

Preferences and Productivity in Job Matching: Theory and Empirics from Internal Labor Markets*

Bo Cowgill

Columbia Business School

Jonathan M.V. Davis

University of Oregon

B. Pablo Montagnes

Emory University

Patryk Perkowski

Columbia Business School

January 27, 2022

Abstract

A principal often needs to match agents together to perform coordinated tasks. However, agents can quit or slack off if they dislike their match. We study two approaches for matching agents, both widely-used in practice: Centralized assignment by firm leaders, and self-organization through market-like mechanisms. Our model connects the choice of method to the degree of specialization in a firm's workforce, the firm's production technology, worker/CEO information asymmetry, incentive alignment, and firm size. We then study these topics using data from a large organization's internal labor market. Centrally assigned matches are highly productive. Using the organization's preferred metric, the optimal match is 36% more productive than randomly assigned matches within job categories. By contrast, preference-based matches (using deferred acceptance) are only 3% better than random, but are ranked more favorably by the workforce. A key driver of results is the degree of assortative matching: The self-organized match is positively assortative, and the most productive match is negatively assortative.

JEL Classification: M5, D47, J4.

Keywords: Internal labor markets, assortative matching, assignment mechanisms, team formation, matching.

*The authors thank Bob Gibbons, Russ Coff, John Hatfield, Zoey Jiang, Fuhito Kojima, Scott Kominers, Jin Li, John Morgan, Tayfun Sonmez, Alex Teytelboym and Brian Wu for helpful comments, as well as seminar participants listed in *additional acknowledgements* (8.0.0.1). Cowgill thanks the Kauffman Foundation. An earlier version of this paper was called, "Matching for Strategic Organizations."

1 Introduction

Principals are often required to coordinate agents by matching. For example, workers and managers in companies must be assigned to each other to perform tasks (Gardenfors, 1973; Roth et al., 1993). In these settings, agents’ preferences are a critical consideration. A worker whose preferences are violated could quit or become demotivated. However, worker preferences are not the *only* consideration. Workers may not fully internalize their organization’s success. Left to their own devices, agents may pursue other goals without knowing (or sufficiently caring) about the principal’s performance.

This paper is about how organizations navigate these tensions in matching. We study two approaches to matching agents, both widely-used in practice: Centralized assignment by firm leaders, and self-organization through market-like mechanisms. Since 2000, a growing number of companies have adopted market-based systems for internal assignments. These systems allow workers and managers to self-organize based entirely on their preferences. An ecosystem of venture-backed software startups exists to facilitate these market-based matching systems,¹ with clients around the Fortune 500 and a combined market capitalization of over \$400M.² Some of these systems explicitly borrow mechanisms developed in the market design literature.³

Our model examines the strengths and weaknesses of these market- or preference- driven systems as compared to centrally-planned hierarchical assignment (“command-and-control”). We derive boundary conditions for when each method is better for a principal facing a matching problem. Our theory model connects this choice to key firm characteristics, including the degree of specialization among the workforce, the production technology, worker/CEO information asymmetry, incentive alignment, and firm size. Our model makes predictions about what types of organizations will adopt either approach. We can also use it to study how important secular trends in the labor market affect team formation, internal organization, and principal-agent problems inside firms.

¹These startups Gloat (<http://gloat.com>), Fuel 50 (<http://fuel50.com>), and Hitch Works (<https://hitch.works/>).

²Gloat’s clients including Wal-Mart, PepsiCo, Vanguard, Unilver, and Nestlé. For Gloat’s valuation, see <https://techcrunch.com/2021/06/16/gloat-raises-57m-to-reinvent-the-internal-job-board/>.

³For an overview of market design, see Roth (2015). We review a list of organizations using market design tools for internal labor markets in Section 2. The list includes Google (Cowgill and Koning, 2018), the US Army (Greenberg et al., 2020; Davis et al., 2020), Teach For America (Davis, 2016), and the International Monetary Fund (Barron and Vardy, 2005).

We then turn to a large organization for an empirical case study of the forces in our model. The structure of the organizations’ internal labor market – and particularly its use of the deferred acceptance algorithm – allows us to use detailed ranking data by workers and managers about each other. We also use the firm’s preferred metric for estimating the productivity of each match.

Our results show a relatively high degree of worker-specialization, and thus matching can be incredibly productive. Using our company’s preferred measure, the optimal match is 36% more productive than randomly assigned matches within job categories. This is a multiple-standard deviation improvement. To achieve the same improvement by training or replacement, the bottom 75% of workers perform as well as the 75th percentile employee. Despite these productivity improvements, workers and managers rank their assignments in this match as approximately as good (bad) as randomly assigned partners.

We then contrast these results with matches using worker/division preferences, implemented using the deferred acceptance algorithm. Preference-driven matches are much less productive: only 3% more valuable to the firm than the average random match. However, workers and managers rank their assignments better by an average of $\approx$ two slots, and are 85 percentage points less likely to be assigned to a role they ranked as tied for last place (but better than quitting).

Why do preference-respecting perform so poorly? In our setting, maximizing firm productivity results in slightly *negative* assortative matching: The firm’s best managers are not necessarily matched with the firm’s best workers (Becker, 1973; Shimer and Smith, 2000; Bandiera et al., 2007; Lazear et al., 2015; Adhvaryu et al., 2020; Kambhampati and Segura-Rodriguez, 2020). This result is driven by opportunity costs: Assigning a star worker to a top manager “wastes” the worker contributions on projects that will easily succeed without elite talent. In our setting, we thus find that disassortative matching is preferable to the firm.

By contrast, deferred acceptance produces greater assortative matching as high-productivity workers and high-productivity managers rank each other favorably. Although pleasing workers can improve performance, several bits of evidence suggest that the magnitude of productivity loss from the DA match are not worth the benefits of retention and other benefits of DA. Despite the drawbacks we highlight, we cannot necessarily characterize the adoption of preference-based assignment as a mistake in our setting, at least without knowing more about the firm’s objectives and costs. Some benefits of these programs may

be subtle or hard to measure.

Although we study only two mechanisms (command-and-control and), the principal’s problem has a variety of possible solutions.

We highlight ways this is distinct from other settings

The remainder of this paper is organized as follows. Section 2 briefly discusses how our paper relates to prior literature. Section 3 presents a model of matching with organizational objectives and constraints. Section 4 describes our empirical setting and Sections C and ?? review our data and econometric specifications. We present empirical findings in Section 6 and conclude in Section 8.

2 Literature Review

The question of how to allocate resources and talent within an organization has been central in the economic and management literature (Boudreau, 2004; Harris et al., 1982). A key constraint faced by firms is the dispersion of relevant information across participants in an organization and incentives for agents or divisions to misreport their private information (Ortega, 2001; Arya et al., 2002; Arya and Mittendorf, 2006; Harris et al., 1982; Christensen and Knudsen, 2010; Groves and Loeb, 1979; Lovejoy, 2006). These problems are similar to those in markets without prices. Coase (1937) famously noted that: “the distinguishing mark of the firm is the suppression of the price mechanism” (p. 389). The similarity of intrafirm labor markets to two-sided markets has been recognized since at least Gardenfors (1973). Our paper is related to several strands of prior literature listed below.

2.0.0.1 Internal labor markets. (Baker et al., 1994; Baker and Holmstrom, 1995). The prior literature on internal labor markets is mostly concerned with vertical transfers, promotions, or other instances where productivity growth has led a worker to desire new responsibilities and pay. By contrast, our paper is about horizontal transfers without promotions. A small literature directly and indirectly examines assortative matching inside firms (Bandiera et al., 2007; Lazear et al., 2015; Adhvaryu et al., 2020; Kambhampati and Segura-Rodriguez, 2020). Most of this literature examines a single vertical dimension of match quality: whether “good” workers should be matched with “good” managers. Our

theory and empirics have space for both vertical and horizontal dimensions of quality, and allow us to measure their relative importance to workers, managers and the firm. In our empirics, we find that workers' and managers' preferences for each other show few signs of vertical differentiation. They instead show strongly idiosyncratic person-specific horizontal preferences.

2.0.0.2 Coordination within organizations. Our paper is related to decentralized management structures ([Groves and Loeb, 1979](#); [Christensen and Knudsen, 2010](#); [Arya et al., 2002](#); [Ortega, 2001](#)), resource allocation ([Arya and Mittendorf, 2006](#); [Harris et al., 1982](#); [Gardenfors, 1973](#)) and coordination within large organizations ([Kouvelis and Lariviere, 2000](#); ?). Prior work analyzes optimal delegation when agents may have different objectives than the principal [Holmstrom \(1984\)](#).⁴

Outside of academia, a few companies have embraced self-organization in internal management (for example, Valve, Zappos, and Morning Star [Martela \(2019\)](#)) and have attracted significant attention from management scholars ?.⁵ Our paper is conceptually related to these management practices, although we are focused on a particular stylized implementation.

2.0.0.3 Matching through market design. Our paper relates to the applied market design literature for matching ([Roth and Sotomayor, 1990](#); [Roth and Peranson, 1999](#); [Abdulkadiroglu and Sönmez, 2013](#)). In this literature, most markets are composed of autonomous agents who are free to match outside of a centralized mechanism. As a result, these papers almost always restrict attention to mechanisms which determine matches using agents' preferences, and often restricts attention to mechanisms which yield pairwise stable matches in the sense that no unmatched pair of agents would prefer to be matched together over their assigned match. As we motivate in our theoretical section (3), orga-

⁴More recently, other scholars compare centralized and decentralized coordination when self-interested division managers have private information about local conditions and communicate strategically with cheap talk [Alonso et al. \(2008\)](#). In their model, decentralization may promote better coordination especially when coordination is relatively important. When a central authority has decision rights and coordination is relatively important, managers strategically exaggerate local conditions in order to manipulate the centralized decision. In contrast, when managers hold decision rights and coordination is important, they optimally share more information with other managers. Our paper highlights the risks associated with using decentralized assignment mechanisms when market participants have different objectives than the organization.

⁵Gary Hamel's article in *Harvard Business Review*, "First, Let's Fire All the Managers" ([Hamel, 2011](#)), is a popular example of this discussion.

nizational settings allow a centralized agent with a separate objective to block matches. Despite this, centralized organizational leaders do have incentives to anticipate worker preferences (insofar as workers can choose to quit the firm entirely or relax effort).

Internal labor markets appear to be a popular and growing application of standard market design tools. These applications have become popular as scholars and practitioners have realized the untapped value of internal mobility (Bidwell, 2011; Cappelli and Keller, 2014; Benson and Rissing, 2020). Market design tools appear to be a principled way to structure internal hiring and assignments. Internal markets explicitly using market design tools have appeared at Google (Cowgill and Koning, 2018), Teach for America (Davis, 2016), the United States Military Academy (Sonmez and Switzer, 2013; Sonmez, 2013), the US Army (Greenberg et al., 2020; Davis et al., 2020), the International Monetary Fund (Barron and Vardy, 2005), and others.⁶ In July 2021, a report by Georgetown and Harvard Universities (?) indicated that the US State Department is collaborating with the National Resident Matching Program (NRMP) “to design, execute, and evaluate a pilot use of a preference-matching algorithm” in assigning American foreign service officers to overseas diplomatic posts.

In addition, firms such as McKinsey (Bryan and Joyce, 2007) and Deloitte (Erickson et al., 2018) have adopted formal internal talent marketplaces and advocated them for clients. There are also several venture-backed startups that sell internal talent marketplace software to Fortune 500 clients, including Gloat (<http://gloat.com>), Fuel 50 (<http://fuel50.com>), and Hitch Works (<https://hitch.works/>). Although the microstructures behind these marketplaces are not disclosed in detail, the descriptions of the applications above suggest that respecting participants’ preferences (the key property in our theoretical results) is a guiding principle.

⁶Other examples include Cowgill et al. (2020), which studies a large Asian bank with “millions of customers, billions of dollars in assets and in revenues, and thousands of employees” that uses deferred acceptance to match and re-allocate workers, managers and projects. The United States Military Academy uses a cumulative offer mechanism to assign cadets to branches (Sonmez and Switzer, 2013; Sonmez, 2013). The International Monetary Fund uses the deferred acceptance algorithm to assign new economists to research teams (Barron and Vardy, 2005).

3 Model

3.1 Preliminaries

An organization consists of a principal (a CEO) and two types of agents: Workers and divisions (the latter led by middle managers). Although the workers and managers are differentiated, we refer to them collectively as “agents,” “participants,” or “the workforce.” There are I workers $i \in \{1, \dots, I\}$ and J divisions $j \in \{1, \dots, J\}$. The CEO wants to match workers and divisions to maximize the output V of the firm. For a private company, V could represent total firm profits, but for non-profits or governments V could represent other, potentially abstract objectives. Both workers and divisions can both be matched with a single member of the other side of the market. All potential worker-division matches yield a match-specific output v_{ij} . The CEO’s utility is equal to the sum of the productivity of all matches.

Without any other constraints, the CEO can find this match and implement it using the [Kuhn-Munkres](#) linear programming algorithm ([Kuhn 1955](#), a.k.a. the “Hungarian algorithm”). This solution finds the set of assignments that maximizes the sum of all v_{ij} s. However, most CEOs are not so unconstrained. Workers and divisions may have preferences over their match partners. They may quit or slack off if assigned to a match they dislike. There are many ways to model how a CEO could be penalized for violating workers’ or managers’ preferences. Below, we develop a model in which unhappy agents exit the company and create retention problems. Unhappy workers could also exert lower effort, or demand higher wages for compensating differentials, or create other problems with similar negative payoffs to the CEO. Analogous issues arise for divisions. Although an entire division is unlikely to quit, the manager leading them can.

To formalize the workforce’s preferences, we say that worker i ’s utility from matching with division j is μ_j^i and division j ’s utility from this match is analogous u_i^j . We also make two substantive assumptions.

Assumption 1 (Quits Reduce Output). *A match is acceptable to a worker or division if it yields greater utility than their outside option $\underline{u}$. If either the worker or the division finds the match unacceptable, the match yields zero output for the CEO.*

Assumption 2 (Private Information). *The CEO knows the productivity of all pairs of workers and managers (v_{ij}), but knows only the probability that each match is acceptable or not.*

The scale of the problem grows very quickly with the size of the organization. There are $N!$ ways to match N workers and N divisions. Each agent has $N!$ possible rankings over all members of the other side of the market. Therefore, there are $(N!)^{2N}$ possible combinations of preference profiles across all agents in the market.

To handle this volume, we provide some additional notation. Let $a = 1, 2, \dots, N!$ index all possible ways to assign N workers and N managers. Let $R(a)$ for a proposed assignment equal its retention rate, or the proportion of the N paired assignments where one or both sides will quit (resulting in zero output). The CEO does not know $R(a)$, but may have beliefs about it.

In addition, $V(a)$ is equal to the sum of all productivities v_{ij} that are assigned in match a , assuming that no worker quits. Formally, $V(a) = \sum_{i=1}^I \sum_{j=1}^J \alpha_{ij}(a) v_{ij}$, where $\alpha_{ij}(a)$ is equal to 1 if worker i and division j are matched in assignment a , and 0 otherwise. Across all $N!$ actions, $V(a)$ has a discrete frequency distribution $\mathcal{V}$ with finite support.

$\mathcal{V}$ is an expression of the firm's production function and output technology. The support of $\mathcal{V}$ measures how sensitive a firm's output is to changes in matching. A firm whose $\mathcal{V}$ places all mass on one point will have the same output, as long as nobody quits, irrespective of how workers are assigned. A $\mathcal{V}$ with wider support represents a firm whose output depends heavily on matching. We will later show that the properties of $\mathcal{V}$ affect the CEO's choice of assignment mechanism.

3.2 The CEO's Problem

The risk-neutral CEO's problem is to select a matching mechanism $\mathcal{M}$ to maximize the expected productivity of all matches:

$$\begin{aligned} \max_{\mathcal{M}} U &= \sum_{i=1}^I \sum_{j=1}^J \alpha_{ij}(\mathcal{M}) v_{ij} P(i \rightarrow j \text{ acceptable} | \mathcal{M}), \\ \text{s.t.: } \sum_{i=1}^I \alpha_{ij}(\mathcal{M}) &= 1, \quad \sum_{j=1}^J \alpha_{ij}(\mathcal{M}) = 1. \end{aligned}$$

We contrast the CEO's problem from other two-sided matching problems that appear to be similar. The key property of our setup is that participants' propensity for quitting (or slacking, etc) is private information. Without access to this information, the CEO could use

prior beliefs about quits, and incorporate these into payoffs into assignments. Workers in this case would have no input, outside of their influence on the CEO's priors. Because of the lack of input or choice, we call this "command and control" (or "CC") in our analysis below.

If the CEO's priors are uninformative, the CEO may consider steps to gather workers' private information. This would give workers input and allow the CEO to better anticipate the assignments the trigger quits (and avoid them). However, this task faces an elicitation problem: Workers and divisions may be tempted to misrepresent their preferences in order to manipulate their assignment.

In other settings, market- and mechanism- designers have developed strategy-proof mechanisms for eliciting preferences. However, the CEO's problem is distinct from the market designer's. We highlight two key reasons: First, the CEO has preferences over the assignment. In typical applications of stable matching, the match institution is broadly indifferent towards outcomes (subject to stability constraints).⁷ Second, although CEOs cannot stop workers from exiting, they do have the power to stop unwanted matches (even if both worker and division care to proceed). For this reason, typical stability considerations may be less critical for the CEO.

Because of these differences, the CEO has the ability (and the incentive) both to commandeer matching, and as well as to solicit input. However, it is possible for the CEO to do both too much or too little of either. It is possible to become too concerned with worker happiness, at the expense of firm output. Similarly, it is possible to be too insensitive to workforce preferences so that workers quit. As such, the CEO may want to engage in tradeoffs between i) the benefits of incorporating private information, ii) the information loss and elicitation challenges of incorporating it, and iii) the coordination benefits of central assignment. In this paper, we study two specific institutions at the extremes: Centralized command-and-control and full delegation.

⁷For example, the National Medical Residency Match does not attempt to match agents according to its own views about which doctors should work where, but rather in a way that pleases doctors and hospitals (subject to stability constraints).

3.3 Additional Assumptions: Matching Protocols and Preferences

As specified above, the CEO’s problem is highly general. To proceed with our analysis, we add some additional structure to the CEO’s options. We begin with command-and-control, or CC. To implement CC, the CEO must calculate the highest payoff assignment, balancing the potential payoffs $V(a)$ of any assignment against the amount (and distribution) of quits. Appendix A.1 shows how to calculate this solution generically using the Kuhn-Munkres method.

For delegation, several mechanisms exist for matching with preferences. We study one in particular:

Assumption 3 (Deferred Acceptance). *The CEO implements delegated matching using the deferred acceptance algorithm (Gale and Shapley, 1962; Roth, 2008a).*

The DA model gives us a tractable model of delegated matches that is strategy proof (for the proposing side, and both sides in many models of large markets)⁸. DA is also widely-used in practice (Roth, 2008a), including the applied setting of our paper. Our results will be framed in terms of workers-proposing DA, but we will highlight where the proposing side matters. While we proceed using DA, many of our results are likely more general than DA. We suggest a few cases where our results seem to be particularly detached from DA, and may apply to a wider variety of preference-driven mechanisms.

