

Due Diligence: Endogenous Information Acquisition and Offer Quality in Search and Matching

Joseph Briggs^a Andrew Caplin^b Mateusz Giezek^b

Daniel Martin^c Christopher Tonetti^{d,*}

August 2026

^aGoldman Sachs & Co. LLC, 200 West Street, New York, NY 10282, USA

^bDepartment of Economics, New York University, 19 West 4th Street, New York, NY 10012, USA

^cDepartment of Economics, University of California, Santa Barbara, 2127 North Hall, Santa Barbara, CA 93106-9210, USA

^dStanford Graduate School of Business, 655 Knight Way, Stanford, CA 94305-7298, USA

Abstract. In many search markets, offers must be accepted or rejected before their quality is known. We build a search and matching model in which suppliers privately choose offer quality and searchers choose costly due diligence before accepting. How carefully searchers screen determines the mix of high- and low-quality offers. This composition channel is absent when information and offer quality are exogenous. In the static model, improving low-quality offers makes accepting one less bad, so searchers screen less carefully and suppliers post more low-quality offers. The mix worsens enough that more searchers remain unmatched and their welfare falls. Raising the value of remaining unmatched actually reduces the number of unmatched searchers, because searchers screen more carefully and suppliers post more high-quality offers. In the dynamic economy the composition channel survives, and a classic market tightness force can reinforce or overturn it. A planner who treats information and offer quality as exogenous can subsidize entry where a tax is optimal.

Keywords: Due diligence; rational inattention; search; matching; offer quality; learning

JEL classification: D82; D83; E24; J64

1 Introduction

Many important search markets require decisions before the value of an offer is fully known. A worker may observe a wage but not the reliability of the schedule, the quality of management, the likelihood of remote work arrangements, the true value of benefits, or the prospects for advancement. A buyer may observe a price but not the quality of the house, car, financial product, or service being purchased. A trader may observe counterparties and prices while remaining uncertain about asset quality. In such markets, accepting an offer is preceded by due diligence: costly effort to learn enough about the offer to decide whether to accept it or reject it and keep searching.

This due diligence margin is increasingly visible. For example, modern job search runs through postings and platforms that disclose only part of the relevant offer, workers place substantial value on the quality of nonwage job characteristics (Mas and Pallais, 2017; Maestas, Mullen, Powell, von Wachter, and Wenger, 2023), and what postings do reveal changes application behavior and beliefs (Batra, Michaud, and Mongey, 2023; Belot, Kircher, and Muller, 2022). Searchers thus face two linked problems: offers differ in quality, and learning about that quality is itself costly.

This paper studies the equilibrium consequences of that link. We build a search and matching model in which firms privately choose the quality of the offers they make, matched searchers choose how much information to acquire before accepting or rejecting, and prices or wages are determined by bargaining after acceptance. The key feature of the model is that due diligence affects not only which offers are accepted, but also which offers firms choose to create. A searcher who inspects offers more carefully changes the return to posting a low-quality offer. Due diligence is therefore not merely a screening technology with private returns. It is an equilibrium determinant of offer quality.

We model due diligence using rational inattention. A matched searcher chooses an information structure before deciding whether to accept or reject the offer, paying an information cost proportional to the expected reduction in Shannon entropy. This approach gives searchers flexible control over what they learn while preserving tractability. It also delivers an economically useful separation between the posterior beliefs that govern searcher behavior and the conditional acceptance probabilities that govern firm payoffs. This separation allows us to characterize equilibria analytically, first in a static version of the model and then in a dynamic Diamond–Mortensen–Pissarides search and matching environment with vacancy creation.

The interior equilibrium has a simple logic. Ex ante identical firms post both high- and low-quality offers. Searchers conduct costly due diligence because offer quality is uncertain. Firms are indifferent between posting the two qualities because high-quality offers are more likely to be accepted, while low-quality offers are more profitable conditional on acceptance. The equilibrium level of due diligence and the equilibrium composition of offers are therefore jointly determined.

Comparative statics depend on how a change in fundamentals affects both the searcher's information problem and the firm's incentive to supply either quality.

The paper's main contribution is this composition mechanism. Costly due diligence changes the profitability of supplying different qualities, and the induced change in offer composition can reverse comparative statics that would be immediate in a model with fixed offer quality. Interventions affecting low-quality jobs, high-quality jobs, or outside options therefore cannot be evaluated solely from their direct effect on payoffs. Their equilibrium effect depends on how they change searchers' due diligence and firms' choices of offer quality.

The composition mechanism overturns several natural intuitions. For example, improving low-quality offers might seem mechanically good for workers. In our model it makes them worse off. When the low-quality offer becomes less bad, the loss from accepting one by mistake shrinks. The worker screens less carefully, and firms respond by posting more low-quality offers. In the static model, this equilibrium composition response lowers welfare and raises the number of unmatched searchers. In the dynamic model the result carries over exactly. Improving low-quality offers raises unemployment, lowers the value of every worker, and lowers worker welfare.

Improving high-quality offers has a different effect. When the high-quality job becomes more valuable, mistakenly rejecting one costs more, so searchers accept on weaker evidence. Low-quality offers slip through due diligence more easily, and firms respond by posting more of them. In the static model the direct effect dominates the worse mix, so fewer searchers remain unmatched and worker welfare rises. In the dynamic model, acceptance on weaker evidence also raises market tightness through free entry, and the worse mix can outweigh the tighter market. Improving high-quality offers still raises the value of every worker and raises worker welfare, but unemployment falls if and only if the tightness response outweighs the composition response, a comparison governed by the matching elasticity.

The analysis of unemployment benefits is similarly nonstandard. In the static model, a higher value of being unmatched makes searchers more selective, which makes low-quality offers less attractive to firms. The composition of offers improves enough that the acceptance rate rises and fewer searchers remain unmatched. In the dynamic model, this composition channel remains, but higher unemployment benefits also compress the surplus from employment, lowering market tightness. Unemployment then depends on whether the better offer mix or the slacker market dominates. Higher unemployment benefits always raise the value of every worker, unemployed and employed alike. Aggregate worker welfare, the employment-weighted average of these values, can nevertheless fall. If tightness contracts sharply in a high-employment market, enough workers lose the employment premium that the average falls even though every value rises.

We then study how due diligence and endogenous offer quality affect constrained efficiency. The planner faces the same frictions as the market — it cannot dictate information or offer quality — and has two instruments, a permanent vacancy levy that cannot be conditioned on quality and an

unemployment benefit financed by lump-sum assessments. In the standard exogenous information and offer quality model, the planner would correct entry and be done. With due diligence and endogenous offer quality, both instruments also move what workers learn and what quality firms supply. The optimal levy corrects the Hosios (1990) benchmark by the response of due diligence and offer composition to market tightness, a correction that can reverse the levy’s sign.

1.1 Relation to the Literature

This paper brings together classic search and matching theory (Stigler, 1961; McCall, 1970; Diamond, 1982; Mortensen, 1982; Pissarides, 1985), endogenous information acquisition (Moscarini and Smith, 2001; Sims, 2003), and strategic offer design. The closest intellectual connection is to search models in which match quality is imperfectly observed. In labor markets, Pries and Rogerson (2005) study jobs as both inspection and experience goods, while Moscarini (2005) embeds learning about match quality into equilibrium search and matching. Menzio (2007) studies partially directed search with incomplete information. In asset markets, Guerrieri, Shimer, and Wright (2010), Guerrieri and Shimer (2014), and Guerrieri and Shimer (2018) study adverse selection and price formation in search environments. These papers show how incomplete information about match or asset quality affects liquidity, prices, and matching outcomes. Our model differs by making the information acquired before acceptance endogenous, and by allowing firms to choose offer quality in response to that information.

The paper is also related to work in which firms choose the quality or terms of the offers they make. The closest paper in this respect is Delacroix and Shi (2013), who study pricing and signaling with frictions in a search environment in which offer quality is endogenous. In their model, prices can signal quality and searchers can direct search based on prices. Our environment is different. Firms choose offer quality privately, the matched searcher chooses how much due diligence to conduct, and bargaining takes place after acceptance and resolution of uncertainty. This timing shuts down posted price signaling as the mechanism revealing quality and instead makes due diligence the central informational margin.

Our modeling of due diligence follows the rational inattention literature. The foundational contribution is Sims (2003); Matějka and McKay (2015) develop the discrete choice formulation that makes Shannon rational inattention especially tractable, Caplin, Dean, and Leahy (2022) provide a posterior characterization and generalization of Shannon entropy costs, and Maćkowiak, Matějka, and Wiederholt (2023) survey the literature. We use the posterior formulation because it cleanly separates beliefs after acceptance and rejection from the conditional acceptance probabilities that determine firm payoffs; this is the source of the paper’s analytical tractability. Several recent papers place rational inattention in strategic or market environments: Martin (2017) studies strategic pricing when consumers are rationally inattentive to quality, Ravid (2020) studies bargaining with rational inattention, and Almog and Martin (2024) provide experimental evidence on rational inat-

tention in games. Closest within labor market search is Acharya and Wee (2020), who incorporate rational inattention into hiring decisions and study business cycle implications. Our paper is complementary, as their contribution is quantitative and focuses on firms’ attention in hiring, while ours is theoretical and focuses on searchers’ due diligence, endogenous offer quality, and analytical comparative statics.

The labor market motivation rests on three facts from recent empirical work. First, workers place substantial value on nonwage job characteristics, long central to compensating differentials theory (Rosen, 1986): Mas and Pallais (2017) estimate willingness to pay for alternative work arrangements, and Maestas, Mullen, Powell, von Wachter, and Wenger (2023) show that working conditions are an important component of compensation. Second, posted offers reveal these characteristics incompletely: online job posts in the United States contain very little wage information (Batra, Michaud, and Mongey, 2023), and the attributes that Norwegian job ads do advertise predict employer attractiveness (Audoly, Bhuller, and Reiremo, 2024). Third, searchers respond to what is revealed: posted wages and posting language affect application behavior (Banfi and Villena-Roldan, 2019; Gee, 2019; Marinescu and Wolthoff, 2020; Belot, Kircher, and Muller, 2022), and advertised amenities and wages shift workers’ beliefs about unlisted job attributes (Adrijaan, Balgova, Jager, Jessen, and Sockin, 2026). These facts support treating job quality as payoff relevant, incompletely posted, and worth learning about.

Related disclosure experiments show why partial posted information need not eliminate the due diligence problem: receivers are often not sufficiently skeptical about nondisclosure (Jin, Luca, and Martin, 2021), and senders can use complexity to make truthful information harder to process (Jin, Luca, and Martin, 2022). Screening also matters on the other side of hiring markets: Martin and Marx (2022) develop a test of prejudice that remains valid when hiring decision makers selectively acquire private information. Nor is the mechanism limited to labor markets. Search frictions and asymmetric information jointly shape liquidity and prices in over-the-counter asset markets (Duffie, Gârleanu, and Pedersen, 2007; Guerrieri and Shimer, 2014; Kaya and Kim, 2018), housing markets (Genesove and Han, 2012; Landvoigt, Piazzesi, and Schneider, 2015), and marriage markets, where dating involves costly learning before acceptance (Burdett and Coles, 1997; Smith, 2006). Our framework applies whenever offer acceptance is preceded by costly due diligence and offer suppliers respond strategically to the resulting acceptance rule.

2 Due Diligence and Offer Quality in a Static Setting

The static one-period model isolates the interaction between a matched firm and worker, and shows how endogenous offer quality and information acquisition jointly determine acceptance, offer composition, and market outcomes. We characterize the unique active “interior” equilibrium, in which firms post both high- and low-quality offers and workers sometimes accept and sometimes reject

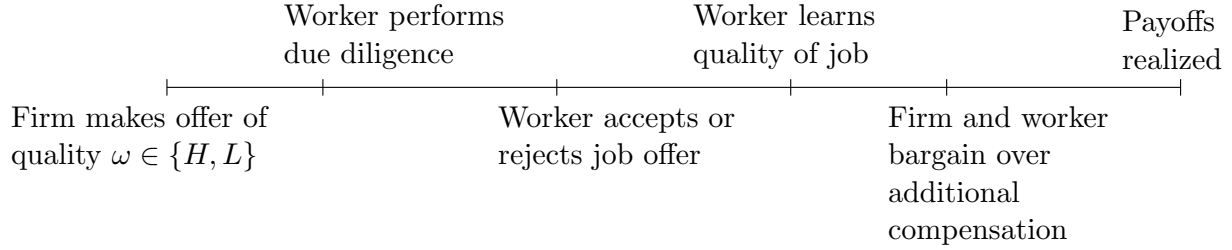


Figure 1: The timeline for the one-period version of the model.

offers. We then analyze comparative statics with respect to match values and the value of being unmatched. We frame the model in terms of a worker searching for a job with uncertain characteristics; the mechanism applies to any market in which a searcher evaluates offers of uncertain endogenous quality before accepting.

2.1 The Interaction Between Firms and Workers

There are two types of agents: a firm that makes a job offer and a worker who decides whether to accept it. There is a continuum of ex ante identical workers and firms, and we describe the behavior of a representative worker and firm. The timeline for their strategic interaction is provided in Figure 1. At the beginning of the period, the firm privately chooses the quality $\omega \in \{H, L\}$ of its offer. The quality of a job offer is determined by characteristics that are under the firm's control and costly for it to improve, while also affecting the worker's utility. Conceptually, job quality is a broad term for components of a job that matter to a worker but are not negotiated during hiring. Crucially, it can be influenced both by characteristics of a job that are known by workers (e.g., base salary, hours worked) and characteristics that are difficult to know without putting in some effort (e.g., office culture, flexibility in working hours, ability to work from home), and we assume that the former are not fully revealing of the latter. Thus, the worker might want to perform due diligence in order to learn more about the job quality before choosing to accept or reject the offer.

Assumption 1 (Contracting). Before due diligence, the firm cannot certify the hidden component of quality. It also cannot commit to compensation that is contingent on that hidden component or on future histories. The additional compensation it negotiates covers only the current period and cannot be committed in advance for later periods. Any terms observed before due diligence are either common across the offers considered in a submarket or are not fully revealing of ω .

Assumption 1 is the contracting foundation for the information friction. It does not require literal wage ignorance: workers may observe a base wage, hours, or other posted terms. What

matters is that some payoff component of quality remains residual, unverifiable, or not contractible at the acceptance stage.¹

Bargaining. We use this interaction to study the strategic interplay between endogenous quality and information acquisition. While this would be possible in a simple model with exogenous surplus splitting (perhaps conditional on job type), it is standard in the search and matching literature to include a bargaining process to determine the split of the surplus of a match (e.g., Diamond, 1982; Mortensen, 1982; Pissarides, 1985; Moscarini, 2005; Pries and Rogerson, 2005). To show that endogenous quality and information acquisition can be adapted to this standard environment, we assume that even though the firm is committed to the quality of an offer, the worker can bargain with the firm for additional compensation after accepting an offer. Because additional compensation is constrained to be nonnegative, the model with no bargaining is nested: when neither quality commands additional compensation, bargaining plays no role.

Firm and Worker Payoffs. If the worker accepts a job offer of quality ω , the firm gains J_ω . If the worker rejects the offer, the firm gets nothing. In this static setting, the firm's acceptance payoff J_ω is just the single-period profit π_ω resulting from the worker's output less the additional negotiated compensation w_ω paid to the worker:

$$J_\omega = \pi_\omega - w_\omega. \quad (1)$$

We study environments in which $\pi_L > \pi_H > 0$: providing a high-quality job is costly for the firm, but can still be profitable if the additional compensation paid to the worker is low enough.

If the worker accepts a job of quality ω , they gain W_ω . If they reject the offer, they get $W_\emptyset$, which can be interpreted as the value of being unemployed. The worker's utility from accepting an offer has two components: the nonnegotiable benefit u_ω from a job of quality ω and the negotiated additional compensation w_ω . The worker's utility from rejecting an offer is the benefit b available to the unemployed:

$$W_\omega = u_\omega + w_\omega, \quad (2)$$

$$W_\emptyset = b. \quad (3)$$

We study environments in which $u_H > b > u_L > 0$: without additional compensation, a good job is better than unemployment benefits and a bad job is worse than either. Thus, if additional compensation for bad jobs is low enough, workers have an incentive to perform due diligence in

¹The model is a theory of endogenous residual quality within pooling submarkets, not a general theory of wage posting under asymmetric information. Disclosure evidence supports the restriction: nondisclosure and complexity limit what receivers learn from observed reports (Jin, Luca, and Martin, 2021, 2022). Proposition 8 in Appendix D.2 makes the formal point: posted terms can index pooling submarkets, the mechanism applies within any cell that pools residual quality, and it disappears only in the limit of perfect separation.

order to distinguish between good and bad jobs. This captures the intuition that unemployed workers might not want to accept any job offer that comes along, but rather try to find a good job.

2.2 Due Diligence

Before accepting or rejecting an offer, the worker can exercise costly due diligence in order to learn about the quality of the job offer. We assume that these costs are increasing in how well informed the worker is about job quality after due diligence is complete. While there are many ways to model costly information acquisition, we use rational inattention theory. Our primary motivation for using this approach is that it produces analytically tractable solutions for our model in both the static and dynamic settings.²

With rational inattention theory, the worker’s uncertainty is measured using Shannon entropy. The worker starts with a prior belief μ that the job is of high quality, which is determined in equilibrium by the probability that the firm makes the job high quality. The entropy of this belief,

$$\mathbb{H}(\mu) = -\mu \ln \mu - (1 - \mu) \ln(1 - \mu), \quad (4)$$

is a nonnegative, symmetric measure of the worker’s uncertainty, strictly positive for $\mu \in (0, 1)$ and maximal at $\mu = 1/2$. The interpretation is straightforward: if μ is close to $1/2$, the worker is highly uncertain about job quality, while μ close to 0 or 1 means they are quite confident about what the quality is, and thus about the action they should take.

Under Shannon mutual information costs, signals that induce the same action can be pooled without changing expected payoffs and while weakly lowering the information cost, so an optimal signal structure uses at most one signal per action. In the active interior region the worker uses both: one signal that “recommends” accepting the offer and another signal that “recommends” rejecting the offer. These signals generate two posterior beliefs: γ is the belief that the job is high quality after receiving the signal that induces acceptance and ν is the belief that the job is high quality after receiving the signal that induces rejection. We refer to the worker’s state-contingent beliefs, γ and ν , as the acceptance and rejection posteriors. Thus, the expression

$$P\mathbb{H}(\gamma) + (1 - P)\mathbb{H}(\nu) \quad (5)$$

is the expected residual entropy after due diligence is complete, where P is the unconditional probability that the worker accepts the offer. Even though the worker acts on these beliefs, uncertainty

²See Maćkowiak, Matějka, and Wiederholt (2023) for a review of the literature, including the advantages and disadvantages of the approach: rational inattention is flexible in the forms information acquisition can take and has foundations in information theory and sequential information acquisition (Morris and Strack, 2019), while it assumes the costs of distinguishing between any two states are the same (Pomatto, Strack, and Tamuz, 2023), and agents may not adjust information acquisition in practice as much as theory suggests (Dean and Neligh, 2023). Appendix D.1 implements the same posterior and choice probability problem using normal signals with a cost function chosen to reproduce entropy reduction.

typically remains, and the posteriors determine how likely the worker is to make each type of mistake: γ gives the chance that the offers they accept are high quality, and ν gives the chance that the offers they reject are high quality.

Choosing the due diligence to perform corresponds to choosing P , γ , and ν , subject to Bayes' rule:

$$\mu = \gamma P + \nu(1 - P). \quad (6)$$

Since the rejection posterior is implied by the pair (γ, P) , we describe due diligence by that pair and collect the Bayes-plausible strategies in the set

$$\mathcal{A}(\mu) := \{(\gamma, P) \in [0, 1]^2 : \gamma P \leq \mu \leq \gamma P + 1 - P\}, \quad (7)$$

the requirement that the implied rejection posterior lie in $[0, 1]$. The information acquired through due diligence is measured by the entropy reduction, prior entropy minus expected posterior entropy:

$$\mathcal{I}(\gamma, P) = \mathbb{H}(\mu) - P\mathbb{H}(\gamma) - (1 - P)\mathbb{H}(\nu). \quad (8)$$

If the worker does not conduct any due diligence, their posteriors equal the prior and the entropy reduction is zero; it is largest when due diligence fully resolves uncertainty ($\gamma = 1$ and $\nu = 0$ imply $\mathbb{H}(\gamma) = \mathbb{H}(\nu) = 0$). It is this reduction in uncertainty that determines the cost of due diligence to the worker: the due diligence cost is $\lambda \mathcal{I}(\gamma, P)$, where $\lambda > 0$, the “unit cost of attention”, linearly scales the entropy reduction.

The posteriors γ and ν describe the worker's beliefs conditional on their action. The same due diligence can equally be described by the probabilities that the worker accepts conditional on the actual job quality, P_H and P_L , and Bayes' rule links the two descriptions:

$$P_H = \frac{\gamma P}{\mu}, \quad P_L = \frac{(1 - \gamma)P}{1 - \mu}, \quad P = \mu P_H + (1 - \mu)P_L. \quad (9)$$

The conditional acceptance probabilities summarize the extent of Type I and Type II errors that the worker can make: the worker can sometimes mistakenly reject high-quality offers, which is given by $1 - P_H$, and the worker can sometimes mistakenly accept low-quality offers, which is given by P_L . The posteriors are the natural description for the worker's information problem, and the conditional acceptance probabilities are the natural description for the firm, whose offer is accepted with probability P_ω when its quality is ω .

2.3 Additional Compensation

We assume that the bargaining between firm and worker for additional compensation follows the Nash bargaining solution, and we follow the standard assumption that the threat point in bargaining is to dissolve the match. We also assume that this bargaining occurs after the job type is revealed to the worker, so there is no asymmetric information about job type during bargaining. This assumption is key to keeping the model tractable. Even with simple bargaining protocols, the interaction between rational inattention and bargaining is both subtle and complex, as demonstrated in Ravid (2020).

We make two additional assumptions about the bargaining process to permit a nondegenerate distribution of endogenous offer quality. First, we do not let workers collect unemployment in the current period if they choose to dissolve the match. This assumption is needed for an active due diligence equilibrium. If a worker could collect unemployment benefits after dissolving the match, any low-quality match that survived bargaining would have to give the worker at least b , eliminating the loss from mistakenly accepting a low-quality offer — and with it the reason to perform due diligence. An interpretation of this assumption is that entering the wage negotiation stage prevents the worker from registering as unemployed and collecting benefits for one period. Second, because the presence of nonnegotiable job benefits introduces the possibility that additional compensation implied by surplus splitting could be negative, we assume that additional compensation cannot be negative. That is, the firm cannot extract money from the worker for the right to work.

With these assumptions, additional compensation is determined by maximizing the following Nash product for a given $\omega \in \{H, L\}$:

$$w_\omega \in \arg \max_{0 \leq w \leq \pi_\omega} (W_\omega(w) - (W_\emptyset - b))^\psi (J_\omega(w))^{1-\psi}, \quad (10)$$

where $W_\omega(w)$ and $J_\omega(w)$ denote the payoffs at candidate compensation w , $\psi \in (0, 1)$ is interpreted as the worker's bargaining power, and the domain $0 \leq w \leq \pi_\omega$ keeps firm profits nonnegative so that the Nash product is well defined. The unconstrained first-order condition for this problem is:

$$\psi (J_\omega(w)) = (1 - \psi) (W_\omega(w) - (W_\emptyset - b)). \quad (11)$$

If the compensation that solves this first-order condition is negative, then the compensation that maximizes the constrained Nash product is 0, so the candidate compensation is

$$w_\omega^0 = \max\{\psi\pi_\omega - (1 - \psi)u_\omega, 0\}, \quad \omega \in \{H, L\}. \quad (12)$$

In our comparative static results, we will indicate the model parameter values under which this additional compensation can influence the strategic interplay between endogenous quality and information acquisition. The object w_ω should be read as additional compensation negotiated

after the worker has decided to pursue the match, not as the entire wage package. Thus the case $w_H^* = 0 < w_L^*$ does not mean that high-quality jobs pay zero total wages. It means that, conditional on the nonnegotiable job attributes already embedded in u_H , high-quality jobs do not need an extra compensating transfer, while low-quality jobs may require one once their quality is revealed. This is the same economic logic as compensating differentials, but applied only to the residual component of compensation that remains negotiable after acceptance.

2.4 Worker and Firm Problems

Given the prior μ and the payoffs $(W_H, W_L, W_\emptyset)$, the worker's due diligence problem is:

$$\max_{(\gamma, P) \in \mathcal{A}(\mu)} P(\gamma W_H + (1 - \gamma)W_L) + (1 - P)W_\emptyset - \lambda \mathcal{I}(\gamma, P). \quad (13)$$

A firm chooses the probability μ of making its offer high quality, taking the worker's behavior as given. Given the conditional acceptance probabilities (P_H, P_L) and the profits (J_H, J_L) , the firm's quality problem is:

$$\max_{\mu \in [0, 1]} \mu P_H J_H + (1 - \mu)P_L J_L. \quad (14)$$

2.5 Equilibrium

To determine market outcomes, we use a Rational Inattention Bayesian Nash Equilibrium (RIBNE) solution concept. In an RIBNE, the equilibrium firm strategy μ^* is the probability that the firm offers a high-quality job. On the other side of the market, the equilibrium worker attention strategy is a pair (γ^*, P^*) , where γ^* is the posterior belief of high quality in equilibrium when accepting offers and P^* is the unconditional probability that the worker accepts a job offer.

Definition 1. *A static Rational Inattention Bayesian Nash Equilibrium (RIBNE) consists of:*

1. *Worker attention strategy $(\gamma^*, P^*) \in [0, 1] \times (0, 1)$ with active due diligence ($\gamma^* \neq \nu^*$)*
2. *Firm offer quality strategy $\mu^* \in (0, 1)$*
3. *Additional compensation $w_H^*, w_L^* \in \mathbb{R}_+$*

such that:

1. *(γ^*, P^*) solves the worker's due diligence problem (13) at prior μ^* and the payoffs implied by (w_H^*, w_L^*) .*
2. *μ^* solves the firm's quality problem (14) at (P_H^*, P_L^*) and the profits implied by (w_H^*, w_L^*) .*

3. Each additional compensation w_ω^* for $\omega \in \{H, L\}$ solves the Nash bargaining problem (10), so that $W_\omega = W_\omega(w_\omega^*)$ and $J_\omega = J_\omega(w_\omega^*)$.

The rejection posterior ν^* is implied by (γ^*, P^*, μ^*) , and the conditional acceptance probabilities P_H^* and P_L^* are given by Bayes' rule (9).

Definition 1 builds the active interior requirements — $\mu^* \in (0, 1)$, $P^* \in (0, 1)$, and $\gamma^* \neq \nu^*$ — directly into the equilibrium concept. Thus, both high- and low-quality jobs are offered, job offers are sometimes accepted and sometimes rejected, and due diligence is active.³

For an interior RIBNE to exist, the Nash bargaining payoffs must put the model in the region where workers want to distinguish job qualities and firms are willing to mix between them. We state this region directly.

Assumption 2 (Static active interiority). The static primitives satisfy

$$\begin{aligned} w_H^0 = 0, \quad W_H(w_H^0) > b > W_L(w_L^0), \\ J_L(w_L^0) > J_H(w_H^0) > 0, \quad \frac{J_L(w_L^0)}{J_H(w_H^0)} < \exp\left(\frac{W_H(w_H^0) - W_L(w_L^0)}{\lambda}\right). \end{aligned} \quad (15)$$

The ordering $W_H(w_H^0) > b > W_L(w_L^0)$ places the value of unemployment between the high- and low-quality worker payoffs after bargaining. This ordering ties the worker's best action to job quality, whatever the offer mix: only high-quality jobs are worth accepting, and due diligence is how the worker tells them apart. The ordering $J_L(w_L^0) > J_H(w_H^0) > 0$ makes both qualities profitable when accepted, and low quality strictly more so. This ranking gives firms a reason to post low-quality offers. A worker performing due diligence accepts high-quality offers more often, whatever quality mix firms offer, so firms are willing to post both qualities only if low-quality offers earn strictly more when accepted. Each ordering is thus necessary for one side of the market on its own, regardless of the other side's strategy.

The bound on the profit ratio, $\frac{J_L(w_L^0)}{J_H(w_H^0)} < \exp\left(\frac{W_H(w_H^0) - W_L(w_L^0)}{\lambda}\right)$, is instead a joint restriction across the two sides of the market, and it is what keeps the two forces balanced at an interior equilibrium. For firms to be willing to post both qualities, workers must accept high-quality offers exactly J_L/J_H times as often as low-quality ones, offsetting the profit advantage of low quality with an acceptance advantage for high quality. Optimal due diligence delivers only limited discrimination. At any prior, the worker accepts high-quality offers fewer than $\exp((W_H - W_L)/\lambda)$ times as often as low-quality ones, because the payoff gap between the qualities, measured in units of the attention cost, caps the discrimination worth paying for. Sharper discrimination comes only with

³The model also admits boundary outcomes in which the mechanism is inactive — for example, workers who expect only low-quality offers and reject without acquiring information, leaving no firm with a profitable deviation. We condition on the active interior equilibrium throughout and do not characterize the boundary outcomes: they fail the active interior requirements by construction, so they are outside the analysis rather than ruled out by it.

rarer acceptance, so the acceptance probability that delivers the ratio firms need is interior exactly when that ratio lies between one and the cap: $J_L > J_H$ places it above one, and the bound places it below the cap.

The equality $w_H^0 = 0$ is what allows these orderings to survive bargaining. The orderings above have the two sides ranking the qualities oppositely — workers prefer high-quality jobs, firms profit more from accepted low-quality jobs — and interior bargaining would eliminate this disagreement. An interior high-quality bargain, $w_H^0 > 0$, would force the low-quality bargain interior as well, since the low-quality job has the higher profit and the lower nonnegotiable payoff ($\pi_L > \pi_H$ and $u_L < u_H$). With bargaining interior for both qualities, worker payoffs and firm profits would be proportional shares of total match surplus, $W_\omega(w_\omega^0) = \psi(\pi_\omega + u_\omega)$ and $J_\omega(w_\omega^0) = (1 - \psi)(\pi_\omega + u_\omega)$, and both sides would agree on which quality is better. The quality workers accept more often would then also be the more profitable one, and firms would post only that quality. The corner $w_H^0 = 0$ blocks this alignment and lets the two sides of the market rank the qualities differently.⁴

The worker's optimal attention strategy depends on the payoffs only through their exponentials in units of the attention cost: $z_\omega := \exp\left(\frac{W_\omega}{\lambda}\right)$ for $\omega \in \{H, L\}$ and $z_\emptyset := \exp\left(\frac{W_\emptyset}{\lambda}\right) = \exp\left(\frac{b}{\lambda}\right)$. Equilibrium objects therefore take their simplest form in these units. Under Assumption 2, $z_H > z_\emptyset > z_L > 0$ and $z_H J_H > z_L J_L$.

