5, \dot{v}_k^6]$, ... the collection of T (where $T \in \mathbb{N} \cup \{\infty\}$) subintervals of $[\underline{v}_k, r_k^h]$ in which $z_1(v_k) < z_2(v_k)$. Because $\int_{\underline{v}_k}^{r_k^h} z_1(v_k) dF_k(v_k) \geq \int_{\underline{v}_k}^{r_k^h} z_2(v_k) dF_k(v_k)$, it is clear that $\mathcal{T} \equiv \bigcup_{t=0}^{T-1} [\dot{v}_k^{2t+1}, \dot{v}_k^{2t+2}]$ is a proper subset of $[\underline{v}_k, r_k^h]$. Now construct $\dot{z}_2(\cdot)$ on the domain $[\underline{v}_k, r_k^h]$ so that

- (a) $\dot{z}_2(v_k) = z_1(v_k) < z_2(v_k)$ for all $v_k \in \mathcal{T}$
- (b) $z_2(v_k) \leq \dot{z}_2(v_k) = \alpha z_1(v_k) + (1 - \alpha)z_2(v_k) \leq z_1(v_k)$, where $\alpha \in [0, 1]$ for all $v_k \in [\underline{v}_k, r_k^h] \setminus \mathcal{T}$,
- (c) $\int_{[\underline{v}_k, r_k^h] \setminus \mathcal{T}} \{\dot{z}_2(v_k) - z_2(v_k)\} dF_k(v_k) = \int_{\mathcal{T}} \{z_2(v_k) - z_1(v_k)\} dF_k(v_k)$.

Because $\int_{\underline{v}_k}^{r_k^h} z_1(v_k) dF_k(v_k) \geq \int_{\underline{v}_k}^{r_k^h} z_2(v_k) dF_k(v_k)$, there always exists some $\alpha \in [0, 1]$ such that (b) and (c) hold. From the construction above, $\dot{z}_2(\cdot)$ is weakly increasing and

$$\int_{\underline{v}_k}^{r_k^h} \dot{z}_2(v_k) dF_k(v_k)/F_k(r_k^h) = \int_{\underline{v}_k}^{r_k^h} z_2(v_k) dF_k(v_k)/F_k(r_k^h). \quad (12)$$

This implies that for all weakly concave and weakly increasing functions $g: \mathbb{R} \rightarrow \mathbb{R}$,

$$\begin{aligned} \int_{\underline{v}_k}^{r_k^h} g(z_2(v_k)) dF_k(v_k)/F_k(r_k^h) &\leq \int_{\underline{v}_k}^{r_k^h} g(\dot{z}_2(v_k)) dF_k(v_k)/F_k(r_k^h) \\ &\leq \int_{\underline{v}_k}^{r_k^h} g(z_1(v_k)) dF_k(v_k)/F_k(r_k^h), \end{aligned}$$

where the first inequality follows from the weak concavity of $g(\cdot)$ along with (12), while the second inequality follows from the fact that $\dot{z}_2(v_k) \leq z_1(v_k)$ for all $v_k \in [\underline{v}_k, r_k^h]$ and $g(\cdot)$ is weakly increasing. $\triangleleft$

We are now ready to prove inequality (11). The results above imply that

$$\begin{aligned}
 \int_{\Theta_k^{h-}} \varphi_k^h(v_k) \cdot g_k(|\mathbf{s}'_k(v_k, \sigma_k)|_l) d\Lambda_k &= \int_{\underline{v}_k}^{r_k^h} \varphi_k^h(v_k) \cdot \mathbb{E}_{\tilde{\sigma}_k} [g_k(|\mathbf{s}'_k(v_k, \tilde{\sigma}_k)|_l) | v_k] dF_k(v_k) \\
 &= \int_{\underline{v}_k}^{r_k^h} \varphi_k^h(v_k) \cdot g_k(z_1(v_k)) dF_k(v_k) \\
 &= F_k(r_k^h) \cdot \mathbb{E}[\varphi_k^h(v_k) \cdot g_k(z_1(v_k)) | v_k \leq r_k^h] \\
 &\leq F_k(r_k^h) \cdot \mathbb{E}[\varphi_k^h(v_k) \cdot g_k(z_2(v_k)) | v_k \leq r_k^h] \\
 &= \int_{\underline{v}_k}^{r_k^h} \varphi_k^h(v_k) \cdot g_k(\mathbb{E}_{\tilde{\sigma}_k} [|\hat{\mathbf{s}}_k(v_k, \tilde{\sigma}_k)|_l | v_k]) dF_k(v_k) \\
 &= \int_{\Theta_k^{h-}} \varphi_k^h(v_k) \cdot g_k(|\hat{\mathbf{s}}_k(v_k, \sigma_k)|_l) d\Lambda_k.
 \end{aligned}$$

The first equality follows from changing the order of integration. The second equality follows from the fact that, since $\mathbf{s}'_k(\cdot)$ is implementable, $g_k(|\mathbf{s}'_k(v_k, \sigma_k)|_l)$ is invariant in σ_k except over a countable subset of $[\underline{v}_k, r_k^h]$, as shown in Lemma 2. The first inequality follows from part (i) of Lemma 3. The equality in the fifth line follows again from the fact that, by construction, $\hat{\mathbf{s}}_k(\cdot)$ is implementable, and hence invariant in σ_k except over a countable subset of $[\underline{v}_k, r_k^h]$. The series of equalities and inequalities above establishes (11), as we wanted to show.

Combining (10) with (11) establishes the result that the threshold rule $\hat{\mathbf{s}}_k(\cdot)$ improves upon the original rule $\mathbf{s}'_k(\cdot)$ in terms of the platform's objective, thus proving the result in Claim 1 for the case of markets that satisfy conditions (i.a) and (i.b) in Condition TP-extended.

Next, consider markets satisfying conditions (ii.a) and (ii.b) in Condition TP-extended.

CASE 2. Consider markets in which the following primitive conditions jointly hold for $k = A, B$: (a) the functions $g_k(\cdot)$ are weakly convex; (b) the random variables $\tilde{\sigma}_k$ and $\tilde{v}_k$ are weakly negatively affiliated.

Again, let $\mathbf{s}'_k(\cdot)$ be any (implementable) rule and for any $\theta_k \in \Theta_k^{h-}$, let $\hat{t}_k(v_k)$ be the threshold defined as follows:

- (i) If $|\Theta_l|_l > |\mathbf{s}'_k(\theta_k)|_l \geq |\Theta_l^{h+}|_l > 0$, then let $\hat{t}_k(v_k) = r_l^h$ (note that in this case $r_l^h \in (\underline{v}_l, \bar{v}_l)$).
- (ii) If $|\mathbf{s}'_k(\theta_k)|_l \geq |\Theta_l^{h+}|_l = 0$, then let $\hat{t}_k(v_k) = \bar{v}_l$.
- (iii) If $|\mathbf{s}'_k(\theta_k)|_l = |\Theta_l^{h+}|_l = |\Theta_l|_l$, then $\hat{t}_k(v_k) = \underline{v}_l$.
- (iv) If $0 \leq |\mathbf{s}'_k(\theta_k)|_l < |\Theta_l^{h+}|_l$, then let $\hat{t}_k(v_k)$ be such that

$$|[\hat{t}_k(v_k), \bar{v}_l] \times \Sigma_l|_l = |\mathbf{s}'_k(\theta_k)|_l.$$

Now apply the construction above to $k = A, B$ and consider the matching rule $\hat{\mathbf{s}}_k(\cdot)$ such that

$$\hat{\mathbf{s}}_k(\theta_k) = \begin{cases} \Theta_l^{h+} \cup \{(v_l, \sigma_l) \in \Theta_l^{h-} : \hat{v}_l(v_l) \leq v_k\} & \Leftrightarrow \theta_k \in \Theta_k^{h+} \\ [\hat{v}_k(v_k), \bar{v}_l] \times \Sigma_l & \Leftrightarrow \theta_k \in \Theta_k^{h-}. \end{cases}$$

By construction, $\hat{\mathbf{s}}_k(\cdot)_k$ is monotone and invariant in σ_k and hence implementable. Moreover, we have that $|\hat{\mathbf{s}}_k(\theta_k)|_l \leq |\mathbf{s}'_k(\theta_k)|_l$ for all $\theta_k \in \Theta_k^{h-}$. This implies that, for $k = A, B$,

$$\int_{\Theta_k^{h-}} \varphi_k^h(v_k) \cdot |\hat{\mathbf{s}}_k(v_k, \sigma_k)|_l d\Lambda_k \geq \int_{\Theta_k^{h-}} \varphi_k^h(v_k) \cdot |\mathbf{s}'_k(v_k, \sigma_k)|_l d\Lambda_k. \quad (13)$$

The arguments below show that the new matching rule $\hat{\mathbf{s}}_k(\cdot)$, relative to $\mathbf{s}'_k(\cdot)$, also increases the surplus from the positive $\varphi_k^h(v_k)$ agents, $k = A, B$ (recall that, by assumption, there exists at least one side $k \in \{A, B\}$ for which $\varphi_k^h(v_k) > 0$ for v_k high enough, $h = P, W$). That is, for any side $k \in \{A, B\}$ for which $\Theta_k^{h+} \neq \emptyset$,

$$\int_{\Theta_k^{h+}} \varphi_k^h(v_k) \cdot |\hat{\mathbf{s}}_k(v_k, \sigma_k)|_l d\Lambda_k \geq \int_{\Theta_k^{h+}} \varphi_k^h(v_k) \cdot |\mathbf{s}'_k(v_k, \sigma_k)|_l d\Lambda_k. \quad (14)$$

We start with the following result.