Finally, we add structure to participant preferences.

Assumption 4 (I.I.D. Preferences). *Workers’ utilities μ_j^i and division utilities u_i^j for each match partner are independently drawn from a common distribution G .*

Under Assumption 4, all preference profiles (i.e., ways for workers to rank managers and vice versa) are equally likely. This type of random matching market is a workhorse model in the literature, particularly when applied to deferred acceptance (Knuth, 1976, 1997; Pittel, 1989; Knuth et al., 1990; Pittel, 1992b; Ashlagi et al., 2017). Adopting this formulation allows us to connect with the prior literature, and makes some results more tractable. For example, Assumption 4 means that the probability that both sides in an assigned pair find the match acceptable is $(1 - G(\underline{u}))^2$. We call this the expected *retention rate* of the pair.

⁸See e.g., Rees-Jones (2018, 2017); Immorlica and Mahdian (2005); Kojima and Pathak (2009); Azevedo and Budish (2019).

However, I.I.D. is a very strong assumption; among other things it rules out: i) vertical preferences on either side of the market, ii) correlated preferences *across* the market (mutual attraction or repulsion), and iii) correlations between what the firm wants and what workers/managers like. In a sense, Assumption 4 helps increase the retention benefits of deferred acceptance (and other delegated solutions to the CEO's problem). Because of the lack of vertical preferences, there is little competition between agents on the same side. As a result, agents in DA are more likely to match with top ranked matches (without being crowded out by competitors) and avoid quits. We explore this further in Section 3.8, where we relax Assumption 4 and allow preferences to be correlated.

3.4 Delegation vs Command and Control

We now have the technology to compare CEO payoffs under delegation and command and control. All proofs are in Appendix A.6. From the CEO's perspective, all $N!$ possible matches have the same expected retention rate *ex-ante* under Assumption 4. Noted above, all pairwise matches have a probability $(1 - G(\underline{u}))^2$ of being acceptable to both parties; the average across all possible matches is the same. We denote this retention rate as $\bar{R}$.

CEO Payoffs from Command and Control. Because retention rates are the same for all possible matches, the CEO's choice under command and control will be based entirely on the productivity $V(a)$ of each match. Under command and control, the CEO will pick the match with the highest output out of all $N!$ possible matches.⁹ We call the output of this assignment V_{CC} ; it equals the output of the best assignment if nobody quits. We now need to integrate the possibility of quitting. Because the value of the match is uncorrelated with the expected retention rate (a constant) in our setup, the performance of command and control is the product $\bar{R}V_{CC}$.

CEO Payoffs from Deferred Acceptance. Under Assumption 4, the deferred acceptance algorithm produces all $N!$ potential assignments with equal probability. As a result, the CEO's expected output of DA is $\bar{V}$, the average output over all $N!$ possible matches. However, the merits of DA also depend on retention.

⁹This would be the equivalent of running the [Kuhn-Munkres](#) and selecting the solution.

Lemma 1. *The expected retention rate of the match selected by DA is weakly higher than the average of all matches, $\bar{R}$.*

The proof for this in Appendix A.6 uses a key property of DA:

Lemma 2. *In an $N \times N$ market where workers and divisions rank all possible matches, at most $\sum_{l=1}^k l$ individuals on the proposing side can match with a partner they rank k^{th} or higher.*

In DA, at most 1 worker can get her last choice, at most 3 workers can get their second to last or last choice, at most 6 workers can get their third to last choice or worse, and etc. By eliminating many assignments in which more workers get low-ranked choices, DA raises the retention rate $\bar{R}$. Let R_{DA} equal this improved retention rate. The expected output of DA is equal to $R_{DA}\bar{V}$.

Lemma 1 is an example of a result that is more general than DA, and likely extends to many two-sided matching algorithms that incorporate workforce preferences in some way. To improve retention above $\bar{R}$, a matching approach would simply need to rule out some bad matches for either side.

Delegation vs. Command and Control. Together, this analysis implies CC is expected to yield greater output than DA when the following inequality holds:

$$\bar{R}V_{CC} \geq R_{DA}\bar{V}. \quad (1)$$

Rearranging, we can see that CEO's choice of CC or DA depends on whether improvements in productivity or retention are more important:

$$\frac{V_{CC}}{\bar{V}} \geq \frac{R_{DA}}{\bar{R}}. \quad (2)$$

The left hand side is a property of $\mathcal{V}$, the firm's output technology; it essentially measures how sensitive firm's output is to changes in matching. The right hand side is a function G , the distribution from which workers' and managers' preferences are drawn. Using these, we can derive our first set of results.

3.5 Match-Specific Output: The Benefit of Command-and-Control

Definition 1 (Specialization). *A worker i is unspecialized if the outputs (v_{ij} s) are equal for all possible assignments. The worker exhibits specialization in some tasks if these outputs vary across assignments.*

Proposition 1. *The performance of delegation will equal or exceed that of command-and-control in firms where workers are completely unspecialized.*

The intuition behind Prop. 1 is that rearranging unspecialized workers generates no productivity benefits. If workers are equally productive in all divisions, then the most productive match (V_{CC}) is equal to the average match ($\bar{V}$). All the mass in $\mathcal{V}$ is concentrated in a single point, and the LHS of Equation 2 is equal to 1. By contrast, rearranging an unspecialized workforce could generate retention benefits. As long as baseline quit rates $G(\underline{u})$ are above zero, then the retention of DA will exceed CC.

Many firms have unspecialized workforces. At the extreme are “gig economy” companies such as Uber, Doordash, and competitors. Prop. 1 suggests a reason these firms are organized like markets. Because there is little match-specific productivity, these firms are relatively indifferent about which driver assists which customer. As a result, these firms’ approach to matching offers flexibility to accommodate workers’ preferences around the timing, location, and total amount of driving (Hall and Krueger, 2018).

Non-specialization also appears in more traditional workforces. A variety of empirical papers document a secular trend toward less specialized job assignments.¹⁰ One alleged cause of this is technological change, which made work in all occupations less routine (Autor et al., 2003). A related literature about the effect of information technology on job design suggests IT has moved work away from fixed categorization and towards “multi-tasking” by workers and job rotation.¹¹

Employers who hire generalists emphasize the flexibility, adaptiveness, and robustness that it affords; theory models by Dessein and Santos (2006); Deming (2017, 2021) formalize

¹⁰See Osterman (1994); Ichniowski et al. (1996); Brynjolfsson and Hitt (1998); for Economic Co-operation and Staff (1999); Caroli (2001); Caroli and Van Reenen (2001); Deming (2017, 2021). As Dessein and Santos (2006) note, the biggest management fad of the 1990s, reengineering (Hammer and Champy, 2009) prescribes “combining several jobs into one” and thus “putting back together again the work that Adam Smith and Henry Ford broke into tiny pieces.”

¹¹For examples, see Bresnahan (1999); Lindbeck and Snower (2000); Caroli and Van Reenen (2001); Bloom and Van Reenen (2011).

some of these intuitions. Google, one of the companies adopting deferred acceptance, openly advertises a preference for “smart generalists” in recruiting.¹² A 2019 Atlantic article, *At Work, Expertise is Falling out of Favor*,¹³ describes to generalist hiring within the US military, which has used preference-driven assignment mechanisms for a variety of personnel problems.¹⁴ Although Google’s engineers may be specialized within the overall economy, they are not specialized to any of their employers’ internal tasks. This reduces the gains from matching and makes deferred acceptance more attractive.

Corollary 1. *For command-and-control to outperform delegation, it is necessary (but not sufficient) for the workforce to be specialized.*

Eq. 2 shows why specialization alone is not sufficient. Specialization must also generate enough extra output to justify the retention loss of command-and-control. Both results (Prop 1 and Cor. 1) are driven partly by our I.I.D. assumption, particularly the lack of correlation between each agent’s preferences and their productivity to the firm.

In Section 3.9, we allow workers preferences to be correlated with their output. However, a degree of uncorrelatedness appears elsewhere in many settings. In our empirical setting, we study an internal mobility program created in response to worker unhappiness with their current specialization. Fear of “overspecialization” is a chronic topic of career advice.¹⁵ More generally, misalignment between principles and agents’ is a classic problem in economics of organizations.

Extension 1 (Noisy CEO Beliefs.) Until now, we have assumed that CEOs have perfect knowledge of workers’ match-specific productivities (v_{ij}). In Appendix A.2, we consider an extension in which the CEO observes these *noisily*. Even if workers are specialized, the CEO may be unable to observe the specializations.

Proposition A1 shows that such measurement problems diminish the benefits of command-

¹²See for example statements by Google HR executives at <https://www.cnbc.com/2017/11/28/the-most-important-skill-you-should-have-to-score-a-job-at-google.html> and <https://www.forbes.com/sites/georgeanders/2014/10/21/googles-people-chief-laszlo-bock-explains-how-to-hire-right/>

¹³<https://www.theatlantic.com/magazine/archive/2019/07/future-of-work-expertise-navy/590647/>

¹⁴The Army describes the “Army Talent Alignment Process” as: “a decentralized, regulated, market-style hiring system that aligns officers with jobs based on preferences. <https://talent.army.mil/atap/>

¹⁵For an example, see <https://www.fastcompany.com/3000791/how-specialize-yourself-right-out-job>.

and-control. If the CEO cannot implement a productive match, then the firm is better off enjoying the retention benefits of delegating. One implication of this result is that better monitoring and/or analysis technology can tilt organizations toward command-and-control. If better technology allows CEOs to reduce measurement noise, the match-specific benefits of CC may outweigh delegation.

3.6 Retention: The Benefit of Delegation

Our next set of results are about retention and the right hand side of Equation 2. Let γ be the expected value of G or the expected value of a random match.

Proposition 2. *Assume that $\underline{u} < \gamma$, then as $\underline{u}$ increases, the unconditional quit probability $G(\underline{u})$ increases and the returns to DA are higher.*

The intuition for this result is that if $G(\underline{u})$ is low, few people are at risk of quitting. As a result, the CEO sees few benefits to putting people into arrangements that prioritize their preferences. These workers are happy enough to remain, even in less preferred assignments. By contrast, when $G(\underline{u})$ is high, workers and managers will quit unless they get their top choices. As a result, there are larger returns to arranging workers in ways that give them their preferred options.

$G(\underline{u})$ can vary for reasons with economic interpretations.

Corollary 2 (Workforce-Unfriendly Firms). *Let G' be a leftward shift of G so that G' equals G minus a positive constant. DA is more attractive under G' than G .*

The shift in G' can be interpreted as making the firm less attractive to workers; draws from G' are more likely to lie below $\underline{u}$. Defined this way, workforce-unfriendliness could either be a common trait to an entire industry, or a firm-specific characteristic (possibly part of a strategy to tradeoff retention for other benefits). Workforce-unfriendly firms are not necessarily unsuccessful; they may be highly profitable by keeping costs low. Worker unfriendliness raises the retention benefit of placing workers in attractive jobs, and makes DA more attractive. DA and worker-unfriendliness are complementary; Corollary 2 also shows being an attractive destination for workers is complementary with command-and-control.

Corollary 3 (Outside Options). *DA is more attractive as $\underline{u}$ increases.*

Corollary 3 suggests that even firms that are relatively worker-friendly may find DA attractive if outside options are high enough. One example of this is Google, a company that regularly appears at the top of Forbes’ *Best Places to Work* list and offers workers free gourmet lunch and subsidized massages. Despite this, Google opened an internal talent marketplace in 2014 based on the deferred acceptance algorithm (Cowgill and Koning, 2018).

Corollary 4 (Asymmetric Information). *Let G' be a mean-preserving spread of G , so that G' has the same mean of G but higher variance. DA is more attractive for G' than G .*

Corollary 4 addresses the variance of G . Higher variance corresponds to greater CEO uncertainty about workers’ (privately-known) tastes, and thus an environment of greater asymmetric information. The CEO’s view of expected quitting is higher amid this uncertainty. This increases the expected returns to DA.

Finally, we can model the effects of higher quitting costs. Suppose every successful match has a benefit of $\frac{c}{N}$ in addition to its productivity v_{ij} . This additional benefit $\frac{c}{N}$ can be interpreted as a quitting cost (it is lost if the match is unsuccessful).

Lemma 3 (Quitting Costs). *Higher quitting costs attenuate the relative benefits of CC over DA, $\frac{V_{CC}}{V_{DA}}$.*

The results above offer some suggestive evidence for why internal talent marketplaces have grown popular in particular settings (high turnover, potentially workforce -unfriendly firms, industries with high outside options and/or retention costs). Although the results are discussed in terms of cross-sectional differences (i.e., differences in firms or industries), we can also apply them to intertemporal changes. Internal talent marketplaces have grown in the last decade, particularly during the pandemic. The last decade included record high quitting rates (Prop. 2), both before the pandemic, as well as the subsequent “great resignation.”

During this time, work may have become unattractive either because work becomes worse (Corollary 2, e.g. all-day Zoom meetings, social distancing at work), or because non-work becomes more attractive (Corollary 3, e.g. rising quality of leisure, Aguiar et al. 2021; post-pandemic attractiveness of self-employment, Fazio et al. 2021; Haltiwanger 2021). Tight labor markets increase quitting costs (Lemma 3).

The last decade also featured rising uncertainty including several uncertainty shocks (Altig et al., 2020, culminating with the pandemic). As workers were reallocated to new

tasks under new circumstances, a CEO’s understanding of workers’ satisfaction with potential assignments (Corollary 4) and workers’ productivity in those assignments (Extension 1) may have become less clear. Some commentators (Useem, 2019) allege a secular shift towards employers’ hiring generalists (Proposition 1), as a way of mitigating shocks (such as the ones mentioned above). Firm size (Proposition 3) addressed next) has also increased in recent decades, and a greater percentage of workers are employed in large firms (Hopenhayn et al., 2018).¹⁶

All of these factors increase the attractiveness of delegating matches. Companies offering talent marketplace software have gained traction in the last decade and claim the pandemic has been good for business;¹⁷ market leader Gloat.com’s reached its peak valuation in June 2021 of \$400M.¹⁸

3.7 Firm Size

We now explore how the CC vs DA choice changes with the size of a firm’s workforce, as measured by the number of workers and divisions. Our results here mainly address how the retention benefits of delegation scale with size (the right hand side of Equation 2). We do not model how the left hand side ($V_{CC}/\bar{V}$) about firm output changes with size. We take it to be exogenous and likely known by the CEO. The relationship between size and the productivity benefits to matching is theoretically ambiguous.¹⁹

Our results about quitting depend directly on the size of the organization. In this context, organizational size affects the amount of potential matches each agent has, and the number of potential competitors on the same side of the market.

Proposition 3 (Firm Size). *For sufficiently large markets, the expected retention rate of DA increases with N , the size of the firm, so long as $G(u)$ has a positive density over its full support.*

This result comes out of order statistics for G . In a random market with 250 workers and

¹⁶The post-pandemic growth in self-employment and entrepreneurship may affect this trend (Fazio et al., 2021; Haltiwanger, 2021).

¹⁷<https://gloat.com/blog/walmart-pepsico-nestle-unilever-deloitte-highlight-how-the-internal-talent-marketplace-is-revolutionizing-the-global-workforce/>.

¹⁸<https://techcrunch.com/2021/06/16/gloat-raises-57m-to-reinvent-the-internal-job-board/>.

¹⁹As firms get larger, they may have sophisticated pools of substitutable labor which make production less dependent on any single worker’s match (at least among rank and file workers). Alternatively, at large companies there may still remain high leverage matches even among non-executives.

managers, a worker is almost certain to match with one of her top 60 choices (Pittel, 1992a) and is most likely to match with one of her top ten choices. Because these are independent draws, even the 60th order statistic across 250 draws is almost certain to be above the $\underline{u}$ outside option.

Unless the benefits of CC similarly grow with size, Prop. 3 suggests that DA will be more attractive for large firms. Indeed, many of the adopters of internal market places are large, including the world’s two largest employers, Wal-Mart and the US military. There may be additional reasons for large firms to prefer internal markets besides the logic of Proposition 3.²⁰

Figure 1 visualizes the benefit of DA and how it varies with firm size, and for various base levels of quitting. Similar to Ashlagi et al. (2017), we estimate the probability a worker matches with her k^{th} choice via simulation. The figure plots $R_{DA}/\bar{R}$, the retention benefit of DA. A CEO who knows his firm’s production technology can compare the plotted values to the productivity benefits of CC ($V_{CC}/\bar{V}$) to determine if it should match with CC or DA (Equation 2).

This figure has a few implications. As Proposition 2 indicated, the payoffs to DA are higher if the baseline retention rate is lower (at all market sizes). As Proposition 3 suggests, larger market sizes increase the DA retention rate for all market sizes. In addition, we see a concavity in the result; as firms grow larger, there are declining marginal benefits.

3.8 Correlated Preferences

As noted above, our I.I.D. assumption is strong. It rules out: i) vertical preferences on either side of the market, ii) correlated preferences *across* the market (mutual attraction or repulsion), and iii) correlations between what the firm wants and what workers/divisions like. We now relax each of these restrictions.

²⁰In our conversations with practitioners, we have particularly heard about additional particular reasons that delegation is more popular with large firms. The first is that setting up a DA marketplace involves fixed costs in creating and populating a match website, and training workers and managers about the rules and expectations around the market. Second, as firms’ workforce grow in size, central management’s ability to keep track of match-specific preferences and outputs becomes harder. Algorithmic approaches to this problem such as those in our empirical section may help firms better keep track of the skills and preferences of a larger workforce.

Vertical Preferences. In many real world settings, preferences are vertical meaning that agents on one side of the market agree on the ranking of agents on the other side. Employers may all agree on the best workers, and workers may agree on the best firms. Vertical preference assumptions appear in many academic papers.²¹

To relax our IID assumptions and use vertical preferences, we need to introduce additional concepts. Vertical preferences can apply to one side, leaving the other I.I.D., or to both sides simultaneously. In our formulation, the CEO knows that preferences are vertical for one or both sides. However, the CEO does not know exactly which options are ranked 1st, 2nd, 3rd (... etc.) by the vertical side(s).²² We begin by assuming one side of the market has vertical preferences and the other side of the market never quits.

Proposition 4 (Vertical Preferences). *If one side of the market has vertical preferences and the other side of the market does not quit, then DA is statewise dominated by command-and-control.*

The intuition of this result is that the CEO's benefit from DA comes from improving retention by avoiding matches that are disliked by the proposing side. DA can never match all proposers to their last choice, but CC can. However, if preferences on one side of the market are vertical, then either all workers or divisions have the same first choice, second choice, etc. No matter how these agents are matched, exactly one agent will match with her k^{th} choice for all k . *Ex-ante*, CC and DA have identical expected retention rates.