Theorem 1 (Equilibrium Existence and Uniqueness). *Maintain Assumption 1. An interior RIBNE exists if and only if the static primitives satisfy Assumption 2, and it is unique. The equilibrium is:*

$$w_H^* = w_H^0 = 0, \quad w_L^* = w_L^0, \quad (16)$$

$$P^* = \frac{z_\emptyset(z_H J_H - z_L J_L)}{z_\emptyset(z_H J_H - z_L J_L) + z_H z_L (J_L - J_H)}, \quad P_H^* = \frac{z_H J_H - z_L J_L}{J_H(z_H - z_L)}, \quad \frac{P_L^*}{P_H^*} = \frac{J_H}{J_L}, \quad (17)$$

$$\mu^* = \frac{z_H J_H (z_\emptyset - z_L)}{z_\emptyset(z_H J_H - z_L J_L) + z_H z_L (J_L - J_H)}, \quad \gamma^* = \frac{z_H(z_\emptyset - z_L)}{z_\emptyset(z_H - z_L)}, \quad \nu^* = \frac{z_\emptyset - z_L}{z_H - z_L}. \quad (18)$$

Proof. See Appendix A. □

Only the first row of the displayed characterization uses the bargaining technology. The final two rows depend on the environment only through the payoff values: for any payoffs satisfying $W_H > W_\emptyset > W_L$, $J_L > J_H > 0$, and $z_H J_H > z_L J_L$, they characterize the unique active attention and quality block.

The proof of this theorem builds on the idea that for an interior equilibrium to exist, the firm has to mix between posting low- and high-quality jobs, so it must be indifferent between posting low- and high-quality jobs. Further, for the firm to be indifferent, it must be that $P_H J_H = P_L J_L$. Under Assumption 2, firms earn a higher profit if a low-quality job is filled than if a high-quality

⁴Assumption 2 is stated in terms of the Nash bargaining payoff objects. Corollary 1 in Appendix A translates it into explicit sufficient conditions on primitives, for both the $w_L^* = 0$ and the $w_L^* > 0$ bargaining regime.

job is filled ($J_L > J_H$). For firms to be indifferent to offering these two job qualities, it must then be the case that in equilibrium the probability of filling a low-quality job is lower than that of filling a high-quality job $P_L < P_H$. While there is a continuous set of such P_H and P_L that makes the firm indifferent, just P_H^* and P_L^* are supported in equilibrium. This is because with rational inattention, whenever due diligence is active the worker's optimal posterior beliefs γ^* and ν^* do not depend on their prior, which is given by the firm's mixing rate between high- and low-quality job offers.⁵ But through Bayes' rule, the mixing rate μ and beliefs γ^* and ν^* determine P , P_H , and P_L . We show that there is exactly one mixing rate, μ^* , at which γ^* and ν^* generate conditional acceptance probabilities P_H and P_L that make the firm indifferent; these are P_H^* and P_L^* .

2.6 Equilibrium Counterfactuals

We study two economic consequences of equilibrium behavior: unemployment and welfare. Unemployment is the probability that the worker does not accept the job offer:

$$\mathcal{U} = 1 - P. \quad (19)$$

We define worker welfare as the worker's expected utility in equilibrium, inclusive of the cost of due diligence:

$$\mathcal{W} = P(\gamma W_H + (1 - \gamma)W_L) + (1 - P)W_\emptyset - \lambda \mathcal{I}(\gamma, P). \quad (20)$$

Appendix A.8 analyzes an alternative measure that excludes attention costs to ease comparison with search and matching models that feature no learning cost, and shows that every comparative static sign is the same for both measures.

Every comparative static concerns the unique interior equilibrium of Theorem 1, and can be read as a statement about how (u_L, u_H, b) move the exponentiated payoffs $(z_H, z_L, z_\emptyset)$ and the profit ratio J_L/J_H , each a known, monotone function of parameters. The results are derivative statements that hold at every interior equilibrium, so they also give the direction of large changes provided the equilibrium remains interior along the path.⁶

The market outcomes work through three equilibrium responses: the posted mix of offers, the action-contingent posteriors, and the conditional acceptance rates. We sign these responses first.

Proposition 1. *In the unique interior equilibrium:*

⁵The local invariance of posteriors in rationally inattentive decision making is established by Caplin, Dean, and Leahy (2022), and the usefulness of this property in solving for equilibria is demonstrated by Martin (2017) and Porcher (2025).

⁶Derivatives are not defined at the knife edge $\psi\pi_L = (1 - \psi)u_L$, where the negotiated wage w_L^* hits zero; every equilibrium object is continuous there and the unemployment and welfare signs are the same on both sides, so even large changes that cross this point move unemployment and welfare in the stated directions.

1. an increase in u_L strictly decreases μ^* , γ^* , ν^* , P_H^* , and P_L^* .
2. an increase in u_H strictly decreases μ^* , γ^* , and ν^* , and strictly increases P_H^* and P_L^* .
3. an increase in b strictly increases μ^* , γ^* , and ν^* , and leaves P_H^* and P_L^* exactly unchanged.

Proof. See Appendix A. □

Unemployment and welfare respond as follows.

Proposition 2. *In the unique interior equilibrium:*

1. Increasing the quality of low-quality jobs, u_L , increases unemployment and decreases worker welfare.
2. Increasing the quality of high-quality jobs, u_H , decreases unemployment and increases worker welfare.
3. Increasing the unemployment benefit, b , decreases unemployment and increases worker welfare.

Proof. See Appendix A. □

Table 1 summarizes the static comparative statics.

Outcome	$\Delta u_L > 0$	$\Delta u_H > 0$	$\Delta b > 0$
Unemployment $\mathcal{U} = 1 - P^*$	+	−	−
Welfare $\mathcal{W}$	−	+	+
Fraction high-quality offers μ^*	−	−	+
Acceptance/rejection posteriors γ^*, ν^*	−	−	+
Acceptance rate of H offers P_H^*	−	+	0
Acceptance rate of L offers P_L^*	−	+	0

Notes: (+) denotes an increase, (−) a decrease, (0) no change, and ($\pm$) an ambiguous sign.

Table 1: Comparative statics of the unique interior equilibrium of the static model (Propositions 1 and 2).

Improving low-quality offers (Δu_L). We first examine the comparative statics for changes in the quality of the bad job, u_L .

Firms are indifferent between posting the two job qualities in the interior equilibrium. Good jobs are accepted more often, while bad jobs are more profitable when accepted, so indifference

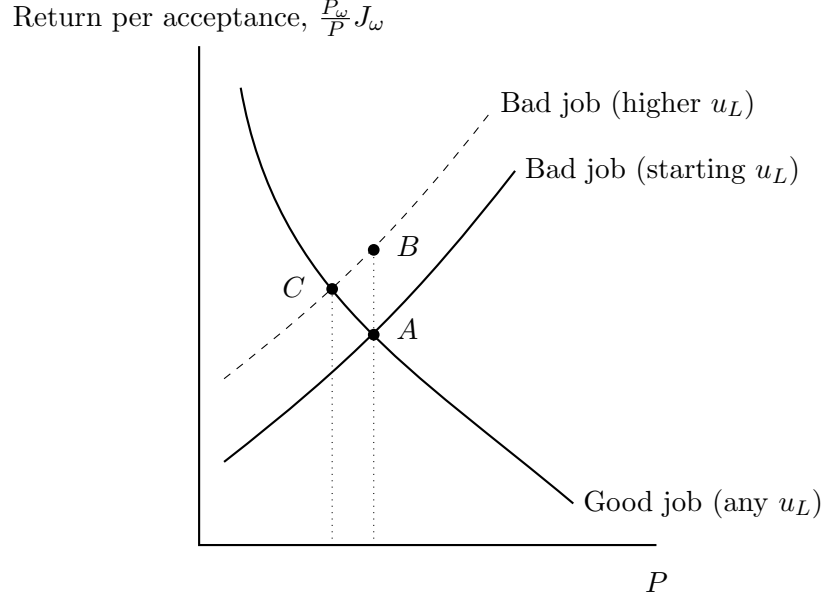


Figure 2: Illustration of changes in the firm's return per acceptance from posting an offer of quality $\omega \in \{H, L\}$, $\frac{P_\omega}{P} J_\omega$, and of the equilibrium acceptance rate P , due to changes in u_L .

trades an acceptance advantage against a profit advantage: $P_H J_H = P_L J_L$. Dividing by the overall acceptance rate P and substituting the worker's optimal conditional acceptance rates gives

$$\frac{P_H}{P} J_H = \frac{P_L}{P} J_L, \quad \text{where} \quad \frac{P_\omega}{P} = \frac{z_\omega}{z_\omega P + z_\emptyset(1 - P)}. \quad (21)$$

Each side of the indifference condition is the return to posting one quality per unit of acceptance. The second equation is the worker's behavior: the worker accepts a quality more readily the larger its exponentiated payoff z_ω relative to the exponentiated unemployment payoff $z_\emptyset$. Figure 2 plots these two returns against P . A worker who accepts more often discriminates less between qualities: P_H rises less than proportionally with P , and P_L rises more than proportionally. The good job curve therefore falls as P rises, and the bad job curve rises. The curves cross once, at point A, the unique interior equilibrium. There the acceptance advantage of good jobs exactly offsets the profit advantage of bad jobs.

An increase in u_L raises the return to posting a bad job at any fixed acceptance rate, because better bad jobs pass due diligence more often. The return to posting a good job at fixed P is unchanged, so posting bad jobs now strictly dominates (point B). Firms respond by shifting the posted mix toward low quality: μ^* falls. A higher u_L also narrows the worker's payoff gap between the two qualities. Both action-contingent posteriors fall, and the worker accepts each quality less often: γ^* , ν^* , P_H^* , and P_L^* all fall. The worse mix and the lower conditional acceptance rates pull the unconditional acceptance rate P down to point C, where firm indifference is restored. Since

unemployment is $1 - P$, improving low-quality jobs raises unemployment. Welfare falls as well: the deterioration in the mix of offers outweighs the direct benefit of better low-quality jobs.

Bargaining determines how the increase in u_L is split. If $w_L^* = 0$, the worker keeps the full increase. If $w_L^* > 0$ both before and after the change, bargaining passes part of it to the firm. In this case, the worker keeps the fraction ψ and the firm's profit J_L rises, which reinforces the shift toward low-quality offers. In either case, the two regimes agree on the direction of counterfactual change.

Improving high-quality offers (Δu_H). The comparative statics for changes in u_H follow a similar logic but move in the opposite directions.

When u_H rises, the stakes in due diligence rise asymmetrically: the loss from mistakenly rejecting a good job grows, while the loss from mistakenly accepting a bad one is unchanged. The worker responds by lowering both the acceptance and rejection posteriors γ^* and ν^* : offers are accepted on weaker evidence, so mistaken rejections of good jobs ($1 - P_H^*$) become rarer while mistaken acceptances of bad jobs (P_L^*) become more frequent. One might have guessed that a more valuable good job would lead firms to post more good jobs. Instead, because bad offers now slip through due diligence more easily, posting them becomes relatively more attractive, and restoring the firms' indifference condition $P_H J_H = P_L J_L$ requires the fraction of high-quality offers μ^* to fall. Making good jobs better makes them scarcer. The direct effect nevertheless dominates in the aggregate: the overall acceptance rate P^* rises, so unemployment falls, and welfare rises.

Raising unemployment benefits (Δb). Next we examine the comparative statics for changes in b , which is the worker's payoff from remaining unmatched. In other models of search, such as the classic model of McCall (1970), an increase in the unemployment payoff leads to an increase in unemployment due to workers becoming more selective about jobs. With due diligence and endogenous offer quality, the equilibrium response is reversed.

The intuition behind the reversal runs through offer composition. A higher b raises the posterior confidence associated with acceptance (γ^* rises) but leaves both conditional acceptance probabilities exactly unchanged in equilibrium. Firms respond by posting more high-quality offers, so the overall acceptance probability increases and unemployment decreases. Looking back at equation (21), an increase in b raises $z_\emptyset$, so the denominator of P_ω/P rises for both qualities. Evaluated at the initial equilibrium acceptance rate P^* , both per acceptance returns shift down. The return to posting a bad job shifts down by more: the two returns are equal at the crossing, and the bad job's denominator is smaller because $W_L < W_H$, so the same rise in $z_\emptyset$ moves it proportionally more. In words, the profitability of posting low-quality jobs decreases more because workers become more worried about mistakenly accepting low-quality offers than mistakenly rejecting high-quality offers.

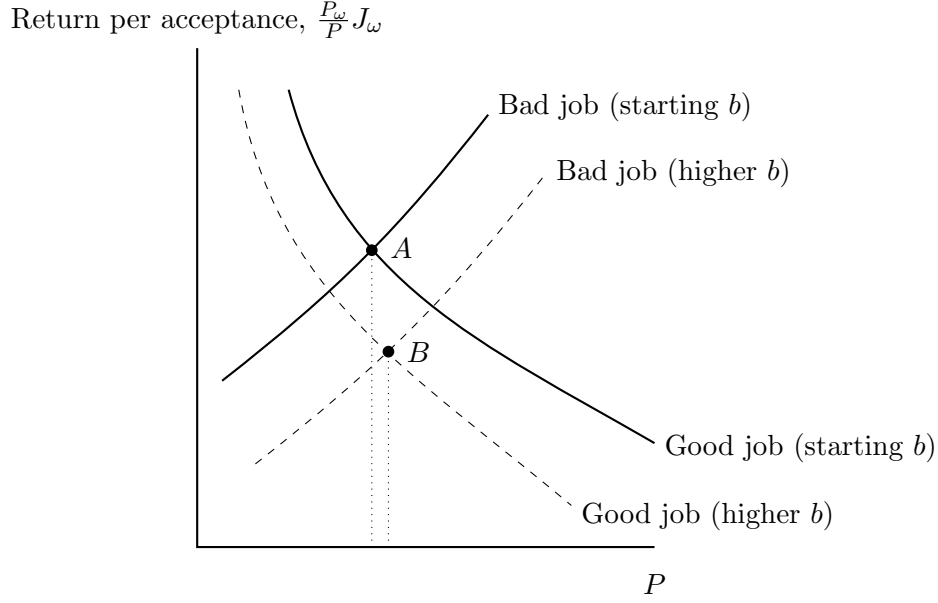


Figure 3: Illustration of changes in the firm's return per acceptance from posting an offer of quality $\omega \in \{H, L\}$, $\frac{P_\omega}{P} J_\omega$, and of the equilibrium acceptance rate P , due to changes in b .

As displayed in Figure 3, both curves drop, with the curve for bad jobs dropping by more near the initial crossing. To restore equilibrium, the unconditional acceptance probability must rise.

The conditional acceptance rates are exactly unchanged. A change in b raises $z_\emptyset$ and leaves z_H , z_L , J_H , and J_L unchanged. The conditional acceptance rates in Theorem 1 do not involve $z_\emptyset$, so they do not move at all. In the static model, unemployment benefits affect acceptance entirely through the composition of offers. The whole increase in P^* comes from firms posting more good jobs, not from workers accepting more readily conditional on quality.

Summary. The model with endogenous information acquisition and offer quality differs in core counterfactual predictions from the standard exogenous information and offer quality model. In the exogenous information and quality model, improving good and bad quality jobs increases welfare and lowers unemployment, while raising unemployment benefits increases unemployment. In contrast, with due diligence and endogenous offer quality, improving low-quality offers raises unemployment and lowers worker welfare, because the offer quality distribution deteriorates. Improving high-quality offers helps workers despite a worse offer quality mix. Raising unemployment benefits lowers unemployment entirely through a better offer mix. With the static mechanism in hand, we turn to the dynamic market.

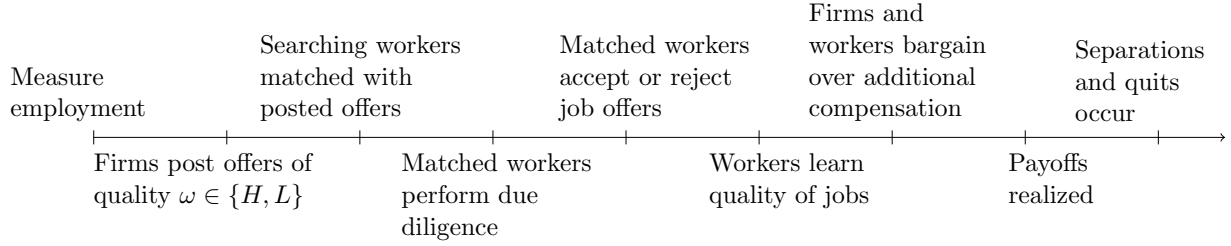


Figure 4: The timeline for the dynamic version of the model.

3 Due Diligence in a Dynamic Setting

The dynamic model embeds the matched firm and worker in a search and matching market. Figure 4 presents the timing of the model. At the beginning of each period, firms enter by posting new job offers. Unemployed workers decide whether to search, and some searching workers are matched with offers. Matched workers perform due diligence and then decide whether to accept the offer given their updated beliefs about job quality. If they accept, the job quality is revealed, surplus bargaining occurs, and then payoffs are realized. Finally, separations and quits occur. A matched worker faces the same due diligence, acceptance, and bargaining protocol as in the static model, but continuation values now enter worker and firm payoffs. What is new is the worker's outside option, the meeting probabilities, and the stocks of employed and unemployed workers, which are now determined in equilibrium through matching and the free entry of vacancies. We solve for the unique interior stationary equilibrium of this dynamic market.

3.1 Search and Matching

There is a unit measure of infinitely lived workers. A portion E_ω of workers are employed in jobs of quality $\omega \in \{H, L\}$ and a portion $\mathcal{U}$ are in the search pool, so that

$$\mathcal{U} + E_H + E_L = 1. \quad (22)$$

A worker who enters the period employed skips that period's search and matching stages: the match renegotiates its additional compensation, they work, payoffs are realized, and they then face the separation and quit stages. Each job ends with exogenous probability $\delta \in (0, 1)$, and, conditional on the job surviving, the worker chooses whether to quit. We measure employment status at the beginning of the period, so either exit ends the job and sends the worker to the next period's search pool. Workers in the search pool can pay a fixed cost $C > 0$ to search for an offer.⁷

⁷For simplicity, employed workers cannot search for new jobs.

Firms enter by paying a fixed cost $F > 0$ to post a job offer that lasts one period. Entry determines the measure of vacancies $\mathcal{V}$ present when job search takes place. As in the static model, each firm privately chooses the quality of its offer.

We use a standard Cobb-Douglas matching technology to determine the probability that searching workers are matched with job offers.⁸ The measure of realized matches is

$$M(\mathcal{U}, \mathcal{V}) = \min \left\{ A\mathcal{V}^{1-\phi}\mathcal{U}^\phi, \mathcal{U}, \mathcal{V} \right\}, \quad A > 0, \quad \phi \in (0, 1). \quad (23)$$

The caps inside the min ensure that there are never more matches than searchers or vacancies. In the interior equilibria we study, the caps do not bind, so $M(\mathcal{U}, \mathcal{V}) = A\mathcal{V}^{1-\phi}\mathcal{U}^\phi$ and ϕ is the elasticity of matches with respect to unemployment.

The firm and worker match probabilities are

$$\alpha_f := \frac{M(\mathcal{U}, \mathcal{V})}{\mathcal{V}}, \quad \alpha_w := \frac{M(\mathcal{U}, \mathcal{V})}{\mathcal{U}}. \quad (24)$$

Defining market tightness $\theta := \mathcal{V}/\mathcal{U}$, the match probabilities become $\alpha_f = A\theta^{-\phi}$ and $\alpha_w = A\theta^{1-\phi}$. A measures matching efficiency, the worker's offer meeting probability rises with tightness with elasticity $1 - \phi$, and the vacancy filling probability falls with elasticity $-\phi$.

3.2 Worker's Problem

Matched workers choose due diligence, accept or reject the offer, and, upon acceptance, negotiate additional compensation after job quality is revealed. The period payoff from a job of quality ω is $u_\omega + w_\omega$ and the period payoff from rejecting an offer is the unemployment benefit b . Acceptance commits the worker to the match for the current period: a worker who accepts cannot search again or collect b during that period. This commitment is what makes mistakenly accepting a bad job costly, and hence due diligence valuable.

Economically, a low-quality job is a bad match discovered during an initial or probationary work period — a misleading schedule, unexpectedly poor management, or hidden benefit limitations, discovered only after the worker has committed to the initial or probationary period. We maintain $u_H > b > u_L > 0$, so the intrinsic flow payoff of a high-quality job lies above unemployment and that of a low-quality job lies below it.

However, because the match can persist into the next period and unemployed workers can choose to search in the next period, the total dynamic payoff from accepting a job of quality $\omega \in \{H, L\}$, W_ω , and the total payoff from rejecting an offer, $W_\emptyset$, differ from the static model due to the addition of continuation values. Workers discount future utility by the factor $\beta \in (0, 1)$, and there are three

⁸See, e.g., Shimer (2005) for why a constant returns Cobb-Douglas matching function describes U.S. data well.

possible continuation values. We define V_H and V_L as the value functions of a worker employed in a job of known high or low quality, and $V_\emptyset$ as the value function of a jobless worker. Based on these value functions, the discount factor β , and the separation rate δ , W_H , W_L , and $W_\emptyset$ are as follows:

$$W_H := u_H + w_H + \beta(1 - \delta)V_H + \beta\delta V_\emptyset, \quad (25)$$

$$W_L := u_L + w_L + \beta(1 - \delta)V_L + \beta\delta V_\emptyset, \quad (26)$$

$$W_\emptyset := b + \beta V_\emptyset. \quad (27)$$

The V 's and W 's sit at different decision nodes. Each W_ω is the value of a committed period of type ω work, entered by accepting an offer or by staying rather than quitting, and $W_\emptyset$ is the value of rejecting an offer. Each V is measured when its holder needs to make a decision: $V_\emptyset$ when the jobless worker chooses whether to search, and V_H and V_L when the employed worker, with payoffs realized and the job having survived the separation shock, chooses whether to quit. Staying commits another period of work at that period's renegotiated compensation, just as acceptance does, and quitting avoids it. The W 's are the same objects the worker's due diligence problem compares in the static model, now each carrying a continuation value.

The continuation values V_H and V_L condition on known quality because, just as in the static model, uncertainty is resolved after acceptance and before bargaining, which preserves the complete information Nash product representation.

Based on these functions, the unemployed worker's problem is:

$$V_\emptyset = \max \left\{ W_\emptyset, \max_{(\gamma, P) \in \mathcal{A}(\mu)} \alpha_w P (\gamma W_H + (1 - \gamma) W_L) + (1 - \alpha_w P) W_\emptyset - \alpha_w \lambda \mathcal{I}(\gamma, P) - C \right\}. \quad (28)$$

The first argument of the outer max is the value of not searching this period. The second is the expected net present value of searching. With probability $\alpha_w P$ the worker meets and accepts an offer, which is worth $\gamma W_H + (1 - \gamma) W_L$ in expectation; otherwise the worker remains jobless, with payoff $W_\emptyset$. A matched worker pays the attention cost $\lambda \mathcal{I}(\gamma, P)$, and every searcher pays the search cost C . When the searching branch is selected, the attention strategy (γ, P) is a maximizer of the inner problem.

Because job quality is known after acceptance, a worker employed in a job of type ω keeps the job when $W_\omega > V_\emptyset$ and quits when $W_\omega < V_\emptyset$, and may mix at equality. Let η_ω denote the probability of quitting. Thus,

$$\eta_\omega \in \arg \max_{\eta \in [0, 1]} (1 - \eta) W_\omega + \eta V_\emptyset, \quad V_\omega = \max\{W_\omega, V_\emptyset\}. \quad (29)$$

3.3 Firm's Problem

In the dynamic model, J_ω is the expected net present value for a firm when a type- ω offer is accepted and production begins:

$$J_H := \pi_H - w_H + \beta(1 - \delta)(1 - \eta_H)J_H, \quad (30)$$

$$J_L := \pi_L - w_L + \beta(1 - \delta)(1 - \eta_L)J_L. \quad (31)$$

We study environments in which the dynamic profits satisfy $J_L > J_H$: firms have an incentive to post low-quality offers even though they are accepted less often.

An entering firm pays the posting cost F , is matched with a searching worker with probability α_f , and chooses the probability μ that its offer is high quality, so the firm's problem is:

$$\max_{\mu \in [0,1]} \alpha_f (\mu P_H J_H + (1 - \mu) P_L J_L) - F. \quad (32)$$

For the firm to post both qualities, it must be indifferent between them:

$$P_H J_H = P_L J_L. \quad (33)$$

With the free entry of firms that post both types of offers, this indifference condition requires that

$$\alpha_f P_H J_H = F = \alpha_f P_L J_L \Rightarrow \quad (34)$$

$$\alpha_f = \frac{F}{P_H J_H} = \frac{F}{P_L J_L}. \quad (35)$$

If the expected payoff to a firm of offering either type of job were greater than the cost of entry, then new firms would enter the market, which would lower the expected return through the matching technology.

Through the matching probabilities α_f and α_w , market tightness plays a key role in equilibrium. If new firms entered the market, tightness θ would rise, which would drive down a firm's expected payoff by lowering the vacancy filling probability and would drive up the value of search for workers by raising the meeting probability.

3.4 Bargaining

Bargaining again occurs after job quality is revealed. Because the firm cannot commit additional compensation in advance (Assumption 1), a match that continues renegotiates each period. The worker first commits to the period, the pair then negotiates that period's additional compensation, and dissolution remains the threat point. A candidate compensation w affects only current-period payoffs. All future negotiations, quit strategies, and market outcomes are held fixed at their

stationary-equilibrium values and are independent of the match's history:

$$W_\omega(w) := u_\omega + w + \beta(1 - \delta)V_\omega + \beta\delta V_\emptyset, \quad J_\omega(w) := \pi_\omega - w + \beta(1 - \delta)(1 - \eta_\omega)J_\omega. \quad (36)$$

The negotiated compensation solves the Nash bargaining problem (10) at these candidate values and the threat point $W_\emptyset - b$, and the equilibrium compensation is a fixed point, $W_\omega = W_\omega(w_\omega^*)$ and $J_\omega = J_\omega(w_\omega^*)$. Because the current compensation is sunk before the next quit decision, the quit strategies do not depend on the candidate.

3.5 Dynamic Equilibrium Definition

A stationary equilibrium of the dynamic infinite-horizon model requires optimality for workers and firms, Bayes-consistent beliefs, market clearing, and stationary employment stocks.

Stationarity of the employment stocks requires that, after separations, quits, and new acceptances, each stock reproduces itself:

$$E_H = (E_H + \mathcal{U}\alpha_w\mu P_H)(1 - \delta)(1 - \eta_H), \quad (37)$$

$$E_L = (E_L + \mathcal{U}\alpha_w(1 - \mu)P_L)(1 - \delta)(1 - \eta_L). \quad (38)$$

Although implied by the employment stocks and the population identity, $\mathcal{U} + E_H + E_L = 1$, the unemployment level can be represented similarly, from this period's searchers who did not form a lasting match together with the employed who separate or quit:

$$\begin{aligned} \mathcal{U} = & \mathcal{U}(1 - \alpha_w P) + (E_H + \mathcal{U}\alpha_w\mu P_H)[1 - (1 - \delta)(1 - \eta_H)] \\ & + (E_L + \mathcal{U}\alpha_w(1 - \mu)P_L)[1 - (1 - \delta)(1 - \eta_L)]. \end{aligned} \quad (39)$$

Definition 2. A Stationary Random Search Rational Inattention Equilibrium (SRSRIE) consists of a worker attention strategy $(\gamma^*, P^*) \in [0, 1] \times (0, 1)$ with active due diligence ($\gamma^* \neq \nu^*$); worker values $V_H^*, V_L^*, V_\emptyset^* \in \mathbb{R}$; worker quit strategies $\eta_H^*, \eta_L^* \in [0, 1]$; a firm offer quality strategy $\mu^* \in (0, 1)$; additional compensation $w_H^*, w_L^* \in \mathbb{R}_+$; and unemployment, vacancy, and employment stocks $\mathcal{U}^* > 0$, $\mathcal{V}^* > 0$, and $E_H^*, E_L^* \in \mathbb{R}_+$, such that:

1. The unemployed worker's problem (28) is solved by $V_\emptyset^*$ and (γ^*, P^*) , with searching strictly optimal, $V_\emptyset^* > W_\emptyset^*$, and the employed worker's problem (29) is solved by V_H^* , V_L^* , η_H^* , and η_L^* .
2. The firm's problem (32) is solved by μ^* , and entry satisfies the free entry condition (34).
3. Each additional compensation level w_ω^* for $\omega \in \{H, L\}$ solves the per-period Nash bargaining problem (10) at the candidate values (36), so that $W_\omega^* = W_\omega(w_\omega^*)$ and $J_\omega^* = J_\omega(w_\omega^*)$.

4. *Beliefs, payoffs, and stocks are consistent at the starred values: the rejection posterior ν^* is implied by (γ^*, P^*, μ^*) , and P_H^* and P_L^* are given by Bayes' rule (9); the payoffs satisfy the value recursions (25)–(27) and (30)–(31); the matching probabilities satisfy (24) at the starred stocks and are interior, $\alpha_f^*, \alpha_w^* \in (0, 1)$; and the stocks satisfy the population identity (22) and the stationarity conditions (37)–(38).*

3.6 Characterizing the Dynamic Equilibrium

We now characterize the interior SRSRIE. The characterization states participation requirements and determines worker values in terms of the reservation flow value $\hat{w} := (1 - \beta)V_\emptyset$, the flow equivalent of the value of unemployment. It encompasses not only unemployment benefits but also the option value of matching, net of search and attention costs; equivalently, $\hat{w}$ is the lowest constant per period payment, paid forever without termination or the option to quit, that an unemployed worker would accept in place of searching.

For this equilibrium to exist, the following restrictions must hold.

Assumption 3 (Dynamic active interiority).

1. $\pi_L - w_L^0 > \frac{\pi_H}{1 - \beta(1 - \delta)}$.
2. $b > u_L + w_L^0$.
3. $\psi\pi_H < (1 - \psi)(u_H - \beta(1 - \delta)\hat{w})$, which ensures $w_H^* = 0$.