LEMMA 5. *Consider the two random variables $z_1, z_2 : [r_k^h, \bar{v}_k] \rightarrow \mathbb{R}_+$ given by $z_1(v_k) \equiv \mathbb{E}_{\tilde{\sigma}_k} [|\hat{\mathbf{s}}_k(v_k, \tilde{\sigma}_k)|_l | v_k]$ and $z_2(v_k) \equiv \mathbb{E}_{\tilde{\sigma}_k} [|\mathbf{s}'_k(v_k, \tilde{\sigma}_k)|_l | v_k]$, where the distribution over $[r_k^h, \bar{v}_k]$ is given by $(F_k^v(v_k) - F_k^v(r_k^h)) / (1 - F_k^v(r_k^h))$. Then z_2 is smaller than z_1 in the monotone convex order.*

PROOF. From (i) the construction of $\hat{\mathbf{s}}_k(\cdot)$, (ii) the assumption of negative affiliation between valuations and salience, (iii) the fact that the measure $F_k(v_k)$ is absolute continuous with respect to the Lebesgue measure, and (iv) Lemma 2, we have that for all $x \in [r_k^h, \bar{v}_k]$,

$$\int_x^{\bar{v}_k} \int_{\Sigma_k} |\hat{\mathbf{s}}_k(v_k, \sigma_k)|_l d\Lambda_k \geq \int_x^{\bar{v}_k} \int_{\Sigma_k} |\mathbf{s}'_k(v_k, \sigma_k)|_l d\Lambda_k,$$

or, equivalently,

$$\int_x^{\bar{v}_k} z_1(v_k) dF_k(v_k) \geq \int_x^{\bar{v}_k} z_2(v_k) dF_k(v_k).$$

The result in the lemma clearly holds if for all $v_k \in [r_k^h, \bar{v}_k]$, $z_1(v_k) \geq z_2(v_k)$. Thus consider the case where $z_1(v_k) < z_2(v_k)$ for some $v_k \in [r_k^h, \bar{v}_k]$ and denote by $[\check{v}_k^1, \check{v}_k^2]$, $[\check{v}_k^3, \check{v}_k^4]$, $[\check{v}_k^5, \check{v}_k^6]$, ... the collection of T (where $T \in \mathbb{N} \cup \{\infty\}$) subintervals of $[r_k^h, \bar{v}_k]$ in which $z_1(v_k) < z_2(v_k)$. Because $\int_{r_k^h}^{\bar{v}_k} z_1(v_k) dF_k(v_k) \geq \int_{r_k^h}^{\bar{v}_k} z_2(v_k) dF_k(v_k)$, it is clear that $\mathcal{T} \equiv \bigcup_{t=0}^{T-1} [\check{v}_k^{2t+1}, \check{v}_k^{2t+2}]$ is a proper subset of $[r_k^h, \bar{v}_k]$. Now construct $\check{z}_2(\cdot)$ on $[r_k^h, \bar{v}_k]$ so that

$$(a) \quad \check{z}_2(v_k) = \alpha z_1(v_k) + (1 - \alpha) z_2(v_k) < z_1(v_k) \text{ for all } v_k \in [r_k^h, \bar{v}_k] \setminus \mathcal{T},$$

(b) $\dot{z}_2(v_k) = z_2(v_k)$ for all $v_k \in \mathcal{T}$,

(c) $\int_{[r_k^h, \bar{v}_k] \setminus \mathcal{T}} \{\dot{z}_2(v_k) - z_2(v_k)\} dF_k(v_k) = \int_{\mathcal{T}} \{z_2(v_k) - z_1(v_k)\} dF_k(v_k)$.

Because $\int_{r_k^h}^{\bar{v}_k} z_1(v_k) dF_k(v_k) \geq \int_{r_k^h}^{\bar{v}_k} z_2(v_k) dF_k(v_k)$, there always exists some $\alpha \in [0, 1]$ such that (b) and (c) hold. From the construction above, $\dot{z}_2(\cdot)$ is weakly increasing and

$$\int_{r_k^h}^{\bar{v}_k} \dot{z}_2(v_k) dF_k(v_k) = \int_{r_k^h}^{\bar{v}_k} z_1(v_k) dF_k(v_k).$$

This implies that for all weakly increasing and weakly convex functions $g : \mathbb{R} \rightarrow \mathbb{R}$,

$$\int_{v_k}^{\bar{v}_k} g(z_2(v_k)) dF_k(v_k) \leq \int_{v_k}^{\bar{v}_k} g(\dot{z}_2(v_k)) dF_k(v_k) \leq \int_{v_k}^{\bar{v}_k} g(z_1(v_k)) dF_k(v_k),$$

where the first inequality follows the fact that $z_2(v_k) \leq \dot{z}_2(v_k)$ for all $v_k \in [r_k^h, \bar{v}_k]$ and $g(\cdot)$ is weakly increasing, while the second inequality follows from the construction of $\dot{z}_2(v_k)$ and the weak convexity of $g(\cdot)$. $\triangleleft$

We are now ready to prove inequality (14). The results above imply that

$$\begin{aligned} \int_{\Theta_k^{h+}} \varphi_k^h(v_k) \cdot g_k(|\mathbf{s}'_k(v_k, \sigma_k)|_l) d\Lambda_k &= \int_{r_k^h}^{\bar{v}_k} \varphi_k^h(v_k) \cdot \mathbb{E}_{\tilde{\sigma}_k} [g_k(|\mathbf{s}'_k(v_k, \tilde{\sigma}_k)|_l) | v_k] dF_k(v_k) \\ &= \int_{r_k^h}^{\bar{v}_k} \varphi_k^h(v_k) \cdot g_k(z_2(v_k)) dF_k(v_k) \\ &= (1 - F_k(r_k^h)) \cdot \mathbb{E}[\varphi_k^h(\tilde{v}_k) \cdot g_k(z_2(\tilde{v}_k)) | v_k \geq r_k^h] \\ &\leq (1 - F_k(r_k^h)) \cdot \mathbb{E}[\varphi_k^h(v_k) \cdot g_k(z_1(v_k)) | v_k \geq r_k^h] \\ &= \int_{r_k^h}^{\bar{v}_k} \varphi_k^h(v_k) \cdot g_k(z_1(v_k)) dF_k(v_k) \\ &= \int_{r_k^h}^{\bar{v}_k} \varphi_k^h(v_k) \cdot g_k(\mathbb{E}_{\tilde{\sigma}_k} [|\hat{\mathbf{s}}_k(v_k, \tilde{\sigma}_k)|_l | v_k]) dF_k(v_k) \\ &= \int_{\Theta_k^{h+}} \varphi_k^h(v_k) \cdot g_k(|\hat{\mathbf{s}}_k(\sigma_k, v_k)|_l) d\Lambda_k. \end{aligned}$$

The first equality follows from changing the order of integration. The second equality follows from the fact that, since $\mathbf{s}'_k(\cdot)$ is implementable, $g_k(|\mathbf{s}'_k(v_k, \sigma_k)|_l)$ is invariant in σ_k except over a countable subset of $[r_k^h, \bar{v}_k]$, as shown in Lemma 2. The first inequality follows from part (ii) of Lemma 3. The equality in the last line follows again from the fact that, by construction, $\hat{\mathbf{s}}_k(\cdot)$ is implementable, and hence invariant over σ_k , except over a countable subset of $[r_k^h, \bar{v}_k]$. The series of equalities and inequalities above establishes (14), as we wanted to show.