Proposition 4 makes the strong assumption that one side of the market has vertical preferences and the other side does not quit. More generally, if one side is vertical and the other is I.D.D., the I.I.D. side could propose. This would possibly improve retention on the I.I.D. side of the market. Our stylized formulation of the problem does not allow the CEO to asymmetrically care more about one side of the market's retention.²³ Appendix A.3 contains a brief extension that allows this, and considers the merits of command-and-control vs. DA (where the proposing side's preferences are I.I.D., and the receiving side's preferences are vertical). We show that the choice depends on an equation very similar to Equation 2 (but focused on the I.I.D. proposing side only).

²¹For example, Agarwal (2015) and Agarwal and Diamond (2017), consider the econometric identification of preferences in a two-sided market under vertical preference assumptions.

²²In practice, this may be easy for a CEO to find out, by (say) bribing a worker to report on his peers' vertical preferences. This would of course decrease the value of DA, insofar as it would allow the CEO to incorporate the benefits of DA (information about quitting) without using DA.

²³In our formulation, the only penalty for quitting is zero productivity in the assignment where one (or both) sides quit. There are no replacement costs, or value retained from a worker who stayed (when their partner quit), or other sources of potential asymmetry between the quitting costs of either side.

Finally, both sides' preferences could be vertical, although still unknown, and uncorrelated with the CEO's preferences. In this case, Appendix Proposition A2 shows that command-and-control statewise dominates DA, irrespective of who proposes.

Semi-Vertical Preferences. The results above show that perfectly vertical preferences eliminates the potential benefits of DA. However, most preferences are not perfectly vertical; they may be highly correlated but not completely. Are the retention benefits of DA monotonically decreasing in workers' preference correlation? Although we do not have analytic results on this question, we present simulation results in Figure 2.²⁴ The figure suggests that the benefits are decreasing, and that the costs of correlation are convex. In a market with 10 workers and divisions, DA is expected to yield an only slightly lower retention rate if the correlation in utilities is 0.75 instead of 0. As the correlation approaches 1, however, the retention benefits quickly converge to 0.

3.8.0.1 Cross-Sided Correlations. Preferences can also be correlated across sides. Workers and managers' preferences may exhibit a high degree of reciprocity. If preferences are correlated across sides, but are otherwise I.I.D. (i.e., no vertical preferences), then the results in our I.I.D. model all hold. None of our results from this model rely on uncorrelatedness across sides. They are mainly driven by the lack of correlation between workforce preferences and production technology.²⁵

3.9 Alignment: Correlation Between CEO & Workforce Preferences

The final way we relax I.I.D. assumptions is to allow alignment between the workforce's preferences and the CEO's. The I.I.D. model proposed Section 3.3 assumed these were independent. In reality, a CEO and his workforce may broadly want very correlated (or negatively correlated) assignments. Workers may care about the CEO's goals if they

²⁴The simulation considers a 10 x 10 market. Workers' preferences are based on latent utilities over matches, all of which are normally distributed with mean zero, variance one, and covariance ρ . Division preferences are I.I.D. We replicate this simulation 2,000 times.

²⁵The addition of cross-sided correlations does not change expected retention rate under DA is the same, and thus Equation 2 is the same. This is partly because of our formulation of the problem, in which the CEO's penalties are the same irrespective of if an agent quits, or if the agent's match partner quits, or both quit. If the penalty for both quitting is higher, then cross-sided preferences would result in a higher degree of both-sides quitting than our standard I.I.D., and the benefits of DA would be higher (since the CEO gets extra credit for avoiding both sides quitting).

have intrinsic motivation for the company's mission. Of, they may be compelled to care through an incentive contract.

To formalize this type of alignment, we introduce two definitions:

Definition 2 (Broad Alignment). *An agent is said to have broadly aligned preferences with the firm if $\frac{\partial \mu_j^i}{\partial V(i-j)} > 0$ where $V(i-j)$ is the firm's total productivity if worker i matches with division j . We say an agent has strictly broadly aligned preferences with the firm if $\mu_j^i = f(V(i-j))$ where $f(\cdot)$ is any strictly monotonic function.*

Broad alignment means that agents gain utility from the company succeeding as a whole, and *strict* broad alignment means that the firm's overall success is the agent's *only* source of utility. Broad alignment would appear to be favorable for the use of deferred acceptance, as it would give workers a vehicle for using private information in a way that's productive for the firm. However, specifying the format of this alignment features some challenges. Broad alignment means that workers care about the output of the organization as a whole. This requires them to have preferences about what other workers do (i.e. preferences over all $N!$ possible ways of assigning all workers to divisions). However, deferred acceptance only permits workers to rank N members of the other side. As such, agents' rankings will be a function of beliefs about what other workers and divisions will do. We analyze the Nash equilibrium of this setup in the following proposition.

Proposition 5. *If workers and divisions have strictly broadly aligned preferences, the CEO's optimal match is the surplus-maximizing Nash equilibrium of deferred acceptance, but is not generally a unique Nash equilibrium. Deferred acceptance can equal the performance of command and control, but cannot exceed it.*

The result exhibits two key aspects of deferred acceptance as a management tool. First, the benefit of deferred acceptances comes from retention. The mechanism helps CEOs by incorporating workers' private information about when they will quit. However, if workers are strictly aligned, there would be no quitting anyway, thus eliminating this benefit. The CEO can implement the optimal match, achieving the same outcome, without delegation.

Second, when the workforce is broadly aligned, multiple equilibrium problems arise. Although Prop. 5 examines an extreme case (*strict* broad alignment) similar multiple equilibria issues would arise if workers and divisions had private values for each match and

were paid a percentage of total firm output. Prop. 5 highlights the coordination aspect of the CEO's problem; even when agents are maximally aligned, multiple equilibria exist. Command-and-control can be seen as an equilibrium selection device.

Broad alignment is a very strong assumption, and in many other settings researchers have found strong firm-wide incentives hard to introduce (Holmström, 1979; Holmstrom, 1982; Oyer, 2004). In practice, firms sometimes face a different form of alignment.

Definition 3 (Narrow Alignment). *An agent is said to have narrowly aligned preferences with the CEO if $\frac{\partial \mu_j^i}{\partial v_{ij}} > 0$. We say an agent has strictly narrowly aligned preferences with the CEO if $\mu_j^i = f(v_{ij})$ where $f(\cdot)$ is any strictly monotonic function.*

Narrow alignment means that agents gain utility from the productivity of their *own* matches (irrespective of their co-workers' matches). This is similar to paying salespeople for their own production, and not the success of the firm as a whole. Under this form of alignment, agents do not have to consider the behavior of others; each agent can rank all N possible match partners considering how productive they would be in the match (individually). However, there are limitations to this form of alignment.

Proposition 6 (Talent Hoarding). *Narrow alignment is not sufficient or to guarantee that DA yields output as high as CC.*

Even if all workers aspire to their most productive uses (individually), the result may be a suboptimal overall match. The intuition for Proposition 6 is about hoarding. Narrow alignment encourages each manager to seek (and retain) the most productive workers, even if these workers would have better use elsewhere in the company or are likely to quit in these jobs.²⁶ Similarly, narrow alignment by workers encourages the reciprocal “project hoarding” by workers seeking out the most productive projects, even if other workers would be more productive in these jobs. In the extreme case, narrow alignment on both sides of the match could produce an assortative match, which is only optimal for output if the production technology is supermodular (Becker, 1973), and whose retention is no better than command and control (because both sides preferences are vertical, Prop. A2).

In more general terms, the problem with narrow alignment is externalities. Even when a worker performs a job well, their assignment may leave others in the company without

²⁶Our setting is motivated by horizontally differentiated placements, but managers could also “hoard” talent by denying promotions to workers (Haeghele, 2021).

productive uses, leading to suboptimal overall output. Maximizing an organization’s performance may require some workers to perform in roles where they are *not* individually their most productive. Through these opportunity costs, each worker’s assignment imposes a *displacement externality* that affects other workers, divisions and the CEO’s objectives. The CEO can incorporate these externalities into decision making; narrowly-aligned agents do not automatically do so.

Prop. 6 shows how aligned preferences are insufficient for preference-respecting mechanisms to perform well. Appendix ?? shows that narrowly-aligned preferences are not *necessary* either. In this example, deferred acceptance leads to an optimal configuration, despite workerteacher preferences *not* being aligned according to our definition.²⁷

Despite these limitations, narrow alignment can be sufficient for DA to perform well for certain types of production technology. For example:

Corollary 5. *Suppose that workers and divisions have strictly narrowly aligned preferences, and that match productivities are the result of a production function that is supermodular in workers’ and divisions’ types. Then the worker-proposing deferred acceptance algorithm selects the CEO-optimal assignment. However, the CEO can select this without deferred acceptance using command-and-control.*

These results show settings where delegation performs well. Unlike our previous results, there is no coordination or multiple equilibria problem. However, like the previous results, there is no quitting under strict narrow alignment. As a result, there is no benefit to incorporating workers’ preferences. The CEO can obtain the same result by command and control, without delegation.

Cor. 5 shows that deferred acceptance and command-and-control sometimes produce the same outcome. In these settings, one may be preferable for reasons outside of our model. Delegating often requires a costly process of soliciting input. The case study of Google’s internal mobility program discusses both infrastructure (an internal website for browsing options and submitting preferences) and training programs to teach workers about the system. However, delegation may also have *benefits* outside the model. A long-

²⁷Of course, this raises the possibility that “negatively aligned” preferences could result in the optimal organizational match, even if preference-respecting mechanisms were used. In Appendix ??, we provide an example that this is possible. In our example, workers and managers want to *avoid* working on the projects where they are individually most productive. Nonetheless, preference-respecting mechanisms produce the optimal assignments.

standing view in organizational psychology and other disciplines is that workers have intrinsic, non-instrumental value for decision rights (Bartling et al. 2014, in addition see the psychology literature about “procedural justice,” Cropanzano 1993; Brockner 1996). Anecdotes from Google suggest that workers value the pageantry and symbolism of workers’ choice, separately from its instrumental value. These issues are outside of our model, but may be an important consideration in practice.

4 Setting and Institutional Details

Our theory model suggests that neither command-and-control nor delegation are optimal for all firms; instead, the CEO’s choice depends on parameters that can differ between firms and industries (and across time). We now study how the forces in our model unfold in an applied setting.

Our data come from a Fortune 500 company that develops software for business clients. The software includes, for example, customized tools for organizing and indexing special files and/or importing and managing content libraries. Employees are organized in teams serving a product and/or client, and each team has a single manager. These teams are the divisions in our theory model, and new (and/or internally mobile) workers need to be matched to them. Prior to the adoption of an internal job marketplace, each participant was assigned to a team indefinitely.

These teams featured a mixture of engineers and non-technical staff. Engineers on these teams held a BS in computer science, and non-technical staff held a BA in a social science, professional, or humanities subject. A single manager oversaw each team. Some products had multiple such teams, each working on different clients or different aspects of the product (with two separate managers).

Our setting features several characteristics appearing in our theoretical setup. The company’s publicly stated recruiting philosophy included both generalists and specialists (Definition 1 and Cor. 1), but overall leaned more towards generalists (Prop. ??). We later measure specialization directly. Workers and managers are relatively skilled and have good outside options (Cor. 3), and replacement costs are high (Prop. 3). The overall retention rate was high compared to industry standards (Prop. 2), but concern about retention led to the adoption of the jobs marketplace.

Workers are paid in part with stock compensation to align incentives with shareholder goals (broad alignment, Definition 2). However, workers and managers are also paid a similar amount in individual bonus based on their own performance (narrow alignment, Definition 3). For the workers in our sample, approximately 40% of their annual compensation is performance based, with about half coming from stock options (broad alignment) and the other half coming from cash performance bonuses. This percentage is higher for managers and more senior workers.

4.1 Internal Mobility

During the sample period, the organization's leadership sought to change the system of indefinite assignment to increase internal mobility. Both workers and team managers had expressed a desire for greater mobility. After employees spent several years in an organization, they sought career growth on a new project. Similarly, some managers sought to update their roster of workers with workers whose skills or interests were better aligned.

The company's leadership embarked on building a market-like structure. All project assignments were given a pre-established "term length," measured in quarters. All workers whose term was ending – including those who were successful in their previous jobs – would go into the market to be reassigned.²⁸ Those who wished to stay in their roles could remain, but only if they were matched again through the market system. Term lengths coordinated internal search around predictable points in the annual calendar in which lots of workers and managers would be searching for new assignments. By aligning the timing, the exchange market was more thick and featured more options for both sides.

To avoid disruption, entry into the market was staggered. The firm aimed for no more than 25% of workers to be on the market at any time, so that the remaining 75% could focus on their day-to-day work. Partly as a result of the staggering, each manager was typically recruiting at most one new worker in any given quarter's match. As a result the match was 1:1. Nearly all matches were unbalanced, and featured an excess of managers hoping to recruit new workers.

Each quarter, eligible participants were given access to a web application through which

²⁸"Term lengths" were created in part to avoid adverse selection in which only bad workers or managers sought new assignments.

they could develop a profile to describe their interests, accomplishments, and skills. Managers also included their job opening and skill requirements. Profiles could be searched or browsed. Although we could not obtain a screenshot for this paper without revealing internal user information, Figure 3 includes a mockup that replicates a hypothetical user's profile page.

Each side then submitted a rank-ordered list of their preferences before a deadline. Figure 4 contains a replica of the submission page. The workers and managers were then matched using the ranked preferences using the workers-propose version of the deferred acceptance algorithm (Assumption 3). The ranking tool also allowed workers to identify managers with whom the worker would rather go unmatched. For workers this would mean quitting, and for managers this meant having one fewer worker. As part of the introduction of match, the firm trained the workforce about deferred acceptance so that all sides could understand how their rankings would be used. The training emphasized that all rankings would be kept strictly confidential. Both workers and managers were privately told that ranking options according to their true preferences was optimal. To assess how the program was understood, the firm anonymously surveyed workers and managers. During the early rounds of the match, participants were anonymously surveyed about whether their rankings reflected their true preferences. 90% of managers confirmed that it did, and the remaining 10% reported feeling "neutral" about whether the reported preferences were their true preferences.

4.2 Market Restrictions

Although the marketplaces broadly welcomed workers on the market to transfer anywhere, the leadership imposed some limitations on the market. First, the market was segmented into (essentially) two submarkets per quarter based on workers specialization (Prop. 1) as either engineers or non-engineers. Workers in one submarket could not bid on the other without special permission. In practice, a few engineers were permitted to bid on non-engineering jobs but never the other way around.

Second, the firm centrally selected which project managers were allowed to enter the market and match with a worker. The match's original design allowed managers to enter the market more freely. However, the firm's executives quickly realized how many managers wanted additional headcount. Without restrictions, high-priority cash-cow projects

could (in principle) lose headcount as workers abandoned them for smaller projects that were cutting-edge, but speculative (and typically lacked a business plan). This can be seen as narrow alignment gone wrong (hoarding, Prop. 6) – workers and managers concerned only with their personal success draw key resources away from critical projects. The company’s executives thus limited the amount and type of projects in the market.

Finally, the firm developed a user-interface cue to nudge workers towards firm-preferred rankings (described more below). We have no way of measuring how segmentation, limited entry and nudging affected preference rankings and final assignments,²⁹ although we assume they pushed them towards the leadership’s objectives. Without these measures, our performance measures of DA might be even worse.

5 Empirical Strategy: Firm Productivity Estimates

Using data from this setting, we aim to measure several of the key properties in our theoretical section. For example:

- How specialized are workers? How does their productivity compare across different matches?
- Do workers and managers prefer to work on assignments that maximize their own productivity? (narrow alignment, Definition 3, Prop. ?? and Prop. ??)
- Do worker/manager preferences exhibit a high degree of vertical preferences (Prop. 4) and/or cross-sided reciprocity (Sec. 3.8.0.1)?
- How assortative/disassortative are the DA and firm-optimal matches?
- How productive is the self-organized match (via DA) vs the leadership-preferred match (Prop. ??)?
- How does the workforce rank the firm-preferred match?

A key input into these questions is the firm’s productivity measures for matches, $V(a)$ in our theory model (particularly V_{CC} and V_{DA} , the outputs under command-and-control and DA). We now lay out our strategy for measuring these productivities. We begin by

²⁹Even without segmentation, non-engineers might not have ranked any engineering jobs. Had they, managers may not have ever ranked them highly enough to have generated a match. Similarly, small speculative projects might not have pulled labor off of established, important projects.

mentioning some important challenges of this problem. We then describe our strategy conceptually, and finally detail our implementation. The next sections contain our analysis plan and results.

Measurement Challenges. Measuring $V(a)$ involves several empirical challenges, but we mention two in particular. The first is “omitted payoffs.” If a match yields strong performance on CEO objectives easily visible to the researcher, how can researchers be certain it improved overall outcomes? Second, two-sided matching faces an enormous set of potential assignments ($N!$), only one of which will be realized. In our smallest market featuring 35 workers, this number is $35!$ or over 2.6×10^{35} . This makes it hard to know how a counterfactual match would have changed outcomes.

A key limitation is that interventions (such as a different matching technique) cannot be randomized at the individual employee level. Changing one employees’ match generates spillovers onto other employees’ matches, and other managers’ matches. The unit of treatment is the entire workforce.³⁰

These challenges are not unique to our setting. For example, empirical papers by [Abdulkadiroğlu et al. \(2020\)](#); [Abdulkadiroglu et al. \(2021\)](#) are about “match quality” in between schools and families. “Match quality” in education is potentially hard to define, and there are a potentially enormous number of match-specific outputs to measure. The researchers overcome these challenges by developing models of match quality based on teacher/student value-added studies ([Todd and Wolpin, 2003](#); [Koedel et al., 2015](#); [Abdulkadiroğlu et al., 2017](#)) developed on observational data.³¹ Treatment effects on match quality are not directly observed through experiments, but inferred through these models.

5.1 Conceptual Approach

The challenges above may appear in many other settings besides matching. Below we describe a more general problem in which researchers are interested in evaluating K options. After, we describe how we operationalized this in our settings. In short, our ap-

³⁰One could sub-divide the market. However, existing matching theory ([Roth, 2008b](#); [Akbarpour et al., 2020](#)) shows the value of larger, thicker markets. If markets are partitioned for measurement reasons, it may destroy the effect we are trying to capture.

³¹In [Abdulkadiroğlu et al. \(2020\)](#), the researchers are able to validate these estimates by comparing a subset of them to value-added estimates using a lottery.

proach involves using modeled outcomes (similar to described above, [Abdulkadiroğlu et al. 2020](#)). However, instead of using prior observational data to train these models, we obtain incentive-compatible estimates of the leadership’s preferences (via direct elicitation). Our approach proceeds with two steps:

Step 1 (Elicitation). *Elicit an objective function for the CEO’s ranking or cardinal utility weights for the K alternatives.*

A full discussion of elicitation techniques is beyond the scope of this paper, but one could use conjoint methods ([Green and Srinivasan, 1978](#)). Should the K options be too many to label individually, researchers could make and justify assumptions for labeling a subset of K , and using this labeled subset to extrapolate labels for the remainder.