Part 1 makes a filled bad job more profitable than a filled good job, giving firms a reason to post both qualities. Part 2 ensures that workers quit bad jobs at the first opportunity. The inequality $b > u_L + w_L^0$ gives $W_L < W_\emptyset$; together with strict search participation ($V_\emptyset > W_\emptyset$), it implies that a worker who accepted a low-quality offer works it for the current period and quits before the next round of matching. Both parts are stated at the low-quality candidate compensation w_L^0 , which may be zero. Part 3 ensures that, at the candidate reservation flow value, the unconstrained Nash transfer for a high-quality match is negative, so the nonnegativity constraint binds and $w_H^* = 0$.

The dynamic model inherits the static attention and quality block. Conditional on searching, the maximand in the search branch of (28) is an increasing affine transformation of the static due diligence objective (13):

$$\begin{aligned} & \alpha_w P (\gamma W_H + (1 - \gamma) W_L) + (1 - \alpha_w P) W_\emptyset - \alpha_w \lambda \mathcal{I}(\gamma, P) - C \\ &= \alpha_w [P (\gamma W_H + (1 - \gamma) W_L) + (1 - P) W_\emptyset - \lambda \mathcal{I}(\gamma, P)] + (1 - \alpha_w) W_\emptyset - C. \end{aligned}$$

Since $\alpha_w > 0$, the attention choice solves the static problem at the dynamic payoffs. Likewise, the entrant's objective (32) is α_f times the static quality objective (14) minus the posting cost,

and since $\alpha_f > 0$ and F does not depend on the offer mix, the quality choice solves the static problem at the dynamic values J_H and J_L . What the dynamic model determines is the payoff levels themselves, together with search participation, quitting, entry, matching, and the stationary stocks.

The worker's due diligence choice depends on the payoffs only through their exponentials in units of the attention cost; the relevant objects are the surplus indices $\hat{z}_\omega := \exp \{(W_\omega - W_\emptyset)/\lambda\}$ for $\omega \in \{H, L\}$, which exponentiate the surplus, at the offer decision, of accepting a type ω offer over rejecting it: $\lambda \ln \hat{z}_\omega = W_\omega - W_\emptyset$. A candidate high-quality surplus index determines due diligence and offer quality through Theorem 1, and corresponds one-for-one to a reservation flow value $\hat{w}$; free entry and reservation consistency each determine a matching probability, and market clearing selects the candidate at which the two agree.

Theorem 2 (Equilibrium Existence and Uniqueness). *Maintain parts 1 and 2 of Assumption 3. The interior SRSRIE is characterized by the following construction.*

1. *Bargaining, profits, quitting, and the implied payoff indices:*

$$\begin{aligned} w_H^* &= 0, & w_L^* &= w_L^0, & \eta_H^* &= 0, & \eta_L^* &= 1, \\ J_H^* &= \frac{\pi_H}{1 - \beta(1 - \delta)}, & J_L^* &= \pi_L - w_L^0, & \hat{z}_L &= \exp \left\{ \frac{u_L + w_L^0 - b}{\lambda} \right\} \in (0, 1). \end{aligned}$$

The high-quality surplus index and the reservation flow value determine each other: at reservation flow value $\hat{w}$,

$$\hat{z}_H = \exp \left\{ \frac{u_H - b + \frac{\beta(1-\delta)(u_H - \hat{w})}{1-\beta(1-\delta)}}{\lambda} \right\}, \quad (40)$$

strictly decreasing in $\hat{w}$, with inverse

$$\hat{w}(\hat{z}_H) = u_H - \frac{(1 - \beta(1 - \delta))(b + \lambda \ln \hat{z}_H - u_H)}{\beta(1 - \delta)}. \quad (41)$$

We parameterize candidates by the surplus index and write $\hat{w}(\hat{z}_H)$ for the implied reservation flow value.

2. *The attention and quality block is inherited from the static equilibrium. An interior candidate satisfies $\hat{z}_H > 1 > \hat{z}_L > 0$, $J_L^* > J_H^* > 0$, and $\hat{z}_H J_H^* > \hat{z}_L J_L^*$, the payoff orderings under which the final two rows of Theorem 1 apply. Under the substitution*

$$(z_H, z_L, z_\emptyset, J_H, J_L) = (\hat{z}_H, \hat{z}_L, 1, J_H^*, J_L^*),$$

they define the candidate maps $\mu(\hat{z}_H)$, $P(\hat{z}_H)$, $P_H(\hat{z}_H)$, $P_L(\hat{z}_H)$, $\gamma(\hat{z}_H)$, and $\nu(\hat{z}_H)$, each interior with active posteriors. In particular, the offer composition and the high-quality acceptance

rate are

$$\mu(\hat{z}_H) = \frac{\hat{z}_H J_H^* (1 - \hat{z}_L)}{\hat{z}_H J_H^* - \hat{z}_L J_L^* + \hat{z}_H \hat{z}_L (J_L^* - J_H^*)}, \quad P_H(\hat{z}_H) = \frac{\hat{z}_H J_H^* - \hat{z}_L J_L^*}{J_H^* (\hat{z}_H - \hat{z}_L)}.$$

3. Free entry and reservation consistency tie the meeting probabilities to the candidate surplus:

$$\alpha_f(\hat{z}_H) = \frac{F}{J_H^* P_H(\hat{z}_H)}, \text{ and}$$

$$\alpha_w(\hat{z}_H) = \frac{C + \hat{w}(\hat{z}_H) - b}{\lambda (P[\gamma \ln \hat{z}_H + (1 - \gamma) \ln \hat{z}_L] - \mathcal{I}(\gamma, P))}, \quad (42)$$

where $\hat{w}(\hat{z}_H)$ is the part 1 inverse (41) and γ , P , and $\mathcal{I}(\gamma, P)$ are the part 2 maps, so that α_w is a function of $\hat{z}_H$ alone; the denominator is the worker's optimized net surplus from a contact. Call a candidate $\hat{z}_H$ admissible if part 3 of Assumption 3 holds at $\hat{w}(\hat{z}_H)$, $u_H > \hat{w}(\hat{z}_H) > b$, $\hat{z}_H J_H^* > \hat{z}_L J_L^*$, and $\alpha_f(\hat{z}_H), \alpha_w(\hat{z}_H) \in (0, 1)$. An interior SRSRIE exists if and only if the market clearing equation

$$\alpha_f(\hat{z}_H)^{1-\phi} \alpha_w(\hat{z}_H)^\phi = A \quad (43)$$

has an admissible solution, and it is then unique, because both α_f and α_w are strictly decreasing in $\hat{z}_H$ on the admissible domain.

4. Every equilibrium object is an explicit function of the admissible solution $\hat{z}_H^*$ and parameters:

$$(\mu^*, P^*, P_H^*, P_L^*, \gamma^*, \nu^*, \alpha_f^*, \alpha_w^*) = (\mu, P, P_H, P_L, \gamma, \nu, \alpha_f, \alpha_w)(\hat{z}_H^*),$$

with the reservation flow value, market tightness, and the worker values

$$\hat{w}^* = \hat{w}(\hat{z}_H^*), \quad \theta^* = \left(\frac{A J_H^* P_H^*}{F} \right)^{\frac{1}{\phi}}, \quad V_\emptyset^* = V_L^* = \frac{\hat{w}^*}{1 - \beta}, \quad V_H^* = V_\emptyset^* + \frac{u_H - \hat{w}^*}{1 - \beta(1 - \delta)}.$$

The rate of entry into lasting high-quality employment — hereafter the high-quality job-finding rate — $f^* := \alpha_w^* \mu^* P_H^*$, which fully determines unemployment through $\mathcal{U}^* = \frac{\delta}{\delta + (1 - \delta)f^*}$, satisfies

$$f^* = \frac{F}{J_H^*} \mu^* \theta^* = A^{\frac{1}{\phi}} \left(\frac{J_H^*}{F} \right)^{\frac{1-\phi}{\phi}} \mu^* (P_H^*)^{\frac{1}{\phi}}. \quad (44)$$

The employment stocks are $E_L^* = 0$ and $E_H^* = 1 - \mathcal{U}^*$, vacancies are $\mathcal{V}^* = \theta^* \mathcal{U}^*$, and welfare is

$$\mathcal{W}^* = V_\emptyset^* + E_H^* (V_H^* - V_\emptyset^*), \quad \text{with } V_H^* - V_\emptyset^* = \frac{u_H - \hat{w}^*}{1 - \beta(1 - \delta)} > 0. \quad (45)$$

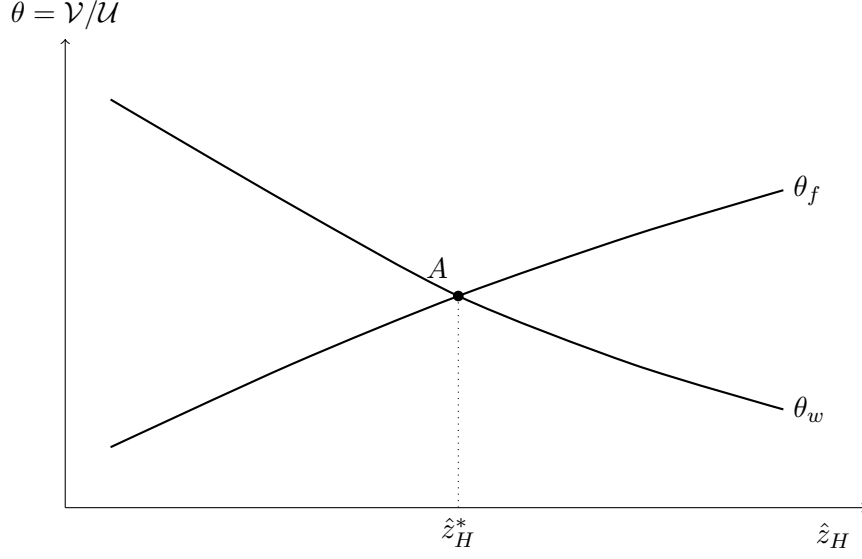


Figure 5: Determination of the interior equilibrium as the unique crossing (point A) of the increasing free entry locus $\theta_f(\hat{z}_H)$ and the decreasing reservation value locus $\theta_w(\hat{z}_H)$, as functions of the high-quality surplus index $\hat{z}_H$.

Proof. See Appendix B, which constructs the candidate maps, verifies every equilibrium condition at an admissible root, proves necessity and uniqueness, and recovers the remaining equilibrium objects. $\square$

Figure 5 visualizes equilibrium determination. For each candidate surplus index $\hat{z}_H$, it draws the market tightness at which firms break even on entry and the tightness that delivers the worker's reservation value.

The free entry locus slopes upward. At a candidate $\hat{z}_H$, the inherited attention and quality block of part 2 determines offer composition and conditional acceptance. A larger surplus index raises the probability a worker accepts a high-quality offer, P_H , so a vacancy breaks even at a lower filling probability, which requires a tighter market: $\theta_f(\hat{z}_H) = (\alpha_f(\hat{z}_H)/A)^{-1/\phi}$ is increasing.

The reservation value locus slopes downward. A larger surplus index corresponds to a lower reservation flow value and a larger optimized net surplus from a contact, so a smaller meeting probability suffices to support the candidate value of unemployment: $\theta_w(\hat{z}_H) = (\alpha_w(\hat{z}_H)/A)^{1/(1-\phi)}$ is decreasing. Because the two schedules move in opposite directions, they cross at most once; the crossing, at A , determines $(\hat{z}_H^*, \theta^*)$.⁹

We now turn to the comparative statics of this equilibrium.

⁹The theorem characterizes the interior equilibrium. The dynamic economy can also have boundary regimes, such as no search, no trade, or rationing outcomes; as in the static model, we condition on the active interior equilibrium — positive search, mixed offer quality, nondegenerate due diligence, and interior matching probabilities — with a primitive test for existence in Corollary 2.

4 Equilibrium Counterfactuals

Despite the richer environment, the comparative statics remain fully analytic and hold at every interior equilibrium. The headline prediction of the static model, that improving low-quality offers u_L raises unemployment and lowers worker welfare, carries over exactly. The predictions for u_H and b , by contrast, interact with a channel that has no static counterpart, the free entry of vacancies. For u_H , a single threshold determines the sign: unemployment falls if and only if the matching elasticity ϕ is below a critical value $\hat{\phi}$. For b , the exact unemployment and welfare signs follow from a few equilibrium statistics, and the same threshold $\hat{\phi}$ gives simple sufficient conditions.

One mechanism organizes the unemployment results that follow. Any change in the environment moves the worker's incentive to tell good offers from bad, and firms respond by changing the mix of offers they post — the composition response of the static model. Because workers quit bad jobs at the first opportunity, only accepted high-quality offers turn into lasting jobs. Free entry then makes the high-quality job-finding rate proportional to the offer mix times market tightness. Two forces therefore drive every unemployment result: how a shock moves the offer mix μ^* (the *composition* channel) and how it moves market tightness θ^* (the *tightness* channel). Composition is the static mechanism. Tightness is new to the dynamic model, and free entry determines whether it reinforces or reverses the composition effect. Worker welfare additionally depends on the state-contingent worker values and the steady state employment rate.

4.1 Unemployment and Worker Welfare

Unemployment and worker welfare need to be redefined for the dynamic setting.

The unemployment rate is the measure of workers searching at the beginning of the period:

$$\mathcal{U} = 1 - E_H - E_L. \quad (46)$$

Recall that E_H and E_L are given in equilibrium by (37) and (38). Workers are ex ante homogeneous, so we measure welfare as the weighted sum of worker values across the three possible employment states:

$$\mathcal{W} = \mathcal{U} V_\emptyset + E_H V_H + E_L V_L. \quad (47)$$

As established in Theorem 2, low-quality jobs last one period and $E_L^* = 0$ at the stock measurement date, so $\mathcal{U} = 1 - E_H$ and equilibrium welfare is a weighted average of $V_\emptyset^*$ and V_H^* .

4.2 Counterfactuals

We study the same three counterfactuals as in the static model: changes in u_L , u_H , and b . Every counterfactual object is a closed-form function of the single unknown $\hat{z}_H^*$ (Theorem 2), and because the per-match block coincides with the static model, everything new in the dynamic comparative statics flows through the equilibrium determination of $\hat{z}_H^*$, the employment stocks, and entry.

Every counterfactual works through the high-quality job-finding rate f^* of (44), defined as the flow of searchers into lasting high-quality employment. A searcher meets an offer with probability α_w^* , the offer is high quality with probability μ^* , and it passes due diligence with probability P_H^* , so $f^* = \alpha_w^* \mu^* P_H^*$. In steady state, the searchers who enter lasting employment replace the employed who separate. With $\eta_H^* = 0$, the stationarity condition (37) rearranges to $(1 - \delta)\mathcal{U}^* f^* = \delta E_H^*$, and substituting $\mathcal{U}^* = 1 - E_H^*$ gives

$$\mathcal{U}^* = \frac{\delta}{\delta + (1 - \delta)f^*}.$$

A fall in unemployment is exactly a rise in f^* .

We work with the flow f^* for our analysis because it decomposes cleanly as a fixed constant times two equilibrium objects,

$$f^* = \frac{F}{J_H^*} \mu^* \theta^*.$$

The posting cost F and the value of a filled high-quality job J_H^* do not respond to u_L , u_H , or b , so f^* tracks the product of the offer mix μ^* and market tightness θ^* .

Furthermore, through free entry, tightness is linked to how often a high-quality offer is accepted. When a high-quality offer is more likely to clear due diligence, a vacancy becomes more valuable, so more firms enter, increasing market tightness. Tightness therefore rises with the acceptance probability of high-quality offers, $\theta^* \propto (P_H^*)^{1/\phi}$. Combining the two, the sign of the change in the high-quality job-finding rate is the sign of the change in $\ln \mu^* + \frac{1}{\phi} \ln P_H^*$. Thus, we analyze the model counterfactuals through the lens of two channels:

- the *composition* channel operates through the fraction μ^* of posted offers that are high quality.
- the *tightness* channel operates through market tightness θ^* . This is driven, through free entry, by the probability P_H^* that a high-quality offer is accepted, with strength $\frac{1}{\phi}$ — a smaller matching elasticity converts a given acceptance response into a larger movement in tightness.

Whenever these two channels move in opposite directions, which one dominates depends on the matching elasticity ϕ . Two statistics make this precise. Along the candidate maps of Theorem 2,

which hold the offer environment $(\hat{z}_L, J_L^*/J_H^*)$ fixed, define the *surplus elasticity of the high-quality job-finding rate*,

$$\varepsilon_f(\hat{z}_H) := \frac{d \ln (\mu(\hat{z}_H) P_H(\hat{z}_H)^{1/\phi})}{d \ln \hat{z}_H} = \frac{d \ln \mu(\hat{z}_H)}{d \ln \hat{z}_H} + \frac{1}{\phi} \frac{d \ln P_H(\hat{z}_H)}{d \ln \hat{z}_H}, \quad (48)$$

and the *critical value of the matching elasticity* that determines its sign,

$$\hat{\phi} := \frac{d \ln P_H(\hat{z}_H)/d \ln \hat{z}_H}{|d \ln \mu(\hat{z}_H)/d \ln \hat{z}_H|} = \left(1 - \frac{J_H^*}{J_L^*}\right) \frac{\hat{z}_H^*}{P^*(\hat{z}_H^* - \hat{z}_L)} > 1 - \frac{J_H^*}{J_L^*}. \quad (49)$$

We write $\varepsilon_f^* := \varepsilon_f(\hat{z}_H^*)$ for the elasticity at the equilibrium root; the threshold $\hat{\phi}$ is already evaluated there, as its closed-form expression shows. An increase in the high-quality surplus index always raises P_H^* (and tightness) and always lowers μ^* , so $\varepsilon_f^* > 0$ if and only if $\phi < \hat{\phi}$. These two statistics organize every comparative static in which the two channels conflict.

Table 2 summarizes the headline results established in the following subsections.¹⁰ Unqualified entries are unconditional signs, the u_H unemployment entry is an exact characterization by the threshold $\hat{\phi}$, and the conditional b entries report simple sufficient regions, with the exact signs for all subregions detailed in Proposition 5. The standalone “ $\pm$ ” entries are the responses to u_L of two intermediate objects, the H acceptance rate and market tightness; neither is signed analytically, and both signs occur (witnessed in Table C.1).

Outcome	$\Delta u_L > 0$	$\Delta u_H > 0$	$\Delta b > 0$
Unemployment $\mathcal{U}^*$	+	– iff $\phi < \hat{\phi}$	– if $\phi \geq \hat{\phi}$, else $\pm$
Value of every worker ($V_\emptyset^* = V_L^*, V_H^*$)	–	+	+
Welfare $\mathcal{W}^*$	–	+	+ if $\phi \geq E_H^* \hat{\phi}$, else $\pm$
Fraction high-quality offers μ^*	–	–	+
Market tightness θ^*	$\pm$	+	–
Acceptance rate of H offers P_H^*	$\pm$	+	–

Notes: (+) denotes an increase, (–) a decrease, (0) no change, and ($\pm$) an ambiguous sign.

Table 2: Counterfactual responses to an increase in u_L , u_H , or b (Propositions 3–5). The b entries provide sufficient regions in terms of $\hat{\phi}$ defined in equation (49), with the exact signs for all subregions detailed in Proposition 5.

4.2.1 Improving Low-quality Offers ($\Delta u_L > 0$)

Suppose we increase u_L without changing any other parameters. Since bad jobs become better, one might expect workers to accept more job offers, leading to fewer unemployed workers. Further,

¹⁰The propositions are stated as derivatives, but the interior equilibrium is unique and moves continuously with the parameters (Theorem 2), so the unqualified signs also give the direction of finite changes along any path that remains interior. Derivatives are one-sided at the kink $\psi\pi_L = (1 - \psi)u_L$, where the negotiated compensation w_L^* hits zero; every equilibrium object is continuous there and the signs are the same on both sides. The governing statistics of the conditional entries can themselves change along such a path, except for the u_H sufficient condition $\phi \leq 1 - J_H^*/J_L^*$, which involves only parameters.

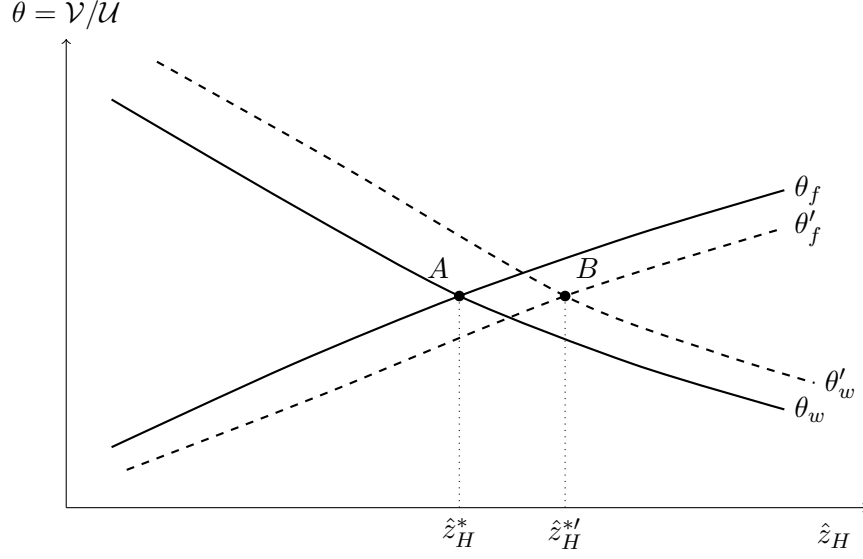


Figure 6: The u_L counterfactual as a shift of the equilibrium loci of Figure 5. Dashed: the loci after an increase in u_L , which moves the equilibrium from A to B . The surplus index $\hat{z}_H^*$ rises; the drawn heights of A and B are schematic, as the effect on tightness θ^* is ambiguous.

since the change raises the flow payoff conditional on accepting a low-quality offer and lowers no other payoff, there is reason to believe worker welfare might increase. The following result shows that the opposite happens.

Proposition 3. *In the interior SRSRIE, an increase in low job quality, u_L :*

1. *strictly increases the high-quality surplus index $\hat{z}_H^*$.*
2. *strictly decreases the fraction of high-quality offers μ^* and both the acceptance and rejection posteriors γ^* and ν^* .*
3. *strictly increases unemployment $\mathcal{U}^*$, because it strictly decreases the high-quality job-finding rate f^* .*
4. *strictly decreases the value of every worker: $V_\emptyset^*$, V_L^* , and V_H^* all fall.*
5. *strictly decreases worker welfare $\mathcal{W}^*$.*

Proof. See Appendix C.3. □

The economics parallels the static result (Proposition 2). The offer composition channel is again decisive, and the dynamic model adds a new margin that reinforces the result. When u_L rises, the loss from mistakenly accepting a bad job rather than remaining unemployed shrinks. The worker's equilibrium acceptance behavior changes — both posteriors fall — and firms respond by shifting

composition toward low-quality offers (μ^* falls, Proposition 3, part 2). When bargaining passes part of the improvement to the firm ($w_L^* > 0$), low-quality matches also become more profitable (J_L^* rises), reinforcing the shift. In the static model this composition response was strong enough to overturn the direct effect of more attractive bad jobs. In the dynamic model it operates even more starkly for the unemployed, because accepted bad jobs are transient. A worker who accepts a low-quality job quits at the first opportunity, so the flow into lasting employment runs entirely through high-quality matches, at rate $f^* = \alpha_w^* \mu^* P_H^*$. The deterioration in offer composition therefore puts direct downward pressure on the high-quality job-finding rate.

In equilibrium the high-quality surplus index $\hat{z}_H^*$ rises (Proposition 3, part 1), because the value of unemployment falls by more than the value of a high-quality job. Figure 6 draws the shift. An increase in u_L moves the reservation value locus up. At a given high-quality surplus, the deterioration in equilibrium offer composition lowers the worker's optimized expected surplus from a contact, so reservation consistency requires a higher meeting rate. The free entry locus shifts down, because P_H falls at a fixed surplus and entry breaks even at a higher matching probability. Both shifts raise the equilibrium surplus, but tightness can move either way, because it tracks the acceptance rate P_H^* through free entry and P_H^* itself faces two opposing pressures. The higher surplus pushes P_H^* up, while the direct movement of the offer environment ($\hat{z}_L, J_L^*/J_H^*$) pushes it down, leaving the total response ambiguous (Table 2).

Whichever way P_H^* moves, however, the tightness response is never strong enough to offset the composition deterioration. The high-quality job-finding rate strictly falls, for any matching elasticity $\phi \in (0, 1)$ and at any interior equilibrium, so unemployment strictly rises (Proposition 3, part 3).

Every worker's equilibrium value falls (Proposition 3, part 4). The unemployed are worse off because the offer quality mix they search over deteriorates (the reservation flow value $\hat{w}^*$ falls). The employed are worse off because their outside option upon separation deteriorates.

The worker welfare result (Proposition 3, part 5) follows from parts 3 and 4. Every value falls, and the steady state distribution of workers shifts from high value employment toward low value unemployment. A weighted average of falling values with weight shifting toward the lowest value must fall.

4.2.2 Improving High-quality Offers ($\Delta u_H > 0$)

Suppose we increase u_H without changing any other parameters. As in the static model, a higher u_H worsens the offer mix. But better good jobs are also accepted more often, which raises market tightness. These two forces push unemployment in opposite directions, so its sign turns on the matching elasticity ϕ . Every worker is nevertheless better off, and worker welfare rises, whether or not unemployment falls.

Proposition 4. *In the interior SRSRIE, an increase in high job quality, u_H :*

1. *strictly increases the high-quality surplus index $\hat{z}_H^*$.*
2. *strictly decreases the fraction of high-quality offers μ^* and the acceptance and rejection posteriors γ^* and ν^* , and strictly increases the acceptance rates P_H^* and P_L^* and market tightness θ^* .*
3. *strictly decreases unemployment $\mathcal{U}^*$ if $\phi < \hat{\phi}$ and strictly increases it if $\phi > \hat{\phi}$, with $\hat{\phi}$ defined in (49). A sufficient condition for unemployment to fall is $\phi \leq 1 - \frac{J_H^*}{J_L^*}$.*
4. *strictly increases the value of every worker: $V_\emptyset^*$, V_L^* , and V_H^* all rise.*
5. *strictly increases worker welfare $\mathcal{W}^*$.*

Proof. See Appendix C.4. □

Since J_H^* , J_L^* , and $\hat{z}_L$ do not depend on u_H , the entire effect on attention, composition, acceptance, and market tightness runs through the equilibrium surplus index $\hat{z}_H^*$. This makes the u_H counterfactual the cleanest way to see how the composition and tightness channels conflict. Among the schedules of Figure 5, an increase in u_H shifts the reservation value locus up and leaves the free entry locus in place, so both the surplus index $\hat{z}_H^*$ (Proposition 4, part 1) and tightness θ^* rise.

The composition mechanism works just as in the static model. A more valuable good job leads workers to accept on weaker evidence, more bad offers slip through due diligence, and firm indifference requires the share of good offers μ^* to fall (Proposition 4, part 2). Making good jobs better again makes them scarcer.

Free entry pushes the other way. Each posted offer is now more likely to be accepted, so vacancies are worth more, and increased entry raises market tightness. Whether unemployment falls then depends on the matching elasticity ϕ (Proposition 4, part 3). When ϕ is small, tightness responds strongly and unemployment falls. When ϕ is large, the worse offer mix dominates and unemployment rises, even though good jobs became better. The threshold is exactly $\hat{\phi}$: unemployment falls exactly when the tightness response exceeds the composition response, which free entry makes equivalent to $\phi < \hat{\phi}$.

Every worker is better off, employed or unemployed, even when unemployment rises (Proposition 4, part 4). Better high-quality jobs raise the return to searching, so the value of unemployment rises, and the employed gain both directly and through a better outside option.

Worker welfare rises either way (Proposition 4, part 5). The gains in both state-contingent worker values dominate any loss from the lower employment rate, for every value of the matching elasticity.

4.2.3 Raising Unemployment Benefits ($\Delta b > 0$)

Suppose we increase the unemployment benefit b . As in the static model, a higher b improves the offer mix, reversing the McCall pickiness intuition. But it also lowers the benefit of working a good job relative to being unemployed, which reduces market tightness. These two forces push unemployment in opposite directions, so its sign turns on the matching elasticity ϕ . Worker welfare is the employment-weighted average of worker values. Every worker's value rises with b , but the share of employment at good jobs can move either way, so welfare too can rise or fall.

Proposition 5. *In the interior SRSRIE, an increase in the unemployment benefit, b :*

1. *strictly decreases the high-quality surplus index $\hat{z}_H^*$.*
2. *strictly increases the fraction of high-quality offers μ^* and the acceptance and rejection posteriors γ^* and ν^* , and strictly decreases the acceptance rates P_H^* and P_L^* and market tightness θ^* .*
3. *strictly decreases unemployment $\mathcal{U}^*$ whenever $\phi \geq \hat{\phi}$. In the remaining region $\phi < \hat{\phi}$, unemployment strictly decreases when $\Sigma^* > \varepsilon_f^* \chi^*$, is locally unchanged when equality holds, and strictly increases when $\Sigma^* < \varepsilon_f^* \chi^*$.*
4. *strictly increases the value of every worker: $V_\emptyset^*$, V_L^* , and V_H^* all rise.*
5. *strictly increases worker welfare $\mathcal{W}^*$ whenever unemployment does not rise, and moreover whenever $\phi \geq E_H^* \hat{\phi}$. When $\phi < E_H^* \hat{\phi}$, both strict signs occur: worker welfare strictly increases if and only if the worker value gains outweigh the employment margin loss,*

$$\underbrace{\lambda \chi^* \left(\frac{1 - \beta(1 - \delta)}{1 - \beta} - E_H^* \right)}_{\text{worker value gains}} > \underbrace{\beta(1 - \delta) (V_H^* - V_\emptyset^*) E_H^* \mathcal{U}^* (\varepsilon_f^* \chi^* - \Sigma^*)}_{\text{employment margin loss}}.$$

Proof. See Appendix C.5. □

A higher unemployment benefit raises the value of remaining unemployed, lowering the benefit of a good job relative to unemployment and, with it, the high-quality surplus index $\hat{z}_H^*$ (Proposition 5, part 1). Among the schedules of Figure 5, the reservation value locus shifts down and the free entry locus shifts up, so market tightness falls. At the lower surplus, firm indifference shifts the offer mix μ^* toward good offers, the posteriors rise, and both acceptance rates fall (Proposition 5, part 2). This is the mirror image of the u_H counterfactual: a better mix in a slacker market, rather than a worse mix in a tighter one.