Combining (13) with (14) establishes that the threshold rule $\hat{s}_k(\cdot)$ improves upon the original rule $\mathbf{s}'_k(\cdot)$ in terms of the platform's objective, thus proving the result in Claim 1 under the conditions in part (ii) of Condition TP-extended. $\square$

PROOF OF PROPOSITION 2. We start with the following lemma, which establishes the first part of the proposition.

LEMMA 6. *Assume Conditions TP and MR hold. For $h = W, P$, the h -optimal matching rule is such that $t_k^h(v_k) = \underline{v}_l$ for all $v_k \in V_k$ if $\Delta_k^h(\underline{v}_k, \underline{v}_l) \geq 0$ and entails separation otherwise.*

PROOF. The proof considers separately the following three different cases.

- First, consider the case where $\varphi_k^h(\underline{v}_k) \geq 0$ for $k = A, B$, implying that $\Delta_k^h(\underline{v}_k, \underline{v}_l) \geq 0$. Because valuations (virtual valuations) are all nonnegative, welfare (profits) is (are) maximized by matching each agent from each side to all agents from the other side, meaning that the optimal matching rule employs a single complete network.
- Next consider the case where $\varphi_k^h(\underline{v}_k) < 0$ for $k = A, B$, so that $\Delta_k^h(\underline{v}_k, \underline{v}_l) < 0$. We then show that, starting from any nonseparating rule, the platform can strictly increase its payoff by switching to a separating rule. To this purpose, let $\hat{\omega}_k^h$ denote the threshold type corresponding to the nonseparating rule so that agents from side k are excluded if $v_k < \hat{\omega}_k^h$ and are otherwise matched to all agents from side l whose valuation is above $\hat{\omega}_l^h$ otherwise.

First suppose that, for some $k \in \{A, B\}$, $\hat{\omega}_k^h > r_k^h$, where we recall that $r_k^h \equiv \inf\{v_k \in V_k : \varphi_k^h(v_k) \geq 0\}$. The platform could then increase its payoff by switching to a separating rule that assigns to each agent from side k with valuation $v_k \geq \hat{\omega}_k^h$ the same matching set as the original matching rule while it assigns to each agent with valuation $v_k \in [r_k^h, \hat{\omega}_k^h]$ the matching set $[\hat{v}_l^\#, \bar{v}_l]$, where $\hat{v}_l^\# \equiv \max\{r_l^h, \hat{\omega}_l^h\}$.

Next, suppose that $\hat{\omega}_k^h < r_k^h$ for both $k = A, B$. Starting from this nonseparating rule, the platform could then increase its payoff by switching to a separating rule $\mathbf{s}_k^\diamond(\cdot)$ such that, for some $k \in \{A, B\}$,²²

$$\mathbf{s}_k^\diamond(v_k) = \begin{cases} [\hat{\omega}_l^h, \bar{v}_l] & \Leftrightarrow v_k \in [r_k^h, \bar{v}_k] \\ [r_l^h, \bar{v}_l] & \Leftrightarrow v_k \in [\hat{\omega}_k^h, r_k^h] \\ \emptyset & \Leftrightarrow v_k \in [\underline{v}_k, \hat{\omega}_k^h]. \end{cases}$$

The new matching rule improves on the original one because it eliminates all matches between agents whose valuations (virtual valuations) are both negative.

Finally, suppose that $\hat{\omega}_k^h = r_k^h$ for some $k \in \{A, B\}$, whereas $\hat{\omega}_l^h \leq r_l^h$ for $l \neq k$. The platform could then do better by switching to the separating rule

$$\mathbf{s}_k^\#(v_k) = \begin{cases} [\hat{\omega}_l^h, \bar{v}_l] & \Leftrightarrow v_k \in [r_k^h, \bar{v}_k] \\ [r_l^h, \bar{v}_l] & \Leftrightarrow v_k \in [\hat{\omega}_k^h, r_k^h] \\ \emptyset & \Leftrightarrow v_k \in [\underline{v}_k, \hat{\omega}_k^h]. \end{cases}$$

By setting the new exclusion threshold $\hat{\omega}_k^\#$ sufficiently close to (but strictly below) r_k^h , the platform increases its payoff. In fact, the marginal benefit of increasing

²²The behavior of the rule on side l is then pinned down by reciprocity.

the quality of the matching sets of those agents from side l whose φ_l^h value is positive more than offsets the marginal cost of getting on board a few more agents from side k whose φ_k^h value is negative, but sufficiently small.²³ Note that for this network expansion to be profitable, it is essential that the new agents from side k that are brought “on board” be matched only to those agents from side l whose φ_l^h value is positive, which requires employing a separating rule.

- Finally, suppose that $\varphi_l^h(\underline{v}_l) < 0 \leq \varphi_k^h(\underline{v}_k)$. First suppose that $\Delta_k^h(\underline{v}_k, \underline{v}_l) \geq 0$ and that the matching rule is different from a single complete network (i.e., $t_k^h(v_k) > \underline{v}_l$ for some $v_k \in V_k$). Take an arbitrary point $v_k \in [\underline{v}_k, \bar{v}_k]$ at which the function $t_k^h(\cdot)$ is strictly decreasing in a right neighborhood of v_k . Consider the effect of a marginal reduction in the threshold $t_k^h(v_k)$ around the point $v_l = t_k^h(v_k)$. This is given by $\Delta_k^h(v_k, v_l)$. Next note that, given any interval $[v'_k, v''_k]$ over which the function $t_k^h(\cdot)$ is constant and equal to v_l , the marginal effect of decreasing the threshold below v_l for any type $v_k \in [v'_k, v''_k]$ is given by $\int_{v'_k}^{v''_k} [\Delta_k^h(v_k, v_l)] dv_k$. Last, note that $\text{sign}\{\Delta_k^h(v_k, v_l)\} = \text{sign}\{\psi_k^h(v_k) + \psi_l^h(v_l)\}$. Under [Condition MR](#), this means that $\Delta_k^h(v_k, v_l) > 0$ for all (v_k, v_l) . The results above then imply that the platform can increase its objective by decreasing the threshold for any type for which $t_k^h(v_k) > \underline{v}_l$, proving that a single complete network is optimal.

Next suppose that $\Delta_k^h(\underline{v}_k, \underline{v}_l) < 0$ and that the platform employs a nonseparating rule. First suppose that such a rule entails full participation (that is, $\hat{\omega}_l^h = \underline{v}_l$ or, equivalently, $t_k^h(\underline{v}_k) = \underline{v}_l$). The fact that $\Delta_k^h(\underline{v}_k, \underline{v}_l) < 0$ implies that the marginal effect of raising the threshold $t_k^h(\underline{v}_k)$ for the lowest type on side k , while leaving the threshold untouched for all other types, is positive. By continuity of the marginal effects, the platform can then improve its objective by switching to a separating rule that is obtained by increasing $t_k^h(\cdot)$ in a right neighborhood of $\underline{v}_k$ while leaving $t_k^h(\cdot)$ untouched elsewhere.

Next consider the case where the original rule excludes some agents (but assigns the same matching set to each agent whose valuation is above $\hat{\omega}_k^h$). From the same arguments as above, for such a rule to be optimal, it must be that $\hat{\omega}_l^h < r_l^h$ and $\hat{\omega}_k^h = \underline{v}_k$, with $\hat{\omega}_l^h$ satisfying the first-order condition

$$\hat{g}_l(\underline{v}_k) \cdot \varphi_l^h(\hat{\omega}_l^h) - \hat{g}_k'(\hat{\omega}_l^h) \cdot \int_{\underline{v}_k}^{\bar{v}_k} \varphi_k^h(v_k) dF_k^v(v_k) = 0.$$

This condition requires that the total effect of a marginal increase of the size of the network on side l (obtained by reducing the threshold $t_k^h(v_k)$ below $\hat{\omega}_l^h$ for all types v_k) be zero. This rewrites as $\int_{\underline{v}_k}^{\bar{v}_k} [\Delta_k^h(v_k, \hat{\omega}_l^h)] dv_k = 0$. Because $\text{sign}\{\Delta_k^h(v_k, \hat{\omega}_l^h)\} = \text{sign}\{\psi_k^h(v_k) + \psi_l^h(\hat{\omega}_l^h)\}$, under [Condition MR](#) this means that there exists a $v_k^\# \in (\underline{v}_k, \bar{v}_k)$ such that $\int_{v_k^\#}^{\bar{v}_k} \Delta_k^h(v_k, \hat{\omega}_l^h) dv_k > 0$. This means that there