Step 2 (Intervention). *Use the elicited objective function in Step 1 to intervene in the organization to encourage favored outcomes.*

In our design, the exclusive purpose of Step 2 is to make Step 1 incentive-compatible. Separately, researchers may want to study the intervention as a natural experiment and measure its effects, that is not part of the strategy we outline here. Through the intervention in Step 2, executives must “live with their choices” ([Ding et al., 2005](#)) in Step 1, and thus face incentives to report their true objectives. The credibility of the strategy clearly depends partly on the strength of the intervention (and how seriously the CEO took the elicitation). The ideal intervention would not only encourage the CEO’s single most preferred option to be realized, but nudge favorable outcomes throughout the entire distribution. This would give the CEO incentives to Step 1 not only to identify the single most preferred option, but to take seriously the ranking of the other $K - 1$. Our application meets this criteria.

Once elicited, the objective function addresses the earlier measurement challenges. If the CEO has payoffs that are “omitted” from typical data, the CEO can incorporate expectations about these in the Step 1 ratings. The resulting function can evaluate any arbitrary assignment among the $N!$, producing the CEO’s $V(a)$ for that match. This includes the DA match and a distribution of random matches. The objective function can also be used to construct the match that would maximize the firm’s objectives.

The procedure above shares some features with conjoint analysis ([Green and Srinivasan, 1978](#); [Ding, 2007](#)), surrogate outcome methodology ([Athey et al., 2019](#); [Prentice, 1989](#)), and

“incentivized resume rating” (an alternative to resume audit studies, [Kessler et al. 2019](#)). Importantly, we do not claim that CEO’s elicited objectives are “correct.” We make no claim that the objectives are the right one for the firm (or other normative distinctions), only that the elicited objectives the elicited objective approximates (“surrogates”) for the leadership’s objectives *ex-ante*. As such, it is useful for assessing tradeoffs between CEO objectives and other constraints.

5.2 Operationalization

In our applied setting, the leadership spotted a need to develop an objective score like the one described above. They were worried about the exact set of issues outlined in the theory section of this paper: A match based only on workforce preferences could produce highly unproductive teams. For example, participants could form matches to facilitate social connections or aspirational jobs, rather than qualified matches based on skills. The firm wanted to develop an *ex-ante* metric of quality they could use to evaluate matches (once they were realized).

As a result, they developed a scoring metric like the one we described in Step 1 to assign a cardinal weight to any pairwise match between a worker and manager/team. The leadership then deployed this metric for inside the web application in order to “nudge” the workforce (Step 2) towards ranking company-favored matches favorably. While each side of the market browsed the other’s profiles, the web application displayed an icon indicating the company’s assessed match quality. The icons not only displayed each user’s top recommended match, but labeled degrees of company approval for each match using color-coded icons. These are visible in our Figure 4 & 3 screenshot mockups. Later, the leadership evaluated the performance of the match based on this same metric.

To be clear, the creation of the firm’s objective score was initiated and executed by the firm itself, without our suggestion or encouragement. The two-step procedure described above describes why this is a useful approach, what challenges it solves, and how it could be replicated elsewhere. The independent development of the score makes it a better representation of the leadership’s priorities, as it shows the firm took this metric seriously. In addition, implementation choices about the metric were made in service of the goals above.

We use the firm’s objective score as part of our analysis of the broader topics of this

paper. Because these scores were shared with workers and managers, they might have influenced rankings. We have no way to know how much this happened. Workers were free to ignore the recommended rankings. Insofar as the scores affected rankings, our results could be interpreted as a *lower bound* of the true level of misalignment between worker preferences and firm objectives.

Appendix ?? contains some implementation details of the firm’s objective score. The scoring algorithm used several input variables, but the most prominent were job skills and requirements, language requirements, and location. Among the many ways to optimize around these variables, the one selected by the firm may have been chosen because it looked good on *other* dimensions. Although one can question the firm’s choices, we encourage readers to think of these scores as a reasonable proxy of the leadership’s objectives.

Data. The objective scores are the final piece of data necessary to study some of the questions motivating our paper. Our data consist of participant characteristics, preferences, objective scores for each pairwise match, and potential assignments under various matching regimes including the workers-propose DA match used by the firm. Appendix C describes all variables in greater detail. To calculate the firm-optimal match, we use the aforementioned [Kuhn-Munkres](#) algorithm that finds the matches that maximize the score, and we include some other matching algorithms for comparison (including random assignment with equal uniform probabilities). Each algorithm is run 50 times resolving ties randomly, generating a distribution of potential assignments.

6 Empirical Results

Table 1 displays summary statistics for our sample. Our data consists of 318 workers applying to 517 divisions/managers across seven submarkets between 2016 and 2017. The average worker went on the market 1.67 times and submitted around three rankings, though some workers ranked up to 26 managers. Meanwhile, managers submitted just over two worker rankings on average.

From here, our analysis is straightforward. We show details below, but summarize our findings first: Our data about workforce rankings suggest preferences are relatively uncorrelated with where the leadership thinks they are individually most productive (i.e.,

little narrow alignment). Workers and managers also have a low degree of vertical preferences. In this respect, worker/manager preferences resemble the I.I.D. model, but with a high degree of cross-sided reciprocity.

Next, our data on the firm’s objective score suggests that workers and managers and managers exhibit a high degree of specialization, and that the firm’s objective function is more submodular than supermodular. Partly for these reasons, the CEO preferred match experiences high gains from using this specialization, and produces a slightly disassortative match. By contrast, the DA match does not prioritize specialization as much (because participants did not rank their specializations highly). Furthermore, the DA match exhibits a slightly positive assortative matching.

As such, the DA output is only slightly more productive than random matching ($\bar{V}$ in our theory model). However, the DA output is significantly better for workers preferences, while the firm-optimal match is approximately as good as random matching. We show these results in detail below, and then discuss their implications at the end.

6.1 Worker/Manager Preferences

Narrow Alignment. Do individual workers and managers rank options according to where the leadership thinks they are most valuable? We answer this question at least in the narrow sense (Prop. ??) in Panel B of Table 2. We run rank-ordered logits of rankings on the firm’s objective score for each assignment (columns 1 and 4), on the agent’s comparative advantage in each match (columns 2 and 5).³² Columns 1–3 display the results for worker rankings while columns 4–6 display them for manager rankings. The results indicate that worker preferences are weakly aligned with the firm’s objective, as workers are more likely to rank managers with whom they have higher match quality. However, manager choices are unrelated to the firm’s objectives. Moreover, the choices of neither workers nor managers are related to their comparative advantage.

Inherently a bad thing

Vertical Preferences We now analyze the extent to which (i) workers on average prefer the same managers and (ii) managers on average prefer the same workers. Prop ?? shows

³²Appendix ?? describes how we calculated comparative advantage.

that this would be relatively bad for deferred acceptance, insofar as such “vertical” preferences would create same-sided competition, prevent workers from obtaining their top choices, and thus reduce the retention (or effort) benefits of DA.

To measure vertical preferences, we create a dataset containing all possible pairs of workers, and calculate the Spearman’s rank-correlation ρ for each pair of workers’ preferences over the same set of managers. We later repeat this analysis from the manager’s side, examining how correlated pairs of managers’ preferences are over the same set of workers. The ρ s in these analyses could in theory vary from -1 to +1.

For both workers and managers, the average ρ is essentially zero (-0.01 and -0.02, respectively). This means that one worker’s preferences are uninformative about another’s; workers have statistically independent views the same set of managers. Our results suggest workers and managers have highly idiosyncratic, personalized, and distinctive set of preferences that do not come into conflict with other members of the same side. Prop ?? suggests this good for deferred acceptance, insofar as the workforce as a whole can obtain high rankings without competing.

Our results thus far somewhat resemble our I.I.D. model insofar as workforce preferences seem uncorrelated both with each other’s preferences, and with the leadership’s (above, at least in a narrow sense).

Cross-Sided Reciprocity Our results above suggest that workers and managers’ tastes are both highly idiosyncratic and personal. We now turn to cross-sided preference correlations. If a worker (idiosyncratically) likes a manager, does the manager like the worker in return (and vice versa)? We measure this by calculating Spearman’s rank-order ρ for cross-sided preferences.

Our results suggest much stronger positive correlations (compared with our same-sided results), but still overall moderate correlations in cross-sided preferences. Among all preferences, the correlation is 0.25 (of a possible range of -1 to +1). Among pairs who ranked each other, the correlation is much higher (0.64).

6.2 Firm Objectives

Specialization. In Definition 1 of our theory section, we say that a worker is unspecialized if they are equally productive in all matches. We can study this empirically, by examining how much a worker’s objective score changes with the assigned partner. Perhaps unsurprisingly, we can statistically reject the null hypothesis that workers are exactly equally productive in all assignments.

However, our data also show of large economic significance. To measure degrees of specialization, we take the difference between the productivity of each worker’s best match and worst match. For ease of interpretation, we standardize these values using the standard deviation of the overall distribution of match qualities; this lets us measure differences in terms of standard deviations from the average match. Figure ?? contains a histogram of these differences. Sure enough, a small minority ($\approx 10\%$) performs roughly the same in all jobs. However, for the vast majority ($\approx 90\%$), the productivity difference between their best and worst job is over 1.5 standard deviations. The median is 2.16 deviations and about 10% of workers differ by 3 standard deviations between their best and worst match.

Vertical Preferences by the Firm. Although there is wide variability in a worker’s (and managers’) productivity (above), we also find strong worker- and manager- fixed effects. Regressions predicting the firm objective score for a match using workers FEs alone have an adjusted R^2 of 0.67. The adjusted R^2 is 0.49 for manager FEs alone and is 0.87 using both FEs together. The distribution of worker and manager fixed effects have the same means, but worker fixed effects are more variable.³³ These numbers suggest the firm has a much stronger notion of “vertical preferences” over workers and managers, even if workers/managers do not have strong vertical preferences over each other. Some workers and managers have high (or low) output on average across all possible match partners.

Complementarity and Production Technology. What the CEO should do to maximize with high/low vertically ranked workers and managers is theoretically ambiguous, and depends on the degree of complementarity in the match production function (Becker, 1973). For supermodular production functions, the most productive match assortatively

³³See Appendix Section E.3 for histograms of the distribution of worker and manager fixed effects).

pairs the best workers with the best managers. For submodular production functions, the most productive match pairs the best manager and the worst manager.

The firm objectives score allows us to measure this directly. Before directly measuring it, we mention some key features of the objective function. These features provide the intuition for our results. First, the firm's objective score algorithm was programmed to have a maximum and a minimum. This will shape some of our results. Because of the upper bound, the performance of the whole will often be "less than the sum of the parts" because performance has an upper bound that cannot be exceeded, no matter how great the parts are. Second, the upper and lower bounds are binding. 28% of all possible pairwise matches are at the maximum, and less than 1% are at the minimum. For the median worker, only 5% of his or her options are at the max, however, 69% of workers have at least one option that will place them at the maximum.³⁴

To analyze the degree of complementarity, we take two approaches. First we analyze all 23K possible matches, using the fixed effects for workers and managers calculated above. For each match, we calculate the difference between the fixed effects and take the absolute value. This measures the average quality distance between the worker and manager in each match. We then use this difference to predict the firm objective score (along with worker and manager fixed effects). A positive coefficient on this variable means that greater differences are better for a match's productivity (conditional on the fixed effects).

Second, we calculate the firm optimal match using the [Kuhn-Munkres](#) algorithm, and evaluate its assortativeness. For comparison, we also calculate several others [X, Y, and Z] and calculate their assortativeness for comparison.

When we analyze our firm's objective function in Panel A of Table 5, we find that match quality is slightly higher under negative associative matching. Column 1 shows that a one standard deviation increase in the interaction between the worker's average productivity (across all matches) and the manager's decreases the firm's objective by 0.21 standard deviations ($p=0.001$). Meanwhile, column 2 shows that the firm's objective score increases as the difference in average quality between workers and managers increases. These results indicate that negative assortative matching is optimal from the firm's perspective in our setting.

³⁴These numbers are 13% and 88% for managers.

We do this in Panel B of Table ??, and find that the.

6.3 DA vs Firm Optimal Matches

Output

Firm Objectives

Assortative Matching

Discussion

7 OLD

7.1 Implementation Details and Limitations

Earlier, we noted that the credibility of our measurement strategy depends on the strength of incentives to report truthfully (and how seriously the firm took the development of the objective score overall). Our applied setting shows how these criteria are fulfilled. Developing the objective scoring algorithm was costly both from design and implementation standpoints. The prominent integration of these scores into the application makes them incentive compatible for the firm. These scores are not *ex-post* cheap-talk labels of match quality. They were strategically designed *ex-ante* to push outcomes towards the firm's objectives.

Objective scores were strategically designed to quantify the firm's objectives and push participants towards them through user-interface engineering. The firm had a wide latitude to develop scores based on any variable or objective, and chose to implement this score after careful deliberation and exploration of alternatives. The developers of this system sampled how various proposed scoring systems would assess samples of potential matches. They fine-tuned the scoring system until it scored outcomes according to the leadership's priorities.

In Appendix B.1, we discuss additional details of this system which we summarize here. Several variables went into the firm objective score, but the most prominent were job skills and requirements, language requirements, and location.

In Appendix B.2, we offer a few caveats to our characterization of the score as the firm’s “objective.” In short: We acknowledge that a firm might want to accomplish more in its “optimal match” than optimizing skills, language and location. However, among the many ways to optimize around these variables, the one selected by the firm may have been chosen because it looked good on *other* dimensions. We encourage readers to think of our “firm objective” variable as a reasonable (if imperfect) proxy of the firm’s objectives.

Finally, one may wonder whether the firm undertook additional measures to influence the match that are outside our data. Our interviews with the company’s leadership and workers suggest nothing like this happened. We discuss this possibility further in Appendix B.3.

Our results about organizational productivity arise from the preferences of workers, managers and the firm. As such, we begin by characterizing these preferences in isolation.

7.2 Tradeoffs: Firm Objectives vs Worker/Management Preferences

We now turn to the central empirical question: How much does the firm sacrifice its objectives when implementing deferred acceptance (compared to the firm-optimal assignment)? Table 7 compares the distribution of firm objective scores obtained by each matching approach algorithm (the Firm optimal, the Manager Draft, the Workers’ Draft, the Managers-Propose DA, the Workers-Propose DA, the Random Assignment and the Firm objective Minimizer).

Panel A contains all matches (including unfilled jobs, which receive a firm objective score of 0), while Panel B limits the sample to only successfully-filled worker-job matches. The results in Table 7 indicate that deferred acceptance depresses match quality compared to the optimal solution. This is true not only on average but also throughout the distribution of matches. In Appendix E.1, we show that the firm optimal solution first-order stochastically dominates deferred acceptance.

Table 4 confirms these results using Equation ???. Deferred acceptance matches scored 15 percentage points lower on the firm objective measure than optimal matches. This dif-

ference is large: using deferred acceptance leads the firm to sacrifice $\approx 24\%$ ($= 0.15/0.62$) of the principal's goals. This magnitude holds regardless of which fixed effects are used, and are similar whether workers or managers propose first (columns 4–6). Moreover, Panel B of Table 5 indicates that through negative assortative matching. Compared to deferred acceptance, the optimal solution is more likely to assign high quality workers with lower quality managers, and vice versa. Meanwhile, random matching does slightly worse than matching through deferred acceptance, and this difference is marginally significant at $p < 0.10$.

We can use the manager and worker fixed effects as a numeraire to measure the size of our effects relative to the quality of workers and managers. Rather than use the optimal match, the firm could try to achieve the same productivity benefit by increasing the quality of hires. However, our results indicate doing so would be difficult. Optimal matching results in a 24% improvement in the firm's objective score over deferred acceptance. However, only 25% of workers and 16% of managers have fixed effects larger than this amount.. Increasing productivity by 24% is about 55% of a standard deviation in the worker fixed effects and 66% of a standard deviation in manager fixed effect, suggesting that achieving the same productivity benefits through hiring would require a large increase in new hire quality.

If the firm were to use the optimal matches instead of those generated by deferred acceptance, how would workers and managers fare? Table 6 examines how optimal matches perform on various measures of worker and manager preferences using Equation ???. Our dependent variables include: (i) a binary indicator whether the worker (manager) was paired with a manager (worker) they did not rank; (ii) the worker's (manager's) ranking of their matched manager (worker); (iii) the worker's (manager's) ranking of their matched manager (worker) with ties randomly resolved through simulation; and (iv) a utility outcome generated from a rank-order logit as described in Section ???. Panel A displays worker losses while Panel B displays manager losses.

The results indicate that matches from the optimal solution perform worse on measures of participant preferences. Workers and managers are much more likely to be matched with a participant they did not rank under the firm-optimal solution. Under deferred acceptance, workers and managers get, on average, their 1.2th and 1.3rd picks, respectively. However, under the optimal solution, workers receive their 3.9th ($= 1.214 + 2.64$) pick and managers get their 3.6th pick on average. Our rank-ordered logit measure indicates

that this causes workers and managers to lose 0.61 and 0.14 standard deviations in utility, respectively. In order for the firm to maximize its objectives, both workers and managers would need to be matched with lower-ranked managers and workers.

8 Conclusion

Since the successful redesign of the National Residency Matching Program ([Roth and Peranson, 1999](#)) the tools of market design have increasingly been adopted as solutions to real world matching problems. In our empirical results, we find optimal matching is incredibly productive for organizations. Using the company’s preferred measure of productivity, the optimal match is 36% more productive than randomly assigned matches within job categories. These results are driven in part by *negative* assortative matching: The firm’s best managers are not necessarily matched with the firm’s best workers.

However, the performance of the deferred acceptance algorithm is much worse. It only narrowly performs better than random matching, in part by producing a sub-optimal match featuring positive assortative matching. In this paper, we argue that results like these arise because that matching problems within an organization are distinct from traditional applications. In particular, organizations have flexibility in choosing an assignment mechanism that best achieves its objectives and may benefit from disregarding agents’ preferences when picking assignments.

As a result, assignments resulting from mechanisms that respect agents’ preferences do not always best serve the organization’s objective, even in situations that seem favorable to the decentralized approach. Although our theoretical results stem from relatively-simple counter-examples, they are not limited to these hypothetical examples. Our data from a real firm that introduced preference-based matching for workers and divisions to illustrate that these tensions have large implications in the real world.

8.0.0.1 Additional Acknowledgements. The authors also thank seminar participants at MIT, Rochester, Carnegie Mellon, University of Minnesota, Stanford (CASBS), University of Hong Kong, and Columbia Business School, as well as participants at the Wharton People and Organizations Conference, the Marketplace Innovations Workshop, the Academy of Management Annual Conference, the Workshop for Information Systems and Economics (WISE), the Strategic Management Society (SMS) Annual Conference, and the 2021 Causal Data Science Meeting.