Whether unemployment falls turns on two opposing forces on the high-quality job-finding rate: the *composition gain* from the improving offer mix and the *premium compression* from the falling

benefit of a good job relative to unemployment,

$$\Sigma^* := \lambda \left. \frac{\partial \ln \mu^*}{\partial b} \right|_{\hat{w}} > 0, \quad \chi^* := \beta(1 - \delta) \left| \frac{d(V_H^* - V_\emptyset^*)}{db} \right| > 0. \quad (50)$$

The composition gain raises the high-quality job-finding rate at fixed acceptance and tightness; the premium compression lowers the acceptance rate of good offers and, through entry, market tightness, and the surplus elasticity ε_f^* carries it into the high-quality job-finding rate.¹¹ The two combine into the total response,

$$\lambda \frac{d \ln f^*}{db} = \Sigma^* - \varepsilon_f^* \chi^*, \quad (51)$$

so unemployment falls exactly when $\Sigma^* > \varepsilon_f^* \chi^*$, as in Proposition 5, part 3. When $\phi \geq \hat{\phi}$ the surplus elasticity is nonpositive, the composition gain necessarily dominates, and unemployment falls: the anti-McCall result of the static model survives whenever the composition channel is at least as strong as tightness. When $\phi < \hat{\phi}$ the composition gain need not dominate (Appendix C.5).

A more generous unemployment benefit makes every worker better off, even though it lowers the reward to being employed (Proposition 5, part 4). The unemployed gain directly. Whenever a searcher meets no offer, or meets one and turns it down, they collect the higher benefit, and the lower reward to employment costs them only in the periods when they accept. A loss borne only upon acceptance does not overturn a gain enjoyed in every other event, so the value of unemployment rises. The employed gain because every job eventually ends. Their pay while employed is unchanged, and the unemployment they return to is more valuable, so the value of holding a good job rises as well, even as its premium over unemployment shrinks.

Welfare (Proposition 5, part 5) moves through two components. A higher benefit raises every worker's value — the *worker value gains*, always positive by Proposition 5, part 4 — but it can lower the employment rate E_H^* that weights the high-quality premium in welfare — the *employment margin loss*. Welfare falls only when the second outweighs the first, which takes two things at once: a contraction in market tightness far stronger than the improvement in the offer mix (ϕ well below $\hat{\phi}$), and a large employed share exposed to the premium compression (E_H^* high). When employment is low enough that $\phi \geq E_H^* \hat{\phi}$, a more generous unemployment benefit raises welfare even if it raises unemployment.

Summary. The three counterfactuals share one structure. Improving low-quality offers degrades the offer mix so strongly that the flow into lasting employment falls regardless of the tightness response. Improving high-quality offers degrades the mix but raises market tightness, so unemployment turns on which force dominates, determined by the matching elasticity relative to $\hat{\phi}$. Raising

¹¹The premium compression $\chi^* = |1 + \lambda d \ln \hat{z}_H^* / db|$ measures how much faster than one-for-one the high-quality surplus index falls, and the composition gain is $\Sigma^* = \frac{J_L^* \hat{z}_L (\hat{z}_H^* - \hat{z}_L)}{(1 - \hat{z}_L) [J_H^* \hat{z}_H^* (1 - \hat{z}_L) + J_L^* \hat{z}_L (\hat{z}_H^* - 1)]}$. Appendix C.5 derives both.

unemployment benefits reverses both movements: the mix improves while the market slackens. In every case free entry determines whether the tightness channel reinforces or overturns the unemployment prediction of the composition channel. Welfare requires the further step through worker values and the employment rate. The composition channel, which works exactly as in the static model, is also what separates these results from the standard exogenous information and offer quality model. In that model the offer mix is a parameter and only the tightness channel operates. With due diligence, the mix responds to what workers choose to learn, and that response is what can reverse the standard predictions.

5 Constrained Efficiency and Policy

We now study constrained efficiency, with a focus on how endogenous information acquisition and offer quality augment the traditional inefficiency in the exogenous information and offer quality benchmark. The constrained planner faces the same frictions as the market: firms and workers meet through the matching technology, workers privately choose due diligence, firms privately choose hidden offer quality, and neither action is contractible. Thus, no policy can dictate the information structure or the offer mix. We give the planner two instruments that change private incentives. A vacancy levy τ_v is paid by every entering firm; because it cannot be conditioned on the hidden quality choice, it moves entry without disturbing the choice between offer qualities. The unemployment benefit b is financed by a common lump-sum assessment and moves the worker's outside option.¹²

The welfare criterion is social surplus, which differs from the worker value $\mathcal{W}^*$ of Section 4. Worker value counts the unemployment benefit as income. Social surplus counts resources: match surplus, search costs, attention costs, and vacancy costs. Total match surplus is $y_\omega := \pi_\omega + u_\omega$ for $\omega \in \{H, L\}$. Negotiated compensation cancels when the worker's utility and the firm's profit are added, and the profit itself is counted. Because u_ω is nonpecuniary, y_ω is a surplus rather than output. The benefit, the levy, and the assessment that balances them are transfers, so they enter social surplus only through the behavior they induce.

Policy does not require a new equilibrium system. Given tightness and the unemployment benefit, reservation consistency and the per-match block of Theorem 2 determine the surplus indices, due diligence, the offer mix, and the acceptance rates. The levy moves tightness through free entry, $F + \tau_v = \alpha_f P_H J_H^*$. Because both offer types pay it, the levy does not enter the choice between

¹²A unit mass of households owns the firms through a diversified portfolio, and the government balances its budget each period with a common assessment $T_t = b\mathcal{U}_t(1 - \alpha_w P) - \tau_v \mathcal{V}_t$, negative values being a rebate; the same rule finances a vacancy subsidy. Benefits are paid to searchers who do not accept an offer (28), hence the base $\mathcal{U}_t(1 - \alpha_w P)$. Payoffs are linear, and the assessment and the dividend are common to every household, independent of employment status and of every action, so subtracting their present value from each value function leaves every payoff difference — and with it search, due diligence, acceptance, quitting, and the bargaining threat points — unchanged.

qualities, and its equilibrium effects run entirely through tightness. The planner therefore chooses two numbers, θ and b , and private behavior fills in the rest.

A contact yields expected match surplus net of the attention cost,

$$\mathcal{S} := \mu P_H y_H + (1 - \mu) P_L y_L - \lambda \mathcal{I}. \quad (52)$$

Each searcher meets an offer with probability α_w and pays the search cost C , and vacancies cost $F\theta$ per searcher. Thus, the planner's objective and the employment transition are

$$\mathcal{P} = \sum_{t \geq 0} \beta^t [E_{H,t} y_H + \mathcal{U}_t (\alpha_w \mathcal{S} - C - F\theta)], \quad E_{H,t+1} = (1 - \delta) (E_{H,t} + \mathcal{U}_t f), \quad (53)$$

where $f = \alpha_w \mu P_H$ is the high-quality job-finding rate at the policy allocation. Accepted low-quality jobs last one period, so they appear in $\mathcal{S}$ but never in the employment stock. The policy is one permanent pair, chosen at date zero given an initial unemployment stock $\mathcal{U}_0$ with the rest of the population in lasting high-quality jobs, and employment then adjusts under the new policy. Constant returns keep the per-match block stationary during that adjustment, so the objective sums exactly (Appendix E):

$$\mathcal{P} = \frac{y_H}{1 - \beta} - \left[\mathcal{U}_0 + \frac{\beta \delta}{1 - \beta} \right] \frac{\Delta}{\Omega}, \quad \Delta := y_H - \alpha_w \mathcal{S} + C + F\theta, \quad \Omega := 1 - \beta(1 - \delta)(1 - f), \quad (54)$$

where Δ is the flow cost of a searcher relative to a lasting job and Ω discounts that cost by how quickly searchers enter lasting jobs. The inherited unemployment stock multiplies the ratio by a positive constant, so the planner minimizes Δ/Ω , and the optimal policy does not depend on $\mathcal{U}_0$. Converting a searcher into a lasting job trades the job's surplus for the searcher's option to keep sampling offers, and the discounted value of that trade, weighted by the share of contacts that become lasting jobs, is

$$\Lambda := \frac{\beta(1 - \delta) \mu P_H \Delta}{\Omega}. \quad (55)$$

Without loss, normalize the initial unemployment stock to be its stationary value, $\mathcal{U}_0 = \mathcal{U}$. Then, differentiating Δ/Ω shows that any policy change p satisfies (Appendix E)

$$\frac{1 - \beta}{\mathcal{U} \alpha_w} \frac{d\mathcal{P}}{dp} = \left[(1 - \phi) (\mathcal{S} + \Lambda) - \frac{F}{\alpha_f} \right] \frac{d \ln \theta}{dp} + \frac{d\mathcal{S}}{dp} + \Lambda \frac{d \ln(\mu P_H)}{dp}. \quad (56)$$

The first term is the standard vacancy creation margin: tightness raises worker contacts with elasticity $1 - \phi$ and costs vacancies. The two remaining terms are this paper's addition. Policy also changes what a contact is worth and how often it becomes a lasting job, because due diligence and offer quality respond.

Each instrument's optimality condition is therefore the standard rule plus a behavioral response. Moving an instrument changes what workers learn and what quality firms post. That response matters for the planner in two ways. It changes the surplus of a contact, $\mathcal{S}$, and it changes the share μP_H of contacts that become lasting high-quality jobs. To make this precise, define the response to each instrument, holding the other fixed,

$$\mathcal{R}_\theta := \left. \frac{\partial \mathcal{S}}{\partial \ln \theta} \right|_b + \Lambda \left. \frac{\partial \ln(\mu P_H)}{\partial \ln \theta} \right|_b, \quad \mathcal{R}_b := \left. \frac{\partial \mathcal{S}}{\partial b} \right|_\theta + \Lambda \left. \frac{\partial \ln(\mu P_H)}{\partial b} \right|_\theta. \quad (57)$$

In each response, the first term is the change in the contact surplus and the second term is the proportional change in the share of contacts that become lasting jobs, valued at Λ .

One further object organizes the results. Consider shifting one unit of offer probability from low to high quality while holding both conditional acceptance rates fixed, the same variation that defines the composition gain Σ^* . The social value of this shift is

$$\mathcal{Q} := P_H \left(y_H - \frac{J_H^*}{J_L^*} y_L \right) - \lambda \left. \frac{\partial \mathcal{I}}{\partial \mu} \right|_{P_H, P_L} + \frac{\Lambda}{\mu}. \quad (58)$$

The first term is the surplus gain. By firm indifference, the added high-quality offers are accepted with probability P_H and the removed low-quality offers were accepted with probability $P_L = (J_H^*/J_L^*) P_H$, so accepted surplus rises by $P_H y_H - P_L y_L$. The second term subtracts the change in the attention cost of a contact as the offer mix improves. The third term is the value of the extra lasting jobs, Λ times the increase $1/\mu$ in their log share.

Proposition 6 (Constrained efficiency and policy assignment). *On the interior equilibrium regime maintained throughout, an optimal permanent policy exists on any compact policy set within the regime and does not depend on the inherited stock; moreover:*

1. *An interior optimum satisfies the tightness condition $F = \alpha_f [(1 - \phi)(\mathcal{S} + \Lambda) + \mathcal{R}_\theta]$ and is implemented by the levy*

$$\tau_v = \alpha_f [P_H J_H^* - (1 - \phi)(\mathcal{S} + \Lambda) - \mathcal{R}_\theta]. \quad (59)$$

Freezing due diligence and offer quality at the same allocation gives the correction $\tau_v^0 := \alpha_f [P_H J_H^ - (1 - \phi)(\mathcal{S} + \Lambda)]$. The true levy differs by $\tau_v - \tau_v^0 = -\alpha_f \mathcal{R}_\theta$, which reverses the levy's sign exactly when $\mathcal{R}_\theta$ and τ_v^0 share a sign and $\alpha_f |\mathcal{R}_\theta| > |\tau_v^0|$.*

2. *At fixed tightness, a higher unemployment benefit improves the offer mix and lowers both conditional acceptance rates: $\left. \frac{\partial \mu}{\partial b} \right|_\theta > 0$, $\left. \frac{\partial P_H}{\partial b} \right|_\theta < 0$, and $\left. \frac{\partial P_L}{\partial b} \right|_\theta < 0$. An interior joint optimum sets $\mathcal{R}_b = 0$ together with part 1, and at any such optimum*

$$\text{sgn}(\mathcal{R}_\theta) = -\text{sgn}(\mathcal{Q}). \quad (60)$$

3. If the vacancy levy is unavailable, then along the zero-levy equilibrium path

$$\frac{1 - \beta}{\mathcal{U} \alpha_w} \frac{d\mathcal{P}}{db} = \mathcal{R}_b - \frac{\tau_v}{\alpha_f} \frac{d \ln \theta^*}{db}, \quad (61)$$

with τ_v the part 1 levy evaluated at the zero-levy allocation and $\frac{d \ln \theta^*}{db} < 0$ by Proposition 5.

Proof. See Appendix E. □

Part 1 of Proposition 6 writes the optimal levy as a standard entry correction plus a response term. The first two terms of the tightness condition are the standard correction, the rule of a planner facing exogenous information and offer quality. With the response absent, the only distortion is matching congestion, and the optimum is the Hosios condition (Hosios, 1990) stated in surplus shares: firms appropriate the share $1 - \phi$ of the social value of a match. Under our bargaining protocol this is not the familiar condition $\psi = \phi$, because the nonnegativity constraint on additional compensation binds for high-quality jobs and firms keep the entire high-quality flow profit regardless of ψ .

The response term $\mathcal{R}_\theta$, defined in (57), is what endogenous due diligence and offer quality add. The contact surplus term is the change in the surplus $\mathcal{S}$ of a contact when the market tightens; the lasting jobs term is the change in the share μP_H of contacts that become lasting high-quality jobs, valued at Λ . A tighter market raises the worker's meeting probability and with it the value of unemployment, so workers screen offers more strictly. The offer mix improves while both conditional acceptance rates fall, and the attention cost of a contact moves with them. $\mathcal{R}_\theta$ is positive when the better offer mix is worth more than the forgone matches, and negative when the forgone matches are worth more.

The lasting jobs term is signed by a threshold already defined. The share μP_H rises with stricter screening exactly when the mix response exceeds the acceptance response, and the threshold $\hat{\phi}$ of (49) is the ratio of the acceptance response to the mix response. Whenever converting searchers into lasting jobs is socially valuable, $\Lambda > 0$, the lasting jobs term is therefore positive exactly when $\hat{\phi} < 1$. In the comparative statics $\hat{\phi}$ was compared to the matching elasticity ϕ , because free entry converts acceptance into tightness with strength $1/\phi$. The planner sets tightness directly, so the amplification is absent and $\hat{\phi}$ is compared to one. The contact surplus term is different. It values offers at y_ω , so its sign turns on the joint surplus levels and the attention cost. Either term can dominate, so the response term can take either sign.

Part 2 gives the unemployment benefit a role that does not exist in the standard model. The benefit is a transfer and enters no firm payoff, so its entire allocative effect runs through due diligence. At fixed tightness, a higher unemployment benefit makes workers reject more readily, which improves the offer mix at the cost of fewer acceptances of both qualities. The unemployment benefit is the planner's route to the offer mix, though not a pure one. It also lowers both acceptance

rates, and the levy moves screening through tightness. The assignment result (60) follows from how the two instruments overlap. Part of what a higher benefit does is exactly what more tightness does. Both raise the value of unemployment, so both make workers screen more strictly, improving the offer mix while lowering both acceptance rates. The remaining effect of the change in unemployment benefits is a pure improvement in the offer mix with both acceptance rates unchanged, and this shift is valued at $\mathcal{Q}$. At a joint optimum the unemployment benefit has been raised until a further increase adds no social value, so the value of the shared screening part and the value of the pure mix shift exactly offset. If the pure mix improvement is still socially valuable there, $\mathcal{Q} > 0$, then the shared part must be socially costly, and the shared part is the response to tightness, $\mathcal{R}_\theta < 0$. By part 1, a negative $\mathcal{R}_\theta$ raises the levy, so the correction leans toward a tax on entry. A coordinated change in the two instruments isolates the pure mix shift, so the pair spans the composition margin that neither moves alone.

Part 3 reads the unemployment benefit counterfactual of Section 4 as second best policy. Without the levy, entry is uncorrected, and each unit of log tightness the benefit induces is worth $-\tau_v/\alpha_f$, the size of the uncorrected wedge. A higher benefit lowers equilibrium tightness. When the missing correction is a tax, $\tau_v > 0$, entry at the zero vacancy levy allocation is locally excessive, so the induced contraction in tightness is a gain and the benefit partly performs as the missing tax. When it is a subsidy, $\tau_v < 0$, the contraction is a loss and the unemployment benefit works against the missing correction.

To highlight the novel elements introduced by due diligence and endogenous offer quality, consider the economy of Appendix E, in which the informed planner's optimal policy cuts the unemployment benefit and taxes entry. A planner who mistakenly believes the economy features exogenous information and an exogenous offer quality distribution would instead leave the benefit unchanged and subsidize entry. With information and offer quality fixed, the benefit is just a transfer that changes no decision, so the mistaken planner has no reason to touch it. The informed planner's benefit cut is also what turns the entry correction into a tax. A lower benefit relaxes screening and raises acceptance, so private entry at the target tightness switches from falling short to overshooting, and the required correction switches from a subsidy to a tax.

Summary. In the standard exogenous information and offer quality model, the planner corrects entry and is done. The unemployment benefit has no work of its own to do. With due diligence and endogenous offer quality, both instruments also move what workers learn and what quality firms supply. The benefit gives the planner access to the offer mix, and the composition response can reverse the sign of the optimal entry correction.

6 Conclusion

Many important search markets require decisions before the value of an offer is fully known, so acceptance is preceded by costly due diligence. This paper builds due diligence and endogenous offer quality into a search and matching model. Due diligence affects not only which offers are accepted but also which offers suppliers choose to create, so the composition of offers is an equilibrium object. This composition channel can reverse comparative statics that would be immediate in a model with fixed offer quality. Interventions affecting offer quality or outside options therefore cannot be evaluated solely from their direct effect on payoffs. Their equilibrium effect depends on how they change searchers' due diligence and suppliers' choices of offer quality.

Credit authorship contribution statement

All authors: Conceptualization; Formal analysis; Methodology; Writing – original draft; Writing – review & editing.

Funding

Andrew Caplin and Daniel Martin acknowledge support from the Alfred P. Sloan Foundation under the grant *Cognitive Economics at Work*.

Declaration of competing interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Data availability

No data were used in the research described in this article. The analytical model, derivations, and numerical parameter values needed to reproduce the reported examples are contained in the article.

Disclaimer

The analysis and conclusions set forth here are those of the authors and do not necessarily reflect the views of The Goldman Sachs Group, Inc.

Declaration of generative AI and AI-assisted technologies in the manuscript preparation process

During the preparation of this work the authors used ChatGPT Codex and Claude Code in order to check proofs and solve the model numerically to verify propositions and provide examples and counterexamples. The authors reviewed and edited all resulting material and take full responsibility for the content of the article.

References

- ACHARYA, S., AND S. L. WEE (2020): “Rational Inattention in Hiring Decisions,” *American Economic Journal: Macroeconomics*, 12(1), 1–40.
- ADRJAN, P., M. BALGOVA, S. JAGER, J. JESSEN, AND J. SOCKIN (2026): “Job Ads as Signals: Evidence from a Priced Amenity and Worker Beliefs,” Working Paper 12794, CESifo.
- ALMOG, D., AND D. MARTIN (2024): “Rational Inattention in Games: Experimental Evidence,” *Experimental Economics*, 27(4), 715–742.
- AMBUEHL, S., A. OCKENFELS, AND C. STEWART (2025): “Who Opt In? Composition Effects and Disappointment from Participation Payments,” *Review of Economics and Statistics*, 107(1), 78–94.
- AUDOLY, R., M. BHULLER, AND T. A. REIREMO (2024): “The Pay and Non-Pay Content of Job Ads,” Staff Report 1124, Federal Reserve Bank of New York, Revised May 2026.
- BANFI, S., AND B. VILLENA-ROLDAN (2019): “Do High-Wage Jobs Attract More Applicants? Directed Search Evidence from the Online Labor Market,” *Journal of Labor Economics*, 37(3), 715–746.
- BATRA, H., A. MICHAUD, AND S. MONGEY (2023): “Online Job Posts Contain Very Little Wage Information,” Working Paper 31984, National Bureau of Economic Research.
- BELOT, M., P. KIRCHER, AND P. MULLER (2022): “How Wage Announcements Affect Job Search—A Field Experiment,” *American Economic Journal: Macroeconomics*, 14(4), 1–67.
- BURDETT, K., AND M. G. COLES (1997): “Marriage and class,” *The Quarterly Journal of Economics*, 112(1), 141–168.
- CAPLIN, A., M. DEAN, AND J. LEAHY (2022): “Rationally Inattentive Behavior: Characterizing and Generalizing Shannon Entropy,” *Journal of Political Economy*, 130(6), 1676–1715.
- DEAN, M., AND N. NELIGH (2023): “Experimental Tests of Rational Inattention,” *Journal of Political Economy*, 131(12), 3415–3461.
- DELACROIX, A., AND S. SHI (2013): “Pricing and signaling with frictions,” *Journal of Economic Theory*, 148(4), 1301–1332.
- DIAMOND, P. A. (1982): “Wage determination and efficiency in search equilibrium,” *The Review of Economic Studies*, 49(2), 217–227.
- DUFFIE, D., N. GÂRLEANU, AND L. H. PEDERSEN (2007): “Valuation in over-the-counter markets,” *The Review of Financial Studies*, 20(6), 1865–1900.

- GEE, L. K. (2019): “The More You Know: Information Effects on Job Application Rates in a Large Field Experiment,” *Management Science*, 65(5), 2077–2094.
- GENESOVE, D., AND L. HAN (2012): “Search and Matching in the Housing Market,” *Journal of Urban Economics*, 72(1), 31–45.
- GUERRIERI, V., AND R. SHIMER (2014): “Dynamic adverse selection: A theory of illiquidity, fire sales, and flight to quality,” *American Economic Review*, 104(7), 1875–1908.
- (2018): “Markets with multidimensional private information,” *American Economic Journal: Microeconomics*, 10(2), 250–74.
- GUERRIERI, V., R. SHIMER, AND R. WRIGHT (2010): “Adverse selection in competitive search equilibrium,” *Econometrica*, 78(6), 1823–1862.
- HOSIOS, A. J. (1990): “On the Efficiency of Matching and Related Models of Search and Unemployment,” *Review of Economic Studies*, 57(2), 279–298.
- JIN, G. Z., M. LUCA, AND D. MARTIN (2021): “Is No News (Perceived As) Bad News? An Experimental Investigation of Information Disclosure,” *American Economic Journal: Microeconomics*, 13(2), 141–173.
- (2022): “Complex Disclosure,” *Management Science*, 68(5), 3236–3261.
- KAYA, A., AND K. KIM (2018): “Trading Dynamics with Private Buyer Signals in the Market for Lemons,” *The Review of Economic Studies*, 85(4), 2318–2352.
- LANDVOIGT, T., M. PIAZZESI, AND M. SCHNEIDER (2015): “The housing market(s) of San Diego,” *American Economic Review*, 105(4), 1371–1407.
- MAĆKOWIAK, B., F. MATĚJKA, AND M. WIEDERHOLT (2023): “Rational Inattention: A Review,” *Journal of Economic Literature*, 61(1), 226–273.
- MAESTAS, N., K. J. MULLEN, D. POWELL, T. VON WACHTER, AND J. B. WENGER (2023): “The Value of Working Conditions in the United States and the Implications for the Structure of Wages,” *American Economic Review*, 113(7), 2007–2047.
- MARINESCU, I., AND R. WOLTHOFF (2020): “Opening the Black Box of the Matching Function: The Power of Words,” *Journal of Labor Economics*, 38(2), 535–568.
- MARTIN, D. (2017): “Strategic pricing with rational inattention to quality,” *Games and Economic Behavior*, 104, 131–145.
- MARTIN, D., AND P. MARX (2022): “A Robust Test of Prejudice for Discrimination Experiments,” *Management Science*, 68(6), 4527–4536.

- MAS, A., AND A. PALLAIS (2017): “Valuing alternative work arrangements,” *American Economic Review*, 107(12), 3722–59.
- MATEJKA, F., AND A. MCKAY (2015): “Rational Inattention to Discrete Choices: A New Foundation for the Multinomial Logit Model,” *American Economic Review*, 105(1), 272–298.
- MCCALL, J. (1970): “Economics of Information and Job Search,” *Quarterly Journal of Economics*, 84(1), 113–126.
- MENZIO, G. (2007): “A theory of partially directed search,” *Journal of political Economy*, 115(5), 748–769.
- MORRIS, S., AND P. STRACK (2019): “The Wald Problem and the Relation of Sequential Sampling and Ex-Ante Information Costs,” *SSRN Electronic Journal*.
- MORTENSEN, D. T. (1982): “Property rights and efficiency in mating, racing, and related games,” *The American Economic Review*, 72(5), 968–979.
- MOSCARINI, G. (2005): “Job matching and the wage distribution,” *Econometrica*, 73(2), 481–516.
- MOSCARINI, G., AND L. SMITH (2001): “The optimal level of experimentation,” *Econometrica*, 69(6), 1629–1644.
- PISSARIDES, C. A. (1985): “Short-run equilibrium dynamics of unemployment, vacancies, and real wages,” *The American Economic Review*, 75(4), 676–690.
- POMATTO, L., P. STRACK, AND O. TAMUZ (2023): “The Cost of Information: The Case of Constant Marginal Costs,” *American Economic Review*, 113(5), 1360–1393.
- PORCHER, C. (2025): “Migration with Costly Information,” *Working paper*.
- PRIES, M., AND R. ROGERSON (2005): “Hiring policies, labor market institutions, and labor market flows,” *Journal of Political Economy*, 113(4), 811–839.
- RAVID, D. (2020): “Ultimatum Bargaining with Rational Inattention,” *American Economic Review*, 110(9), 2948–2963.
- ROSEN, S. (1986): “The theory of equalizing differences,” *Handbook of labor economics*, 1, 641–692.
- SHIMER, R. (2005): “The cyclical behavior of equilibrium unemployment and vacancies,” *American economic review*, 95(1), 25–49.
- SIMS, C. (2003): “Implications of Rational Inattention,” *Journal of Monetary Economics*, 50(3), 665–690.
- SMITH, L. (2006): “The marriage model with search frictions,” *Journal of political Economy*, 114(6), 1124–1144.

STIGLER, G. J. (1961): “The economics of information,” *Journal of political economy*, 69(3), 213–225.

Appendix

A Static Model and the Shared Attention and Quality Block

A.1 Proof of Theorem 1

Throughout the proof, the additional compensation levels are fixed at the Nash bargaining candidates w_H^0, w_L^0 in Assumption 2. The Nash bargaining first-order conditions and the nonnegativity constraint imply these are the unique compensation levels whenever $w_H^0 = 0$ and $W_H > b > W_L$. Thus $J_H = \pi_H - w_H^0$, $J_L = \pi_L - w_L^0$, and z_ω incorporates the bargained compensation. Assumption 2 gives $z_H > z_\emptyset > z_L$, $J_L > J_H > 0$, and $z_H J_H > z_L J_L$. Taking the first-order conditions of the worker's problem or applying Matějka and McKay (2015) directly we get:

$$P_\omega = \frac{z_\omega P}{z_\omega P + z_\emptyset(1 - P)} \quad (\text{A.1})$$

Substituting that into the firm's first-order condition identifies P :

$$\begin{aligned} P_H J_H &= P_L J_L \Rightarrow \\ \frac{z_H J_H P}{z_H P + z_\emptyset(1 - P)} &= \frac{z_L J_L P}{z_L P + z_\emptyset(1 - P)} \Rightarrow \\ P &= \frac{z_\emptyset(z_H J_H - z_L J_L)}{z_\emptyset(z_H J_H - z_L J_L) + z_H z_L (J_L - J_H)} \end{aligned}$$

Because we study the active interior RIBNE, the displayed solution must lie in $(0, 1)$. Given $J_L > J_H$, this holds if and only if $z_H J_H - z_L J_L > 0 \Leftrightarrow \frac{J_L}{J_H} < \frac{z_H}{z_L}$, which is imposed directly in Assumption 2. For $P > 0$, Bayes' rule (or, equivalently, substituting A.1 into the definition of P) implies:

$$1 = \frac{\mu z_H}{z_H P + z_\emptyset(1 - P)} + \frac{(1 - \mu) z_L}{z_L P + z_\emptyset(1 - P)}$$

So the relationship between the firm's probability of posting a good job μ and the P associated with the worker's best response is:

$$\mu = \frac{(z_\emptyset - z_L)(z_H P + z_\emptyset(1 - P))}{z_\emptyset(z_H - z_L)} \quad (\text{A.2})$$

Substituting the interior solution back into A.1 yields:

$$\begin{aligned} P_H &= \frac{z_H - z_L (J_L/J_H)}{z_H - z_L} \\ P_L &= P_H (J_H/J_L) \end{aligned}$$

and substituting it into A.2 renders:

$$\mu = \frac{z_H J_H (z_\emptyset - z_L)}{z_\emptyset (z_H J_H - z_L J_L) + z_H z_L (J_L - J_H)}$$

These match the expressions displayed in Theorem 1, and the posteriors follow from Bayes' rule: $\gamma^* = \mu^* P_H^* / P^*$ and $\nu^* = (\mu^* - \gamma^* P^*) / (1 - P^*)$. Any active interior RIBNE must satisfy the worker first-order conditions, Bayes' rule, and firm indifference used above, so at most one exists.

It does exist. The closed forms are interior under Assumption 2. The ordering $z_L < z_\emptyset < z_H$ gives $\nu^* = \frac{z_\emptyset - z_L}{z_H - z_L} \in (0, 1)$, and $\gamma^* = \frac{z_H}{z_\emptyset} \nu^*$ then satisfies $\nu^* < \gamma^* < 1$, the upper bound because $z_H(z_\emptyset - z_L) < z_\emptyset(z_H - z_L)$ whenever $z_H > z_\emptyset$. The bound $z_H J_H > z_L J_L$ gives $P_H^* > 0$, the ordering $J_L > J_H > 0$ gives $P_H^* < 1$ and $0 < P_L^* = (J_H / J_L) P_H^* < P_H^*$, and the same two conditions put the displayed P^* in $(0, 1)$. Bayes plausibility $\mu^* = \gamma^* P^* + \nu^* (1 - P^*)$ with $P^* \in (0, 1)$ then places μ^* strictly between the two posteriors, so $0 < \nu^* < \mu^* < \gamma^* < 1$. Finally, the worker's first-order conditions are sufficient as well as necessary: the due diligence problem (13) is the Shannon problem, whose solution is the global logit optimum (Matějka and McKay, 2015). Firm indifference holds by construction of P^* , and the compensation levels are the Nash bargaining solution, so the displayed profile is an active interior RIBNE.