²³To see this, note that, starting from $\hat{\omega}_k^\# = r_k^h$, the marginal benefit of decreasing the threshold $\hat{\omega}_k^\#$ is $-\hat{g}_l'(r_k^h) \int_{r_k^h}^{\bar{v}_l} \varphi_l^h(v_l) dF_l^v(v_l) > 0$, whereas the marginal cost is given by $-\hat{g}_k(r_k^h) \cdot \varphi_k^h(r_k^h) f_k^v(r_k^h) = 0$ since $\varphi_k^h(r_k^h) = 0$.

exists a $\omega_l^\# < \hat{\omega}_l^h$ such that the platform could increase its payoff by switching to the separating rule

$$s_k^h(v_k) = \begin{cases} [\omega_l^\#, \bar{v}_l] \Leftrightarrow v_k \in [v_k^\#, \bar{v}_k] \\ [\hat{\omega}_l^h, \bar{v}_l] \Leftrightarrow v_k \in [\underline{v}_k, v_k^\#]. \end{cases}$$

We conclude that a separating rule is optimal when $\Delta_k^h(\underline{v}_k, \underline{v}_l) < 0$. $\triangleleft$

The rest of the proof shows that when, in addition to Conditions TP and MR, $\Delta_k^h(\underline{v}_k, \underline{v}_l) < 0$, then the optimal separating rule satisfies properties (i) and (ii) in the proposition.

To see this, note that the h -optimal matching rule solves the program (which we call the *full program* (P^F))

$$P^F : \quad \max_{\{\omega_k, t_k(\cdot)\}_{k=A,B}} \sum_{k=A,B} \int_{\omega_k}^{\bar{v}_k} \hat{g}_k(t_k(v_k)) \cdot \varphi_k^h(v_k) \cdot dF_k^v(v_k)$$

subject to the following constraints for $k, l \in \{A, B\}$, $l \neq k$,

$$t_k(v_k) = \inf\{v_l : t_l(v_l) \leq v_k\} \quad (15)$$

$$t_k(\cdot) \text{ weakly decreasing} \quad (16)$$

$$\text{and } t_k(\cdot) : [\omega_k, \bar{v}_k] \rightarrow [\omega_l, \bar{v}_l] \quad (17)$$

with $\omega_k \in [\underline{v}_k, \bar{v}_k]$ and $\omega_l \in [\underline{v}_l, \bar{v}_l]$. Constraint (15) is the reciprocity condition, rewritten using the result in Proposition 1. Constraint (16) is the monotonicity constraint required by incentive compatibility. Finally, constraint (17) is a domain–codomain restriction that requires the function $t_k(\cdot)$ to map each type on side k that is included in the network into the set of types on side l that is also included in the network.

Because $\Delta_k^h(\underline{v}_k, \underline{v}_l) < 0$, it must be that $r_k^h > \underline{v}_k$ for some $k \in \{A, B\}$. Furthermore, from the arguments in the proof of Lemma 6 above, at the optimum, $\omega_k^h \in [\underline{v}_k, r_k^h]$. In addition, whenever $r_l^h > \underline{v}_l$, then $\omega_l^h \in [\underline{v}_l, r_l^h]$ and $t_k^h(r_k^h) = r_l^h$. Hereafter, we will assume that $r_l^h > \underline{v}_l$. When this is not the case, then $\omega_l^h = \underline{v}_l$ and $t_k^h(v_k) = \underline{v}_l$ for all $v_k \geq r_k^h$, while the optimal ω_k^h and $t_k^h(v_k)$ for $v_k < r_k^h$ are obtained from the solution to program P_k^F below by replacing r_l^h with $\underline{v}_l$.

Thus assume $\varphi_k^h(\underline{v}_k) < 0$ for $k = A, B$. Program P^F can then be decomposed into the two independent programs P_k^F , $k = A, B$:

$$P_k^F : \quad \max_{\omega_k, t_k(\cdot), t_l(\cdot)} \int_{\omega_k}^{r_k^h} \hat{g}_k(t_k(v_k)) \cdot \varphi_k^h(v_k) \cdot dF_k^v(v_k) + \int_{r_l^h}^{\bar{v}_l} \hat{g}_l(t_l(v_l)) \cdot \varphi_l^h(v_l) \cdot dF_l^v(v_l) \quad (18)$$

subject to $t_k(\cdot)$ and $t_l(\cdot)$ satisfying the reciprocity and monotonicity constraints (15) and (16), along with the constraints

$$t_k(\cdot) : [\omega_k, r_k^h] \rightarrow [r_l^h, \bar{v}_l], \quad t_l(\cdot) : [r_l^h, \bar{v}_l] \rightarrow [\omega_k, r_k^h]. \quad (19)$$

Program P_k^F is not a standard calculus of variations problem. As an intermediate step, we will thus consider the *auxiliary program* (P_k^{Au}), which strengthens constraint (16) and fixes $\omega_k = \underline{v}_k$ and $\omega_l = \underline{v}_l$:

$$P_k^{Au} : \quad \max_{t_k(\cdot), t_l(\cdot)} \int_{\underline{v}_k}^{r_k^h} \hat{g}_k(t_k(v_k)) \cdot \varphi_k^h(v_k) \cdot dF_k^v(v_k) + \int_{r_l^h}^{\bar{v}_l} \hat{g}_l(t_l(v_l)) \cdot \varphi_l^h(v_l) \cdot dF_l^v(v_l) \quad (20)$$

subject to (15),

$$t_k(\cdot), t_l(\cdot) \quad \text{strictly decreasing} \quad (21)$$

$$\text{and } t_k(\cdot) : [\underline{v}_k, r_k^h] \rightarrow [r_l^h, \bar{v}_l], \quad t_l(\cdot) : [r_l^h, \bar{v}_l] \rightarrow [\underline{v}_k, r_k^h] \text{ are bijections.} \quad (22)$$

By virtue of (21), (15) can be rewritten as $t_k(v_k) = t_l^{-1}(v_k)$. Plugging this into the objective function (20) yields

$$\int_{\underline{v}_k}^{r_k^h} \hat{g}_k(t_k(v_k)) \cdot \varphi_k^h(v_k) \cdot f_k^v(v_k) dv_k + \int_{r_l^h}^{\bar{v}_l} \hat{g}_l(t_k^{-1}(v_l)) \cdot \varphi_l^h(v_l) \cdot f_l^v(v_l) dv_l. \quad (23)$$

Changing the variable of integration in the second integral in (23) to $\tilde{v}_l \equiv t_k^{-1}(v_l)$, using the fact that $t_k(\cdot)$ is strictly decreasing and hence differentiable almost everywhere, and using the fact that $t_k^{-1}(r_l^h) = r_k^h$ and $t_k^{-1}(\bar{v}_l) = \underline{v}_k$, the auxiliary program can be rewritten as

$$P_k^{Au} : \quad \max_{t_k(\cdot)} \int_{\underline{v}_k}^{r_k^h} \{ \hat{g}_k(t_k(v_k)) \cdot \varphi_k^h(v_k) \cdot f_k^v(v_k) - \hat{g}_l(v_k) \cdot \varphi_l^h(t_k(v_k)) \cdot f_l^v(t_k(v_k)) \cdot t_k'(v_k) \} dv_k \quad (24)$$

subject to $t_k(\cdot)$ being continuous, strictly decreasing, and satisfying the boundary conditions

$$t_k(\underline{v}_k) = \bar{v}_l \quad \text{and} \quad t_k(r_k^h) = r_l^h. \quad (25)$$

Consider now the *relaxed auxiliary program* (P_k^R) that is obtained from P_k^{Au} by dispensing with the condition that $t_k(\cdot)$ be continuous and strictly decreasing, and instead allowing for any measurable control $t_k(\cdot) : [\underline{v}_k, r_k^h] \rightarrow [r_l^h, \bar{v}_l]$ with bounded subdifferential that satisfies the boundary condition (25).