References

- Abdulkadiroglu, Atila and Tayfun Sönmez**, “Matching Markets: Theory and Practice,” in “Advances in Economics and Econometrics: 10th World Congress, Volume 1: Economic Theory,” New York: Cambridge University Press, 2013.
- Abdulkadiroğlu, Atila, Joshua D Angrist, Yusuke Narita, and Parag A Pathak**, “Research design meets market design: Using centralized assignment for impact evaluation,” *Econometrica*, 2017, 85 (5), 1373–1432.
- , **Parag A Pathak, Jonathan Schellenberg, and Christopher R Walters**, “Do parents value school effectiveness?,” *American Economic Review*, 2020, 110 (5), 1502–39.
- Abdulkadiroglu, Atila, Umut M Dur, and Aram Grigoryan**, “School Assignment by Match Quality,” Technical Report, National Bureau of Economic Research 2021.
- Adhvaryu, Achyuta, Vittorio Bassi, Anant Nyshadham, and Jorge A Tamayo**, “No line left behind: Assortative matching inside the firm,” Technical Report, National Bureau of Economic Research 2020.
- Agarwal, Nikhil**, “An Empirical Model of the Medical Match,” *American Economic Review*, 2015, 105 (7), 1939–1978.
- **and William Diamond**, “Latent indices in assortative matching models,” *Quantitative Economics*, 2017, 8, 685–728.
- Aguiar, Mark, Mark Bils, Kerwin Kofi Charles, and Erik Hurst**, “Leisure luxuries and the labor supply of young men,” *Journal of Political Economy*, 2021, 129 (2), 337–382.
- Akbarpour, Mohammad, Shengwu Li, and Shayan Oveis Gharan**, “Thickness and information in dynamic matching markets,” *Journal of Political Economy*, 2020, 128 (3), 783–815.
- Alonso, Ricardo, Wouter Dessein, and Niko Matouschek**, “When Does Coordination Require Centralization?,” *The American Economic Review*, 2008, 98 (1), 145–179.
- Altig, Dave, Scott Baker, Jose Maria Barrero, Nicholas Bloom, Philip Bunn, Scarlet Chen, Steven J Davis, Julia Leather, Brent Meyer, Emil Mihaylov et al.**, “Economic

- uncertainty before and during the COVID-19 pandemic," *Journal of Public Economics*, 2020, 191, 104274.
- Arya, Anil and Brian Mittendorf**, "Project assignments when budget padding taints resource allocation," *Management science*, 2006, 52 (9), 1345–1358.
- , **Jonathan Glover, and Bryan R Routledge**, "Project assignment rights and incentives for eliciting ideas," *Management Science*, 2002, 48 (7), 886–899.
- Ashlagi, Itai, Yash Kanoria, and Jacob D Leshno**, "Unbalanced random matching markets: The stark effect of competition," *Journal of Political Economy*, 2017, 125 (1), 69–98.
- Athey, Susan, Raj Chetty, Guido W Imbens, and Hyunseung Kang**, "The surrogate index: Combining short-term proxies to estimate long-term treatment effects more rapidly and precisely," Technical Report, National Bureau of Economic Research 2019.
- Autor, David H, Frank Levy, and Richard J Murnane**, "The skill content of recent technological change: An empirical exploration," *The Quarterly journal of economics*, 2003, 118 (4), 1279–1333.
- Azevedo, Eduardo M and Eric Budish**, "Strategy-proofness in the large," *The Review of Economic Studies*, 2019, 86 (1), 81–116.
- Baker, George and Bengt Holmstrom**, "Internal labor markets: Too many theories, too few facts," *The American Economic Review*, 1995, 85 (2), 255–259.
- , **Michael Gibbs, and Bengt Holmstrom**, "The internal economics of the firm: Evidence from personnel data," *The Quarterly Journal of Economics*, 1994, 109 (4), 881–919.
- Bandiera, Oriana, Iwan Barankay, and Imran Rasul**, "Incentives for managers and inequality among workers: Evidence from a firm-level experiment," *The Quarterly Journal of Economics*, 2007, 122 (2), 729–773.
- Barron, Greg and Felix Vardy**, "The Internal Job Market of the IMF ' s Economist Program," *IMF Staff Papers*, 2005, 52 (3), 410–429.
- Bartling, Björn, Ernst Fehr, and Holger Herz**, "The intrinsic value of decision rights," *Econometrica*, 2014, 82 (6), 2005–2039.
- Becker, Gary S**, "A theory of marriage: Part I," *Journal of Political economy*, 1973, 81 (4), 813–846.

- Benson, Alan and Ben A Rissing**, “Strength from within: Internal mobility and the retention of high performers,” *Organization Science*, 2020, 31 (6), 1475–1496.
- Bidwell, Matthew**, “Paying more to get less: The effects of external hiring versus internal mobility,” *Administrative Science Quarterly*, 2011, 56 (3), 369–407.
- Bloom, Nicholas and John Van Reenen**, “Human resource management and productivity,” in “Handbook of labor economics,” Vol. 4, Elsevier, 2011, pp. 1697–1767.
- Boudreau, John W**, “50th Anniversary Article: Organizational behavior, strategy, performance, and design in management science,” *Management Science*, 2004, 50 (11), 1463–1476.
- Bresnahan, Timothy F**, “Computerisation and wage dispersion: an analytical reinterpretation,” *The economic journal*, 1999, 109 (456), 390–415.
- Brockner, Joel**, “Understanding the interaction between procedural and distributive justice: The role of trust,” 1996.
- Bryan, Lowell L and Claudia I Joyce**, *Mobilizing minds: creating wealth from talent in the 21st-century organization*, McGraw-Hill, 2007.
- Brynjolfsson, Erik and Lorin M Hitt**, “Information technology and organizational design: evidence from micro data,” *Manuscript, MIT*, 1998.
- Cappelli, Peter and JR Keller**, “Talent management: Conceptual approaches and practical challenges,” *Annu. Rev. Organ. Psychol. Organ. Behav.*, 2014, 1 (1), 305–331.
- Caroli, Eve**, “New technologies, organizational change and the skill bias: what do we know,” *Technology and the future of European employment*, 2001, pp. 259–292.
- **and John Van Reenen**, “Skill-biased organizational change? Evidence from a panel of British and French establishments,” *The Quarterly Journal of Economics*, 2001, 116 (4), 1449–1492.
- Christensen, Michael and Thorbjørn Knudsen**, “Design of decision-making organizations,” *Management Science*, 2010, 56 (1), 71–89.
- Coase, R.H.**, “The Nature of the Firm,” *Economica*, 1937, 4 (16), 386–405.
- Cowgill, Bo and Rembrand Koning**, “Matching Markets for Googlers,” 2018.

- , **Zoë Cullen, Anh Nguyen, and Ricardo Perez-Truglia**, “Gender Segregation in the Workplace,” *Working paper*, 2020.
- Cropanzano, Russell Ed**, *Justice in the workplace: Approaching fairness in human resource management.*, Lawrence Erlbaum Associates, Inc, 1993.
- Davis, Jonathan, Damon Jones, and Kyle Greenberg**, “An Experimental Evaluation of a Matching Market Mechanism: Evidence from Army Officer Labor Markets,” *AEA Registry*, 2020.
- Davis, Jonathan M.V.**, “Deferred Acceptance Mechanisms Can Improve Match Quality: Quasi-Experimental Evidence from a Teach for America Pilot,” 2016.
- Deming, David J**, “The growing importance of social skills in the labor market,” *The Quarterly Journal of Economics*, 2017, 132 (4), 1593–1640.
- , “The Growing Importance of Decision-Making on the Job,” Technical Report, National Bureau of Economic Research 2021.
- Dessein, Wouter and Tano Santos**, “Adaptive organizations,” *Journal of political Economy*, 2006, 114 (5), 956–995.
- Ding, Min**, “An incentive-aligned mechanism for conjoint analysis,” *Journal of Marketing Research*, 2007, 44 (2), 214–223.
- , **Rajdeep Grewal, and John Liechty**, “Incentive-aligned conjoint analysis,” *Journal of marketing research*, 2005, 42 (1), 67–82.
- Erickson, R, D Moulton, and B Cleary**, “Are you overlooking your greatest source of talent,” *Deloitte Review*, (23), 2018, pp. 42–51.
- Fazio, Catherine E, Jorge Guzman, Yupeng Liu, and Scott Stern**, “How is COVID Changing the Geography of Entrepreneurship? Evidence from the Startup Cartography Project,” Technical Report, National Bureau of Economic Research 2021.
- for Economic Co-operation, Organisation and Development Staff**, *OECD Employment Outlook, June 1999*, OECD, 1999.
- Gale, David and Lloyd S . Shapley**, “College Admissions and the Stability of Marriage,” *The American Mathematical Monthly*, 1962, 69 (1), 9–15.

- Gardenfors, Peter**, "Assignment problem based on ordinal preferences," *Management Science*, 1973, 20 (3), 331–340.
- Green, Paul E and Venkatachary Srinivasan**, "Conjoint analysis in consumer research: issues and outlook," *Journal of consumer research*, 1978, 5 (2), 103–123.
- Greenberg, Kyle, Mark Crow, and Carl Wojtaszek**, "Winning in the Marketplace: How Officers and Units Can Get the Most Out of the Army Talent Alignment Process," *Modern War Institute*, 2020.
- Groves, Theodore and Martin Loeb**, "Incentives in a divisionalized firm," *Management Science*, 1979, 25 (3), 221–230.
- Haeghele, Ingrid**, "Talent Hoarding in Organizations," 2021.
- Hall, Jonathan V and Alan B Krueger**, "An analysis of the labor market for Uber's driver-partners in the United States," *Ilr Review*, 2018, 71 (3), 705–732.
- Haltiwanger, John C**, "Entrepreneurship During the COVID-19 Pandemic: Evidence from the Business Formation Statistics," Technical Report, National Bureau of Economic Research 2021.
- Hamel, Gary**, "First, let's fire all the managers," *Harvard Business Review*, 2011, 89 (12), 48–60.
- Hammer, Michael and James Champy**, *Reengineering the Corporation: Manifesto for Business Revolution*, A, Zondervan, 2009.
- Harris, Milton, Charles H Kriebel, and Artur Raviv**, "Asymmetric information, incentives and intrafirm resource allocation," *Management Science*, 1982, 28 (6), 604–620.
- Holmström, Bengt**, "Moral hazard and observability," *The Bell journal of economics*, 1979, pp. 74–91.
- Holmstrom, Bengt**, "Moral hazard in teams," *The Bell Journal of Economics*, 1982, pp. 324–340.
- , "On the theory of delegation," in Marcel Boyer Kihlstrom and Richard E., eds., *Bayesian Models in Economic Theory*, Elsevier Science Ltd, 1984, pp. 115–141.

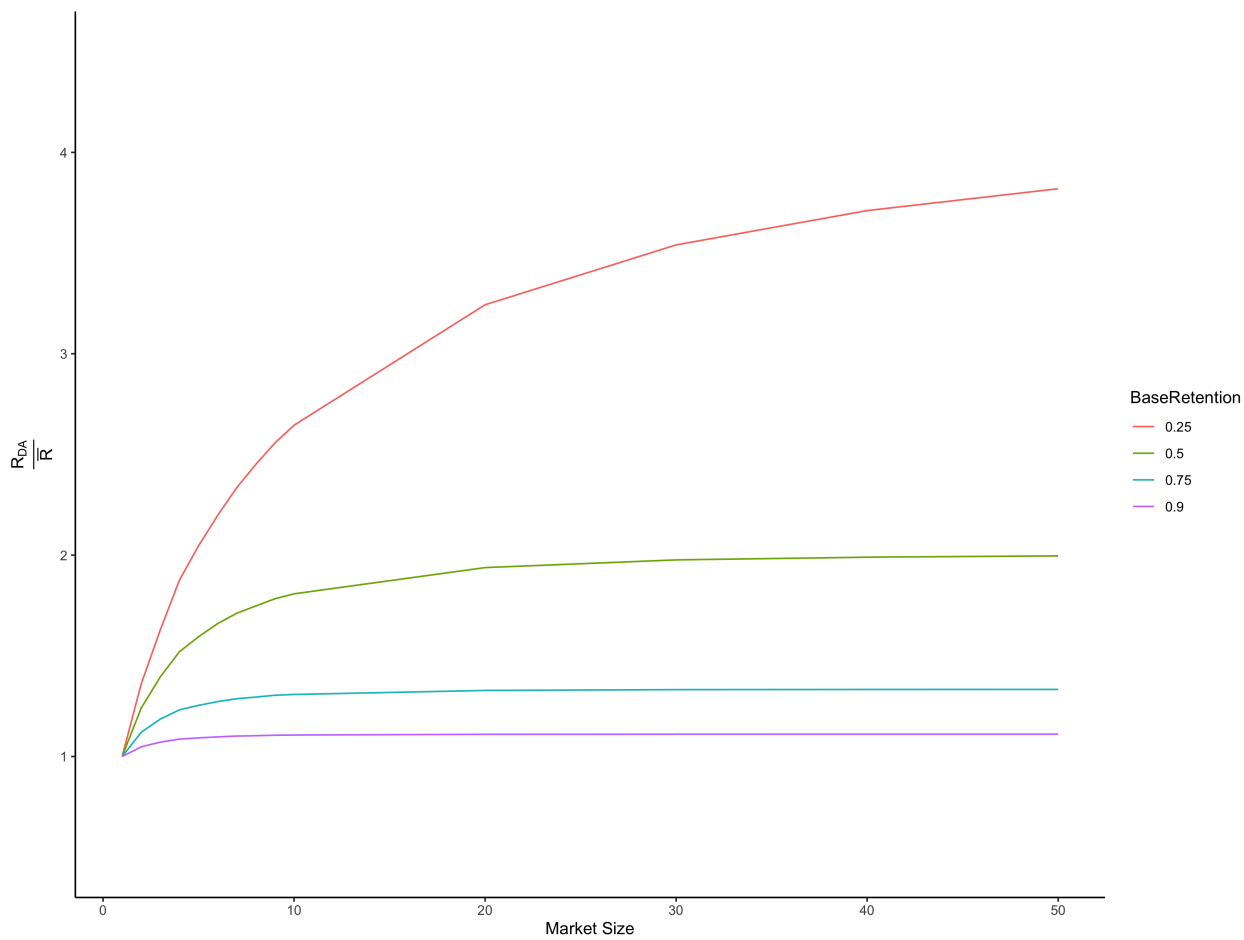
- Hopenhayn, Hugo, Julian Neira, and Rish Singhania**, “From population growth to firm demographics: Implications for concentration, entrepreneurship and the labor share,” Technical Report, National Bureau of Economic Research 2018.
- Ichniowski, Casey, Thomas A Kochan, David Levine, Craig Olson, and George Strauss**, “What works at work: Overview and assessment,” *Industrial Relations: A Journal of Economy and Society*, 1996, 35 (3), 299–333.
- Immorlica, Nicole and Mohammad Mahdian**, “Marriage, Honesty, and Stability,” in “Proceedings of the Sixteenth Annual ACM-SIAM Symposium on Discrete Algorithms” SODA ’05 Society for Industrial and Applied Mathematics USA 2005, p. 53–62.
- Kambhampati, Ashwin and Carlos Segura-Rodriguez**, “The Optimal Assortativity of Teams Inside the Firm,” *Available at SSRN 3531123*, 2020.
- Kessler, Judd B, Corinne Low, and Colin D Sullivan**, “Incentivized resume rating: Eliciting employer preferences without deception,” *American Economic Review*, 2019, 109 (11), 3713–44.
- Knuth, Donald E, Rajeev Motwani, and Boris Pittel**, “Stable husbands,” *Random Structures & Algorithms*, 1990, 1 (1), 1–14.
- Knuth, Donald Ervin**, *Mariages stables et leurs relations avec d’autres problemes combinatoires: introduction a l’analyse mathematique des algorithmes-*, Les Presses de l’Université de Montréal, 1976.
- , *Stable marriage and its relation to other combinatorial problems: An introduction to the mathematical analysis of algorithms*, Vol. 10, American Mathematical Soc., 1997.
- Koedel, Cory, Kata Mihaly, and Jonah E Rockoff**, “Value-added modeling: A review,” *Economics of Education Review*, 2015, 47, 180–195.
- Kojima, Fuhito and Parag A Pathak**, “Incentives and stability in large two-sided matching markets,” *American Economic Review*, 2009, 99 (3), 608–27.
- Kouvelis, Panagiotis and Martin A Lariviere**, “Decentralizing cross-functional decisions: Coordination through internal markets,” *Management Science*, 2000, 46 (8), 1049–1058.
- Kuhn, Harold W**, “The Hungarian method for the assignment problem,” *Naval research logistics quarterly*, 1955, 2 (1-2), 83–97.

- Lazear, Edward P, Kathryn L Shaw, and Christopher T Stanton**, "The value of bosses," *Journal of Labor Economics*, 2015, 33 (4), 823–861.
- Lindbeck, Assar and Dennis J Snower**, "Multitask learning and the reorganization of work: From tayloristic to holistic organization," *Journal of labor economics*, 2000, 18 (3), 353–376.
- Lovejoy, William S**, "Optimal mechanisms with finite agent types," *Management Science*, 2006, 52 (5), 788–803.
- Martela, Frank**, "What makes self-managing organizations novel? Comparing how Weberian bureaucracy, Mintzberg's adhocracy, and self-organizing solve six fundamental problems of organizing," *Journal of Organization Design*, 2019, 8 (1), 1–23.
- Ortega, Jaime**, "Job rotation as a learning mechanism," *Management Science*, 2001, 47 (10), 1361–1370.
- Osterman, Paul**, "How common is workplace transformation and who adopts it?," *Ilr Review*, 1994, 47 (2), 173–188.
- Oyer, Paul**, "Why do firms use incentives that have no incentive effects?," *The Journal of Finance*, 2004, 59 (4), 1619–1650.
- Pittel, Boris**, "The average number of stable matchings," *SIAM Journal on Discrete Mathematics*, 1989, 2 (4), 530–549.
- , "On Likely Solutions of a Stable Marriage Problem," *The Annals of Applied Probability*, 1992, May, 358–401.
- , "On likely solutions of a stable matching problem," in "Proceedings of the third annual ACM-SIAM symposium on Discrete algorithms" 1992, pp. 10–15.
- Prentice, Ross L**, "Surrogate endpoints in clinical trials: definition and operational criteria," *Statistics in medicine*, 1989, 8 (4), 431–440.
- Rees-Jones, Alex**, "Mistaken play in the deferred acceptance algorithm: Implications for positive assortative matching," *American Economic Review*, 2017, 107 (5), 225–29.
- , "Suboptimal behavior in strategy-proof mechanisms: Evidence from the residency match," *Games and Economic Behavior*, 2018, 108, 317–330.