For the converse, take any active interior RIBNE. We first show $w_H^* = 0$. The Nash bargaining problem gives $w_\omega^* = w_\omega^0$, so $w_H^* > 0$ requires $\psi \pi_H > (1 - \psi) u_H$. Suppose first that also $w_L^* > 0$. Then $W_\omega = \psi(\pi_\omega + u_\omega)$ and $J_\omega = (1 - \psi)(\pi_\omega + u_\omega)$ are proportional, so $W_H \geq W_L$ exactly as $J_H \geq J_L$; active due diligence requires distinct posteriors, hence $W_H \neq W_L$. By Bayes' rule $P_H \geq P_L$ exactly as $\gamma^* \geq \nu^*$, and the worker's optimum has $\gamma^* > \nu^*$ when $W_H > W_L$ and $\gamma^* < \nu^*$ when $W_L > W_H$. In either case the higher payoff job is both accepted more often and more profitable, so $P_H J_H = P_L J_L$ fails — a contradiction. Suppose instead $w_L^* = 0$. Then $\psi \pi_L \leq (1 - \psi) u_L < (1 - \psi) u_H < \psi \pi_H$ gives $\pi_L < \pi_H$, contradicting the environment ordering $\pi_L > \pi_H$. Hence $w_H^* = 0$, and $W_H = u_H > b$ by the environment ordering. If $W_L \geq b$, the worker could deviate to accepting every offer without acquiring information: this never loses in either state, gains $W_H - b > 0$ on any high-quality offer that was being rejected, and saves the attention cost of any strategy with distinct conditional acceptance rates, so it is strictly profitable whenever $P^* < 1$ — a contradiction. Hence $W_H > b > W_L$. The worker's optimal attention strategy then has distinct posteriors $\gamma^* > \nu^*$, so $P_H > P_L > 0$. Since $w_H^* = 0$, $J_H = \pi_H > 0$, and firm mixing with both qualities used requires $P_H J_H = P_L J_L$, so $J_L > J_H > 0$. Finally, the formula for the interior unconditional acceptance probability gives $z_H J_H > z_L J_L$, equivalently $J_L / J_H < z_H / z_L$. The Nash bargaining problem gives $w_\omega^* = w_\omega^0$, so the active interior equilibrium satisfies Assumption 2.

A.2 Primitive Sufficient Conditions for the Static Model

Corollary 1 (Primitive sufficient conditions). *Each of the following primitive regions implies Assumption 2, and hence the unique active interior RIBNE of Theorem 1.*

1. If $w_L^* = 0$, it is sufficient that

$$\begin{aligned} \psi\pi_H &\leq (1 - \psi)u_H, & \psi\pi_L &\leq (1 - \psi)u_L, \\ u_H &> b > u_L, & \pi_L &> \pi_H > 0, & \frac{\pi_L}{\pi_H} &< \exp\left(\frac{u_H - u_L}{\lambda}\right). \end{aligned}$$

2. If $w_L^* > 0 = w_H^*$, it is sufficient that

$$\begin{aligned} \psi\pi_H &\leq (1 - \psi)u_H, & \psi\pi_L &> (1 - \psi)u_L, \\ u_H &> b > \psi(\pi_L + u_L), & (1 - \psi)(\pi_L + u_L) &> \pi_H > 0, \end{aligned}$$

and

$$\frac{(1 - \psi)(\pi_L + u_L)}{\pi_H} < \exp\left(\frac{u_H - \psi(\pi_L + u_L)}{\lambda}\right).$$

Proof. In the first case, $W_H = u_H$, $W_L = u_L$, $J_H = \pi_H$, and $J_L = \pi_L$. In the second, $W_H = u_H$, $W_L = \psi(\pi_L + u_L)$, $J_H = \pi_H$, and $J_L = (1 - \psi)(\pi_L + u_L)$. Substituting these expressions into Assumption 2 gives exactly the displayed primitive inequalities. $\square$

Lemma 1 below shows that the attention and quality block depends on the environment only through the payoff and profit values, so the role of bargaining is exactly to place those values in the region of Assumption 2. When $w_H^* = 0 < w_L^*$, a compensation profile with the required payoff ordering $W_H > W_\emptyset > W_L$ and profit ordering $J_L > J_H$ exists if and only if $\psi(\pi_L + u_L) < \min\{\pi_L - \pi_H + u_L, b\}$, the two bounds being $J_L > J_H$ and $W_L < b$ evaluated at $W_L = \psi(\pi_L + u_L)$ and $J_L = (1 - \psi)(\pi_L + u_L)$.

A.3 The Normalized Per-Match Block

Both models share the same per-match problem, and every proof below works in the same normalized coordinates. Fix payoff values $W_H > W_\emptyset > W_L$ and firm values $J_L > J_H > 0$, and let

$$x := \exp\left\{\frac{W_H - W_\emptyset}{\lambda}\right\}, \quad z := \exp\left\{\frac{W_L - W_\emptyset}{\lambda}\right\}, \quad \rho := \frac{J_L}{J_H}, \quad s := \rho z, \quad (\text{A.3})$$

so that $x > 1 > z > 0$ and $\rho > 1$, and the bound on the profit ratio, $J_L/J_H < z_H/z_L$, is $s < x$. In the static model $(x, z) = (z_H/z_\emptyset, z_L/z_\emptyset)$, since $W_\emptyset = b$. In the dynamic model $W_\emptyset$ is the value of rejecting the current offer, so $(x, z) = (\hat{z}_H, \hat{z}_L)$ and (ρ, s) are built from the equilibrium firm values J_H^* and J_L^* . Write

$$Q := x(1 - z) + s(x - 1), \quad L := \ln \frac{x}{s} > 0, \quad \ell_Q := \ln \left(\frac{Q}{s(x - z)} \right), \quad (\text{A.4})$$

and note the following identities, each verified by direct expansion:

$$Q = (x - z) + (x - 1)(s - z) > x - z \quad (\text{A.5})$$

$$Q - s(x - z) = (1 - z)(x - s) > 0 \quad (\text{A.6})$$

$$xQ - (x - s)(x - z) = x^2(s - z) + z(x - s) > 0 \quad (\text{A.7})$$

By (A.6), $\ell_Q > 0$.

Lemma 1 (Fixed payoff attention and quality block). *Fix constants $W_H, W_\emptyset, W_L, J_H$, and J_L satisfying $x > 1 > z > 0$, $\rho > 1$, and $s < x$ in the notation (A.3). Then there is a unique active interior profile that solves the worker's due diligence problem (13) and the firm's quality problem (14) simultaneously, with the conditional acceptance probabilities given by Bayes' rule (9) and the rejection posterior implied by Bayes plausibility. It is given by the final two rows of Theorem 1 and, in normalized coordinates, reads*

$$\mu = \frac{x(1 - z)}{Q}, \quad P = \frac{x - s}{Q}, \quad P_H = \frac{x - s}{x - z}, \quad P_L = \frac{P_H}{\rho}, \quad \gamma = \frac{x(1 - z)}{x - z}, \quad \nu = \frac{1 - z}{x - z}. \quad (\text{A.8})$$

Proof. The proof of Theorem 1 in Appendix A.1 treats the payoff values as fixed constants and uses only the displayed orderings, so its attention and quality argument applies verbatim; (A.8) is that solution divided through by $z_\emptyset$. $\square$

The following log derivatives are obtained by direct differentiation of (A.8):

$$\frac{\partial \ln \mu}{\partial x} = -\frac{s}{xQ}, \quad \frac{\partial \ln P_H}{\partial x} = \frac{s - z}{(x - s)(x - z)}, \quad (\text{A.9})$$

$$\frac{\partial \ln \mu}{\partial s} = -\frac{x - 1}{Q}, \quad \frac{\partial \ln P_H}{\partial s} = -\frac{1}{x - s}, \quad (\text{A.10})$$

$$\left. \frac{\partial \ln \mu}{\partial z} \right|_\rho = -\frac{s(x - 1)}{zQ(1 - z)}, \quad \left. \frac{\partial \ln P_H}{\partial z} \right|_\rho = -\frac{x(s - z)}{z(x - s)(x - z)}, \quad (\text{A.11})$$

where derivatives “at fixed ρ ” move $s = \rho z$ together with z . Two combinations recur. The radial response of the offer mix,

$$\Sigma := - \left(x \frac{\partial \ln \mu}{\partial x} + z \frac{\partial \ln \mu}{\partial z} \Big|_{\rho} \right) = \frac{s}{Q} + \frac{s(x-1)}{Q(1-z)} = \frac{s(x-z)}{Q(1-z)} > 0, \quad (\text{A.12})$$

measures how much the offer mix improves when both surpluses fall by the same proportion; it is the composition gain of Proposition 5. And $x \left| \frac{\partial \ln \mu}{\partial x} \right| = s/Q$, which by (A.5) and Lemma 4(i) below satisfies $s \ln x \leq s(x-1) \leq Q$.

Lemma 2 (Conditional choice probability form). *Fix a prior $\mu \in (0, 1)$. On the interior domain $P \in (0, 1)$, the worker’s due diligence problem (13), stated over Bayes-plausible posterior strategies $(\gamma, P) \in \mathcal{A}(\mu)$, is equivalent to the same problem stated over conditional acceptance probabilities (P_H, P_L) : the two parameterizations describe the same set of feasible interior information structures and assign each structure the same objective value, so they have the same interior optima.*

Proof. Bayes’ rule (9) maps (γ, P) to $P_H = \gamma P / \mu$ and $P_L = (1 - \gamma)P / (1 - \mu)$, with inverse $P = \mu P_H + (1 - \mu)P_L$ and $\gamma = \mu P_H / P$; Bayes plausibility then gives $\nu = (\mu - \gamma P) / (1 - P)$. The map is a bijection onto

$$\{(P_H, P_L) \in [0, 1]^2 : 0 < \mu P_H + (1 - \mu)P_L < 1\},$$

which restricts to $(P_H, P_L) \in (0, 1)^2$ on the fully interior subdomain. Note that $P_H = P_L$ is feasible: it is the no information strategy, against which the active optimum is compared below. Active due diligence is the further condition $P_H \neq P_L$, not part of feasibility. The map leaves the expected payoff $\mu P_H W_H + (1 - \mu)P_L W_L + (1 - P)W_\emptyset$ unchanged by construction, and it leaves the information cost unchanged, because the mutual information between the state and the acceptance decision has the two equivalent expressions

$$\mathbb{H}(\mu) - P \mathbb{H}(\gamma) - (1 - P) \mathbb{H}(\nu) = \mathbb{H}(P) - \mu \mathbb{H}(P_H) - (1 - \mu) \mathbb{H}(P_L), \quad (\text{A.13})$$

the first conditioning the four joint probabilities $\{\gamma P, (1 - \gamma)P, \mu - \gamma P, 1 - \mu - (1 - \gamma)P\}$ on the acceptance decision and the second conditioning the same four on the state. Both sides equal $\mathbb{H}(\mu) + \mathbb{H}(P) - \mathbb{H}(\text{joint})$, where $\mathbb{H}(\text{joint})$ is the entropy of that common distribution, and both are therefore nonnegative. Hence the two problems have the same objective on the same feasible set. $\square$

Lemma 3 (Value of the fixed payoff attention problem). *At prior $\mu \in (0, 1)$, let P be the optimal unconditional acceptance probability in (13) at payoff indices $(x, 1, z)$. The optimized surplus over rejecting, in units of the attention cost, is*

$$d = \mu \ln(1 - P + Px) + (1 - \mu) \ln(1 - P + Pz), \quad (\text{A.14})$$

and in payoff units it is $D = \lambda d$. Evaluated at the block (A.8),

$$d(x, z, s) = \mu \ln x + (1 - \mu) \ln s + \ln \left(\frac{x - z}{Q} \right) = \mu L - \ell_Q. \quad (\text{A.15})$$

Proof. Write the problem in conditional acceptance probabilities, which Lemma 2 shows is equivalent. In those coordinates the surplus over rejecting is

$$d = \sum_{\omega} \mu_{\omega} \left[P_{\omega} \ln z_{\omega} - P_{\omega} \ln \frac{P_{\omega}}{P} - (1 - P_{\omega}) \ln \frac{1 - P_{\omega}}{1 - P} \right], \quad (z_H, z_L) = (x, z),$$

with $\mu_H = \mu$ and $\mu_L = 1 - \mu$. The two subtracted terms are the statewise Kullback–Leibler divergences between the conditional and unconditional acceptance distributions; averaged over states they sum to the mutual information, which is the right side of (A.13). The optimum satisfies the logit rule (A.1) with $z_{\emptyset} = 1$, namely $P_{\omega} = P z_{\omega} / (1 - P + P z_{\omega})$, so $P_{\omega} / P = z_{\omega} / (1 - P + P z_{\omega})$ and $(1 - P_{\omega}) / (1 - P) = 1 / (1 - P + P z_{\omega})$. Substituting, the bracket collapses to $[P_{\omega} + (1 - P_{\omega})] \ln(1 - P + P z_{\omega}) = \ln(1 - P + P z_{\omega})$, which is (A.14).

At the block (A.8), $1 - P + P x = 1 + P(x - 1) = [Q + (x - s)(x - 1)] / Q = x(x - z) / Q$, by the expansion $Q + (x - s)(x - 1) = x(1 - z) + x(x - 1) = x(x - z)$, and $1 - P + P z = 1 - P(1 - z) = [Q - (x - s)(1 - z)] / Q = s(x - z) / Q$ by (A.6). Taking logs and collecting the common factor gives (A.15). $\square$

Differentiating (A.15) directly,

$$\frac{\partial d}{\partial x} = \frac{(1 - z)[(x - s)Q - s(x - z)L]}{Q^2(x - z)}, \quad \frac{\partial d}{\partial s} = -\frac{x(1 - z)(x - 1)L}{Q^2}, \quad (\text{A.16})$$

$$\left. \frac{\partial d}{\partial z} \right|_{\rho} = \frac{(x - 1)[(x - s)Qz - xs(x - z)L]}{zQ^2(x - z)}. \quad (\text{A.17})$$

Lemma 4. For all $t > 1$, (i) $\ln t < t - 1$, (ii) $\ln t < \frac{t-1}{\sqrt{t}}$, and (iii) $\ln t > \frac{2(t-1)}{t+1}$.

Proof. All three expressions vanish at $t = 1$. Differencing gives (ii), since $\frac{d}{dt} \left[\frac{t-1}{\sqrt{t}} - \ln t \right] = \frac{(\sqrt{t}-1)^2}{2t^{3/2}} > 0$. The same argument gives (iii), since $\frac{d}{dt} \left[\ln t - \frac{2(t-1)}{t+1} \right] = \frac{(t-1)^2}{t(t+1)^2} > 0$. Part (i) is standard. $\square$

Lemma 5 (Properties of the optimized contact surplus). *On the interior region $x > 1 > z > 0$, $z < s < x$:*

(a) $d > 0$.

(b) $\frac{\partial d}{\partial x} > 0$.

(c) $\frac{\partial d}{\partial s} < 0$ (so d is decreasing in ρ at fixed z).

$$(d) \left. \frac{\partial d}{\partial z} \right|_{\rho} < 0.$$

Proof. (a) λd is the maximized value of the worker's attention problem with payoffs $\lambda \ln x > 0$, $\lambda \ln z < 0$, and rejection payoff 0, evaluated at the prior μ . Acquiring no information and rejecting every offer is feasible and yields exactly 0. Since the solution to the attention problem is unique and interior ($P \in (0, 1)$ because $z < s < x$), the optimal strategy differs from this feasible strategy, so d is strictly positive.

(b) We must show $(x - s)Q > s(x - z)L$. By Lemma 4(ii) with $t = x/s$, $L < \frac{x-s}{\sqrt{xs}}$, so $s(x - z)L < (x - z)(x - s)\sqrt{s/x} < (x - z)(x - s) \leq Q(x - s)$, where the last step is (A.5).

(c) Immediate from (A.16) since $L > 0$.

(d) We must show $(x - s)Qz < xs(x - z)L$. By Lemma 4(iii) with $t = x/s$, $L > \frac{2(x-s)}{x+s}$, so it suffices that $z(x + s)Q \leq 2xs(x - z)$. Define $q(s) := 2xs(x - z) - z(x + s)Q$. As a function of s , q is a concave quadratic (the coefficient on s^2 is $-z(x - 1) < 0$), with $q(z) = z(x - z)^2 > 0$ and $q(x) = 2x^2(1 - z)(x - z) > 0$. Hence $q > 0$ on $[z, x]$. $\square$

Part (d) is worth pausing over: raising z at a fixed profitability ratio ρ makes bad jobs *less* bad, yet strictly *lowers* the optimized contact surplus. The direct payoff gain is more than undone by the deterioration in the composition of offers, since μ falls as z rises.

The next lemma isolates the behavior of the key objects under a *radial* displacement of the surpluses, in which both log surpluses over rejection fall by the same amount.

Lemma 6 (Radial bounds). *On the interior region, with the z -derivatives taken at fixed ρ :*

$$(i) \left. x \frac{\partial \ln P_H}{\partial x} + z \frac{\partial \ln P_H}{\partial z} \right|_{\rho} = 0.$$

$$(ii) \left. x \frac{\partial d}{\partial x} + z \frac{\partial d}{\partial z} \right|_{\rho} = P - \frac{sx(x-z)L}{Q^2} = P - L\mu\Sigma < 1.$$

Proof. (i) $P_H = \frac{x-s}{x-z}$ is homogeneous of degree zero in (x, z, s) , and the displacement $(dx, dz, ds) \propto (x, z, s)$ is exactly the radial one. Moving z at fixed ρ moves $s = \rho z$ proportionally. Equivalently, multiply the expressions in (A.9) and (A.11) by x and z and add. (ii) Multiplying the expressions in (A.16) and (A.17) by x and z and adding, the coefficients collapse using $x(1 - z) + z(x - 1) = x - z$ and $x(1 - z) + x(x - 1) = x(x - z)$:

$$\left. x \frac{\partial d}{\partial x} + z \frac{\partial d}{\partial z} \right|_{\rho} = \frac{(x - s)Q - sx(x - z)L}{Q^2} = P - \frac{sx(x - z)L}{Q^2} < P < 1.$$

The second form follows from $\mu\Sigma = sx(x - z)/Q^2$ by (A.8) and (A.12). $\square$

Part (ii) bounds the radial derivative of the optimized contact surplus above by one: a radial contraction of the surpluses can lower d by at most one-for-one. With probability $1 - P$ the worker rejects the offer and is insulated from the contraction, and the second term provides further insulation — it can be strong enough that d rises under the contraction. That term has a direct reading. By the envelope theorem in P , the value (A.14) responds to the prior at rate $\frac{\partial d}{\partial \mu} = \ln(1 - P + Px) - \ln(1 - P + Pz) = L$, so $L\mu\Sigma$ is the composition gain Σ of (A.12) weighted by the worker's own valuation of a better offer mix.

A.4 Proof of Proposition 1

Work in normalized variables: $x := z_H/z_0$, $z := z_L/z_0$, $\rho := J_L/J_H$, $s := \rho z$, and $Q := x(1 - z) + s(x - 1)$. In these units the equilibrium objects of Theorem 1 read $\mu^* = x(1 - z)/Q$, $P^* = (x - s)/Q$, $P_H^* = (x - s)/(x - z)$, $P_L^* = P_H^*/\rho$, $\gamma^* = x(1 - z)/(x - z)$, and $\nu^* = (1 - z)/(x - z)$, and Assumption 2 gives $x > 1 > z > 0$, $\rho > 1$, and $z < s < x$. The three shocks move (x, z, ρ) as follows. The shock to u_L raises z at rate $\frac{z}{\lambda} \frac{dW_L}{du_L} > 0$ and weakly raises ρ . The effect on ρ is strict, at rate $\frac{1 - \psi}{J_H}$, iff $w_L^* > 0$. The shock to u_H raises x only. The shock to b lowers x and z equiproportionally ($\frac{dx}{db} = -\frac{x}{\lambda}$, $\frac{dz}{db} = -\frac{z}{\lambda}$) and leaves ρ unchanged, because a change in b moves neither w_ω^* nor J_ω .

Parts 1 and 2 then follow from the elementary partial derivatives of the closed forms, collected as (A.9)–(A.11) in Appendix A.3:

$$\frac{\partial \ln \mu^*}{\partial x} < 0, \quad \left. \frac{\partial \ln \mu^*}{\partial z} \right|_\rho < 0, \quad \frac{\partial \ln \mu^*}{\partial \rho} < 0, \quad \frac{\partial \ln P_H^*}{\partial x} > 0, \quad \left. \frac{\partial \ln P_H^*}{\partial z} \right|_\rho < 0, \quad \frac{\partial \ln P_H^*}{\partial \rho} < 0,$$

together with $\frac{\partial \gamma^*}{\partial x}, \frac{\partial \gamma^*}{\partial z}, \frac{\partial \nu^*}{\partial x}, \frac{\partial \nu^*}{\partial z} < 0$ (and γ^*, ν^* do not depend on ρ), and $P_L^* = P_H^*/\rho$.

For part 3, a rise in b by $\Delta b > 0$ scales (x, z) by the common factor $\varepsilon := e^{-\Delta b/\lambda} < 1$ at fixed ρ . $P_H^* = \frac{x - \rho z}{x - z}$ and $P_L^* = P_H^*/\rho$ are homogeneous of degree zero in (x, z) at fixed ρ , hence exactly invariant. Equivalently, z_0 appears nowhere in their expressions in Theorem 1. The posteriors rise: $\gamma^*(\varepsilon x, \varepsilon z) = \frac{x(1 - \varepsilon z)}{x - z}$ is strictly decreasing in ε , and likewise $\nu^*(\varepsilon x, \varepsilon z) = \frac{1 - \varepsilon z}{\varepsilon(x - z)}$. And μ^* rises because the radial derivative of $\ln \mu^*$ is strictly negative:

$$x \frac{\partial \ln \mu^*}{\partial x} + z \frac{\partial \ln \mu^*}{\partial z} \Big|_\rho = -\frac{s}{Q} \left[1 + \frac{x - 1}{1 - z} \right] = -\frac{s(x - z)}{Q(1 - z)} < 0$$

(this is the object $-\Sigma$ that reappears in Proposition 5).

A.5 Proof of Proposition 2, Part 1 (u_L)

Recall that we solve for P from the firm's first-order condition:

$$\frac{z_H J_H}{z_H P + z_\emptyset(1 - P)} = \frac{z_L J_L}{z_L P + z_\emptyset(1 - P)}$$

Since $z_H > z_\emptyset$, the denominator of the left side increases in P , so the left side decreases in P . Likewise, since $z_L < z_\emptyset$, the right side increases in P . It is readily verified that the left side increases in z_H and the right side in z_L and in J_L . Now suppose W_L decreases. Then, z_L decreases and so does the right side. Moreover, since W_L and J_L are both increasing functions of u_L , a decrease in u_L weakly decreases J_L as well, which lowers the right side further. When $w_L^* > 0$, $W_L = \psi(\pi_L + u_L)$ and $J_L = (1 - \psi)(\pi_L + u_L)$. When $w_L^* = 0$, $W_L = u_L$ and $J_L = \pi_L$. To restore equality, the right side needs to increase and the left side to decrease. Both of these effects are associated with an increase in P . Thus $\mathcal{U} = 1 - P$ decreases.

The welfare claim of the proposition — for the attention-inclusive measure — is proven in Appendix A.8. Here we show that gross worker welfare $\mathcal{W}^G := \mathcal{W} + \lambda \mathcal{I}(\gamma^*, P^*)$ falls as well. By Bayes' rule, $\mathcal{W}^G = \mu P_H W_H + (1 - \mu) P_L W_L + (1 - P) W_\emptyset$. Since $W_\omega = \lambda \ln z_\omega$ and $W_\emptyset = b = \lambda \ln z_\emptyset$, and since $\mu P_H + (1 - \mu) P_L = P$, the $\ln z_\emptyset$ terms collect exactly:

$$\frac{\mathcal{W}^G - b}{\lambda} = \mu P_H \ln \left(\frac{z_H}{z_\emptyset} \right) + (1 - \mu) P_L \ln \left(\frac{z_L}{z_\emptyset} \right),$$

which can be expressed as:

$$\begin{aligned} & P \left[\frac{z_H(z_\emptyset - z_L)}{z_\emptyset(z_H - z_L)} \ln \left(\frac{z_H}{z_\emptyset} \right) + \frac{z_L(z_H - z_\emptyset)}{z_\emptyset(z_H - z_L)} \ln \left(\frac{z_L}{z_\emptyset} \right) \right] = \\ & P \left[\frac{z_H z_\emptyset \ln \left(\frac{z_H}{z_\emptyset} \right) + z_L z_H \ln \left(\frac{z_L}{z_H} \right) - z_L z_\emptyset \ln \left(\frac{z_L}{z_\emptyset} \right)}{z_\emptyset(z_H - z_L)} \right] \equiv P(z_L) \mathcal{F}(z_L) \end{aligned}$$

Note that $\mathcal{F}(z_L) > 0$ in the interior equilibrium: the bracketed term is the worker's expected log surplus at the acceptance posterior, and with $P > 0$ obedience of the acceptance recommendation requires it to be strictly positive. From the previous part we know that $P'(z_L) < 0$. Thus showing that $\mathcal{F}'(z_L) < 0$ is sufficient for $\mathcal{W}^G$ to be decreasing in z_L . We have:

$$\mathcal{F}'(z_L) = \frac{(z_H - z_\emptyset) \left[z_H \ln \left(\frac{z_L}{z_H} \right) + z_H - z_L \right]}{z_\emptyset(z_H - z_L)^2}$$

Since $z_H > z_\emptyset > z_L > 0$, it is enough to sign:

$$\mathcal{K}(z_H, z_L) = z_H \ln \left(\frac{z_L}{z_H} \right) + z_H - z_L$$

Note that $\mathcal{K}(z_\emptyset, z_\emptyset) = 0$. Also:

$$\begin{aligned}\frac{\partial \mathcal{K}}{\partial z_H} &= \ln \left(\frac{z_L}{z_H} \right) < 0 \\ \frac{\partial \mathcal{K}}{\partial z_L} &= \frac{z_H - z_L}{z_L} > 0\end{aligned}$$

Thus for any $z_H > z_\emptyset > z_L > 0$ we have that:

$$\mathcal{K}(z_H, z_L) < \mathcal{K}(z_\emptyset, z_L) < \mathcal{K}(z_\emptyset, z_\emptyset) = 0$$

The first inequality follows because $\mathcal{K}$ is decreasing in its first argument and $z_H > z_\emptyset$, and the second because $\mathcal{K}$ is increasing in its second argument and $z_L < z_\emptyset$. Thus we have that $\mathcal{F}'(z_L) < 0$, so $\frac{\partial}{\partial z_L} \mathcal{W}^G < 0$ implying $\frac{\partial}{\partial W_L} \mathcal{W}^G < 0$ as required (and the effect of u_L on J_L only reinforces this, since it further decreases P while leaving $\mathcal{F}$ unchanged).

The same decomposition delivers the π_L direction claimed in Proposition 7. When $w_L^* = 0$, an increase in π_L raises J_L only, which lowers P and leaves $\mathcal{F}$ unchanged; when $w_L^* > 0$, it raises W_L and J_L at the same rates as u_L , the case just treated. In both regimes $\mathcal{W}^G$ falls.

A.6 Proof of Proposition 2, Part 2 (u_H)

Work in the normalized variables defined in the proof of Proposition 1. An increase in u_H raises x (at rate $\frac{dx}{du_H} = \frac{x}{\lambda}$, since $W_H = u_H$ and $w_H^* = 0$) and leaves z and ρ unchanged.

For unemployment, $P^* = (x-s)/Q$, and direct expansion gives the identity $Q - (x-s)(1-z+s) = s(s-z)$, so

$$\frac{\partial \ln P^*}{\partial x} = \frac{1}{x-s} - \frac{1-z+s}{Q} = \frac{Q - (x-s)(1-z+s)}{(x-s)Q} = \frac{s(s-z)}{(x-s)Q} > 0.$$

Hence P^* strictly rises and $\mathcal{U} = 1 - P^*$ strictly falls.

The attention-inclusive welfare claim is proven in Appendix A.8. For the gross measure, write $\mathcal{W}^G - b = P^* [\gamma^* W_H + (1 - \gamma^*) W_L - b]$ and define $h(x) := P^*(x) [\gamma^*(x) \ln x + (1 - \gamma^*(x)) \ln z]$ with $\gamma^* = \frac{x(1-z)}{x-z}$, so that $\mathcal{W}^G - b = \lambda h(x)$. Then

$$h'(x) = \frac{\partial P^*}{\partial x} [\gamma^* \ln x + (1 - \gamma^*) \ln z] + P^* \left[\frac{\gamma^*}{x} + \frac{\partial \gamma^*}{\partial x} \ln \frac{x}{z} \right].$$

The first term is strictly positive: $\frac{\partial P^*}{\partial x} > 0$ as above, and the bracket is the worker's surplus from accepting at posterior γ^* , which is strictly positive in the interior equilibrium. For the second term,

$\frac{\gamma^*}{x} = \frac{1-z}{x-z}$ and $\frac{\partial \gamma^*}{\partial x} = -\frac{z(1-z)}{(x-z)^2}$, so

$$P^* \left[\frac{\gamma^*}{x} + \frac{\partial \gamma^*}{\partial x} \ln \frac{x}{z} \right] = P^* \frac{1-z}{x-z} \left[1 - \frac{z \ln(x/z)}{x-z} \right] \geq 0,$$

since $z \ln(x/z) \leq z \left(\frac{x}{z} - 1 \right) = x - z$ by $\ln t \leq t - 1$. Hence $h'(x) > 0$ and $\mathcal{W}^G$ strictly rises.