LEMMA 7. *The program P_k^R admits a piecewise absolutely continuous maximizer $\tilde{t}_k(\cdot)$.*

PROOF. Program P_k^R is equivalent to the optimal control problem

$$\mathcal{P}_k^R : \quad \max_{y(\cdot)} \int_{\underline{v}_k}^{r_k^h} \{ \hat{g}_k(x(v_k)) \cdot \varphi_k^h(v_k) \cdot f_k^v(v_k) - \hat{g}_l(v_k) \cdot \varphi_l^h(x(v_k)) \cdot f_l^v(x(v_k)) \cdot y(v_k) \} dv_k$$

subject to

$$\begin{aligned} x'(v_k) &= y(v_k) \quad \text{a.e.,} & x(\underline{v}_k) &= \bar{v}_l, & x(r_k^h) &= r_l^h \\ y(v_k) &\in [-K, +K] \quad \text{and} & x(v_k) &\in [r_l^h, \bar{v}_l], \end{aligned}$$

where K is a large number. Program $\mathcal{P}_k^R$ satisfies all the conditions of the Filipov–Cesari theorem (see Cesari 1983). By that theorem, we know that there exists a measurable function $y(\cdot)$ that solves $\mathcal{P}_k^R$. By the equivalence of P_k^R and $\mathcal{P}_k^R$, it then follows that P_k^R admits a piecewise absolutely continuous maximizer $\tilde{t}_k(\cdot)$. $\triangleleft$

LEMMA 8. Consider the function $\eta(\cdot)$ implicitly defined by

$$\Delta_k^h(v_k, \eta(v_k)) = 0. \quad (26)$$

Let $\tilde{v}_k \equiv \inf\{v_k \in [\underline{v}_k, r_k^h] : (26) \text{ admits a solution}\}$. The solution to P_k^R is given by

$$\tilde{t}_k(v_k) = \begin{cases} \bar{v}_l & \text{if } v_k \in [\underline{v}_k, \tilde{v}_k] \\ \eta(v_k) & \text{if } v_k \in (\tilde{v}_k, r_k^h]. \end{cases} \quad (27)$$

PROOF. From Lemma 7, we know that P_k^R admits a piecewise absolutely continuous solution. Standard results from calculus of variations then imply that such a solution $\tilde{t}_k(\cdot)$ must satisfy the Euler equation at any interval $I \subset [\underline{v}_k, r_k^h]$ where its image $\tilde{t}_k(v_k) \in (r_l^h, \bar{v}_l)$. The Euler equation associated with program P_k^R is given by (26). Condition MR ensures that (i) there exists a $\tilde{v}_k \in [\underline{v}_k, r_k^h]$ such that (26) admits a solution if and only if $v_k \in [\tilde{v}_k, r_k^h]$, (ii) that at any point $v_k \in [\tilde{v}_k, r_k^h]$ such solution is unique and given by $\eta(v_k) = (\psi_l^h)^{-1}(-\psi_k^h(v_k))$, and (iii) that $\eta(\cdot)$ is continuous and strictly decreasing over $[\tilde{v}_k, r_k^h]$.

When $\tilde{v}_k > \underline{v}_k$, (26) admits no solution at any point $v_k \in [\underline{v}_k, \tilde{v}_k]$, in which case $\tilde{t}_k(v_k) \in \{r_l^h, \bar{v}_l\}$. Because $\varphi_k^h(v_k) < 0$ for all $v_k \in [\underline{v}_k, \tilde{v}_k]$ and because $\hat{g}_k(\cdot)$ is decreasing, it is then immediate from inspecting the objective (24) that $\tilde{t}_k(v_k) = \bar{v}_l$ for all $v_k \in [\underline{v}_k, \tilde{v}_k]$.

It remains to show that $\tilde{t}_k(v_k) = \eta(v_k)$ for all $v_k \in [\tilde{v}_k, r_k^h]$. Because the objective function in P_k^R is not concave in (t_k, t_k') for all v_k , we cannot appeal to standard sufficiency arguments. Instead, using the fact that the Euler equation is a necessary optimality condition for interior points, we will prove that $\tilde{t}_k(v_k) = \eta(v_k)$ by arguing that there is no function $\hat{t}_k(\cdot)$ that improves on $\tilde{t}_k(\cdot)$ and such that $\hat{t}_k(\cdot)$ coincides with $\tilde{t}_k(\cdot)$ except on an interval $(v_k^1, v_k^2) \subseteq [\tilde{v}_k, r_k^h]$ over which $\hat{t}_k^h(v_k) \in \{r_l^h, \bar{v}_l\}$.

To see that this is true, fix an arbitrary $(v_k^1, v_k^2) \subseteq [\tilde{v}_k, r_k^h]$ and consider the problem that consists in choosing optimally a step function $\hat{t}_k(\cdot) : (v_k^1, v_k^2) \rightarrow \{r_l^h, \bar{v}_l\}$. Because step functions are such that $\hat{t}_k'(v_k) = 0$ at all points of continuity and because $\varphi_k^h(v_k) < 0$ for all $v_k \in (v_k^1, v_k^2)$, it follows that the optimal step function is given by $\hat{t}_k(v_k) = \bar{v}_l$ for all $v_k \in (v_k^1, v_k^2)$. Notice that the value attained by the objective (24) over the interval (v_k^1, v_k^2) under such a step function is zero. Instead, an interior control $t_k(\cdot) : (v_k^1, v_k^2) \rightarrow (r_l^h, \bar{v}_l)$ over the same interval with derivative

$$t_k'(v_k) < \frac{\hat{g}_k(t_k(v_k)) \cdot \varphi_k^h(v_k) \cdot f_k^v(v_k)}{\hat{g}_l(v_k) \cdot \varphi_l^h(t_k(v_k)) \cdot f_l^v(t_k(v_k))}$$

for all $v_k \in (v_k^1, v_k^2)$ yields a strictly positive value. This proves that the solution to P_k^R must indeed satisfy the Euler equation (26) for all $v_k \in [\tilde{v}_k, r_k^h]$. Together with the property established above that $\tilde{t}_k(v_k) = \bar{v}_l$ for all $v_k \in [\underline{v}_k, \tilde{v}_k]$, this establishes that the

unique piecewise absolutely continuous function that solves P_k^R is the control $\tilde{t}_k(\cdot)$ that satisfies (27). $\triangleleft$

Denote by $\max\{P_k^R\}$ the value of program P_k^R (i.e., the value of the objective (24) evaluated under the control $\tilde{t}_k^h(\cdot)$ defined in Lemma 8). Then denote by $\sup\{P_k^{Au}\}$ and $\sup\{P_k^F\}$ the supremum of programs P_k^{Au} and P_k^F , respectively. Note that we write $\sup$ rather than $\max$ as, a priori, a solution to these problems might not exist.

LEMMA 9. *We have $\sup\{P_k^F\} = \sup\{P_k^{Au}\} = \max\{P_k^R\}$.*

PROOF. Clearly, $\sup\{P_k^F\} \geq \sup\{P_k^{Au}\}$ for P_k^{Au} is more constrained than P_k^F . Next note that $\sup\{P_k^F\} = \sup\{\hat{P}_k^F\}$, where $\hat{P}_k^F$ coincides with P_k^F except that ω_k is constrained to be equal to $\underline{v}_k$ and $t_k(\underline{v}_k)$ is constrained to be equal to $\bar{v}_l$. This follows from the fact that excluding types below a threshold ω'_k gives the same value as setting $t_k(v_k) = \bar{v}_l$ for all $v_k \in [\underline{v}_k, \omega'_k)$. That $\sup\{\hat{P}_k^F\} = \sup\{P_k^{Au}\}$ then follows from the fact that any pair of measurable functions $t_k(\cdot)$, $t_l(\cdot)$ satisfying conditions (15), (16), and (19) with $\omega_k = \underline{v}_k$ and $t_k(\underline{v}_k) = \bar{v}_l$ can be approximated arbitrarily well in the L^2 norm by a pair of functions satisfying conditions (15), (21), and (22). That $\max\{P_k^R\} \geq \sup\{P_k^{Au}\}$ follows from the fact that P_k^R is a relaxed version of P_k^{Au} . That $\max\{P_k^R\} = \sup\{P_k^{Au}\}$ in turn follows from the fact that the solution $\tilde{t}_k^h(\cdot)$ to P_k^R can be approximated arbitrarily well in the L^2 norm by a function $t_k(\cdot)$ that is continuous and strictly decreasing. $\triangleleft$

From the results above, we are now in a position to exhibit the solution to P_F^k . Let $\omega_k^h = \tilde{v}_k$, where $\tilde{v}_k$ is the threshold defined in Lemma 8. Next for any $v_k \in [\tilde{v}_k, r_k^h]$, let $t_k^h(v_k) = \tilde{t}_k(v_k)$, where $\tilde{t}_k(\cdot)$ is the function defined in Lemma 8. Finally, given $t_k^h(\cdot) : [\omega_k^h, r_k^h] \rightarrow [r_l^h, \bar{v}_l]$, let $t_l^h(\cdot) : [r_l^h, \bar{v}_l] \rightarrow [\omega_k^h, r_k^h]$ be the unique function that satisfies (15). It is clear that the triple $\omega_k^h, t_k^h(\cdot), t_l^h(\cdot)$ constructed in this way satisfies conditions (15), (16), and (19), and is therefore a feasible candidate for program P_k^F . It is also immediate that the value of the objective (18) in P_k^F evaluated at $\omega_k^h, t_k^h(\cdot), t_l^h(\cdot)$ is the same as $\max\{P_k^R\}$. From Lemma 9, we then conclude that $\omega_k^h, t_k^h(\cdot), t_l^h(\cdot)$ is a solution to P_k^F .