- Roth, Alvin E.**, “Deferred acceptance algorithms: history, theory, practice, and open questions,” *International Journal of Game Theory*, jan 2008, 36 (3-4), 537–569.
- Roth, Alvin E.**, “What have we learned from market design?,” *Innovations: Technology, Governance, Globalization*, 2008, 3 (1), 119–147.
- Roth, Alvin E.**, *Who Gets What – and Why*, Mariner Books, 2015.
- Roth, Alvin E and Elliott Peranson**, “The Redesign of the Matching Market for American Physicians: Some Engineering Aspects of Economic Design,” *The American Economic Review*, 1999, 89 (4), 748–780.
- Roth, Alvin E. and Marilda Sotomayor**, *Two-Sided Matching: A Study in Game-Theoretic Modeling and Analysis*, Cambridge, UK: Cambridge University Press, 1990.
- Roth, Alvin E, Uriel G Rothblum, and John H Vande Vate**, “Stable Matchings, Optimal Assignments, and Linear Programming,” *Mathematics of Operations Research*, 1993, 18 (4), 803–828.
- Shimer, Robert and Lones Smith**, “Assortative matching and search,” *Econometrica*, 2000, 68 (2), 343–369.
- Sonmez, Tayfun**, “Bidding for Army Career Specialties : Improving the ROTC Branching Mechanism,” *Journal of Political Economy*, 2013, 121 (1), 186–219.
- **and Tobias B Switzer**, “Matching With (Branch-of-Choice) Contracts at the United States Military Academy,” *Econometrica*, 2013, 81 (2), 451–488.
- Todd, Petra E and Kenneth I Wolpin**, “On the specification and estimation of the production function for cognitive achievement,” *The Economic Journal*, 2003, 113 (485), F3–F33.
- Useem, Jerry**, “At work, expertise is falling out of favor,” *The Atlantic*, 2019, pp. 56–65.

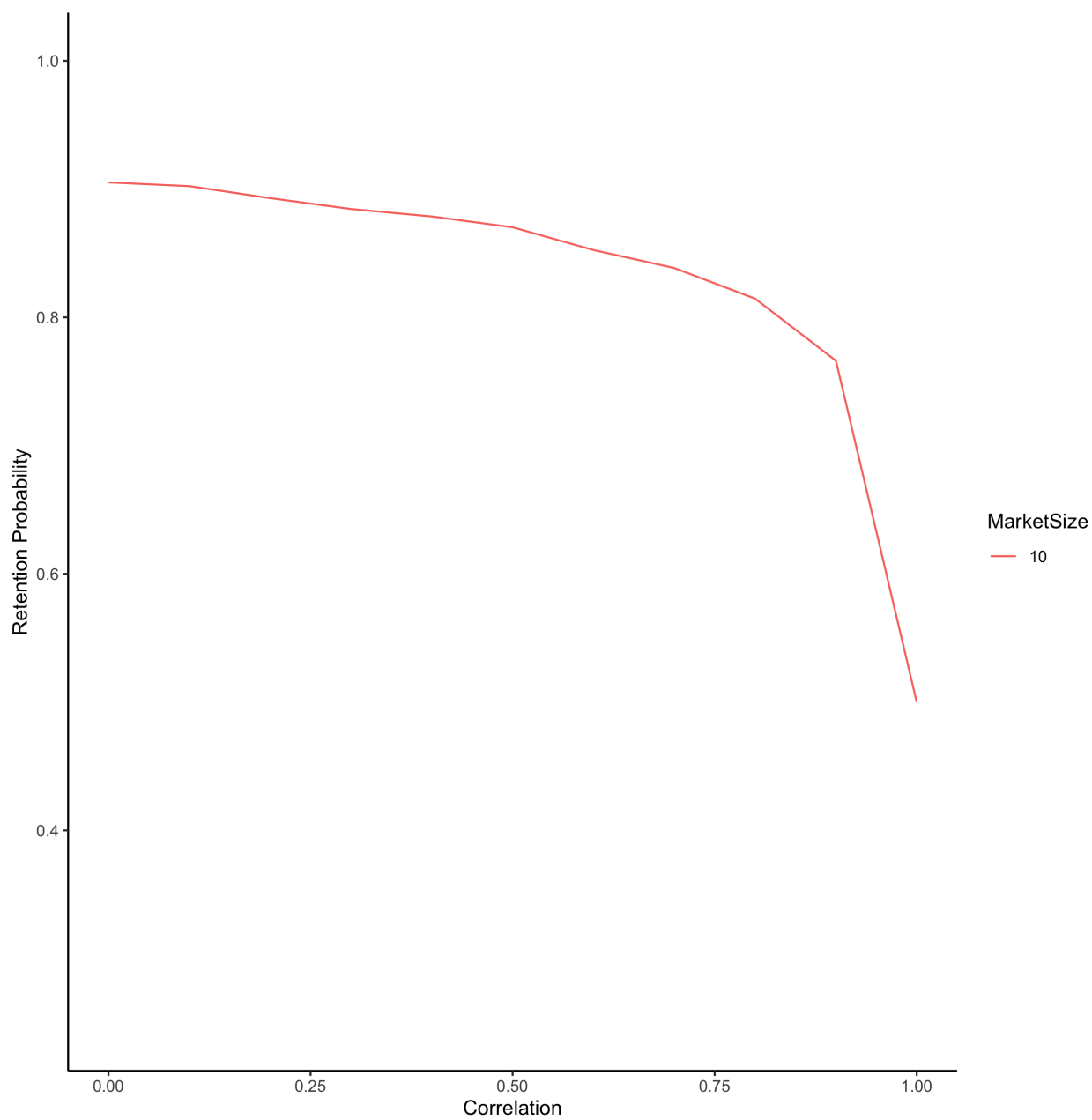
Tables and Figures

Figure 1: DA Retention Benefits and Market Size



Notes.

Figure 2: DA Retention Benefits with Preference Correlation



Notes.

Table 1: **Summary statistics**

	Workers N=318	Managers N=517
	(1)	(2)
Choices ranked		
Mean	2.66	2.22
Min	0	0
Max	26	11
# of times on the market		
Mean	1.67	1
Min	1	1
Max	4	1
Firm objective score		
Mean	0.58	0.54
Min	0	0
Max	1	1
Comparative advantage		
Mean	0.02	0.02
Min	0	0
Max	1	1

Notes: This table examines summary statistics for our data.

Table 2: **Preference data analysis**

Panel A: Correlation in preferences

	Mean	Std.Dev	Min	P25	Median	P75	Max
Worker preferences	-0.01	0.08	-0.20	-0.02	-0.02	-0.02	1.00
Manager preferences	-0.02	0.05	-0.22	-0.03	-0.02	-0.02	1.00
Worker and manager preferences							
All pairs	0.25						
Mutually-ranked pairs	0.64						

Panel B: Preference alignment

	Worker ranking of manager		Manager ranking of worker	
Firm objective score	0.017*	0.017*	0.0024	0.0027
	(0.0093)	(0.0093)	(0.0060)	(0.0060)
Comparative advantage, worker		-0.015		
		(0.015)		
Comparative advantage, manager			-0.0066	-0.0074
			(0.017)	(0.017)
Pseudo-R2	0.000	0.000	0.000	0.000
Observations	23361	23361	23361	23361

Notes: This table examines our preference data. Panel A displays the distribution of pair-wise Spearman correlations. It calculates the correlation of (i) worker preferences over managers, (ii) manager preferences over worker, and (iii) worker and manager preferences for each other. Panel B examines the alignment between worker/manager preferences and the firm's objectives. Columns 1–3 display the results of a rank-ordered logit of each worker's ranking on the firm's objective score (column 1), their comparative advantage (column 2), and both (column 3). Columns 4–6 display the results of a rank-ordered logit of each manager's ranking on the firm's objective score (column 4), their comparative advantage (column 5), and both (column 6). Columns 1–3 include worker fixed effects and robust standard errors clustered at the worker level, while and columns 4–6 include manager fixed effects with robust standard errors clustered at the manager level.

*** $p < 0.01$ ** $p < 0.05$ * $p < 0.10$

Table 3: Firm objective score by matching algorithm

	Mean	Std.Dev	Min	P25	Median	P75	Max
Firm Optimal	0.75	0.22	0.21	0.57	0.76	1.00	1.00
Manager Draft	0.58	0.31	0.00	0.32	0.53	1.00	1.00
Worker Draft	0.57	0.30	0.00	0.33	0.53	0.98	1.00
Managers-Propose DA	0.57	0.31	0.00	0.32	0.53	0.99	1.00
Workers-Propose DA	0.57	0.31	0.00	0.32	0.53	0.99	1.00
Random Assignment	0.55	0.32	0.01	0.29	0.48	0.99	1.00
Firm Objective (minimizer)	0.34	0.32	0.00	0.12	0.22	0.36	1.00

Notes: This table displays the average firm objective score for matches generated by various algorithms. Because there were more managers on the market than workers, some managers are unmatched. This table includes only complete manager worker/pairs (i.e., dropping managers who were not paired to workers). All matching algorithms left the same number of managers unmatched, but differed in their composition. A similar table including unmatched managers (containing a score of zero) is accessible in Table E1.

Table 4: Effects on Firm Objectives

	Firm objective score					
Deferred Acceptance	-0.15*** (0.022)	-0.15*** (0.022)	-0.15*** (0.022)			
Random Assignment	-0.17*** (0.026)	-0.17*** (0.026)	-0.17*** (0.026)	-0.17*** (0.026)	-0.17*** (0.026)	-0.17*** (0.026)
Managers-Propose DA				-0.15*** (0.022)	-0.15*** (0.022)	-0.15*** (0.022)
Workers-Propose DA				-0.15*** (0.022)	-0.15*** (0.022)	-0.15*** (0.022)
R^2	0.819	0.448	0.919	0.819	0.448	0.919
Observations	129250	129250	129250	129250	129250	129250
Fixed effect	Worker	Manager	Both	Worker	Manager	Both
P-values:						
DA = Random	0.059	0.059	0.059	0.002	0.002	0.002
DA (MP) = Random				0.059	0.058	0.059
DA (WP) = Random				0.059	0.059	0.059

Notes: This table displays the results of a regression of the firm objective score on indicators for the match being generated through deferred acceptance or random matching. The excluded category is the optimal match. Each regression includes robust standard errors clustered at the simulation-submarket level. The bottom of the table displays p-values that test the difference in coefficients on deferred acceptance versus random assignment. The average firm objective score from the optimal matches is 0.62.

*** $p < 0.01$ ** $p < 0.05$ * $p < 0.10$

Table 5: **Assortative Matching**

Panel A: All possible worker-manager matches

	Firm Objective Score (Norm)	Firm Objective Score (Norm)
Interaction	-0.21*** (0.0057)	
Abs[Worker - Manager Quality]		0.22*** (0.0066)
Worker Fixed Effects	Y	Y
Manager Fixed Effects	Y	Y
R^2	0.89	0.89
Observations	23361	23361

Panel B: Difference in worker-manager quality (absolute value), by matching algorithm

	Mean	Std.Dev	Min	P25	Median	P75	Max
Firm Optimal	0.75	0.73	0.00	0.27	0.47	1.05	3.31
Random Assignment	0.69	0.61	0.00	0.28	0.48	0.96	2.98
Manager Draft	0.68	0.62	0.00	0.27	0.44	0.99	2.93
Worker Draft	0.66	0.60	0.00	0.27	0.43	0.97	2.87
Workers-Propose DA	0.66	0.60	0.00	0.27	0.44	0.95	2.87
Managers-Propose DA	0.66	0.60	0.00	0.26	0.44	0.95	2.87
Firm Objective (minimizer)	0.65	0.62	0.00	0.27	0.43	0.89	3.25

Notes: This table investigates assortative matching in our sample. Panel A uses all possible worker-manager matches. Column 1 regresses the firm objective score on the interaction between worker and manager quality, while column 2 regresses the firm objective score on the absolute value of the difference between worker and manager quality. Worker and manager quality are normalized to have mean 0 and standard deviation one. Both columns include fixed effects for workers and managers. Panel B examines the average absolute value of the difference between worker and manager quality generated by various matching algorithms.

Table 6: **Worker and manager losses from firm optimal match**

Panel A: Worker losses

	Manager unranked by worker	Worker ranking of manager	Simulated ranking of manager	Worker utility, z-score
Firm optimal	0.85*** (0.027)	2.22*** (0.40)	25.6*** (2.68)	-0.61*** (0.065)
R^2	0.792	0.555	0.589	0.520
Observations	12840	12840	12840	12720
DA mean	0.063	1.173	1.403	0.770
Fixed effect	Worker	Worker	Worker	Worker

Panel B: Manager losses

	Worker unranked by manager	Manager ranking of worker	Simulated ranking of worker	Manager utility, z-score
Firm optimal	0.70*** (0.043)	0.73** (0.24)	16.7*** (2.63)	-0.14* (0.067)
R^2	0.768	0.423	0.564	0.458
Observations	15510	15510	15510	15300
DA mean	0.065	8.221	9.273	-0.660
Fixed effect	Manager	Manager	Manager	Manager

Notes: This table examines the effect of using the optimal matching system on worker and manager preferences. Panel A displays the results of regression of four measures of worker preferences on an indicator for the optimal match, with worker fixed effects. Panel B displays the results of regression of four measures of manager preferences on an indicator for the optimal match, with manager fixed effects. The regression drops worker-manager pairs obtained from random matching, so the comparison group is matching from deferred acceptance. Each regression uses robust standard errors clustered at the simulation-submarket level. Some workers did not rank managers and vice-versa. In column 2, we impute these as workers being indifferent among these choices (if a worker submitted N rankings, we impute the non-ranked managers as the worker's N+1 choice. In column 3, we randomly assign preferences over unranked choices. Column 5 uses a standardized utility measure from a rank-ordered logit, and drops some workers/managers who only submitted one ranking.

*** $p < 0.01$ ** $p < 0.05$ * $p < 0.10$

Table 7: **Worker/Manager Ranking of Assignment, by matching algorithm**

Panel A: Workers' Rankings of Managers

	Mean	Std.Dev	Min	P25	Median	P75	Max
Worker Draft	1.15	0.49	1.00	1.00	1.00	1.00	5.08
Workers-Propose DA	1.17	0.64	1.00	1.00	1.00	1.00	8.00
Managers-Propose DA	1.17	0.64	1.00	1.00	1.00	1.00	8.00
Manager Draft	1.95	1.75	1.00	1.00	1.00	2.00	16.34
Random Assignment	3.36	2.46	1.00	2.00	2.00	4.00	22.40
Firm Optimal	3.39	2.28	1.00	2.00	2.00	4.00	16.00
Firm Objective (minimizer)	3.48	2.38	1.00	2.00	2.00	4.00	15.58

Panel B: Managers' Rankings of Workers

	Mean	Std.Dev	Min	P25	Median	P75	Max
Managers-Propose DA	1.19	0.65	1.00	1.00	1.00	1.00	6.32
Workers-Propose DA	1.19	0.65	1.00	1.00	1.00	1.00	6.32
Worker Draft	1.32	0.82	1.00	1.00	1.00	1.00	7.32
Manager Draft	1.62	1.19	1.00	1.00	1.00	2.00	10.20
Random Assignment	3.00	1.80	1.00	2.00	2.00	3.79	11.50
Firm Optimal	3.04	1.79	1.00	2.00	2.00	4.00	11.90
Firm Objective (minimizer)	3.15	1.82	1.00	2.00	2.66	4.00	11.22

Notes: This table displays the average firm objective score for matches generated by various algorithms. Because there were more managers on the market than workers, some managers are unmatched. This table includes only complete manager worker/pairs (i.e., dropping managers who were not paired to workers). All matching algorithms left the same number of managers unmatched, but differed in their composition. A similar table including unmatched managers (containing a score of zero) is accessible in Table [E1](#).

Table 8: **Heterogeneity in impact of DA on firm objectives**

Panel A: Worker characteristics

	Firm objective score					
Deferred Acceptance	-0.18*** (0.024)	-0.18*** (0.027)	-0.21*** (0.014)	-0.19*** (0.030)	-0.39*** (0.050)	-0.14 (0.087)
X Average Ranking (Worker)	0.033 (0.024)					
X % Ranked By Managers		-0.020 (0.011)				
X Above-Median			0.062 (0.033)			
X Comparative Advantage				0.29 (0.19)		
X Objective Score Mean					0.39*** (0.064)	
X Objective Score StdDev						-0.18 (0.40)
R^2	0.779	0.777	0.779	0.777	0.804	0.776
Observations	64200	64200	64200	64200	64200	64200
Covariate	Worker average ranking	Ranked by X% of managers	Above- median ranking	Comparative advantage	Objective score, mean	Objective score, std. dev.
Fixed effect	Worker	Worker	Worker	Worker	Worker	Worker

Panel B: Manager characteristics

	Firm objective score					
Deferred Acceptance	-0.15*** (0.019)	-0.15*** (0.023)	-0.16*** (0.015)	-0.17*** (0.026)	-0.14 (0.081)	0.073 (0.061)
X Average Ranking (Manager)	0.031 (0.021)					
X % Ranked By Workers		0.0084 (0.011)				
X Above-Median			0.025 (0.045)			
X Comparative Advantage				0.80** (0.32)		
X Objective Score Mean					-0.016 (0.19)	
X Objective Score StdDev						-0.94** (0.28)
R^2	0.615	0.614	0.614	0.620	0.613	0.621
Observations	77550	77550	77550	77550	77550	77550
Covariate	Manager average ranking	Ranked by X% of workers	Above- median ranking	Comparative advantage	Objective score, mean	Objective score, std. dev.
Fixed effect	Manager	Manager	Manager	Manager	Manager	Manager

Notes: This table examines heterogeneity in the impact of deferred acceptance on firm objectives. It displays the results of regressions of the firm objective score on indicators for the match being generated through deferred acceptance, and then interactions between this indicator and various worker (Panel A) and manager (Panel B) characteristics. The excluded category is the optimal match. Panel A includes worker fixed effects while Panel B includes manager fixed effects. Each regression includes robust standard errors clustered at the simulation-submarket level. Column 1 uses a normalized z-score of worker/manager rankings (where positive = more preferred by the other side of the market).