A.7 Proof of Proposition 2, Part 3 (b)

We solve for P from the firm's first-order condition, which implies:

$$z_H J_H [z_L P + z_\emptyset (1 - P)] = z_L J_L [z_H P + z_\emptyset (1 - P)]$$

Note that a change in b leaves w_ω^* , and hence J_ω and z_ω , unchanged. Only $z_\emptyset$ increases. The terms in square brackets increase in $W_\emptyset$ at the same rate. However, as in the interior equilibrium $z_H J_H > z_L J_L$, the left side increases at a faster rate than the right side at any fixed P . Since $z_L < z_\emptyset < z_H$, the first square bracket decreases in P and the second square bracket increases in P , so the left side decreases and the right side increases in P . Thus P has to increase in $W_\emptyset$ to keep both sides equal, which implies that $\mathcal{U} = 1 - P$ decreases in $W_\emptyset$.

For welfare, the attention-inclusive claim is proven in Appendix A.8; here we show that gross worker welfare $\mathcal{W}^G = \mathcal{W} + \lambda \mathcal{I}(\gamma^*, P^*)$ rises as well. Recall that:

$$\mathcal{W}^G = P[\gamma W_H + (1 - \gamma)W_L] + (1 - P)W_\emptyset$$

In proving that unemployment decreases we have shown that P increases in b . Now consider:

$$\begin{aligned} \gamma W_H + (1 - \gamma)W_L &= \lambda [\gamma \ln z_H + (1 - \gamma) \ln z_L] \\ &= \lambda \left[\frac{z_H}{z_\emptyset} \frac{z_\emptyset - z_L}{z_H - z_L} \cdot \ln z_H + \left(1 - \frac{z_H}{z_\emptyset} \frac{z_\emptyset - z_L}{z_H - z_L} \right) \cdot \ln z_L \right] \end{aligned}$$

The derivative of this term with respect to $z_\emptyset$ is

$$\frac{\partial}{\partial z_\emptyset} [\gamma W_H + (1 - \gamma)W_L] = \frac{\lambda z_H z_L \ln \left(\frac{z_H}{z_L} \right)}{z_\emptyset^2 (z_H - z_L)} > 0$$

For $P > 0$ we clearly need $\gamma W_H + (1 - \gamma)W_L > W_\emptyset$. Thus an increase in both P and $\gamma W_H + (1 - \gamma)W_L$ is sufficient for an increase in $\mathcal{W}^G$ if $W_\emptyset$ does not decrease. Since it, in fact, increases, so does $\mathcal{W}^G$.

A.8 Worker Value Including the Disutility of Learning

Both baseline welfare measures are attention inclusive. The static $\mathcal{W}$ is the equilibrium value of the worker's due diligence problem (13), and the dynamic $\mathcal{W}$ aggregates equilibrium value functions, in which workers internalize every cost they incur, including the disutility of due diligence $\lambda\mathcal{I}$ and the search cost C . This subsection records the closed form of the static measure and analyzes the corresponding *gross* measure, which excludes the resource cost of attention, as in search and matching models that feature no learning cost. No comparative static conclusion in the static model depends on the distinction, and the result is fully analytic. Define static *gross* worker welfare by adding back the equilibrium attention cost:

$$\mathcal{W}^G := \mathcal{W} + \lambda\mathcal{I}(\gamma^*, P^*) = P^* (\gamma^* W_H + (1 - \gamma^*) W_L) + (1 - P^*) W_\emptyset.$$

Proposition 7. *In the unique interior RIBNE of the static model, worker welfare admits the closed-form expression*

$$\mathcal{W} = b + \lambda d(x, z, s), \quad d(x, z, s) = \mu^* \ln x + (1 - \mu^*) \ln s + \ln \left(\frac{x - z}{Q} \right), \quad (\text{A.18})$$

in the normalized coordinates (A.3) of the static model, and its comparative static signs are

$$\frac{\partial \mathcal{W}}{\partial u_L} < 0, \quad \frac{\partial \mathcal{W}}{\partial \pi_L} < 0, \quad \frac{\partial \mathcal{W}}{\partial u_H} > 0, \quad \frac{\partial \mathcal{W}}{\partial b} > 0.$$

Gross worker welfare $\mathcal{W}^G$ moves in the same four directions.

Proof. Closed-form expression. By definition,

$$\mathcal{W} - W_\emptyset = P^* [\gamma^* (W_H - W_\emptyset) + (1 - \gamma^*) (W_L - W_\emptyset)] - \lambda\mathcal{I}(\gamma^*, P^*),$$

which is the value of the worker's attention problem at payoff indices $(x, 1, z)$ and the equilibrium prior μ^* . In the static model $W_H = u_H$, $W_L = u_L + w_L^*$, $W_\emptyset = b$, $J_H = \pi_H$, and $J_L = \pi_L - w_L^*$, and the interior equilibrium satisfies $x > 1 > z > 0$ and $z < s < x$, so Lemma 3 applies and gives (A.18).

Comparative statics. The parameters enter (A.18) only through (x, z, s) and the additive b .

- u_H enters only through x , with $\frac{dx}{du_H} = \frac{x}{\lambda}$ (as $W_H = u_H$ and $w_H^* = 0$). Hence $\frac{\partial \mathcal{W}}{\partial u_H} = x \frac{\partial d}{\partial x} > 0$ by Lemma 5(b).
- u_L and π_L enter only through (z, s) . If $w_L^* > 0$, then $W_L = \psi(\pi_L + u_L)$ and $J_L = (1 - \psi)(\pi_L + u_L)$, so an increase in either u_L or π_L raises z at rate $\frac{\psi z}{\lambda}$ and raises ρ at rate $\frac{1 - \psi}{J_H}$. If $w_L^* = 0$, an increase in u_L raises z only, and an increase in π_L raises ρ only. In every case the

displacement of (z, s) is a nonnegative combination, with at least one strict, of the fixed ρ z direction and the pure s direction, along both of which d strictly falls by Lemma 5(c)–(d). Hence $\frac{\partial \mathcal{W}}{\partial u_L} < 0$ and $\frac{\partial \mathcal{W}}{\partial \pi_L} < 0$.

- b enters through the additive term and radially: $\frac{dx}{db} = -\frac{x}{\lambda}$ and $\frac{dz}{db} = -\frac{z}{\lambda}$ at fixed ρ , since a change in b moves neither w_L^* nor ρ . Hence, by Lemma 6(ii),

$$\frac{\partial \mathcal{W}}{\partial b} = 1 - \left(x \frac{\partial d}{\partial x} + z \frac{\partial d}{\partial z} \Big|_{\rho} \right) = 1 - P^* + \frac{sx(x-z)L}{Q^2} \geq 1 - P^* > 0. \quad (\text{A.19})$$

The gross directions are established in Appendices A.5–A.7. □

Two aspects of this result are worth emphasizing. First, the worker value conclusions of the static model are driven by the equilibrium response of offer composition, not by the resource cost of attention: the distinction between the gross and attention-inclusive measures changes no sign. Indeed, for the counterintuitive u_L result the attention-inclusive proof is *shorter* than the gross one. An increase in u_L (or π_L) lowers the worker's equilibrium problem value directly through the deterioration of the offer pool, with no auxiliary argument needed. Second, the effect of the unemployment payoff takes a particularly transparent form: (A.19) shows that an increase in b erodes all employment surpluses one-for-one, but the worker is exposed to that erosion only on the accepted branch, of probability P^* , while with probability $1 - P^*$ due diligence ends in rejection and the worker simply pockets the higher b .

B Dynamic Equilibrium Characterization

Two value dates carry through the proof. The offer stage values W_ω and $W_\emptyset$ are the payoffs of accepting a type ω offer and of rejecting it, evaluated when the offer is in hand. The state values V_ω and $V_\emptyset$ are evaluated when their holder needs to make a decision: $V_\emptyset$ before the jobless worker searches, and V_ω , with payoffs realized and the separation shock survived, before the worker's quit decision; the employment stocks are measured immediately after that decision. At a quit node, staying leads to another committed production period whose compensation is set by that period's negotiation. The values are linked by $W_\emptyset = b + \beta V_\emptyset$, so rejecting means collecting the benefit and returning to the search pool, and by the reservation flow value $\hat{w} = (1 - \beta)V_\emptyset$. The static model is the special case $W_\emptyset = b$; the dynamic model is not.

B.1 Proof of Theorem 2

1. We first show that every interior SRSRIE has $\eta_H = 0$ and $\eta_L = 1$. Strict search participation gives $V_\emptyset > W_\emptyset$. From the recursions (25)–(27), $W_H - W_\emptyset = u_H + w_H - b + \beta(1 - \delta)(V_H - V_\emptyset) > 0$,

since $u_H > b$, $w_H \geq 0$, and $V_H \geq V_\emptyset$. Next, $W_L < W_\emptyset$: if $W_L \geq W_\emptyset$, accepting every offer without acquiring information would weakly improve the payoff in each state and save the strictly positive attention cost of the active strategy (active due diligence, $\gamma \neq \nu$, implies $\mathcal{I}(\gamma, P) > 0$), contradicting the optimality of (γ, P) with $P \in (0, 1)$. Since $V_\emptyset > W_\emptyset > W_L$, the employed worker's problem (29) gives $\eta_L = 1$ and $V_L = V_\emptyset$. Next, the searching branch of (28) evaluated at the optimum gives $V_\emptyset - W_\emptyset = \alpha_w (P [\gamma(W_H - W_\emptyset) + (1 - \gamma)(W_L - W_\emptyset)] - \lambda \mathcal{I}(\gamma, P)) - C$. The term in parentheses is at most $P(W_H - W_\emptyset)$, which is below $W_H - W_\emptyset$ since $P < 1$; with $\alpha_w < 1$ and $C > 0$, it follows that $V_\emptyset - W_\emptyset < W_H - W_\emptyset$, so $W_H > V_\emptyset$, and (29) gives $\eta_H = 0$ and $V_H = W_H$. Finally, $W_H > W_\emptyset > W_L$ implies $\gamma > \nu$ and hence $P_H > P_L$, so interior firm mixing and free entry give $P_H J_H = P_L J_L > 0$ and therefore $J_L > J_H > 0$. With the quitting regime established, equations (25)–(27) become:

$$\begin{aligned} W_H &= w_H + u_H + \beta(1 - \delta)W_H + \beta\delta V_\emptyset \\ W_L &= w_L + u_L + \beta V_\emptyset \\ W_\emptyset &= b + \beta V_\emptyset \end{aligned}$$

These simplify to:

$$\begin{aligned} W_H &= \frac{w_H + u_H + \beta\delta V_\emptyset}{1 - \beta(1 - \delta)} \\ W_L &= w_L + u_L + \beta V_\emptyset \end{aligned}$$

By $\eta_H = 0$ and $\eta_L = 1$ we also get:

$$\begin{aligned} J_H &= \pi_H - w_H + \beta(1 - \delta)J_H \Leftrightarrow J_H = \frac{\pi_H - w_H}{1 - \beta(1 - \delta)} \\ J_L &= \pi_L - w_L \end{aligned}$$

These quit strategies are the continuations in the per-period bargaining problems below; the current candidate compensation is sunk before the next quit decision, so they apply at every candidate.

2. Let us solve for the negotiated compensation. The per-period Nash bargaining problems are

$$w_\omega = \arg \max_{0 \leq w \leq \pi_\omega} [W_\omega(w) - (W_\emptyset - b)]^\psi [J_\omega(w)]^{1-\psi} \text{ for } \omega = H, L,$$

where, under the established quitting regime, the candidate values (36) read

$$\begin{aligned} W_H(w) &= u_H + w + \beta(1 - \delta)W_H + \beta\delta V_\emptyset, & J_H(w) &= \pi_H - w + \beta(1 - \delta)J_H, \\ W_L(w) &= u_L + w + \beta V_\emptyset, & J_L(w) &= \pi_L - w, \end{aligned}$$

with the continuation values W_H , J_H , and $V_\emptyset$ at their stationary equilibrium levels (by the previous step, $V_H = W_H$ and $V_L = V_\emptyset$). Each candidate value is affine in w with slope ± 1 , so each Nash product is strictly log-concave on its positive surplus domain and the first order or corner condition identifies the global per-period solution. The first-order condition for $\omega = H$ is

$$\psi J_H(w) = (1 - \psi) (W_H(w) - \beta V_\emptyset).$$

At the stationary solution, $W_H = W_H(w_H)$ and $J_H = J_H(w_H)$, and the recursions of the previous step give $W_H = \frac{w_H + u_H + \beta \delta V_\emptyset}{1 - \beta(1 - \delta)}$ and $J_H = \frac{\pi_H - w_H}{1 - \beta(1 - \delta)}$; substituting, the stationary compensation solves

$$\psi \left[\frac{\pi_H - w}{1 - \beta(1 - \delta)} \right] = (1 - \psi) \left[\frac{w + u_H + \beta \delta V_\emptyset}{1 - \beta(1 - \delta)} - \beta V_\emptyset \right]$$

We can now solve for w_H :

$$w_H = \psi \pi_H - (1 - \psi) [u_H - \beta(1 - \beta)(1 - \delta) V_\emptyset]$$

Similarly, the first-order condition for w_L is:

$$\psi [\pi_L - w] = (1 - \psi) [w + u_L]$$

which gives

$$w_L = \psi \pi_L - (1 - \psi) u_L$$

as in the static case.

3. When the respective nonnegativity constraint is slack, substituting the unconstrained solutions back into W_H and W_L gives:

$$W_H = \frac{\psi \pi_H - (1 - \psi) [u_H - \beta(1 - \beta)(1 - \delta) V_\emptyset] + u_H + \beta \delta V_\emptyset}{1 - \beta(1 - \delta)}$$

$$W_L = \psi \pi_L - (1 - \psi) u_L + u_L + \beta V_\emptyset$$

which simplify to:

$$W_H = \frac{\psi(\pi_H + u_H) + [(1 - \psi)(1 - \beta)(1 - \delta) + \delta] \beta V_\emptyset}{1 - \beta(1 - \delta)}$$

$$W_L = \psi(\pi_L + u_L) + \beta V_\emptyset$$

W_H can also be expressed as:

$$W_H = \frac{\psi(\pi_H + u_H) + [1 - (1 - \delta)(\psi + (1 - \psi)\beta)]\beta V_\emptyset}{1 - \beta(1 - \delta)}$$

Substituting into J_H and J_L yields:

$$J_H = (1 - \psi) \left[\frac{\pi_H + u_H - \beta(1 - \beta)(1 - \delta)V_\emptyset}{1 - \beta(1 - \delta)} \right]$$

$$J_L = (1 - \psi)(\pi_L + u_L)$$

4. The Nash problem is solved on $0 \leq w_\omega \leq \pi_\omega$, but the upper bound cannot bind in an active interior SRSRIE: free entry requires $\alpha_f P_\omega J_\omega = F > 0$, hence $J_\omega > 0$ and $w_\omega < \pi_\omega$. Only the nonnegativity constraint can bind, so

$$w_H = \max \{ \psi \pi_H - (1 - \psi)[u_H - \beta(1 - \beta)(1 - \delta)V_\emptyset], 0 \}$$

$$w_L = \max \{ \psi \pi_L - (1 - \psi)u_L, 0 \}$$

The low-quality candidate is therefore the static candidate, $w_L = w_L^0$, and Nash optimality at w_L^0 — the first-order condition when $w_L^0 > 0$, the corner condition $\psi \pi_L \leq (1 - \psi)u_L$ when $w_L^0 = 0$ — implies, in both regimes,

$$\psi (\pi_L - w_L^0) \leq (1 - \psi) (u_L + w_L^0).$$

It remains to exclude $w_H > 0$. If the high-quality constraint is slack, the first-order condition reads $\psi J_H = (1 - \psi)(W_H - \beta V_\emptyset)$, and the ordering $W_H > W_\emptyset$ established above gives $W_H - \beta V_\emptyset > b$, so $\psi J_H > (1 - \psi)b$. Part 1 of Assumption 3 gives $\pi_L - w_L^0 > \frac{\pi_H}{1 - \beta(1 - \delta)} \geq \frac{\pi_H - w_H}{1 - \beta(1 - \delta)} = J_H$, and part 2 gives $b > u_L + w_L^0$. Chaining, $\psi(\pi_L - w_L^0) > \psi J_H > (1 - \psi)b > (1 - \psi)(u_L + w_L^0)$, which contradicts the displayed Nash optimality condition. Hence $w_H = 0$ in both low-quality compensation regimes.

5. The bargaining regime is therefore $w_H^* = 0$ and $w_L^* = w_L^0$, with

$$J_H = \frac{\pi_H}{1 - \beta(1 - \delta)}, \quad J_L = \pi_L - w_L^0.$$

The conditions defining this regime are the maintained assumptions, evaluated at the candidate value of unemployment $V_\emptyset = \hat{w}/(1 - \beta)$. The profit ordering $J_L > J_H$ is part 1 of Assumption 3. The payoff ordering $W_\emptyset > W_L$ is part 2, $b > u_L + w_L^0$. The binding constraint $w_H^* = 0$ is part 3, $\psi \pi_H < (1 - \psi)(u_H - \beta(1 - \delta)\hat{w})$. The remaining ordering $W_H > W_\emptyset$ reads $u_H > (1 - \beta(1 - \delta))b + \beta(1 - \delta)\hat{w}$, whose right side is a strict convex combination of b and $\hat{w}$ and is therefore implied by the admissibility requirement $u_H > \hat{w} > b$ of Theorem 2. The witness equilibria in Table C.1 satisfy all four.

The surplus indices of the body, $\hat{z}_\omega = \exp\{(W_\omega - W_\emptyset)/\lambda\}$, therefore read

$$\hat{z}_L = \exp\left\{\frac{u_L + w_L^0 - b}{\lambda}\right\} < 1 < \hat{z}_H = \exp\left\{\frac{(u_H - b) + \frac{\beta(1-\delta)(u_H - \hat{w})}{1-\beta(1-\delta)}}{\lambda}\right\},$$

with $\hat{w} = (1-\beta)V_\emptyset$. Only $\hat{z}_H$ depends on the unknown $\hat{w}$, and it does so strictly monotonically, so we may index candidates by $\hat{z}_H$ and recover

$$\hat{w}(\hat{z}_H) = u_H - k(b + \lambda \ln \hat{z}_H - u_H), \quad \frac{d\hat{w}(\hat{z}_H)}{d\hat{z}_H} = -\frac{k\lambda}{\hat{z}_H} < 0, \quad k := \frac{1 - \beta(1 - \delta)}{\beta(1 - \delta)}.$$

The per-match block can therefore be solved at each candidate $\hat{z}_H$, with the implied offer composition $\mu(\hat{z}_H)$ recovered from the shared block; the candidate surplus index is the only unknown.

6. Restrict attention to the active candidate domain $\hat{z}_H J_H > \hat{z}_L J_L$. The worker's attention problem at payoffs $(\lambda \ln \hat{z}_H, \lambda \ln \hat{z}_L, 0)$ is the static due diligence problem (13) with exponentiated payoffs $(z_H, z_\emptyset, z_L) = (\hat{z}_H, 1, \hat{z}_L)$, and the zero profit conditions for the two offer types,

$$\alpha_f P_H J_H = F = \alpha_f P_L J_L,$$

require the same indifference condition $P_H J_H = P_L J_L$ as in the static model. On this domain the payoff orderings of Lemma 1 hold: $\hat{z}_H > 1 > \hat{z}_L > 0$ by the ordering in the previous step and $J_L > J_H > 0$ by part 1 of Assumption 3. Lemma 1 therefore yields the candidate maps $\mu(\hat{z}_H)$, $P(\hat{z}_H)$, $P_H(\hat{z}_H)$, $P_L(\hat{z}_H)$, $\gamma(\hat{z}_H)$, and $\nu(\hat{z}_H)$ of part 2 of Theorem 2, each in $(0, 1)$. The resulting prior lies strictly inside the interval

$$\frac{1 - \hat{z}_L}{\hat{z}_H - \hat{z}_L} < \mu(\hat{z}_H) < \frac{\hat{z}_H(1 - \hat{z}_L)}{\hat{z}_H - \hat{z}_L} :$$

the lower inequality is equivalent to $\hat{z}_H J_H > \hat{z}_L J_L$ and the upper to $J_L > J_H$. An active equilibrium with both offer types posted therefore requires $J_L > J_H$; we do not characterize the boundary outcomes that may arise when $J_L \leq J_H$.

7. Free entry for either offer type gives the same filling probability,

$$\alpha_f(\hat{z}_H) = \frac{F}{J_L P_L(\hat{z}_H)} = \frac{F}{J_H P_H(\hat{z}_H)} = \frac{F(\hat{z}_H - \hat{z}_L)}{\hat{z}_H J_H - \hat{z}_L J_L},$$

which is strictly decreasing on the active domain, since

$$\alpha'_f(\hat{z}_H) \sim \left(\hat{z}_H - \frac{J_L}{J_H} \hat{z}_L\right) - (\hat{z}_H - \hat{z}_L) = \left(1 - \frac{J_L}{J_H}\right) \hat{z}_L \sim 1 - \frac{J_L}{J_H} < 0,$$

and falls to F/J_H as $\hat{z}_H \rightarrow \infty$. Hence $F < J_H$ is necessary for a filling probability below one, as free entry also requires directly: the posting cost is paid whether or not the offer is accepted, and not even high-quality offers are always accepted. Admissibility restricts candidates to those with $\alpha_f(\hat{z}_H) \in (0, 1)$, and we do not characterize the boundary regimes outside that range.

8. Reservation consistency determines the meeting probability. At a candidate $\hat{z}_H$ the searching branch of (28) holds with equality at $\hat{w}(\hat{z}_H)$, which says that the flow cost of searching, C plus the forgone reservation flow value, equals the expected gain from a contact,

$$C + \hat{w}(\hat{z}_H) - b = \alpha_w(\hat{z}_H) D(\hat{z}_H), \quad D(\hat{z}_H) := P[\gamma \lambda \ln \hat{z}_H + (1 - \gamma) \lambda \ln \hat{z}_L] - \lambda \mathcal{I}(\gamma, P),$$

the optimized net surplus from a contact, in which $\mathcal{I}(\gamma, P)$ is the entropy reduction at the candidate's own prior $\mu(\hat{z}_H)$. Rearranging gives (42),

$$\alpha_w(\hat{z}_H) = \frac{n(\hat{z}_H)}{D(\hat{z}_H)}, \quad n(\hat{z}_H) := C + \hat{w}(\hat{z}_H) - b.$$

Lemma 3 evaluates the contact surplus in closed form, $D = \lambda d$, and Lemma 5(b) shows d is strictly increasing in $\hat{z}_H$: raising the reward to employment in the high state raises the optimized surplus from a contact. Since n is strictly decreasing, α_w is strictly decreasing as well.

9. On the selected interior region $\alpha_f(\hat{z}_H), \alpha_w(\hat{z}_H) \in (0, 1)$, neither quantity cap in M binds, so $M = A\mathcal{V}^{1-\phi}\mathcal{U}^\phi$ and the two matching probabilities each define a corresponding ratio of vacancies to unemployment,

$$\theta_f(\hat{z}_H) = \left[\frac{\alpha_f(\hat{z}_H)}{A} \right]^{-\frac{1}{\phi}}, \quad \theta_w(\hat{z}_H) = \left[\frac{\alpha_w(\hat{z}_H)}{A} \right]^{\frac{1}{1-\phi}}.$$

Boundary matching regimes, in which a quantity cap binds and one side of the market is matched for sure, are outside the scope of our analysis.

10. Market clearing requires the two implied ratios to agree. Setting $\theta_f(\hat{z}_H) = \theta_w(\hat{z}_H)$ is equivalent to the scalar equation (43): raising $\alpha_f = A\theta^{-\phi}$ to the power $1 - \phi$, raising $\alpha_w = A\theta^{1-\phi}$ to the power ϕ , and multiplying eliminates θ . Conversely, given a solution $\hat{z}_H^*$ of (43), setting $\theta^* = (\alpha_f(\hat{z}_H^*)/A)^{-1/\phi}$ recovers market clearing.
11. An admissible solution of (43) is an equilibrium. Evaluate every candidate map at $\hat{z}_H^*$. Each condition of Definition 2 then holds. Lemma 1 makes (γ^*, P^*) optimal in the attention problem at the payoff indices $(\hat{z}_H^*, 1, \hat{z}_L)$ and leaves the entrant indifferent between qualities, $P_H^* J_H^* = P_L^* J_L^*$, so mixing at μ^* is a best response. Reservation consistency is the definition of α_w , so the searching branch of (28) holds with equality at $\hat{w}(\hat{z}_H^*)$, and admissibility gives $\hat{w}^* > b$: search is strictly preferred to permanent joblessness. The bargaining and quitting

analysis above gives $w_H^* = 0$, $w_L^* = w_L^0$, and $\eta_H^* = 0 < \eta_L^* = 1$, with $u_H > \hat{w}^*$ delivering $V_H^* > V_\emptyset^*$, and the value recursions (25)–(27) and (30)–(31) hold by construction of the payoff indices. Free entry is the definition of α_f . Admissibility puts $\alpha_f^*, \alpha_w^* \in (0, 1)$, so neither quantity cap in M binds, the matching probabilities (24) are those of the smooth technology, and both sides face the common tightness θ^* . Because $\eta_L^* = 1$, no worker is employed in a low-quality job at the measurement date, so $E_L^* = 0$ solves (38); the stocks displayed below solve (37) and the population identity (22).

Conversely, every interior SRSRIE satisfies the quitting, bargaining, and attention conditions established above, so it lies on the candidate maps at its own surplus index, and its common market tightness makes that index an admissible solution of (43). An interior SRSRIE therefore exists if and only if (43) has an admissible solution. The equivalence is stated on the strict bargaining regime maintained by part 3 of Assumption 3. At the knife edge where the unconstrained high-quality transfer is exactly zero, $w_H^* = 0$ still obtains and the same construction applies with the weak inequality.

12. The solution is unique. Free entry $\alpha_f = F/(J_H^* P_H)$ is strictly decreasing in $\hat{z}_H$ because $\frac{\partial \ln P_H}{\partial x} > 0$ by (A.9), and $\alpha_w = n/(\lambda d)$ is strictly decreasing because $n'(x) = -k\lambda/x < 0$ while $\frac{\partial d}{\partial x} > 0$ by Lemma 5(b). Hence the left side of (43) is strictly decreasing on the admissible domain and (43) has at most one admissible solution.
13. Every remaining equilibrium object is an explicit function of the root. Free entry $\alpha_f P_H J_H^* = F$ gives $\theta^* = (\alpha_f/A)^{-1/\phi} = (AJ_H^* P_H^*/F)^{1/\phi}$, and $\hat{w} = (1-\beta)V_\emptyset^*$ gives the values: with $w_H^* = 0$ and $\eta_H^* = 0$, $V_H^* = \frac{u_H + \beta\delta V_\emptyset^*}{1-\beta(1-\delta)}$, so

$$V_H^* - V_\emptyset^* = \frac{u_H - (1-\beta)V_\emptyset^*}{1-\beta(1-\delta)} = \frac{u_H - \hat{w}^*}{1-\beta(1-\delta)} > 0$$

in the interior equilibrium. Market clearing gives the unemployment rate and, from the ratio of vacancies to unemployment, the vacancy stock,

$$\mathcal{U}^* = \frac{\delta}{\delta + (1-\delta)f^*}, \quad f^* = \alpha_w(\hat{z}_H^*)\mu(\hat{z}_H^*)P_H(\hat{z}_H^*), \quad \mathcal{V}^* = \theta^*\mathcal{U}^*.$$

Substituting $\alpha_w^* = A(\theta^*)^{1-\phi}$ and then θ^* reduces the high-quality job-finding rate to the two forms in (44):

$$\begin{aligned} f^* &= A(\theta^*)^{1-\phi}\mu^*P_H^* = A \cdot \frac{\theta^*}{(\theta^*)^\phi} \cdot \mu^*P_H^* = \frac{F}{J_H^*P_H^*} \theta^* \mu^*P_H^* = \frac{F}{J_H^*} \mu^*\theta^* \\ &= A^{\frac{1}{\phi}} \left(\frac{J_H^*}{F} \right)^{\frac{1-\phi}{\phi}} \mu^*(P_H^*)^{\frac{1}{\phi}}, \end{aligned}$$

using $A(\theta^*)^{-\phi} = \alpha_f^* = F/(J_H^*P_H^*)$ in the third equality. Finally $\eta_L^* = 1$ implies $E_L^* = 0$, so $\mathcal{W}^* = E_H^*V_H^* + \mathcal{U}^*V_\emptyset^*$ with $E_H^* + \mathcal{U}^* = 1$, which is the welfare expression of part 4.

The construction turns into an operational existence test once the candidate maps are written out in primitives.

Corollary 2 (Primitive test for interior existence). *Fix primitive parameters satisfying parts 1 and 2 of Assumption 3, and let*

$$\rho = \frac{J_L^*}{J_H^*}, \quad \hat{z}_L = \exp \left\{ \frac{u_L + w_L^0 - b}{\lambda} \right\}, \quad \hat{w}(x) = u_H - \frac{(b + \lambda \ln x - u_H)(1 - \beta(1 - \delta))}{\beta(1 - \delta)}.$$

For each $x > \max\{1, \rho \hat{z}_L\}$ define $\alpha_f(x)$ and $\alpha_w(x)$ by the formulas in Theorem 2 and (42), replacing $\hat{z}_H$ by x . Let $\mathcal{X}$ be the set of such x for which $\alpha_f(x) \in (0, 1)$, $\alpha_w(x) \in (0, 1)$, the induced $\mu(x), P(x), P_H(x), P_L(x)$ all lie in $(0, 1)$, $u_H > \hat{w}(x) > b$ (so search is strictly worthwhile), and part 3 of Assumption 3 holds at the candidate $\hat{w}(x)$. Define the primitive equilibrium map

$$A^{\text{impl}}(x) := \alpha_f(x)^{1-\phi} \alpha_w(x)^\phi. \quad (\text{B.1})$$

If $\mathcal{X}$ contains an interval on which A^{impl} is continuous and strictly decreasing, and the matching efficiency parameter A lies strictly between the two endpoint values of A^{impl} on that interval, then there exists a unique interior SRSRIE. Conversely, every interior SRSRIE has $\hat{z}_H^* \in \mathcal{X}$ and satisfies $A^{\text{impl}}(\hat{z}_H^*) = A$.