Applying the construction above to $k = A, B$ and combining the solution to program P_A^F with the solution to program P_B^F then gives the solution $\{\omega_k^h, t_k^h(\cdot)\}_{k \in \{A, B\}}$ to program P_F .

By inspection, it is easy to see that the corresponding rule is maximally separating. Furthermore, from the arguments in Lemma 8, one can easily verify that there is exclusion at the bottom on side k (and no bunching at the top on side l) if $\tilde{v}_k > \underline{v}_k$ and bunching at the top on side l (and no exclusion at the bottom on side k) if $\tilde{v}_k = \underline{v}_k$. By the definition of $\tilde{v}_k$, in the first case, there exists a $v'_k > \underline{v}_k$ such that $\Delta_k^h(v'_k, \bar{v}_l) = 0$ or, equivalently, $\psi_k^h(v'_k) + \psi_l^h(\bar{v}_l) = 0$. Condition MR along with the fact that $\text{sign}\{\Delta_k^h(v_k, v_l)\} = \text{sign}\{\psi_k^h(v_k) + \psi_l^h(v_l)\}$ then implies that $\Delta_k^h(\underline{v}_k, \bar{v}_l) = \Delta_l^h(\bar{v}_l, \underline{v}_k) < 0$. Hence, whenever $\Delta_k^h(\underline{v}_k, \bar{v}_l) = \Delta_l^h(\bar{v}_l, \underline{v}_k) < 0$, there is exclusion at the bottom on side k and no bunching at the top on side l . Symmetrically, $\Delta_l^h(\underline{v}_l, \bar{v}_k) = \Delta_k^h(\bar{v}_k, \underline{v}_l) < 0$ implies that there is exclusion at the bottom on side l and no bunching at the top on side k , as stated in the proposition.

Next, consider the case where $\bar{v}_k = \underline{v}_k$. In this case there exists a $\eta(\underline{v}_k) \in [r_l^h, \bar{v}_l]$ such that $\Delta_k^h(\underline{v}_k, \eta(\underline{v}_k)) = 0$ or, equivalently, $\psi_k^h(\underline{v}_k) + \psi_l^h(\eta(\underline{v}_k)) = 0$. Assume first that $\eta(\underline{v}_k) < \bar{v}_l$. By [Condition MR](#), it then follows that $\psi_k^h(\underline{v}_k) + \psi_l^h(\bar{v}_l) > 0$ or, equivalently, that $\Delta_k^h(\underline{v}_k, \bar{v}_l) = \Delta_l^h(\bar{v}_l, \underline{v}_k) > 0$. Hence, whenever $\Delta_k^h(\underline{v}_k, \bar{v}_l) = \Delta_l^h(\bar{v}_l, \underline{v}_k) > 0$, there is no exclusion at the bottom on side k and bunching at the top on side l . Symmetrically, $\Delta_l^h(\underline{v}_l, \bar{v}_k) = \Delta_k^h(\bar{v}_k, \underline{v}_l) > 0$ implies that there is bunching at the top on side k and no exclusion at the bottom on side l , as stated in the proposition.

Next, consider the case where $\eta(\underline{v}_k) = \bar{v}_l$. In this case $\omega_k^h = \underline{v}_k$ and $t_k^h(\underline{v}_k) = \bar{v}_l$. This is the knife-edge case where $\Delta_k^h(\underline{v}_k, \bar{v}_l) = \Delta_l^h(\bar{v}_l, \underline{v}_k) = 0$ in which there is neither bunching at the top on side l nor exclusion at the bottom on side k . $\square$

PROOF OF PROPOSITION 4. Hereafter, we use the caret ($\hat{\cdot}$) notation for all variables in the mechanism $\hat{M}^P$ corresponding to the new salience function $\hat{\sigma}_k(v_k)$ and continue to denote the variables in the mechanism M^P corresponding to the original function $\sigma_k(v_k)$ without annotation. By definition, we have that $\hat{\psi}_k^P(v_k) \geq \psi_k^P(v_k)$ for all $v_k \leq r_k^P$ while $\hat{\psi}_k^P(v_k) \leq \psi_k^P(v_k)$ for all $v_k \geq r_k^P$. Recall, from the arguments in the proof of [Proposition 2](#), that for any $v_k < \omega_k^P$, $\Delta_k^P(v_k, \bar{v}_l) < 0$ or, equivalently, $\psi_k^P(v_k) + \psi_l^P(\bar{v}_l) < 0$, whereas for any $v_k \in (\omega_k, r_k^P]$, $t_k^P(v_k)$ satisfies $\psi_k^P(v_k) + \psi_l^P(t_k^P(v_k)) = 0$. The ranking between $\hat{\psi}_k^P(\cdot)$ and $\psi_k^P(\cdot)$, along with the strict monotonicity of these functions then implies that $\hat{\omega}_k^P \leq \omega_k^P$ and, for any $v_k \in [\omega_k^P, r_k^P]$, $\hat{t}_k^P(v_k) \leq t_k^P(v_k)$. Symmetrically, because $\hat{\psi}_k^P(v_k) + \psi_l^P(v_l) < \psi_k^P(v_k) + \psi_l^P(v_l)$ for all $v_k > r_k^P$, all v_l , we have that $\hat{t}_k^P(v_k) \geq t_k^P(v_k)$ for all $v_k > r_k^P$. This completes the proof of part (i) in the proposition.

Next consider part (ii). The result in part (i) implies that $|\hat{\mathbf{s}}_k(v_k)|_l \geq |\mathbf{s}_k(v_k)|_l$ if and only if $v_k \leq r_k^P$. Using (3), note that for all types with valuation $v_k \leq r_k^P$,

$$\Pi_k(v_k; \hat{M}^P) = \int_{\underline{v}_k}^{v_k} |\hat{\mathbf{s}}_k(x)|_l dx \geq \Pi_k(v_k; M^P) = \int_{\underline{v}_k}^{v_k} |\mathbf{s}_k(x)|_l dx.$$

Furthermore, since $|\hat{\mathbf{s}}_k(v_k)|_l \leq |\mathbf{s}_k(v_k)|_l$ for all $v_k \geq r_k^P$, there exists a threshold type $\hat{v}_k > r_k^P$ (possibly equal to $\bar{v}_k$) such that $\Pi_k(v_k; \hat{M}^P) \geq \Pi_k(v_k; M^P)$ if and only if $v_k \leq \hat{v}_k$, which establishes part (ii) in the proposition. $\square$

PROOF OF COROLLARY 1. Let $y_k(v_k) \equiv |\mathbf{s}_k^P(v_k)|_l$ denote the quality of the matching set that each agent with valuation v_k obtains under the original mechanism, and let $\hat{y}_k(v_k) \equiv |\hat{\mathbf{s}}_k^P(v_k)|_l$ denote the corresponding quality under the new mechanism. Using (3), for any $q \in y_k(V_k) \cap \hat{y}_k(V_k)$, i.e., for any q offered both under M^P and $\hat{M}^P$,

$$\begin{aligned} \rho_k^P(q) &= y_k^{-1}(q)q - \int_{\underline{v}_k}^{y_k^{-1}(q)} y_k(v) dv \quad \text{and} \\ \hat{\rho}_k^P(q) &= \hat{y}_k^{-1}(q)q - \int_{\underline{v}_k}^{\hat{y}_k^{-1}(q)} \hat{y}_k(v) dv, \end{aligned}$$

where $y_k^{-1}(q) \equiv \inf\{v_k : y_k(v_k) = q\}$ is the generalized inverse of $y_k(\cdot)$ and $\hat{y}_k^{-1}(q) = \inf\{v_k : \hat{y}_k(v_k) = q\}$ is the corresponding inverse for $\hat{y}_k(\cdot)$. We thus have that

$$\rho_k^P(q) - \hat{\rho}_k^P(q) = \int_{\underline{v}_k}^{y_k^{-1}(q)} [\hat{y}_k(v) - y_k(v)] dv + \int_{y_k^{-1}(q)}^{\hat{y}_k^{-1}(q)} [\hat{y}_k(v) - q] dv.$$