*** $p < 0.01$ ** $p < 0.05$ * $p < 0.10$

Table 9: **Heterogeneity in impact of optimal match on worker utility**

Panel A: Worker ranking of manager

	Worker ranking of matched manager					
Firm optimal	2.27*** (0.17)	2.07*** (0.32)	2.41*** (0.64)	2.13*** (0.38)	3.39*** (0.70)	2.37* (1.13)
X Average Ranking (Worker)	-0.59*** (0.15)					
X % Ranked By Managers		1.34*** (0.35)				
X Above-Median			-0.41 (0.76)			
X Comparative Advantage				2.28* (1.10)		
Objective Score Mean					-2.17** (0.73)	
Objective Score StdDev						-0.84 (4.08)
R^2	0.579	0.742	0.559	0.559	0.579	0.556
Observations	64200	64200	64200	64200	64200	64200
Covariate	Worker average ranking	Ranked by X% of managers	Above- median ranking	Comparative advantage	Objective score, mean	Objective score, std. dev.
Fixed effect	Worker	Worker	Worker	Worker	Worker	Worker

Panel B: Worker is matched to unranked manager

	Manager unranked by worker					
Firm optimal	0.85*** (0.029)	0.86*** (0.024)	0.82*** (0.047)	0.85*** (0.030)	0.77*** (0.062)	0.91*** (0.11)
X Average Ranking (Worker)	0.0044 (0.034)					
X % Ranked By Managers		-0.064** (0.018)				
X Above-Median			0.067 (0.050)			
X Comparative Advantage				-0.11 (0.12)		
Objective Score Mean					0.15 (0.090)	
Objective Score StdDev						-0.32 (0.51)
R^2	0.790	0.796	0.791	0.790	0.791	0.790
Observations	64200	64200	64200	64200	64200	64200
Covariate	Worker average ranking	Ranked by X% of managers	Above- median ranking	Comparative advantage	Objective score, mean	Objective score, std. dev.
Fixed effect	Worker	Worker	Worker	Worker	Worker	Worker

Notes: This table examines heterogeneity in the impact of the optimal match on two measures of worker utility: average ranking of matched manager (Panel A), and a binary indicator whether the worker was matched to a manager she did not rank (Panel B). I will add a third measure of estimated utility once I get more covariates. It displays the results of regressions of these measures on indicators for the optimal match, and then interactions between this indicator and various worker characteristics. The excluded category is the match obtained via deferred acceptance. All regressions include worker fixed effects, with robust standard errors clustered at the simulation-submarket level.

*** $p < 0.01$ ** $p < 0.05$ * $p < 0.10$

Table 10: Heterogeneity in impact of optimal match on manager utility

Panel A: Manager ranking of worker

	Manager ranking of matched worker					
Firm optimal	0.69 (0.44)	0.34 (0.34)	0.47 (0.40)	1.48** (0.49)	11.0* (5.20)	3.94 (2.64)
X Average Ranking (Manager)	0.56 (0.37)					
X % Ranked By Workers		1.97 (1.38)				
X Above-Median			0.51 (0.95)			
X Comparative Advantage				-23.2** (9.14)		
X Objective Score Mean					-19.5 (12.1)	
X Objective Score StdDev						-13.7 (10.2)
R^2	0.398	0.404	0.398	0.402	0.412	0.399
Observations	77550	77550	77550	77550	77550	77550
Covariate	Manager average ranking	Ranked by X% of workers	Above- median ranking	Comparative advantage	Objective score, mean	Objective score, std. dev.
Fixed effect	Manager	Manager	Manager	Manager	Manager	Manager

Panel B: Manager is matched to unranked worker

	Worker unranked by manager					
Firm optimal	0.70*** (0.050)	0.71*** (0.045)	0.65*** (0.082)	0.71*** (0.042)	0.63*** (0.11)	0.81*** (0.14)
X Average Ranking (Manager)	0.026 (0.068)					
X % Ranked By Workers		-0.023 (0.020)				
X Above-Median			0.11 (0.089)			
X Comparative Advantage				-0.18 (0.14)		
X Objective Score Mean					0.13 (0.20)	
X Objective Score StdDev						-0.44 (0.53)
R^2	0.764	0.765	0.767	0.764	0.764	0.765
Observations	77550	77550	77550	77550	77550	77550
Covariate	Manager average ranking	Ranked by X% of workers	Above- median ranking	Comparative advantage	Objective score, mean	Objective score, std. dev.
Fixed effect	Manager	Manager	Manager	Manager	Manager	Manager

Notes: This table examines heterogeneity in the impact of the optimal match on two measures of manager utility: average ranking of matched worker (Panel A) and a binary indicator whether the manager was matched to a worker she did not rank (Panel B). I will add a third measure of estimated utility once I get more covariates. It displays the results of regressions of these measures on indicators for the optimal match, and then interactions between this indicator and various manager characteristics. The excluded category is the match obtained via deferred acceptance. All regressions include manager fixed effects, with robust standard errors clustered at the simulation-submarket level.

*** $p < 0.01$ ** $p < 0.05$ * $p < 0.10$

Appendix: For Online Publication Only

A Theory Appendix

A.1 Formalization of Command and Control using the [Kuhn-Munkres Algorithm](#)

Any proposed match between worker i and division j has probability $P(i \rightarrow j \text{ acceptable})$ of quitting, and payoff of v_{ij} , for an expected value of $v_{ij}P(i \rightarrow j \text{ acceptable})$. Because the principal's objective in Equation ?? is to find matches that maximize the sum of these, the problem is equivalent to the assignment problem ([Kuhn, 1955](#)), even after integrating beliefs about quitting. It is thus solvable through the [Kuhn-Munkres](#) linear programming algorithm (a.k.a. the "Hungarian algorithm").

A.2 Noisy CEO Beliefs

Assumption 5 (Noisy CEO Beliefs). *Suppose that match specific productivities v_{ij} are drawn I.I.D. from a distribution F with mean $\bar{v}$. CEO's do not directly observe each v_{ij} draw from F , but instead observe each v_{ij} with noise; i.e. they observe $v'_{ij} = v_{ij} + \epsilon$, where $\epsilon \sim \mathcal{N}(0, \sigma^2)$ and drawn independently for each possible match.*

Proposition A1. *As the noise in CEO's observations of match productivity (σ^2) increases, the returns to command-and-control decrease.*

Proof. The CEO's prior belief about each match's productivity is F . For each observation of v'_{ij} , the CEO updates beliefs about the match's value using Bayes' rule.

As σ^2 goes to infinity, the signal is increasingly uninformative and the CEO places more weight on his prior, which has a mean of $\bar{v}$ for all applications. $\square$

A.3 Non-Vertical Side Proposing

As discussed in the main text, our stylized formulation of the problem does not allow the CEO to asymmetrically care more about one side of the market's retention. In our formulation, the only penalty for quitting is zero productivity in the assignment where one (or both) sides quit. There are no replacement costs, or value retained from a worker who stayed (when their partner quit), or other sources of potential asymmetry between the quitting costs of either side. Here, we allow the CEO to care about both sides quitting.

Assumption 6. Assume that both sides quitting is bad, and that one side has vertical preferences and the other has I.I.D. preferences.

Under Assumption 6, DA may return higher results if the side without vertical preferences proposes. The returns to command-and-control do not change.

A.4 Both Sides have Vertical Preferences

Proposition A2. If both sides have vertical preferences, then command-and-control statewise dominates DA, irrespective of who proposes.

Proof. If either side of the market has vertical preferences, then regardless of which of the $N!$ matches is selected, exactly one agent will match with her 1st, 2nd, $\dots$, N th most preferred division. As a result, under I.I.D. preferences the retention rate for each side will be $1 - G(\underline{u})$, and the retention rate for DA matches will be $(1 - G(\underline{u}))^2 = \bar{R}$. This is the same as the retention rate for CC, however CC can achieve a higher productivity match if $\mathcal{V}$ isn't flat. $\square$

A.5 Narrow Alignment not Sufficient

To show this, we utilize a counterexample. Suppose a 2×2 firm contains payoffs as below.

	Division A	Division B
Worker 1	x	1
Worker 2	1	0

This means that matching Worker 1 with Division A yields a v_{ij} yields x units of productivity. Suppose that workers and divisions rank their match partners based on v_{ij} (narrow alignment). If $x < 2$, then the optimal assignment will be {(Worker 1, Division B), (Worker 2, Division A)}, which yields a total output of two.

Suppose that $1 < x < 2$. The optimal assignment is 2 (above). However, the Worker 1 prefers Division A, who reciprocates (this is the consequence of narrow alignment). In a preference-respecting setup, the match will be {(Worker 1, Division A), (Worker 2, Division B)}, which yields an output of x (which is less than the optimal output of two). The expected output of a delegated mechanism is lower, despite workers and division being narrowly aligned.

A.6 Proofs

A.6.1 Proof for Lemma 1

Lemma 1. *The expected retention rate of the match selected by DA is weakly higher than the average of all matches, $\bar{R}$.*

Proof. The retention rate for DA is equal to the retention rate of the proposing side multiplied by the retention rate of the receiving side. The average retention rate $\bar{R}$ is $(1 - G(\underline{u}))^2$. Under Lemma 2, at most one agent on the proposing side is assigned this worst match. As such, the proposers' retention rate is higher in DA. The receiving side is weakly better off in DA. It is generally not better off than average, except in circumstances in which two proposers target the same receiver, and the receiver gets a choice. Because retention is higher for proposers and weakly higher for receivers, the product is higher than the average retention. $\square$

Lemma A1. *Under Assumption 4 (I.I.D. preferences), retention rates for the receiving side are weakly better than random.*

A.6.2 Proof for Lemma 2

Lemma 2. *In an $N \times N$ market where workers and divisions rank all possible matches, at most $\sum_{l=1}^k l$ individuals on the proposing side can match with a partner they rank k^{th} or higher.*

Proof. Without loss generality (because we have assumed workers and divisions are symmetric other than in who makes proposals), consider the worker-proposing algorithm. We would like to track the evolution of the algorithm. Let s_i denote the number of proposals worker i has made up to a particular point in the algorithm. Note that this is worker i 's rank of the division they have most recently proposed to. Define $S = \max_i s_i$ be the highest number of proposals any worker has made. When $S = 1$, all workers have proposed to their most preferred division. When $S = 2$, workers rejected by their first choice have proposed to their 2nd most preferred division. And so on.

Our assumption that workers and divisions rank all possible matches implies all workers and divisions prefer all matches to being unmatched (although they may quit later). At least S divisions (and workers) must have received at least one proposal and therefore be holding a (possibly not final) match. Then at most $N - S$ workers will make proposals at any iteration of the algorithm so there can be at most $N - S$ rejections while S is at this level. This implies at most 1 worker can be rejected when $S = N - 1$ so at most one worker will match with her least preferred division. This implies there can be up to $\sum_{l=1}^k l$ total rejections when $S = k - 1$. This proves the claim.

For example, consider a worker in an $N \times N$ matching market who is proposing to her N^{th} ranked match, she must have been rejected by $N - 1$ divisions. These $N-1$ divisions must be holding matches or else they would have accepted her proposal. The final division must not be holding a match because the other $N-1$ workers are tentatively matched with the other $N-1$ divisions. This worker must not have proposed to the unfilled position yet otherwise it would have to be holding a match. Therefore she will propose to the unfilled position, displacing no other workers, and the algorithm terminates because everyone is matched. $\square$

A.6.3 Proof for Proposition 1

Proposition 1. *The performance of delegation will equal or exceed that of command-and-control in firms where workers are completely unspecialized.*

Proof. If workers are unspecialized under Definition 1, then v_{ij} s are constant across all assignments. For any assignment a of N workers to N jobs will produce the same total output if nobody quits. This output is equal to $V(a) = \sum_{i=1}^I \sum_{j=1}^J \alpha_{ij}(a)$, where $\alpha_{ij}(a)$ is equal to 1 if worker i and division j are matched in assignment a , and 0 otherwise. If all workers are equally productive in all jobs, then v_{ij} equals a constant v_i for the worker i , and $V(a) = \sum_{i=1}^N v_i$ for all a . The most productive match (V_{CC}) is equal to the average match ($\bar{V}$). All the mass in $\mathcal{V}$ is concentrated in a single point, and the LHS of Equation 2 is equal to 1. Lemma 1 shows that the RHS of Equation 2 is equal to or greater than one. $\square$

A.6.4 Proof for Corollary 1

Corollary 1. *For command-and-control to outperform delegation, it is necessary (but not sufficient) for the workforce to be specialized.*

Proof. For command-and-control to outperform delegation, the LHS of Equation 2 needs to exceed the RHS. The right hand side of Equation 2 will be greater than or equal to 1 because of Lemma 1. If workers are not specialized, the LHS of Equation 2 will equal 1 under Proposition 1. $\square$

A.6.5 Proof for Proposition 2

Proposition 2. *Assume that $\underline{u} < \gamma$, then as $\underline{u}$ increases, the unconditional quit probability $G(\underline{u})$ increases and the returns to DA are higher.*

Proof. The returns to DA are the right hand side of Equation 2, $\frac{R_{DA}}{\bar{R}}$. The denominator $\bar{R}$ is $(1 - G(\underline{u}))^2$. When N is sufficiently high, all the utilities associated with DA are greater

than γ and so for $\underline{u} < \gamma$ and N sufficiently high, the retention rate for DA is 1, while it is decreasing in $\underline{u}$ for CC, the so the returns to DA are increasing in $\underline{u}$. If N is not so large that that all utilities under DA are greater than γ , the distribution is sufficiently right skewed, such that the retention rate is both greater than under CC and declining at a slower rate and so again, the returns to DA are increasing in $\underline{u}$. □

A.6.6 Proofs for Corollaries 2, 3 and 4

Corollaries 2, 3 and 4 are applications of Proposition 2.

Corollary 2 (Workforce-Unfriendly Firms). *Let G' be a leftward shift of G so that G' equals G minus a positive constant. DA is more attractive under G' than G .*

Proof. If G is a leftward shift by a constant, then $G(\bar{u})$, the probability of quitting has increased. DA is more attractive by Proposition 2. □

Corollary 3 (Outside Options). *DA is more attractive as $\underline{u}$ increases.*

Proof. If $\bar{u}$ improves, then $G(\bar{u})$, the probability of quitting has increased because G is increasing in its argument. DA is more attractive by Proposition 2. □

Corollary 4 (Asymmetric Information). *Let G' be a mean-preserving spread of G , so that G' has the same mean of G but higher variance. DA is more attractive for G' than G .*

Proof. If G' is a mean preserving spread of G , then $G'(\underline{u})$ is greater than $G(\underline{u})$ because greater probability mass falls below $\underline{u}$. DA is more attractive by Proposition 2. □

A.6.7 Proof for Lemma 3

Lemma 3 (Quitting Costs). *Higher quitting costs attenuate the relative benefits of CC over DA, $\frac{V_{CC}}{V_{DA}}$.*

Proof. To see this, assume every successful match has a benefit of $\frac{c}{N}$ in addition to its productivity v_{ij} . We can interpret $\frac{c}{N}$ as a quitting cost because it is lost if the match is unsuccessful. The ratio $\frac{V_{CC}}{V_{DA}}$ is decreasing in c :

$$\frac{\partial}{\partial c} \left[\frac{V_{CC} + c}{V_{DA} + c} \right] = \frac{V_{DA} - V_{CC}}{(V_{DA} + c)^2}. \quad (3)$$

This is negative whenever $V_{CC} > V_{DA}$. Therefore, increasing the costs of quits makes DA relatively more appealing than CC. $\square$

A.6.8 Proof for Proposition 3

Proposition 3 (Firm Size). *For sufficiently large markets, the expected retention rate of DA increases with N , the size of the firm, so long as $G(u)$ has a positive density over its full support.*

Proof. Let Q_i be worker i 's rank of her assigned division in a stable matching, $\mathcal{M}$. Define $\mathcal{Q} = \max_i Q_i(\mathcal{M})$. Pittel (1992a), Theorem 6.1 shows that $P(\mathcal{Q} \leq (2+a)\log^2 n) \geq 1 - O(n^{-c})$ for all $c < c(a)$ where $c(a) = 2a[3 + (4a+9)^{1/2}]^{-1}$. In words, this says it is almost certain that the least preferred match of a worker in an $n \times n$ random matching market using working-proposing deferred acceptance will be ranked less than $(2+a)\log^2 n$. To be conservative, we will use $3\log^2 n$. This suggests the worker who receives the least preferred match in the entire market is almost certain to match with utility draw at least in the p_n percentile, where $p_n = \frac{n-3\log^2 n}{n}$. This is increasing for $n > 7$:

$$\begin{aligned} p'_n &= \frac{(n - 6\log n - n + 3\log^2 n)}{n^2} \\ &= \frac{3(\log^2 n - 2\log n)}{n^2}, \end{aligned}$$

Where $\log^2 n > 2\log n$ if $n > 7$. This shows that the person who randomly gets the worst outcome in DA, and is therefore the most likely to quit, will match with a more preferred match as N grows.

?, Theorem 2 shows that if $x_1, x_2, \dots, x_n$ are drawn from $f(x)$ where $f(\cdot)$ is continuous and does not vanish in the neighborhood of the p^{th} percentile, then the p^{th} quantile is asymptotically normal with mean $F^{-1}(p)$ and variance $\frac{p(1-p)}{n[f(F^{-1}(p))]^2}$. This suggests that the quit probability for a worker with outside option $\underline{u}$ who matches with her p^{th} percentile match is:

$$\begin{aligned} P(Q|p) &= P(\mu^i < \underline{u}) \\ &= \Phi\left(\frac{\underline{u} - F^{-1}(p)}{n[f(F^{-1}(p))]^2}\right). \end{aligned}$$

Together, these two results show that even workers who get the least preferred match in Deferred Acceptance will almost surely match with increasingly preferred positions in percentile terms and the utility from this match will converge to that percentile of the

utility distribution. Therefore, the probability that a worker quits converges to zero as the market size grows.

On the receiving side of the market, [Pittel \(1992a\)](#), Theorem 6.2 shows that the worst match on the receiving side of the market is almost surely to have rank n . Therefore, the least preferred match on the receiving side of the market does not similarly improve with n . However, ? shows the average rank of matches on the receiving side converges to $n/\ln(n)$ whereas the average rank of a random match is $n/2$. This suggests the attrition rate on the receiving side of the market will be lower than in CC. $\square$

A.6.9 Proof for Proposition 4

Proposition 4 (Vertical Preferences). *If one side of the market has vertical preferences and the other side of the market does not quit, then DA is statewise dominated by command-and-control.*

Proof. If one side of the market has vertical preferences then regardless of which of the $N!$ matches is selected, exactly one agent will match with her 1st, 2nd, $\dots$, N th most preferred division. Because we have assumed agents on the potentially non-vertical side of the market do not quit, this implies $R_{DA} = \bar{R}$. Because there is no retention penalty from command and control, DA can do no better than CC but can do worse. $\square$

A.6.10 Proof for Proposition 5

Proposition 5. *If workers and divisions have strictly broadly aligned preferences, the CEO's optimal match is the surplus-maximizing Nash equilibrium of deferred acceptance, but is not generally a unique Nash equilibrium. Deferred acceptance can equal the performance of command and control, but cannot exceed it.*

Proof. Because workers and divisions have strictly broadly aligned preferences, $\mu_j^i = v_i^j = V(a)$ if matching together will yield overall match a given all other workers' and divisions' preference reports. V_{CC} is the output of the optimal match. Therefore it must be the case that $V_{CC} \geq V(a)$ for all $a \in 1, \dots, N!$. Suppose all workers except worker i rank their optimal match first. Then worker i 's best response is to also rank her optimal match, say j^* first because $\mu_j^i = V_{CC} \geq V(a)$ for all $a \in 1, \dots, N!$. The same logic applies to division j . This proves that the optimal assignment is a Nash equilibrium.