Proof. Every interior equilibrium is equivalent to a solution of $A^{\text{impl}}(x) = A$, with all other objects pinned down in closed form. The interiority restrictions in the definition of $\mathcal{X}$ are exactly the requirements that the matching probabilities, worker search decision, quitting decisions, and posteriors and choice probabilities be interior. Strict monotonicity gives uniqueness, and the intermediate value theorem gives existence. $\square$

C Proofs for the Dynamic Comparative Statics

C.1 Notation

Throughout this appendix we work in the normalized coordinates (A.3) of Appendix A.3, evaluated in the dynamic model: $x := \hat{z}_H > 1$, $z := \hat{z}_L \in (0, 1)$, $\rho := J_L^*/J_H^* > 1$, and $s := \rho z$, with Q , L , and ℓ_Q as in (A.4). Unstarred x is a candidate coordinate and $x^* := \hat{z}_H^*$ is the market clearing root; a prime denotes the corresponding object at a shifted equilibrium. The attention and quality objects that part 2 of Theorem 2 inherits are those of Lemma 1, their derivatives are (A.9)–(A.17), and the optimized surplus from a contact is $D = \lambda d$ with d given by (A.15). The dynamic model adds

the free entry and reservation consistency maps of Theorem 2,

$$\alpha_f(x) = \frac{F}{J_H^* P_H(x)}, \quad \alpha_w(x) = \frac{n(x)}{\lambda d(x)}, \quad n(x) := C + \hat{w}(x) - b,$$

$$\hat{w}(x) = u_H - k(b + \lambda \ln x - u_H), \quad k := \frac{1 - \beta(1 - \delta)}{\beta(1 - \delta)}.$$

Derivatives “at fixed ρ ” move $s = \rho z$ together with z .

C.2 Responses of the high-quality surplus

Lemma 7 (Responses of the high-quality surplus). *Define*

$$\Phi(x) := (1 - \phi) \ln \alpha_f(x) + \phi \ln \alpha_w(x),$$

so that the interior equilibrium is $\Phi(\hat{z}_H^*) = \ln A$. Then $\Phi' < 0$, and at the interior equilibrium:

- (a) $\frac{\partial \Phi}{\partial u_H} = \frac{\phi(1+k)}{n} > 0$, so $\frac{d\hat{z}_H^*}{du_H} > 0$.
- (b) $\frac{\partial \Phi}{\partial u_L} > 0$, so $\frac{d\hat{z}_H^*}{du_L} > 0$.
- (c) $\frac{\partial \Phi}{\partial b} < 0$, so $\frac{d\hat{z}_H^*}{db} < 0$.

Proof. Since $\alpha_f = F/(J_H^* P_H)$ and $\alpha_w = n/D$,

$$\Phi'(x) = -(1 - \phi) \frac{\partial \ln P_H}{\partial x} + \phi \left[\frac{n'(x)}{n(x)} - \frac{\partial \ln d}{\partial x} \right] < 0, \quad (\text{C.1})$$

because $\frac{\partial \ln P_H}{\partial x} > 0$, $n' = -k\lambda/x < 0$ (with $n > 0$ since $\alpha_w > 0$), and $\frac{\partial d}{\partial x} > 0$ with $d > 0$ by Lemma 5. By the implicit function theorem, $\frac{d\hat{z}_H^*}{dt} = -\frac{\partial \Phi / \partial t}{\Phi'}$ has the sign of $\frac{\partial \Phi}{\partial t}$.

- (a) u_H enters only through n : $\frac{\partial n}{\partial u_H} = 1 + k$.

(b) An increase in u_L moves z and s through $z = e^{(u_L + w_L^0 - b)/\lambda}$ and $s = \rho z$ with $\rho = \frac{(\pi_L - w_L^0)(1 - \beta(1 - \delta))}{\pi_H}$. When $w_L^0 > 0$, bargaining shares the improvement: $\frac{dz}{du_L} = \frac{\psi}{\lambda} z$ and $\frac{d\rho}{du_L} = \frac{(1 - \psi)(1 - \beta(1 - \delta))}{\pi_H} > 0$. When $w_L^0 = 0$, it accrues to the worker alone: $\frac{dz}{du_L} = \frac{z}{\lambda}$ and $\frac{d\rho}{du_L} = 0$. In both regimes the displacement of (z, s) is $\frac{dz}{du_L} > 0$ times the fixed- ρ z -direction (in which $ds = \frac{s}{z} dz$) plus $z \frac{d\rho}{du_L} \geq 0$ times the pure s -direction. In the pure s -direction, $\frac{\partial \Phi}{\partial s} = -(1 - \phi) \frac{\partial \ln P_H}{\partial s} - \phi \frac{\partial \ln d}{\partial s} > 0$ since both partial derivatives are negative by (A.10) and Lemma 5(c). In the fixed- ρ z -direction, $\frac{\partial \Phi}{\partial z}|_\rho = -(1 - \phi) \frac{\partial \ln P_H}{\partial z}|_\rho - \phi \frac{\partial \ln d}{\partial z}|_\rho > 0$ since both partial derivatives are negative by (A.11) and Lemma 5(d). A combination with nonnegative weights, at least one strictly positive, is positive.

(c) b enters through z (with $\frac{dz}{db} = -\frac{z}{\lambda} < 0$, at fixed ρ) and through n (with $\frac{\partial n}{\partial b} = -(1+k)$, since $\frac{\partial \hat{w}}{\partial b} = -k$). Hence

$$\frac{\partial \Phi}{\partial b} = \underbrace{\left[-(1-\phi) \frac{\partial \ln P_H}{\partial z} \Big|_{\rho} - \phi \frac{\partial \ln d}{\partial z} \Big|_{\rho} \right]}_{>0} \cdot \left(-\frac{z}{\lambda} \right) - \frac{\phi(1+k)}{n} < 0.$$

□

Two facts serve more than one of the three proofs that follow. First, aggregate worker value $\mathcal{W}^* = (1 - E_H^*)V_{\emptyset}^* + E_H^*V_H^*$ obeys the differencing identity

$$\Delta \mathcal{W}^* = (1 - E_H^*) \Delta V_{\emptyset}^* + E_H^* \Delta V_H^* + \Delta E_H^* (V_H^* - V_{\emptyset}^*), \quad (\text{C.2})$$

where a prime on a starred object denotes its value at the shifted equilibrium. Primes on functions of the candidate coordinate, such as Φ' and n' below, are ordinary derivatives. Since $V_H^* > V_{\emptyset}^*$, aggregate worker value moves in the common direction of the two values and the employment stock whenever all three move together. Second, the employment premium is bounded by the high-quality surplus,

$$V_H^* - V_{\emptyset}^* = (1+k)g = \lambda \ln \hat{z}_H^* - (\hat{w}^* - b) \leq \lambda \ln \hat{z}_H^*, \quad g := b + \lambda \ln \hat{z}_H^* - u_H > 0, \quad (\text{C.3})$$

which follows from $\hat{w}(x) = u_H - kg$ and $\hat{w}^* \geq b$ in the interior equilibrium, where search is worthwhile. Because $k > 0$, it also gives $g < \lambda \ln \hat{z}_H^*$.

C.3 Proof of Proposition 3

We first isolate the principal technical step.

Lemma 8. *For $t \in \{s, z\}$, define (all functions evaluated on the interior region, with the z -derivatives taken at fixed ρ):*

$$A_t := d \cdot \left[\frac{\partial \ln \mu}{\partial t} \frac{\partial \ln P_H}{\partial x} - \frac{\partial \ln \mu}{\partial x} \frac{\partial \ln P_H}{\partial t} \right] + \left[\frac{\partial \ln P_H}{\partial t} \frac{\partial d}{\partial x} - \frac{\partial \ln P_H}{\partial x} \frac{\partial d}{\partial t} \right],$$

$$B_t := \left[\frac{\partial \ln \mu}{\partial t} + \frac{\partial \ln P_H}{\partial t} \right] \frac{\partial d}{\partial x} - \left[\frac{\partial \ln \mu}{\partial x} + \frac{\partial \ln P_H}{\partial x} \right] \frac{\partial d}{\partial t}.$$

Then $A_s < 0$, $A_z < 0$, $B_s < 0$, and $B_z < 0$ everywhere on the interior region.

Proof. Substituting (A.9)–(A.17) and (A.15) and simplifying, the terms involving L cancel identically in A_s and A_z (because $dQ = x(1-z)L - Q\ell_Q$ and the coefficient multiplying L in the second bracket exactly offsets the one generated by d), and each inequality reduces to a single logarithm

inequality:

$$A_s < 0 \iff x(1-z)(x-s) > G \ell_Q, \quad G := s(x-z) + x(x-1)(s-z), \quad (\text{C.4})$$

$$A_z < 0 \iff (1-z)(x-s) > s(x-z) \ell_Q, \quad (\text{C.5})$$

$$B_s < 0 \iff x(x-z)(x-s) > [xs(x-z) - z(x-1)(x-s)] L, \quad (\text{C.6})$$

$$B_z < 0 \iff xs(s-z)(x-z) L < \frac{(x-1)z(x-s)^2 Q}{x} + (s-z)(x-s)[x^2(1-z) + sz(x-1)]. \quad (\text{C.7})$$

Proof of (C.5). By Lemma 4(i) with $t = \frac{Q}{s(x-z)} > 1$ and identity (A.6), $\ell_Q < \frac{Q}{s(x-z)} - 1 = \frac{(1-z)(x-s)}{s(x-z)}$, which is exactly (C.5).

Proof of (C.4). Using the same bound on ℓ_Q , it suffices that $x(1-z)(x-s) \geq G \cdot \frac{(1-z)(x-s)}{s(x-z)}$, i.e. that $xs(x-z) \geq G$. And indeed $xs(x-z) - G = (x-1)[s(x-z) - x(s-z)] = (x-1)z(x-s) > 0$.

Proof of (C.6). If the bracket is nonpositive the claim is trivial. Otherwise, drop the negative term $-z(x-1)(x-s)$ and use Lemma 4(ii): $xs(x-z)L < xs(x-z)\frac{x-s}{\sqrt{xs}} = \sqrt{xs}(x-z)(x-s) < x(x-z)(x-s)$.

Proof of (C.7). By Lemma 4(ii) and then the AM-GM inequality $\sqrt{xs} \leq \frac{x+s}{2}$,

$$xs(s-z)(x-z)L < \sqrt{xs}(s-z)(x-z)(x-s) \leq \frac{1}{2}(x+s)(x-s)(s-z)(x-z).$$

So it suffices to show $E := 2(x-1)z(x-s)Q + 2x(s-z)[x^2(1-z) + sz(x-1)] - x(x+s)(s-z)(x-z) \geq 0$. Direct expansion gives the factorization $E = (x-s)W$ with

$$W := s(x^2 - 3xz + 2z) + xz(x + z - 2),$$

which is linear in s with $W|_{s=z} = 2z(x-1)(x-z) > 0$ and $W|_{s=x} = x(x-z)^2 > 0$. Hence $W > 0$ on $[z, x]$ and $E > 0$. $\square$

Proof of Proposition 3. Recall from Theorem 2 that u_L enters only through (z, s) , which it moves as in Lemma 7(b), and leaves A, F, J_H^*, C, u_H, b unchanged. Throughout the proof, ∂_y denotes the derivative along that displacement of (z, s) at fixed x ; by Lemma 7(b) it is a nonnegative combination, with at least one weight strictly positive, of the fixed- ρ z -derivative and the pure s -derivative.

Part 1. $\frac{d\hat{z}_H^*}{du_L} > 0$ is Lemma 7(b).

Part 2. Taking total derivatives of μ^* gives

$$\frac{d \ln \mu^*}{du_L} = \underbrace{\frac{\partial_y \ln \mu}{du_L}}_{<0 \text{ by (A.10),(A.11)}} + \underbrace{\frac{\partial \ln \mu}{\partial x}}_{<0} \cdot \underbrace{\frac{dz_H^*}{du_L}}_{>0} < 0.$$

For the posteriors, γ and ν depend only on (x, z) , with $\frac{\partial \gamma}{\partial z} = \frac{x(1-x)}{(x-z)^2} < 0$, $\frac{\partial \nu}{\partial z} = \frac{1-x}{(x-z)^2} < 0$, and both decreasing in x . Since both z and x^* rise, both posteriors fall.

Part 3. By (44),

$$\frac{d \ln f^*}{du_L} = \left[\partial_y \ln \mu + \frac{1}{\phi} \partial_y \ln P_H \right] + \left[\frac{\partial \ln \mu}{\partial x} + \frac{1}{\phi} \frac{\partial \ln P_H}{\partial x} \right] \frac{dz_H^*}{du_L},$$

with $\frac{dz_H^*}{du_L} = \frac{\partial_y \Phi}{|\Phi'|}$. Multiplying through by $|\Phi'| > 0$ and using (C.1), a direct rearrangement gives (the terms in $\frac{1-\phi}{\phi}$ cancel)

$$|\Phi'| \frac{d \ln f^*}{du_L} = (1-\phi) \mathcal{J}_1 + \phi \mathcal{J}_2 + \mathcal{J}_3 + \kappa (\phi \partial_y \ln \mu + \partial_y \ln P_H), \quad (\text{C.8})$$

where $\kappa := \frac{k\lambda}{x n} > 0$, and (writing ∂_x for $\partial/\partial x$)

$$\begin{aligned} \mathcal{J}_1 &:= \partial_y \ln \mu \partial_x \ln P_H - \partial_x \ln \mu \partial_y \ln P_H, & \mathcal{J}_2 &:= \partial_y \ln \mu \partial_x \ln d - \partial_x \ln \mu \partial_y \ln d, \\ \mathcal{J}_3 &:= \partial_y \ln P_H \partial_x \ln d - \partial_x \ln P_H \partial_y \ln d. \end{aligned}$$

The κ -term in (C.8) is strictly negative, since $\partial_y \ln \mu < 0$ and $\partial_y \ln P_H < 0$. The remaining terms are linear in ϕ , so it suffices to show they are negative at $\phi = 0$ and $\phi = 1$, i.e. that $\mathcal{J}_1 + \mathcal{J}_3 < 0$ and $\mathcal{J}_2 + \mathcal{J}_3 < 0$. Multiplying by $d > 0$, and decomposing the u_L displacement into the pure s -direction and the fixed- ρ z -direction with nonnegative weights as in Lemma 7(b), these are exactly nonnegative combinations, with at least one weight strictly positive, of A_s, A_z and of B_s, B_z respectively, all strictly negative by Lemma 8. Hence $\frac{d \ln f^*}{du_L} < 0$, and $\mathcal{U}^* = \frac{\delta}{\delta + (1-\delta)f^*}$ strictly increases.

Part 4. $\hat{w}(x)$ is strictly decreasing in x and does not depend on u_L directly, so $\hat{w}^*$ falls. Hence $V_\emptyset^* = V_L^* = \hat{w}^*/(1-\beta)$ falls, and $V_H^* = \frac{u_H + \beta \delta V_\emptyset^*}{1-\beta(1-\delta)}$ falls.

Part 5. By Part 4, $V_\emptyset^*$ and V_H^* both fall. By Part 3, E_H^* falls. The differencing identity (C.2), with all three terms now strictly negative, gives $\Delta \mathcal{W}^* < 0$. $\square$

C.4 Proof of Proposition 4

Proof. Part 1. $\frac{dz_H^*}{du_H} > 0$ is Lemma 7(a).

Part 2. With z, s fixed, P_H^* rises by (A.9), $P_L^* = P_H^*/\rho$ rises, $\theta^* = (AJ_H^*P_H^*/F)^{1/\phi}$ rises, μ^* falls by (A.9), and $\gamma = \frac{x(1-z)}{x-z}$, $\nu = \frac{1-z}{x-z}$ fall since $\frac{\partial \gamma}{\partial x} = -\frac{z(1-z)}{(x-z)^2} < 0$ and $\frac{\partial \nu}{\partial x} = -\frac{1-z}{(x-z)^2} < 0$.

Part 3. By (44) with A, F, J_H^*, z, s fixed,

$$\frac{d \ln f^*}{du_H} = \left[\frac{\partial \ln \mu}{\partial x} + \frac{1}{\phi} \frac{\partial \ln P_H}{\partial x} \right] \frac{dz_H^*}{du_H},$$

and using (A.9), the bracket is positive iff

$$\begin{aligned} \frac{1}{\phi} \cdot \frac{s-z}{(x-s)(x-z)} &> \frac{s}{xQ} \\ \iff \phi &< \frac{(s-z)xQ}{s(x-s)(x-z)} = \left(1 - \frac{1}{\rho}\right) \frac{x}{P(x-z)} = \hat{\phi}, \end{aligned}$$

where we used $P = (x-s)/Q$ and $z/s = 1/\rho$. The bound $\hat{\phi} > 1 - \frac{1}{\rho}$ follows because $\frac{x}{P(x-z)} = \frac{xQ}{(x-s)(x-z)} > 1$ by (A.7). Since unemployment is strictly decreasing in f^* , unemployment falls iff $\phi < \hat{\phi}$.

Part 4. Since z, s, J_H^*, J_L^* are unaffected by u_H, α_f and D move only through x . At the new equilibrium $x' > x^*$ we have $\alpha_f(x') < \alpha_f(x^*)$, so (43) forces $\alpha'_w > \alpha_w^*$. Since $D(x') > D(x^*)$ by Lemma 5(b), $n = \alpha_w D$ strictly increases, i.e. $\hat{w}^*$ strictly increases. Then $V_\emptyset^* = V_L^* = \hat{w}^*/(1-\beta)$ rises, and $V_H^* = \frac{u_H + \beta \delta V_\emptyset^*}{1-\beta(1-\delta)}$ rises since both u_H and $V_\emptyset^*$ rise.

Part 5. If $\phi \leq \hat{\phi}$, the result follows from Parts 3 and 4 by the differencing identity (C.2): E_H^* weakly rises, both values strictly rise, and $V_H^* > V_\emptyset^*$, so $\Delta \mathcal{W}^* > 0$. For the general case, including $\phi > \hat{\phi}$ where E_H^* falls, we compute the derivative exactly. Denote $\Psi := \lambda \frac{d \ln \hat{z}_H^*}{du_H} > 0$, $\kappa := \frac{k\lambda}{xn} > 0$, $X := (1-\phi) \frac{\partial \ln P_H}{\partial x} + \frac{\phi}{d} \frac{\partial d}{\partial x} > 0$, and recall $1+k = \frac{1}{\beta(1-\delta)}$ and the surplus elasticity of the high-quality job-finding rate (48), which in this notation reads $\varepsilon_f^* = x \frac{\partial \ln \mu}{\partial x} + \frac{x}{\phi} \frac{\partial \ln P_H}{\partial x}$ at $x = \hat{z}_H^*$. From $\hat{w}(x) = u_H - k(b + \lambda \ln x - u_H)$:

$$\frac{dV_\emptyset^*}{du_H} = \frac{(1+k) - k\Psi}{1-\beta}, \quad V_H^* - V_\emptyset^* = (1+k)g, \quad \frac{d(V_H^* - V_\emptyset^*)}{du_H} = (1+k)(\Psi - 1),$$

and, since $\frac{d \ln E_H^*}{d \ln f^*} = \mathcal{U}^*$, $\frac{dE_H^*}{du_H} = E_H^* \mathcal{U}^* \varepsilon_f^* \Psi / \lambda$. Moreover $\Psi = \frac{\phi(1+k)\kappa}{k|\Phi'|}$ and $1 - \frac{k}{1+k} \Psi = \frac{X}{|\Phi'|}$ by (C.1). Assembling the total derivative of $\mathcal{W}^* = V_\emptyset^* + E_H^*(V_H^* - V_\emptyset^*)$ and multiplying by $\beta(1-\delta)|\Phi'| > 0$:

$$\beta(1-\delta)|\Phi'| \frac{d\mathcal{W}^*}{du_H} = X \left(\frac{1}{1-\beta} - E_H^* \right) + E_H^* \frac{\phi\kappa}{k} \left[1 + \mathcal{U}^* \frac{V_H^* - V_\emptyset^*}{\lambda} \varepsilon_f^* \right]. \quad (\text{C.9})$$

The first term is strictly positive since $E_H^* < 1 < \frac{1}{1-\beta}$. If $\varepsilon_f^* \geq 0$ (i.e. $\phi \leq \hat{\phi}$) the bracket is at least one. If $\varepsilon_f^* < 0$, the bracket exceeds $1 - \mathcal{U}^* > 0$ by three bounds: the premium bound (C.3); $s \ln x \leq Q$ from Appendix A.3; and, on this branch only, $|\varepsilon_f^*| \leq x \left| \frac{\partial \ln \mu}{\partial x} \right| = \frac{s}{Q}$, since $\varepsilon_f^* = -\frac{s}{Q} + \frac{x}{\phi} \frac{\partial \ln P_H}{\partial x}$ with the second term positive, so that $\varepsilon_f^* < 0$ forces $|\varepsilon_f^*|$ below the first term.

Combining, $\mathcal{U}^* \frac{V_H^* - V_\emptyset^*}{\lambda} |\varepsilon_f^*| \leq \mathcal{U}^* \cdot \ln x \cdot \frac{s}{Q} \leq \mathcal{U}^* < 1$. Hence both terms in (C.9) are positive and $\frac{d\mathcal{W}^*}{du_H} > 0$ for every $\phi \in (0, 1)$. $\square$

C.5 Proof of Proposition 5 and counterexamples

For reference, the premium compression of Proposition 5 satisfies the proved identity, and the composition gain is, explicitly,

$$\chi^* = \left| 1 + \lambda \frac{d \ln \hat{z}_H^*}{db} \right|, \quad \Sigma^* = \frac{\rho \hat{z}_L (\hat{z}_H^* - \hat{z}_L)}{Q(\hat{z}_H^*) (1 - \hat{z}_L)}, \quad (\text{C.10})$$

with $\rho = J_L^*/J_H^*$ and $Q(\hat{z}_H^*) = \hat{z}_H^*(1 - \hat{z}_L) + \rho \hat{z}_L(\hat{z}_H^* - 1)$, so that Σ^* is the composition gain in the surplus indices $(\hat{z}_H^*, \hat{z}_L)$ and firm values (J_H^*, J_L^*) of the body. The second expression equals the defining derivative (50). At a fixed reservation flow value, a unit increase in b lowers both $\lambda \ln \hat{z}_H$ and $\lambda \ln \hat{z}_L$ by one. P_H and P_L are ratios of the exponentiated surpluses, hence invariant to a common scaling of $(\hat{z}_H, \hat{z}_L)$, so

$$\lambda \frac{\partial \ln \mu^*}{\partial b} \Big|_{\hat{w}} = - \left(\frac{\partial \ln \mu}{\partial \ln \hat{z}_H} + \frac{\partial \ln \mu}{\partial \ln \hat{z}_L} \right) = \frac{\rho \hat{z}_L (\hat{z}_H^* - \hat{z}_L)}{Q(\hat{z}_H^*) (1 - \hat{z}_L)},$$

where the second equality is a direct computation. The same invariance gives $\lambda \frac{\partial \ln f^*}{\partial b} \Big|_{\hat{w}} = \Sigma^*$, which is how the composition gain enters the decomposition (51).

$$\frac{dV_\emptyset^*}{db} = \frac{1 - \beta(1 - \delta)}{1 - \beta} \left| \frac{d(V_H^* - V_\emptyset^*)}{db} \right| \quad \text{and} \quad \frac{dV_H^*}{db} = \frac{\beta\delta}{1 - \beta} \left| \frac{d(V_H^* - V_\emptyset^*)}{db} \right|, \quad (\text{C.11})$$

$$\frac{d\mathcal{W}^*}{db} = \underbrace{\left(\frac{1 - \beta(1 - \delta)}{1 - \beta} - E_H^* \right) \left| \frac{d(V_H^* - V_\emptyset^*)}{db} \right|}_{\text{net value gains} > 0} + \underbrace{(V_H^* - V_\emptyset^*) \frac{dE_H^*}{db}}_{\text{employment margin}}, \quad (\text{C.12})$$

$$\mathcal{T}^* := \lambda \chi^* \left(\frac{1 - \beta(1 - \delta)}{1 - \beta} - E_H^* \right) + \beta(1 - \delta) (V_H^* - V_\emptyset^*) E_H^* \mathcal{U}^* (\Sigma^* - \varepsilon_f^* \chi^*). \quad (\text{C.13})$$

Rearranging, $\mathcal{T}^* > 0$ is exactly the worker value gains against employment margin loss inequality displayed in part 5 of Proposition 5.

Proof of Proposition 5. Part 1. Lemma 7(c).

Part 2. A rise in b lowers z at fixed ρ ($\frac{dz}{db} = -z/\lambda$) and lowers x^* . Then

$$\frac{d \ln \mu^*}{db} = \underbrace{\frac{\partial \ln \mu}{\partial z} \Big|_{\rho}}_{<0} \cdot \underbrace{\frac{dz}{db}}_{<0} + \underbrace{\frac{\partial \ln \mu}{\partial x}}_{<0} \cdot \underbrace{\frac{d \hat{z}_H^*}{db}}_{<0} > 0,$$

and similarly γ and ν , which are decreasing in both x and z , strictly rise. For P_H^* , the same decomposition with $\lambda \frac{d \ln \hat{z}_H^*}{db} = -(1 + \chi^*)$ (established in Part 3, which does not use Part 2) and Lemma 6(i) gives

$$\lambda \frac{d \ln P_H^*}{db} = z \left| \frac{\partial \ln P_H}{\partial z} \right|_{\rho} - x \frac{\partial \ln P_H}{\partial x} (1 + \chi^*) = - \left(x \frac{\partial \ln P_H}{\partial x} \right) \chi^* < 0 :$$

the direct improvement in screening (z falls, so P_H tends to rise) is exactly offset by the radial component of the surplus decline, leaving only the excess compression χ^* . Then $P_L^* = P_H^*/\rho$ and $\theta^* = (AJ_H^* P_H^*/F)^{1/\phi}$ strictly fall as well.

Part 3. From $\hat{w}(x) = u_H - k(b + \lambda \ln x - u_H)$ and $\frac{d \hat{z}_H^*}{db} = \Phi_b/|\Phi'|$ (writing $\Phi_b \equiv \frac{\partial \Phi}{\partial b} < 0$),

$$\frac{d \hat{w}^*}{db} = -k \left(1 + \lambda \frac{d \ln \hat{z}_H^*}{db} \right) = -k - \frac{k\lambda}{x} \cdot \frac{\Phi_b}{|\Phi'|} = k \frac{\lambda |\Phi_b| - x |\Phi'|}{x |\Phi'|} =: k \chi^*,$$

which also shows $\chi^* = - \left(1 + \lambda \frac{d \ln \hat{z}_H^*}{db} \right)$, so the two definitions of χ^* in the theorem agree once we establish $\chi^* > 0$ (then $\chi^* = |1 + \lambda d \ln \hat{z}_H^*/db|$, and since $V_H^* - V_{\emptyset}^* = \frac{u_H - \hat{w}^*}{1 - \beta(1 - \delta)}$, $\frac{d(V_H^* - V_{\emptyset}^*)}{db} = -\frac{k\chi^*}{1 - \beta(1 - \delta)} = -(1 + k)\chi^*$, giving $\chi^* = \beta(1 - \delta) \left| \frac{d(V_H^* - V_{\emptyset}^*)}{db} \right|$). Collecting the expressions from the proof of Lemma 7,

$$\lambda |\Phi_b| - x |\Phi'| = (1 - \phi) \left[z \left| \frac{\partial \ln P_H}{\partial z} \right|_{\rho} - x \frac{\partial \ln P_H}{\partial x} \right] + \frac{\phi}{d} \left[z \left| \frac{\partial d}{\partial z} \right|_{\rho} - x \frac{\partial d}{\partial x} \right] + \frac{\phi \lambda (1 + k)}{n} - \frac{\phi \lambda k}{n}.$$

The first bracket vanishes by Lemma 6(i). Interiority gives $\alpha_w^* = n/(\lambda d) < 1$, i.e. $\lambda/n > 1/d$, so

$$\lambda |\Phi_b| - x |\Phi'| > \frac{\phi}{d} \left[z \left| \frac{\partial d}{\partial z} \right|_{\rho} - x \frac{\partial d}{\partial x} + 1 \right] = \frac{\phi}{d} \left[1 - \left(x \frac{\partial d}{\partial x} + z \left| \frac{\partial d}{\partial z} \right|_{\rho} \right) \right] \geq \frac{\phi}{d} (1 - P) > 0$$

by Lemma 6(ii). Hence $\chi^* > 0$.

Decomposing the total derivative of $\ln f^*$ and substituting $\lambda \frac{d \ln \hat{z}_H^*}{db} = -(1 + \chi^*)$ established above:

$$\begin{aligned} \lambda \frac{d \ln f^*}{db} &= -z \left[\frac{\partial \ln \mu}{\partial z} \Big|_{\rho} + \frac{1}{\phi} \frac{\partial \ln P_H}{\partial z} \Big|_{\rho} \right] + \left[x \frac{\partial \ln \mu}{\partial x} + \frac{x}{\phi} \frac{\partial \ln P_H}{\partial x} \right] \lambda \frac{d \ln \hat{z}_H^*}{db} \\ &= z \left| \frac{\partial \ln \mu}{\partial z} \right|_{\rho} + \frac{x}{\phi} \frac{\partial \ln P_H}{\partial x} - \varepsilon_f^* (1 + \chi^*) = z \left| \frac{\partial \ln \mu}{\partial z} \right|_{\rho} + x \left| \frac{\partial \ln \mu}{\partial x} \right| - \varepsilon_f^* \chi^* = \Sigma^* - \varepsilon_f^* \chi^*, \end{aligned}$$

where the second line uses Lemma 6(i) ($\frac{z}{\phi}|\partial_z \ln P_H| = \frac{x}{\phi}\partial_x \ln P_H$) and $\varepsilon_f^* = \frac{x}{\phi} \frac{\partial \ln P_H}{\partial x} - x \left| \frac{\partial \ln \mu}{\partial x} \right|$, and, by (A.9) and (A.11),

$$\Sigma^* = z \left| \frac{\partial \ln \mu}{\partial z} \right| + x \left| \frac{\partial \ln \mu}{\partial x} \right| = \frac{s(x-1)}{Q(1-z)} + \frac{s}{Q} = \frac{s(x-z)}{Q(1-z)} > 0.$$

Since unemployment is strictly decreasing in f^* , it strictly falls iff $\Sigma^* > \varepsilon_f^* \chi^*$. If $\phi \geq \hat{\phi}$ then $\varepsilon_f^* \leq 0$ and the condition holds automatically. For $\phi < \hat{\phi}$ the condition can fail: column (ii) of Table C.1 records an interior equilibrium, satisfying all conditions of Theorem 2, with $\Sigma^* < \varepsilon_f^* \chi^*$.