From the results in [Proposition 4](#), we know that $[y_k(v_k) - \hat{y}_k(v_k)][v_k - r_k^P] \geq 0$ with $y_k(r_k^P) = \hat{y}_k(r_k^P)$. Therefore, for all $q \in y_k(V_k) \cap \hat{y}_k(V_k)$, with $q \leq y_k(r_k^P) = \hat{y}_k(r_k^P)$,

$$\begin{aligned} \rho_k^P(q) - \hat{\rho}_k^P(q) &= \int_{\underline{v}_k}^{y_k^{-1}(q)} [\hat{y}_k(v) - y_k(v)] dv - \int_{\hat{y}_k^{-1}(q)}^{y_k^{-1}(q)} [\hat{y}_k(v) - q] dv \\ &= \int_{\underline{v}_k}^{\hat{y}_k^{-1}(q)} [\hat{y}_k(v) - y_k(v)] dv + \int_{\hat{y}_k^{-1}(q)}^{y_k^{-1}(q)} [q - y_k(v)] dv \\ &\geq 0, \end{aligned}$$

whereas for $q \geq y_k(r_k^P) = \hat{y}_k(r_k^P)$,

$$\begin{aligned} \rho_k^P(q) - \hat{\rho}_k^P(q) &= \int_{\underline{v}_k}^{r_k^P} [\hat{y}_k(v) - y_k(v)] dv + \int_{r_k^P}^{y_k^{-1}(q)} [\hat{y}_k(v) - y_k(v)] dv \\ &\quad + \int_{y_k^{-1}(q)}^{\hat{y}_k^{-1}(q)} [\hat{y}_k(v) - q] dv \\ &= \rho_k^P(y_k(r_k^P)) - \hat{\rho}_k^P(y_k(r_k^P)) + \int_{r_k^P}^{y_k^{-1}(q)} [\hat{y}_k(v) - y_k(v)] dv \\ &\quad + \int_{y_k^{-1}(q)}^{\hat{y}_k^{-1}(q)} [\hat{y}_k(v) - q] dv \\ &= \rho_k^P(y_k(r_k^P)) - \hat{\rho}_k^P(y_k(r_k^P)) + \left(\int_{r_k^P}^{\hat{y}_k^{-1}(q)} \hat{y}_k(v) dv - \hat{y}_k^{-1}(q)q \right) \\ &\quad - \left(\int_{r_k^P}^{y_k^{-1}(q)} y_k(v) dv - y_k^{-1}(q)q \right). \end{aligned}$$

Integrating by parts, using the fact that $y_k(r_k^P) = \hat{y}_k(r_k^P)$, and changing variables, we have that

$$\begin{aligned} &\left(\int_{r_k^P}^{\hat{y}_k^{-1}(q)} \hat{y}_k(v) dv - \hat{y}_k^{-1}(q)q \right) - \left(\int_{r_k^P}^{y_k^{-1}(q)} y_k(v) dv - y_k^{-1}(q)q \right) \\ &= \left(r_k^P \hat{y}_k(r_k^P) - \int_{r_k^P}^{\hat{y}_k^{-1}(q)} v \frac{d\hat{y}_k(v)}{dv} dv \right) - \left(r_k^P y_k(r_k^P) - \int_{r_k^P}^{y_k^{-1}(q)} v \frac{dy_k(v)}{dv} dv \right) \\ &= - \int_{y_k(r_k^P)}^q (\hat{y}_k^{-1}(z) - y_k^{-1}(z)) dz. \end{aligned}$$

Because $\hat{y}_k^{-1}(z) \geq y_k^{-1}(z)$ for $z > y_k(r_k^P)$, we then conclude that the price differential $\rho_k^P(q) - \hat{\rho}_k^P(q)$, which is positive at $q = y_k(r_k^P) = \hat{y}_k(r_k^P)$, declines as q grows above $y_k(r_k^P)$. Going back to the original notation, it follows that there exists $\hat{q}_k > |\mathbf{s}_k^P(r_k^P)|_l = |\hat{\mathbf{s}}_k^P(r_k^P)|_l$ (possibly equal to $|\hat{\mathbf{s}}_k^P(\bar{v}_k)|_l$) such that $\hat{\rho}_k^P(q) \leq \rho_k^P(q)$ if and only if $q \leq \hat{q}_k$. This establishes the result. $\square$

APPENDIX B

This appendix complements the discussion in [Section 3.1](#) by exhibiting an example where threshold rules fail to be optimal when salience is nonincreasing and preferences are strictly concave.

EXAMPLE 7 (Sub-optimality of threshold rules III). Agents from sides A and B have their valuations drawn uniformly from $V_A = [0, 1]$ and $V_B = [-2, 0]$, respectively. The salience of side- A agents is constant and normalized to 1, i.e., $\sigma_A(v_A) \equiv 1$ for all $v_A \in V_A$, while the salience of the side- B agents is given by

$$\sigma_B(v_B) = \begin{cases} 1 & \text{if } v_B \in [-1, 0] \\ 8 & \text{if } v_B \in [-2, -1]. \end{cases}$$

Preferences for matching intensity are linear on side B (that is, g_B is the identity function), whereas preferences on side A are given by the concave function²⁴

$$g_A(x) = \min\left\{x, \frac{1}{2}\right\}.$$

In this environment, the welfare-maximizing threshold rule is described by threshold function $t_A(v) = t_B(v) = -v$, with exclusion types $\omega_A = 0$ and $\omega_B = -1$, as can be easily verified from [Proposition 2](#). Total welfare under the optimal threshold rule is $\frac{1}{12}$. Now consider the following nonthreshold rule, which we describe by its side- A correspondence:

$$\mathbf{s}_A(v_A) = \begin{cases} [-\frac{9}{8}, -1] & \text{if } v_A \in [\frac{3}{4}, 1] \\ [-1, 0] & \text{if } v_A \in [\frac{1}{2}, \frac{3}{4}] \\ [-v_A, 0] & \text{if } v_A \in [0, \frac{1}{2}]. \end{cases}$$

It is easy to check that this rule is implementable. Total welfare under this rule equals $\frac{3}{32} > \frac{1}{12}$. $\diamond$

The matching rules in this example are illustrated in [Figure 3](#).

Intuitively, the reason why threshold rules fail to be optimal (and segmentation occurs) is that they fail to maximize the benefits of cross-subsidization. Agents from side B with valuation $v_B \in [-2, -1]$ are more expensive but significantly more attractive than

²⁴That the function g_A has a kink simplifies the computations but is not important for the result; the sub-optimality of threshold rules clearly extends to an environment identical to the one in the example but where the function g_A is replaced by a sufficiently close smooth concave approximation.

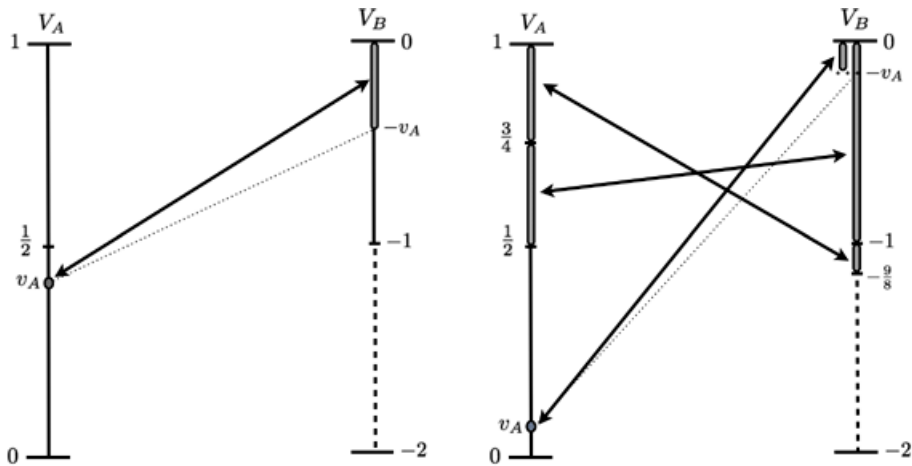


FIGURE 3. The welfare-maximizing rule among those with a threshold structure (left) and the welfare-improving nonthreshold rule (right) from Example 7.

agents with valuation $v_B \in [-1, 0]$. That salience is decreasing in valuations (weakly on side A , strictly on side B) per se does not make threshold rules suboptimal. Indeed, as established in Proposition 1, were preferences for matching intensity weakly convex on both sides, threshold rules would maximize welfare. Under concavity, however, once the high-valuation agents from side A interact with the high-valuation agents from side B (those with $v_B > -1$), they no longer benefit from interacting with agents from side B whose valuation is low (those with $v_B \leq -1$). This is inefficient, for those side- B agents with a low valuation are in fact the most attractive ones from the eyes of the side- A agents. More efficient cross-subsidization (and hence higher welfare) can be achieved by matching high-valuation agents from side A *only* to low-valuation agents from side B (segmented matching).