Now, suppose there is another Nash equilibrium that yields output $\tilde{V} < V_{CC}$. If all workers except i rank their match corresponding to this alternative equilibrium first, then it is a best response for i to also rank this alternative equilibrium first by our assumption that it is a Nash equilibrium. The same logic applies to divisions.

Together, these two results prove that there are possibly Nash equilibria that yield the same output as command and control but there are possibly others that do not.

□

A.6.11 Proof for Corollary 5

Corollary 5. *Suppose that workers and divisions have strictly narrowly aligned preferences, and that match productivities are the result of a production function that is supermodular in workers' and divisions' types. Then the worker-proposing deferred acceptance algorithm selects the CEO-optimal assignment. However, the CEO can select this without deferred acceptance using command-and-control.*

Proof. Narrow alignment implies deferred acceptance will yield positive assortative matching. It is well known that assortative matching maximizes output when production is super-modular (Becker, 1973). □

B Firm Objective Score: Implementation Details and Discussion

B.1 Additional Implementation Details

Several variables went into the firm objective score, but the most prominent were job skills and requirements, language requirements, and location. Using the aforementioned structured skill taxonomy, the company developed a distance metric between the skills sought on the manager's job requisition, and the skills listed on the worker's profile. Workers and managers' self-reported skills and requirements were reviewed by a centralized third party for accuracy. Matches on specific skills (or the absence of a match) was weighted differently depending on the type of skill and how the skill was characterized in the manager's requisition and on the worker's profile. These weightings were hand-tuned by the central management to nudge workers and managers towards the outcomes preferred by the firm.

B.2 Limitations

We offer a few caveats to our characterization of the score as the firm's "objective." In short: We acknowledge that a firm might want to accomplish more in its "optimal match" than optimizing skills, language and location. However, among the many ways to optimize around these variables, the one selected by the firm may have been chosen because it looked good on *other* dimensions.

In addition, incorporating other dimensions more directly would face measurement problems. For example, a firm may want to consider a worker’s personality fit with colleagues and/or intrinsic motivation for the task. Unfortunately, these variables may not be measured credibly by anyone, much less by a centralized manager for all possible combinations of worker/manager pairs in a scalable way. Workers and managers themselves may not accurately know these values. If they did, collecting data about these characteristics would be subject to strategizing (as workers may pose to be unmotivated to force an assignment towards their preferences). Notably, the firm made no effort to capture these measures in its nudge. A firm-optimal match that utilizes this data may be unachievable in practice.

As such, we encourage readers to think of our “firm objective” variable as a reasonable (if imperfect) proxy of the firm’s objectives, subject to the firm’s information constraints.

B.3 Additional Interventions

Finally, One may wonder whether the firm undertook additional measures to influence the match that are outside our data. For example: They may have intervened through private channels to encourage certain matches, possibly in spite of a low match score. This would affect our characterization of the scores as the “objective.” If this were true, the choice of a match score may have represented only *part* of the firm’s objectives; the other part coming through offline interventions invisible to us researchers.

Although this is possible, our interviews with the company have suggested nothing like this happened. The central managers representing the firm’s broader interests are several layers of management removed from the rank-and-file workers and managers in the match. Centralized managers have little information on which to base these interventions. To the contrary, central managers reported feeling nervous about the outcome of the match, compared it to a “leap of faith” and desired better tools to incorporate worker preferences without fully losing control.

B.4 Bounded Support of Objective Function

Was it a mistake for the firm to implement a maximum? Is it realistic to feel that worker/manager pairings have a maximum and minimum potential quality? One can imagine scenarios in which a star worker assigned to a particular project could raise the ceiling of the project indefinitely high by (say) pitching new lines of business with the client, proposing a blockbuster merger, preventing a blockbuster merger, or creating a new software patent licensed to the client for billions of dollars.¹

¹In a sense, any of these outcomes are bad matches because the worker should be in a new division creating new products, or possibly *becoming* the manager.

However, it does seem plausible for there to be a maximum in settings like ours. In our setting, each assigned worker will become a rank-and-file member of an existing team of 5-12 people. The teams each have an existing leader, road-map and relationship with a client. The ability of the worker to significantly alter the course of the project is potentially large, but plausible limited on both sides. In addition, the firm's objective score could have been (in principle) unbounded. As we discussed above, the properties of this scoring system reflected some deliberate design.

C Data

Our data consist of four tables about participant characteristics, preferences, pairwise match scores and potential assignments under various matching regimes.

Participant Characteristics. For both workers and managers, we have a set of individual covariates. Some of the manager characteristics are associated with the job he/she is hiring for. Note that some workers and managers are observed multiple times, as they finish one assignment and enter another one over our sample period.

Participant Preferences. For each worker and manager, we have a complete set of ordinal preferences over the opposite side of the market. For workers and managers who appear in the market multiple times, we have multiple sets of preferences. Note: We describe our as "complete," although a large portion of possible matches were left unranked, or essentially tied for last place. We code such labels as expressing indifferences.² Given our setting has more managers than workers, some managers will not be matched to workers. Non-matches are coded as the least desirable choice for participants based on their inputs.³

Pairwise Match Scores. For all possible pairs of workers and managers, we have several numeric scores characterizing the match. The most important in our analysis is the *firm objective score*. This is (roughly) the firm's estimate of how productive each match would be for the firm's goals. The firm would prefer the match that maximizes the sum of total pairwise scores across all assignments.

Utilizing these pairwise productivity scores, we then calculate two additional variables about each pair. These measures utilize notions of comparative advantage and apply

²For example, if a worker ranked seven out of a total of ten managers, we code the three non-ranked managers as choice 8. We also run a version where we randomize participant preferences among unranked choices.

³As described above, all participants were given the option to rank "prefer to be unmatched" ahead of some least-preferred. No participant chose to utilize this option.

them to both sides. For each possible match, we calculate how much more productivity the match yields, compared to the workers' next most productive match. We do the same for managers, comparing each match to the manager's next most productive match.

Actual and Simulated Assignments. Finally, our data includes the assignments generated by the worker-proposed match used by the company. We have also re-run this algorithm 50 independent times, resolving indifferences and ties randomly, to generate a distribution of potential assignments. In addition, we also use the preference data to generate counterfactual assignments based on a variety of other assignment methods. Specifically, we examine the managers-propose DA algorithm which is nearly always identical to the workers-propose. We also examine completely random assignments, as well as random serial dictatorship (?) led by the workers and managers (we call these the "workers' draft" and the "managers' draft"). Our goal using these latter two is to help us evaluate non-parametrically how much our results are driven by workers and managers preferences by examining matching based only on one side's preferences. We similarly run each of these algorithms 50 times, resolving ties randomly. When running matches using participants preferences, we utilize the same random draws to break ties across all different assignment algorithms.

Finally, we calculate the "firm-optimal match," which maximizes the firm's objective (as described above). To calculate this, we utilize the aforementioned [Kuhn-Munkres](#) algorithm (aka the "Hungarian algorithm," [Kuhn 1955](#)). This algorithm identifies the set of assignments which maximize a numeric objective. We apply this algorithm to find the set of worker/manager assignments that maximize the total firm objective score (described above). Because this sometimes involves resolving ties between equally productive pairwise matches, we run this algorithm 50 times. We also run a version of the [Kuhn-Munkres](#) algorithm that minimizes the total firm objective to provide a lower bound of the firm's objective function in our data (while still assigning all workers).⁴

D Screenshots of Internal Marketplace (Mockups)

E Additional Empirical Analysis

⁴One possible way to minimize the firm's objective might be to leave them all unmatched; our implementation of the objective-minimizer forbade this.

Figure 3: Profile Screen

myMatchEN▼Amy Smith ▼

InstructionsMy ProfileMy PositionsAll PositionsCandidates

← **Amy Smith (Engineering Manager)**
My teams work on Hooli Banking and Search.
At Hooli Since: 2015



Previous Projects: Hooli Apps, Nucleus, hooli.xyz.

- I wrote the compiler behind hooli.xyz.
- Founding engineer on compression project.

Education: UCLA (BS/MS, Computer Science, '05),
MBA (Berkeley, 2014)

Skills: Java; C++; SQL, HooliReduce, software architecture.

- I was a software engineer focused on databases before entering management at Hooli.

Previous Employers: Microsoft (2005-2015).

Open Positions

- ✓ • **Software Engineer (Hooli Banking):**
 - My team is hiring a dedicated engineer to automate loan applications.
 - [Full description here](#); please email me if you may be a fit!

Home | Contact | FAQs & Help

Notes.

Figure 4: Ranking Screen

myMatch

EN  Amy Smith ▼

Instructions

My Profile

My Positions

All Positions

Candidates

← Role: Software Engineer (Hooli Banking)

RANKED		UNRANKED		EXCLUDED	
(1)  Lila 	 Jessica 	 Kenneth 	(X)  Gabriel 		
(2)  Derrick 	 Michael 	 Jason 	(X)  Joan 		
(3)  Jennifer 	 John 	 Gloria 	(X)  Joseph 		
(4)  Anita 	 Ashley 	 Maria 			
(5)  Isaac 	 Ryan 	 William 			
(6)  Shawn 	 Cynthia 	 Victoria 			
	 Abigail 	 Thomas 			
	 Richard 	 Sarah 			

[Home](#) | [Contact](#) | [FAQs & Help](#)

Notes.

Table E1: Firm objective score by matching algorithm

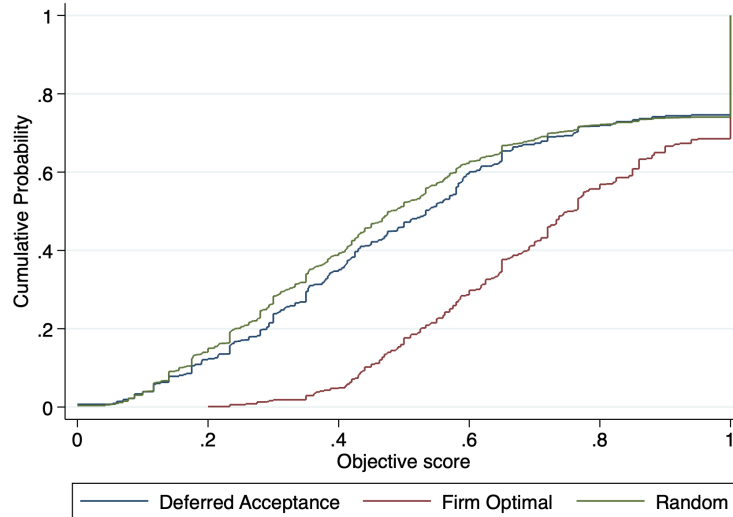
	Mean	Std.Dev	Min	P25	Median	P75	Max
Firm Optimal	0.75	0.22	0.21	0.57	0.76	1.00	1.00
Manager Draft	0.58	0.31	0.00	0.32	0.53	1.00	1.00
Worker Draft	0.57	0.30	0.00	0.33	0.53	0.98	1.00
Managers-Propose DA	0.57	0.31	0.00	0.32	0.53	0.99	1.00
Workers-Propose DA	0.57	0.31	0.00	0.32	0.53	0.99	1.00
Random Assignment	0.55	0.32	0.01	0.29	0.48	0.99	1.00
Firm Objective (minimizer)	0.34	0.32	0.00	0.12	0.22	0.36	1.00

Notes: This table displays the average firm objective score for matches generated by various algorithms. Because there were more managers on the market than workers, some managers were unmatched. This table includes all managers, labeling unmatched managers an objective score of zero. All matching algorithms left the same number of managers unmatched, but differed in their composition. A similar table including only complete manager/worker pairs is accessible in Table 7.

E.1 Distribution of match quality by algorithm

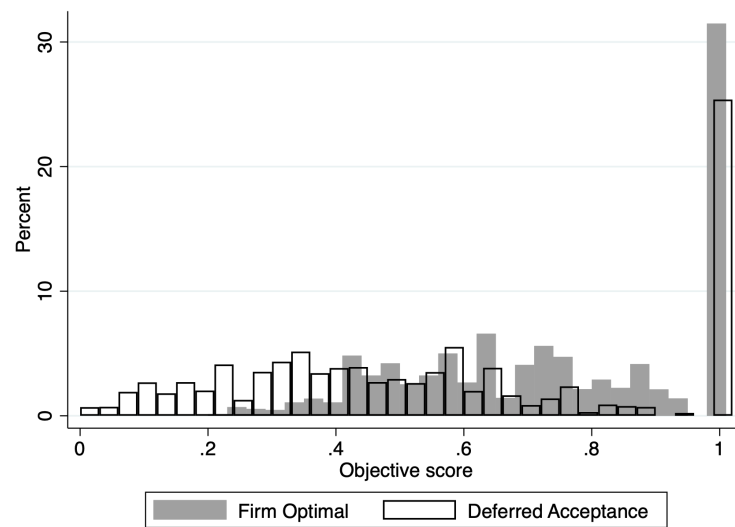
Figure 5 displays a cumulative distribution plot of match quality by algorithm for successful worker-manager matches (i.e., we drop instances when a manager is not matched to any worker). Meanwhile, Figure 6 displays the corresponding histogram.

Figure 5: **Cumulative distribution plot of match quality by algorithm**



Notes: This figure displays a cumulative distribution plot of match quality by the firm-optimal (Kuhn-Munkres) algorithm, deferred acceptance, and random matching.

Figure 6: Histogram of match quality by algorithm

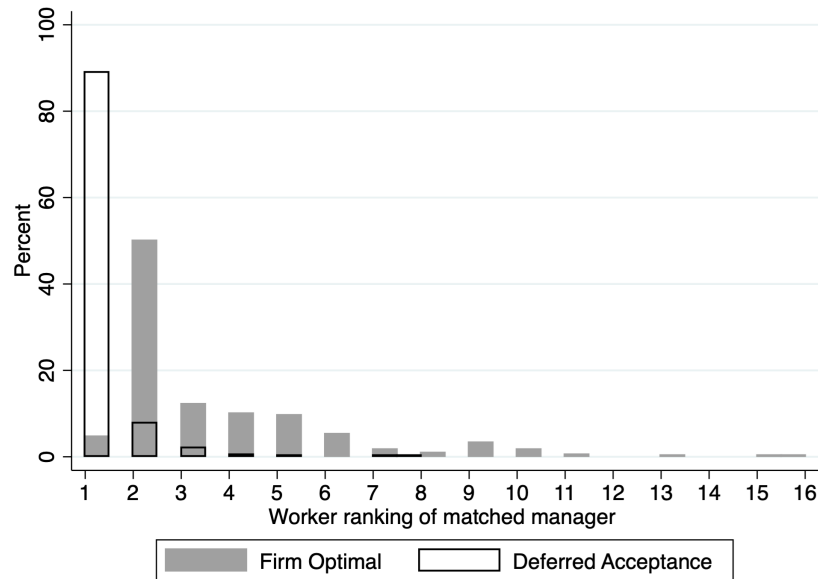


Notes: This figure displays a histogram of match quality by the firm-optimal ([Kuhn-Munkres](#)) algorithm versus deferred acceptance.

E.2 Distribution of match quality by algorithm

Figure 7 displays a histogram of each worker's ranking of their matched manager by algorithm. Figure 8 displays the corresponding plot for manager ranking of their matched workers, excluding managers who are not matched to a worker.

Figure 7: **Histogram of worker ranking of matched manager by algorithm**



Notes: This figure displays a histogram of worker rankings of their matched manager for deferred acceptance and the firm-optimal ([Kuhn-Munkres](#)) algorithm.

Figure 8: Histogram of worker ranking of matched manager by algorithm

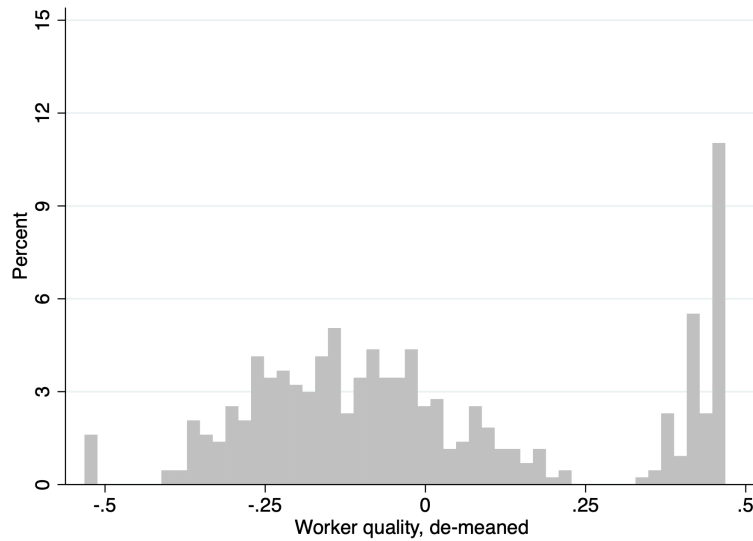


Notes: This figure displays a histogram of manager rankings of their matched worker for deferred acceptance and the firm-optimal ([Kuhn-Munkres](#)) algorithm. The figure drops managers who were not matched to a worker.

E.3 Distribution of worker and manager quality

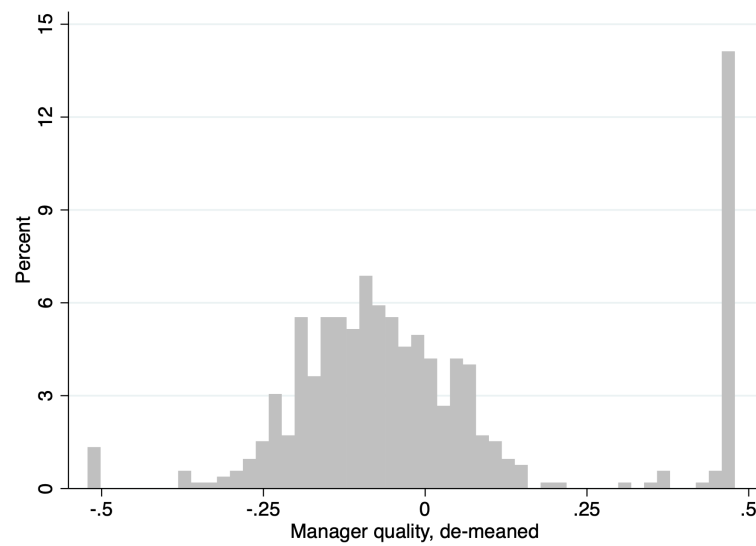
Figure 9 displays a histogram of worker quality. We take the average firm objective score for each worker across all managers and de-mean this measure. Figure 10 displays the corresponding plot for manager quality.

Figure 9: **Histogram of worker quality (de-meaned)**



Notes: This figure displays a histogram of worker quality. We take the average firm objective score for each worker across all managers, de-mean this measure, and plot the distribution across all workers.

Figure 10: **Histogram of manager quality (de-meaned)**



Notes: This figure displays a histogram of manager quality. We take the average firm objective score for each manager across all workers, de-mean this measure, and plot the distribution across all managers.