Part 4. By Part 3, $\chi^* > 0$, so $\hat{w}^*$ strictly rises, and hence $V_\emptyset^* = V_L^* = \hat{w}^*/(1-\beta)$ strictly rises, with $\frac{dV_\emptyset^*}{db} = \frac{k\chi^*}{1-\beta} = \frac{1-\beta(1-\delta)}{1-\beta}(1+k)\chi^*$ (using $\frac{k}{1+k} = 1-\beta(1-\delta)$), the first equality in (C.11). Finally $V_H^* = \frac{u_H + \beta\delta V_\emptyset^*}{1-\beta(1-\delta)}$ gives $\frac{dV_H^*}{db} = \frac{\beta\delta}{1-\beta(1-\delta)} \frac{dV_\emptyset^*}{db} = \frac{\beta\delta}{1-\beta}(1+k)\chi^* > 0$, the second equality in (C.11).

Part 5. Differentiating $\mathcal{W}^* = V_\emptyset^* + E_H^*(V_H^* - V_\emptyset^*)$ and substituting the value derivatives from Part 4,

$$\frac{d\mathcal{W}^*}{db} = \left[\frac{1-\beta(1-\delta)}{1-\beta} - E_H^* \right] (1+k)\chi^* + (V_H^* - V_\emptyset^*) \frac{dE_H^*}{db},$$

which is (C.12) since $(1+k)\chi^* = \left| \frac{d(V_H^* - V_\emptyset^*)}{db} \right|$. The first term is strictly positive because $\frac{1-\beta(1-\delta)}{1-\beta} \geq 1 > E_H^*$. Indeed it is at least $\mathcal{U}^*(1+k)\chi^*$. If unemployment does not rise, then $\frac{dE_H^*}{db} \geq 0$ and aggregate worker value strictly rises. Substituting $\frac{dE_H^*}{db} = E_H^* \mathcal{U}^* \frac{d \ln f^*}{db} = \frac{E_H^* \mathcal{U}^*}{\lambda} (\Sigma^* - \varepsilon_f^* \chi^*)$ and $V_H^* - V_\emptyset^* = \frac{g}{\beta(1-\delta)}$ with $g := b + \lambda \ln \hat{z}_H^* - u_H$, and multiplying by $\frac{\lambda}{1+k} > 0$, gives exactly $\mathcal{T}^*$ of (C.13), so $\text{sign}\left(\frac{d\mathcal{W}^*}{db}\right) = \text{sign}(\mathcal{T}^*)$.

For the sufficient condition, write

$$\mathcal{T}^* = \chi^* \left[\lambda \left(\frac{1-\beta(1-\delta)}{1-\beta} - E_H^* \right) - g E_H^* \mathcal{U}^* \varepsilon_f^* \right] + g E_H^* \mathcal{U}^* \Sigma^*.$$

The last term is strictly positive, so it suffices that the bracket be nonnegative. Its first piece is at least $\lambda \mathcal{U}^*$. For the second, use $g < \lambda \ln x$ from (C.3), together with $x \left| \frac{\partial \ln \mu}{\partial x} \right| = s/Q$ and $s \ln x \leq Q$ from Appendix A.3 and $\frac{x}{\phi} \frac{\partial \ln P_H}{\partial x} = \frac{\hat{\phi}}{\phi} \cdot \frac{s}{Q}$ (by the definition (49) of $\hat{\phi}$), to get

$$g \varepsilon_f^* \leq g \cdot \frac{x}{\phi} \frac{\partial \ln P_H}{\partial x} = \frac{g s}{Q} \cdot \frac{\hat{\phi}}{\phi} < \lambda \ln x \cdot \frac{s}{Q} \cdot \frac{\hat{\phi}}{\phi} \leq \lambda \frac{\hat{\phi}}{\phi}.$$

Hence $g E_H^* \mathcal{U}^* \varepsilon_f^* < \lambda E_H^* \mathcal{U}^* \frac{\hat{\phi}}{\phi} \leq \lambda \mathcal{U}^*$ whenever $\phi \geq E_H^* \hat{\phi}$, so $\mathcal{T}^* > 0$. For $\phi < E_H^* \hat{\phi}$, columns (ii) and (iii) of Table C.1 record interior equilibria with each aggregate worker value sign. $\square$

The unqualified signs hold at every interior equilibrium, so no example can contradict them. The conditional signs and the two unsigned u_L responses are different: each is claimed to take both

values inside the admissible region, and that claim needs witnesses. Table C.1 records five, each an exact interior equilibrium at which every condition of Theorem 2 was verified.

	(i)	(ii)	(iii)	(iv)	(v)		(i)	(ii)	(iii)	(iv)	(v)
β	0.96	0.95	0.9132	0.9656	0.9365	u_L	0.3	0.5	0.2684	0.1255	0.7477
δ	0.1	0.05	0.0365	0.03351	0.2171	π_L	1	3	2.9671	3.689	5.323
ψ	0.3	0.3	0.3082	0.05007	0.7064	π_H	0.1125	0.0796	0.0216	0.06684	0.1782
λ	1	1	1	0.04328	0.08111	b	1.0831	1.7431	2.0839	0.2108	4.339
C	0.05	0.0124	0.0232	0.001469	0.008767	u_H	1.2031	1.895	2.4077	0.2212	4.404
F	0.4	0.35	0.063	0.5891	0.469	A	0.422	0.5244	0.5554	0.4256	0.6797
ϕ	0.72	0.25	0.235	0.3602	0.1245						
ρ	1.10	3	12.4	3.62	2.67	$\hat{z}_L$	0.50	0.50	0.337	0.633	0.536
$\hat{z}_H^*$	2.0	4.0	12.7	29.9	18.3	μ^*	0.645	0.308	0.146	0.142	0.256
P_H^*	0.967	0.714	0.689	0.9433	0.9496	$\hat{\phi}$	0.13	1.98	6.39	2.07	1.27
$\mathcal{U}^*$	0.3108	0.4064	0.337	0.5447	0.7523	$\mathcal{U}^*$ after	0.3133	0.4236	0.352	0.5521	0.7822
$\mathcal{W}^*$	28.280	36.044	25.917	6.213	68.434	$\mathcal{W}^*$ after	28.434	36.099	25.914	6.211	68.400

Table C.1: Parameter witnesses for the conditional signs and the unsigned u_L responses: (i) unemployment rises with u_H ; (ii) unemployment rises with b while aggregate worker value rises; (iii) aggregate worker value falls with b ; (iv) the acceptance rate of high-quality offers rises with u_L ; (v) it falls with u_L . Each column is an exact interior equilibrium satisfying every condition of Theorem 2. The shock is $\Delta u_H = 0.01$ in column (i), $\Delta b = 0.01$ in columns (ii) and (iii), and $\Delta u_L = 0.01$ in columns (iv) and (v). Reported equilibrium objects are rounded, and $E_H^* = 1 - \mathcal{U}^*$ because $E_L^* = 0$.

In column (i), $\phi > \hat{\phi}$, so unemployment rises with u_H , even as every worker's value rises. Changing only the matching technology to $\phi = 0.1 < \hat{\phi}$ and $A = 0.48857$, which leaves the equilibrium allocation unchanged, reverses the unemployment response, so both signs occur.

In column (ii), $\phi < \hat{\phi}$ and $\Sigma^* < \varepsilon_f^* \chi^*$: a more generous unemployment benefit improves offer composition yet raises unemployment, because the contraction in the high-quality surplus and in market tightness dominates. Every worker's value rises, and $\mathcal{T}^* > 0$, so aggregate worker value rises as well.

In column (iii), $\phi < E_H^* \hat{\phi}$ and $\mathcal{T}^* < 0$: the benefit raises every worker's value and improves offer composition, yet the loss of employment premia outweighs those gains and aggregate worker value falls.

Columns (iv) and (v) witness the two u_L responses that Proposition 3 does not sign: P_H^* rises from 0.94329 to 0.94344 in column (iv) and falls from 0.9496 to 0.9481 in column (v), and market tightness moves with P_H^* through free entry, $\theta^* \propto (P_H^*)^{1/\phi}$, so both tightness signs occur as well. In both columns unemployment rises and aggregate worker value falls, as parts 3 and 5 require.

D Robustness and Scope

D.1 Implementing the Same Due Diligence Problem with Normal Signals

The posterior and choice probability representation can be implemented with normally distributed signals under a suitable cost function. This is an implementation result, not robustness to arbitrary normal signal costs. Assume the latent signal ξ is drawn from $N(1, \sigma^2)$ if the job is high quality and from $N(0, \sigma^2)$ if it is low quality. Due diligence consists of choosing, before the signal is realized, both the variance σ^2 and a threshold $\bar{\xi}$; the worker then observes only whether ξ lies above or below $\bar{\xi}$, and accepts if and only if it lies above. The pair $(\sigma^2, \bar{\xi})$ is a commitment to a binary information structure, and we verify obedience below. The worker's problem is

$$\max_{\sigma^2 > 0, \bar{\xi} \in \mathbb{R}} P[\gamma W_H + (1 - \gamma) W_L] + (1 - P) W_\emptyset - \lambda \mathcal{I}^N(\sigma^2, \bar{\xi}), \quad (\text{D.1})$$

where (P, γ, ν) are the objects induced by $(\sigma^2, \bar{\xi})$ and $\mathcal{I}^N$ is the entropy reduction of the binary observation,

$$\mathcal{I}^N(\sigma^2, \bar{\xi}) = \mathbb{H}(\mu) - P \mathbb{H}(\gamma) - (1 - P) \mathbb{H}(\nu). \quad (\text{D.2})$$

Writing Φ_N for the standard normal distribution function, the conditional acceptance probabilities are $P_H = 1 - \Phi_N((\bar{\xi} - 1)/\sigma)$ and $P_L = 1 - \Phi_N(\bar{\xi}/\sigma)$, with $P = \mu P_H + (1 - \mu) P_L$, $\gamma = \mu P_H / P$, and $\nu = \mu(1 - P_H)/(1 - P)$ by Bayes' rule (9).

The implementation result is that any pair $1 > P_H > P_L > 0$ is generated by exactly one choice of $(\sigma^2, \bar{\xi})$. Writing $c_H := \Phi_N^{-1}[1 - P_H] = (\bar{\xi} - 1)/\sigma$ and $c_L := \Phi_N^{-1}[1 - P_L] = \bar{\xi}/\sigma$, the ordering $P_H > P_L$ is $c_L > c_H$, and the unique solution is

$$\bar{\xi} = \frac{c_L}{c_L - c_H}, \quad \sigma = \frac{1}{c_L - c_H}.$$

The thresholded normal family is therefore onto the set of strictly ordered conditional acceptance probabilities. The cost (D.2) depends on $(\sigma^2, \bar{\xi})$ only through the induced (P_H, P_L) , and Lemma 2 restates the posterior problem (13) in those same coordinates, so the two problems have the same objective on this set. At the interior rational inattention optimum the posteriors are distinct, $\gamma > \nu$, and Bayes' rule orders the conditional acceptance probabilities, so the two optima coincide and are implemented by a unique finite pair $(\sigma^2, \bar{\xi})$. If the firm's problem and the wage setting process are as in Definition 1, the firm's first-order condition is again $P_H J_H = P_L J_L$ and the resulting equilibrium is characterized by Theorem 1. Obedience holds at the optimum: the acceptance posterior strictly

prefers acceptance and the rejection posterior strictly prefers rejection,¹³ so a worker who could disregard the recommendation after observing the binary outcome would not do so.

Two features delimit the result. It relies on the worker observing only the *binary* outcome. Were the continuous realization ξ observed instead, the appropriate information cost would be the expected entropy reduction over the full posterior distribution, $\mathbb{H}(\mu) - \mathbb{E}_\xi [\mathbb{H}(\Pr(H | \xi))]$, which strictly exceeds (D.2) because the binary outcome is a coarsening of ξ ; and sequential rationality would pin the threshold down at

$$\bar{\xi} = \frac{1}{2} - \sigma^2 \mathcal{G}(\mu), \quad \mathcal{G}(\mu) := \ln \left[\frac{\mu}{1-\mu} \left(\frac{W_H - W_\emptyset}{W_\emptyset - W_L} \right) \right],$$

the cutoff at which the posterior $\Pr(H | \xi)$ leaves the worker indifferent. In that variant σ^2 is the only ex ante choice, the model must be solved numerically as in Ambuehl, Ockenfels, and Stewart (2025), and its equilibrium does not in general coincide with the rational inattention equilibrium, because the sequentially rational cutoff generally differs from the threshold that replicates the rational inattention acceptance probabilities. The result also depends on the binary mutual information cost (D.2): other normal signal cost functions generally produce different optimal precision choices and need not reproduce the rational inattention equilibrium.

D.2 Posted Terms, Contracting, and Partial Signaling

This subsection records a simple posted-terms extension. The purpose is not to solve a full signaling game with arbitrary contracts, but to show that the paper's mechanism survives whenever posted terms leave a pooling or partially pooling residual quality problem.

Let firms first post a publicly observed term $\tau \in T$. A worker who observes τ has posterior belief μ_τ that the offer is high quality. Conditional on τ , the worker then chooses due diligence and an acceptance decision as in the baseline model. Let $W_\omega(\tau)$ and $J_\omega(\tau)$ denote the worker and firm continuation payoffs in the submarket indexed by τ , after all contracting and compensation are determined. We do not model how τ maps into these payoffs; in particular, if compensation is renegotiated after acceptance, part of a posted transfer can be offset in bargaining. The residual quality component ω remains hidden and noncontractible.

Proposition 8 (Pooling posted terms). *Fix a posted term τ that is attached with positive probability to offers of both qualities, so that $\mu_\tau \in (0, 1)$. Suppose the continuation payoffs $W_\omega(\tau)$ and $J_\omega(\tau)$ satisfy the strict interior conditions of Theorem 1. Then the continuation equilibrium conditional on observing τ is exactly the unique active interior RIBNE of Theorem 1 with W_ω and J_ω replaced*

¹³In the normalized coordinates (A.3) these statements read $x(1-z)\ln x + z(x-1)\ln z > 0$ and $(1-z)\ln x + (x-1)\ln z < 0$ for $x > 1 > z > 0$. The second follows from $\ln z \leq z-1$: $(1-z)\ln x + (x-1)\ln z \leq (1-z)[\ln x - (x-1)] < 0$. The first follows from $\ln z \geq 1 - \frac{1}{z}$: $x(1-z)\ln x + z(x-1)\ln z \geq (1-z)[x\ln x - (x-1)] > 0$.

by $W_\omega(\tau)$ and $J_\omega(\tau)$. In particular, the equilibrium offer composition inside the pooling cell satisfies

$$P_H(\tau)J_H(\tau) = P_L(\tau)J_L(\tau), \quad (\text{D.3})$$

and the comparative static composition mechanism survives inside the cell.

Proof. Conditional on observing a pooling term τ , the worker faces the same binary due diligence problem (13), with payoffs $W_\omega(\tau)$ and prior μ_τ . The firm indifference condition inside the pooling cell is likewise the same condition as in the baseline model with firm values $J_\omega(\tau)$. The proof of Theorem 1 therefore applies verbatim. $\square$

The same argument gives a partial signaling result. Let a posted term induce a partition of offer types into information cells. In any cell that separates quality perfectly, due diligence about quality is unnecessary and the mechanism is absent. In any cell that pools both qualities and satisfies the interior conditions, Proposition 8 applies cell by cell. Thus posted terms weaken the mechanism only to the extent that they separate residual quality. They do not eliminate it in pooling or partial pooling regions.

The contracting restriction of Assumption 1 is about what can be contracted, not about what can be posted. If high-quality firms could write fully enforceable dynamic or contingent contracts that low-quality firms could not profitably mimic, those contracts would serve as separating signals, and in the limit of perfect separation workers would not need to acquire costly information at all. The model is therefore not a description of attributes that can be cheaply certified or fully priced through enforceable contracts. It describes residual quality: job attributes, product characteristics, or service reliability that are payoff relevant but difficult to verify, enforce, or summarize in a posted term before acceptance. The intermediate cases are the relevant ones. Posted terms may be coarse, noisy, constrained, or partially pooling; firms may have limited commitment; and some dimensions of quality may be learned only after costly inspection or experience. In those environments the composition channel survives, though the exact formulas would have to be recomputed for the richer contract space.

The mechanism also does not require a literal absence of posted wages. Suppose firms in a submarket can post a common noncontingent transfer τ before due diligence, while the residual quality component remains hidden and noncontractible. Then the same analysis applies after replacing worker flow payoffs by $u_\omega + \tau$ and firm flow profits by $\pi_\omega - \tau$, with post-acceptance Nash compensation recomputed under the transformed primitives, provided those primitives satisfy the interiority assumptions. Directed search across posted terms can be read as workers first choosing among such submarkets and then facing the due diligence problem within the chosen one. What matters for our results is that posted terms do not fully separate the residual quality types inside the relevant submarket.

E Constrained Efficiency

Proof of Proposition 6. Throughout, the economy is on the interior equilibrium regime maintained in the text, and every derivative incorporates the private responses of due diligence and offer quality.

E.1 Implementable allocations

With the levy, free entry reads $F + \tau_v = \alpha_f (\mu P_H J_H^* + (1 - \mu) P_L J_L^*)$. Both offer types pay τ_v , so the firm indifference condition $P_H J_H^* = P_L J_L^*$ is unchanged, entry reduces to $F + \tau_v = \alpha_f P_H J_H^*$, and the levy spans tightness at any fixed b . Given (θ, b) : the unemployment benefit determines $\hat{z}_L$; reservation consistency (42) at $\alpha_w = A\theta^{1-\phi}$ determines $\hat{z}_H$; and the per-match block of Theorem 2 determines $(\mu, P_H, P_L, \gamma, \nu, \mathcal{I})$. Constant returns keep this block stationary while the stock transitions: a permanent policy holds θ constant, vacancies scale with the pool, $\mathcal{V}_t = \theta \mathcal{U}_t$, so the matching probabilities, and with them the whole per-match allocation, are the same in every period of the transition. The planner therefore chooses a permanent pair (θ, b) from a compact set within the maintained interior regime.

E.2 The exact objective and the master derivative

Using $E_{H,t} = 1 - \mathcal{U}_t$ and the definition of Δ in (54), period surplus is $E_{H,t} y_H + \mathcal{U}_t (\alpha_w \mathcal{S} - C - F\theta) = y_H - \mathcal{U}_t \Delta$, so the objective is

$$\mathcal{P} = \frac{y_H}{1 - \beta} - \sum_{t \geq 0} \beta^t \mathcal{U}_t \Delta.$$

The transition in (53) is linear in the stock, $\mathcal{U}_{t+1} = \delta + (1 - \delta)(1 - f)\mathcal{U}_t$, so multiplying by β^{t+1} and summing gives

$$\sum_{t \geq 0} \beta^t \mathcal{U}_t = \frac{\mathcal{U}_0 + \beta\delta/(1 - \beta)}{1 - \beta(1 - \delta)(1 - f)},$$

which is (54). The inherited stock enters only through the positive factor $\mathcal{U}_0 + \beta\delta/(1 - \beta)$, so maximizing $\mathcal{P}$ is minimizing Δ/Ω for every inherited stock; the ratio is continuous on any compact policy set within the maintained regime, so an optimum exists, and the set of optima is independent of $\mathcal{U}_0$. This proves the preamble of Proposition 6.

A permanent policy changes (Δ, f) once and for all. Differentiating $\mathcal{P}$ at the stationary inherited stock $\mathcal{U} = \delta/(1 - (1 - \delta)(1 - f))$, where $\mathcal{U}_0 + \beta\delta/(1 - \beta) = \mathcal{U}(1 - \beta(1 - \delta)(1 - f))/(1 - \beta)$, gives

$$d\mathcal{P} = -\frac{\mathcal{U}}{1 - \beta} d\Delta + \frac{\beta(1 - \delta)\mathcal{U}\Delta}{(1 - \beta)(1 - \beta(1 - \delta)(1 - f))} df.$$

Substituting $-d\Delta = \alpha_w [d\mathcal{S} + ((1 - \phi)\mathcal{S} - F/\alpha_f) d\ln\theta]$, which uses $d\alpha_w = (1 - \phi)\alpha_w d\ln\theta$ and $F d\theta = \alpha_w(F/\alpha_f) d\ln\theta$, together with $df = \alpha_w\mu P_H [(1 - \phi) d\ln\theta + d\ln(\mu P_H)]$, gives (56).

E.3 Part 1: the vacancy margin

Setting (56) to zero in the direction $d\ln\theta$ at fixed b gives $(1 - \phi)(\mathcal{S} + \Lambda) - F/\alpha_f + \mathcal{R}_\theta = 0$, the tightness condition of part 1. The levy that implements the optimal tightness is $\tau_v = \alpha_f P_H J_H^* - F$, and substituting the tightness condition for F gives (59). A planner who holds $(\mu, P_H, P_L, \mathcal{I})$ fixed when differentiating sets $\mathcal{R}_\theta = 0$ and obtains τ_v^0 ; the difference is $-\alpha_f \mathcal{R}_\theta$, and the two levies differ in sign exactly when $\tau_v^0 (\tau_v^0 - \alpha_f \mathcal{R}_\theta) < 0$, that is, when $\mathcal{R}_\theta$ and τ_v^0 share a sign and $\alpha_f |\mathcal{R}_\theta| > |\tau_v^0|$. This proves part 1.

The reversal occurs at an optimum of the informed problem, not only in the same-allocation inference of part 1. Take the economy of column (ii) of Table C.1 with matching technology $\phi = 0.809$ and $A = 0.38785$ and high-quality flow profit $\pi_H = 0.06$. Minimizing Δ/Ω over the admissible policies in the search box $\theta \in [0.310, 142.6]$, the matching-interior range, and $b \in [1.05, 2.79]$ gives an interior optimum $(\theta^*, b^*) = (0.4447, 1.6718)$: every condition of Theorem 2 holds with slack, both first-order conditions hold with residuals below 10^{-7} , the Hessian of Δ/Ω is positive definite in $(\ln\theta, b)$, and multistart searches and a 301×301 grid find no better admissible policy. At the optimum, $\mathcal{R}_\theta = -0.0223 < \tau_v^0/\alpha_f = -0.0071 < 0$, so $\tau_v = +0.0113$ and $\tau_v^0 = -0.0053$: the levy is a tax while the correction inferred with the responses frozen is a subsidy. Local versions of both sign patterns occur at zero-levy interior equilibria, with the part 1 formula read as the correction the missing levy would apply there: the same matching variant at the unmodified column (ii) parameters gives $\tau_v^0 = -0.0093$ and $\tau_v = +0.0095$, and column (i) with posting cost $F = 0.22$ gives $\tau_v^0 = +0.0061$ and $\tau_v = -0.0050$.

E.4 A mistaken planner

Fix the zero-levy equilibrium at the prevailing benefit as the reference market, $(\tau_v, b) = (0, b^r)$, with allocation $(\mu^r, P_H^r, P_L^r, \mathcal{I}^r)$. The mistaken planner models what it observes as fixed primitives: an entrant draws a high-quality offer with probability μ^r ; every matched worker receives a binary signal that recommends acceptance with probability P_H^r for high quality and P_L^r for low quality, at the fixed attention cost $\lambda \mathcal{I}^r$ per contact; workers follow the signal wherever obedience is strictly optimal. The calibrated model reproduces every level of the reference market, including expected

entry profits, because firm indifference gives $P_H^r J_H^* = P_L^r J_L^*$. Its contact surplus $\mathcal{S}^X := \mu^r P_H^r y_H + (1 - \mu^r) P_L^r y_L - \lambda \mathcal{I}^r$ and the share of contacts yielding high-quality jobs, $\mu^r P_H^r$, are constants, so its objective reduces, exactly as in (54), to a ratio in tightness alone. Write $d = 1 - \beta(1 - \delta)$, $g = \beta(1 - \delta)A\mu^r P_H^r$, $B = y_H + C$, and $H = A\mathcal{S}^X$, so that

$$q^X(\theta) = \frac{B - H\theta^{1-\phi} + F\theta}{d + g\theta^{1-\phi}}.$$

Multiplying the numerator of $dq^X/d\theta$ by $\theta^\phi > 0$ gives $Fd\theta^\phi + Fg\phi\theta - (1 - \phi)(Hd + gB)$, strictly increasing in θ , so the interior root of the first-order condition is the unique global tightness optimum of the perceived problem. The benefit moves nothing in the perceived model — quality, the signal, and its cost are primitives — so the mistaken planner leaves $b = b^r$.

The mistaken planner implements its optimum θ^X through its own free entry condition, recommending the levy $\tau_v^M = \alpha_f(\theta^X) P_H^r J_H^* - F$. Once the levy is imposed, tightness solves the actual free entry condition at the endogenous acceptance rate. In the witness economy at $b^r = 1.7431$: the reference tightness is $\theta^r = 0.3530$; the perceived optimum is $\theta^X = 0.4500$, implemented by the subsidy $\tau_v^M = -0.0624$; realized tightness is $\theta^D = 0.4480$; and the informed optimum is the pair above, an entry tax with a lower benefit, attaining strictly higher welfare in the true economy. Signal obedience, active search, and the maintained branch hold at every reported allocation. The sign contrast is a property of the policy pair: at the status-quo benefit, the informed entry correction alone is also a subsidy ($\theta = 0.4174$, $\tau_v = -0.0452$); the tax arises jointly with the benefit choice, which relaxes screening, raises acceptance, and converts private entry at the target tightness from falling short to overshooting.

E.5 Part 2: responses at fixed tightness and the assignment result

Work in the normalized coordinates (A.3), $x = \hat{z}_H$ and $z = \hat{z}_L$, and let $d(x, z)$ be the optimized contact surplus of Lemma 3 with the offer mix adjusting to (x, z) through firm indifference. Reservation consistency at fixed α_w is (42),

$$\alpha_w \lambda d(x, z) = C + \hat{w} - b, \quad \hat{w} = u_H - k(b + \lambda \ln x - u_H), \quad (\text{E.1})$$

with $\lambda \ln z = u_L + w_L^0 - b$, so that $\lambda d \ln z / db = -1$.

Lemma 9 (The surplus index falls more than one-for-one). *At fixed tightness the high-quality surplus index satisfies*

$$m_b := \lambda \left. \frac{d \ln x}{db} \right|_\theta < -1.$$

Proof. Differentiating (E.1) with respect to b at fixed α_w and collecting terms,

$$m_b (\alpha_w x d_x + k) = \alpha_w z d_z - (1 + k) \quad \implies \quad 1 + m_b = \frac{\alpha_w (x d_x + z d_z) - 1}{\alpha_w x d_x + k}.$$

The denominator is positive because $d_x > 0$ by Lemma 5(b). The numerator is negative by the radial bound of Lemma 6(ii), which gives $x d_x + z d_z = P - L\mu\Sigma^* < P < 1 \leq 1/\alpha_w$, the last step because $\alpha_w \in (0, 1)$. Hence $m_b < -1$. $\square$

The fixed tightness responses in part 2 follow. First, $\lambda d \ln \mu / db|_\theta = \frac{\partial \ln \mu}{\partial \ln x} m_b - \frac{\partial \ln \mu}{\partial \ln z} > -\frac{\partial \ln \mu}{\partial \ln x} - \frac{\partial \ln \mu}{\partial \ln z} = \Sigma^* > 0$, using $\frac{\partial \ln \mu}{\partial \ln x} < 0$ from (A.9) and $m_b < -1$. Second, P_H depends on (x, z) only through x/z , so $\lambda d \ln P_H / db|_\theta = \frac{\partial \ln P_H}{\partial \ln x} (m_b + 1) < 0$, using $\frac{\partial \ln P_H}{\partial \ln x} > 0$ (A.9). Third, $P_L = (J_H^*/J_L^*) P_H$ inherits the sign of P_H .

For the assignment result, write policy directions as changes in $(\ln x, \ln z)$ per unit of the policy. The unemployment benefit at fixed tightness moves $(m_b, -1)/\lambda$. Tightness at fixed benefit moves $(e_\theta, 0)$ with $e_\theta := \frac{\partial \ln x}{\partial \ln \theta}|_b < 0$, because a higher θ raises α_w and α_w is strictly decreasing in the surplus index (Theorem 2). The ray direction $-(1, 1)/\lambda$ is the mechanical part of the benefit. These satisfy

$$\left(\frac{m_b}{\lambda}, -\frac{1}{\lambda} \right) = c(e_\theta, 0) + \left(-\frac{1}{\lambda}, -\frac{1}{\lambda} \right), \quad c = \frac{m_b + 1}{\lambda e_\theta} > 0,$$

the sign by Lemma 9 and $e_\theta < 0$. The per contact response $\mathcal{R}(v) := d\mathcal{S} + \Lambda d \ln(\mu P_H)$ evaluated along a direction v is linear in v , with $\mathcal{R}_\theta$ and $\mathcal{R}_b$ its values in the two policy directions. Along the ray both conditional acceptance rates are unchanged, because they depend on (x, z) only through x/z , so only the offer mix moves and

$$\mathcal{R}(\text{ray}) = \left[P_H y_H - P_L y_L - \lambda \frac{\partial \mathcal{I}}{\partial \mu} \Big|_{P_H, P_L} + \frac{\Lambda}{\mu} \right] d\mu = \mathcal{Q} d\mu, \quad d\mu = \frac{\mu \Sigma^*}{\lambda} db > 0,$$

using $P_L y_L = (J_H^*/J_L^*) P_H y_L$ in (58). At an interior joint optimum $\mathcal{R}_b = 0$, so linearity across the decomposition gives $c \mathcal{R}_\theta = -(\mu \Sigma^*/\lambda) \mathcal{Q}$, and since $c > 0$ and $\mu \Sigma^*/\lambda > 0$, $\text{sgn}(\mathcal{R}_\theta) = -\text{sgn}(\mathcal{Q})$. This proves part 2.

E.6 Part 3: the unemployment benefit without the levy

Along the zero-levy equilibrium path, tightness responds to the unemployment benefit, and (56) gives

$$\frac{1 - \beta}{\mathcal{U} \alpha_w} \frac{d\mathcal{P}}{db} = \left[(1 - \phi)(\mathcal{S} + \Lambda) - \frac{F}{\alpha_f} + \mathcal{R}_\theta \right] \frac{d \ln \theta^*}{db} + \mathcal{R}_b.$$

At zero levy, free entry gives $F/\alpha_f = P_H J_H^*$, so the bracket equals $-\tau_v/\alpha_f$ with τ_v the part 1 levy (59) evaluated at the zero-levy allocation, which is (61). Proposition 5 gives $d \ln \theta^*/db < 0$. This proves part 3. $\square$