REFERENCES

- Allard, Nicholas W. (2008), "Lobbying is an honorable profession: The right to petition and the competition to be right." *Stanford Law & Policy Review*, 19, 23–68. [1006]
- Ambrus, Attila and Rossella Argenziano (2009), "Asymmetric networks in two-sided markets." *American Economic Journal: Microeconomics*, 1, 17–52. [1030]
- Arnott, Richard and John Rowse (1987), "Peer group effects and educational attainment." *Journal of Public Economics*, 32, 287–305. [1009]
- Becker, Gary S. (1973), "A theory of marriage: Part 1." *Journal of Political Economy*, 81, 813–846. [1010]
- Bedre-Defolie, Özlem and Emilio Calvano (2013), "Pricing payment cards." *American Economic Journal: Microeconomics*, 5, 206–231. [1010]
- Board, Simon (2009), "Monopoly group design with peer effects." *Theoretical Economics*, 4, 89–125. [1009]

- Caillaud, Bernard and Bruno Jullien (2001), “Competing cybermediaries.” *European Economic Review*, 45, 797–808. [1030]
- Caillaud, Bernard and Bruno Jullien (2003), “Chicken & egg: Competition among intermediation service providers.” *Rand Journal of Economics*, 34, 309–328. [1030]
- Cesari, Lamberto (1983), *Optimization Theory and Applications: Problems With Ordinary Differential Equations* (A. V. Balakrishnan, ed.), volume 17 of Applications of Mathematics. Springer-Verlag, New York. [1044]
- Crémer, Jacques and Richard P. McLean (1988), “Full extraction of the surplus in Bayesian and dominant strategy auctions.” *Econometrica*, 56, 1247–1257. [1013]
- Damiano, Ettore and Hao Li (2007), “Price discrimination and efficient matching.” *Economic Theory*, 30, 243–263. [1009, 1010, 1012, 1021]
- Damiano, Ettore and Hao Li (2008), “Competing matchmaking.” *Journal of the European Economic Association*, 6, 789–818. [1032]
- Eeckhout, Jan and Philipp Kircher (2010), “Sorting and decentralized price competition.” *Econometrica*, 78, 539–574. [1010]
- Ellison, Glenn and Drew Fudenberg (2003), “Knife-edge or plateau: When do market models tip?” *Quarterly Journal of Economics*, 118, 1249–1278. [1030]
- Epplé, Dennis and Richard E. Romano (1998), “Competition between private and public schools, vouchers, and peer-group effects.” *American Economic Review*, 88, 33–62. [1009]
- Evans, David and Richard Schmalensee (2010), “Failure to launch: Critical mass in platform businesses.” Working Paper, SSRN 1353502. [1030]
- Helsley, Robert W. and William C. Strange (2000), “Social interactions and the institutions of local government.” *American Economic Review*, 90, 1477–1490. [1009]
- Hoppe, Heidrun C., Benny Moldovanu, and Emre Ozdenoren (2011), “Coarse matching with incomplete information.” *Economic Theory*, 47, 75–104. [1029]
- Johnson, Terence R. (2013), “Matching through position auctions.” *Journal of Economic Theory*, 148, 1700–1713. [1009, 1010]
- Jones, Stephen (1805), *The Spirit of the Public Journals for 1799: Being an Impartial Selection of the Most Exquisite Essays and Jeux d’Esprits, Particularly Prose, That Appear in the Newspapers and Other Publications*, Vol. 3. James Ridgway, York Street, London. [1006]
- Jullien, Bruno (2012), “Two-sided b to b platforms.” In *The Oxford Handbook of Digital Economy* (Martin Peitz and Joel Waldfogel, eds.). Oxford University Press, New York. [1012]
- Jullien, Bruno and Alessandro Pavan (2014), “Platform pricing under dispersed information.” Working Paper, Center for Mathematical Studies in Economics and Management Science 1568R. [1032]

Kaiser, Ulrich and Minjae Song (2009), “Do media consumers really dislike advertising?: An empirical assessment of the role of advertising in print media markets.” *International Journal of Industrial Organization*, 27, 292–301. [1011]

Kaiser, Ulrich and Julian Wright (2006), “Price structure in two-sided markets: Evidence from the magazine industry.” *International Journal of Industrial Organization*, 24, 1–28. [1011]

Kang, Karam and Hye Young You (2016), “Lobbyists as matchmakers in the market for access.” Unpublished, Carnegie Mellon University and Vanderbilt University. [1008, 1012]

Lazear, Edward P. (2001), “Educational production.” *Quarterly Journal of Economics*, 116, 777–803. [1009]

Lee, Robin S. (2013), “Vertical integration and exclusivity in platform and two-sided markets.” *American Economic Review*, 103, 2960–3000. [1010]

Lee, Robin S. (2014), “Competing platforms.” *Journal of Economics and Management Strategy*, 23, 507–526. [1032]

Legros, Patrick and Andrew F. Newman (2002), “Monotone matching in perfect and imperfect worlds.” *Review of Economic Studies*, 69, 925–942. [1010]

Legros, Patrick and Andrew F. Newman (2007), “Beauty is a beast, frog is a prince: Assortative matching with nontransferabilities.” *Econometrica*, 75, 1073–1102. [1010]

Lucking-Reiley, David and Daniel F. Spulber (2001), “Business-to-business electronic commerce.” *Journal of Economic Perspectives*, 15, 55–68. [1012]

Maskin, Eric S. and John Riley (1984), “Monopoly with incomplete information.” *Rand Journal of Economics*, 15, 171–196. [1008, 1009, 1021]

McAfee, R. Preston (2002), “Coarse matching.” *Econometrica*, 70, 2025–2034. [1029]

Mezzetti, Claudio (2007), “Mechanism design with interdependent valuations: Surplus extraction.” *Economic Theory*, 31, 473–488. [1013]

Milgrom, Paul R. and Robert J. Weber (1982), “A theory of auctions and competitive bidding.” *Econometrica*, 50, 1089–1122. [1033]

Mussa, Michael and Sherwin Rosen (1978), “Monopoly and product quality.” *Journal of Economic Theory*, 18, 301–317. [1008, 1009]

Myerson, Roger B. (1981), “Optimal auction design.” *Mathematics of Operations Research*, 6, 58–73. [1010]

Myerson, Roger B. and Mark A. Satterthwaite (1983), “Efficient mechanisms for bilateral trading.” *Journal of Economic Theory*, 29, 265–281. [1021]

Rayo, Luis (2013), “Monopolistic signal provision.” *B.E. Journal of Theoretical Economics*, 13, 27–58. [1009]

- Rysman, Marc (2009), “The economics of two-sided markets.” *Journal of Economic Perspectives*, 23, 125–143. [1010]
- Seymour, John Barton (1928), *The British Employment Exchange*. PS King & Son, London. [1006]
- Shimer, Robert and Lones Smith (2000), “Assortative matching and search.” *Econometrica*, 68, 343–369. [1010]
- Strausz, Roland (2006), “Deterministic versus stochastic mechanisms in principal-agent models.” *Journal of Economic Theory*, 128, 306–314. [1012]
- Weyl, E. Glen (2010), “A price theory of two-sided markets.” *American Economic Review*, 100, 1642–1672. [1010, 1030]
- White, Alexander and E. Glen Weyl (2015), “Insulated platform competition.” Working Paper, SSRN 1694317. [1030]
- Wilson, Robert (1989), “Efficient and competitive rationing.” *Econometrica*, 57, 1–40. [1029]
- Wilson, Robert B. (1993), *Nonlinear Pricing*. Oxford University Press, Oxford. [1009]

Co-editor Johannes Hörner handled this manuscript.

Submitted 2014-6-26. Final version accepted 2015-10-19. Available online 2015-10-21.