

Who Is Shown to Whom?

Sorting, Screening, and Exposure Design on a Two-Sided Matching Platform*

Kei Ikegami[†]

27th September 2026

Abstract

I study how platforms should design recommendations when users also search independently. Using contact and bilateral acceptance data from a Japanese platform matching doctors to short-term jobs, I estimate an equilibrium model with endogenous doctor search. The counterfactuals favor recommendations that complement doctor search and prioritize long-run value. Retaining search makes welfare and the distribution of benefits across posts less sensitive to the recommendation objective, while long-run targeting improves both immediate surplus and long-run value relative to existing recommendations, albeit with uneven gains. The value of redesign also persists as recommendations expand: reallocating recommendations across posts yields a larger gain than doubling recommendation volume and equalizing totals across doctors combined.

JEL Classification Codes: D47; D83; C78; J64; C61; C63; L86.

Keywords: two-sided platforms; sequential search; sorting; screening; intermediation; market design; recommendation systems; doctor–job matching.

1 Introduction

Two-sided matching platforms rarely choose matches directly. Instead, they shape who can match by bringing selected pairs into contact, while participants also search for partners on their own. A recommendation policy therefore operates alongside decentralized search rather than in isolation. Because both sides can reject after contact and participants can adjust their search

*I am grateful to Medical Principle Co. for providing the data and for their thoughtful support in helping me understand the institutional background. I thank Yu Awaya, Michihiro Kandori, Fuhito Kojima, Kyohei Okumura, and Kosuke Uetake for valuable comments. I also thank participants at the ERATO meeting, IIOC, and SSCW for helpful feedback. Earlier versions of this paper circulated under the title “Exposure Design for Two-Sided Platforms.”

[†]Yokohama City University, ikegami.kei.pr@yokohama-cu.ac.jp.

intensity, a rule that creates more matches immediately need not create the most value over time or distribute matching opportunities broadly across the market. The platform consequently faces three linked choices: whether recommendations should complement or replace participants' own search, whether the algorithm should prioritize immediate expected surplus, longer-run value, or broad access across the market, and whether to make more recommendations or reallocate those already being made. This paper asks how a platform should make these choices when participants can respond by changing how intensively they search on their own.

I answer this question using data from a Japanese platform where doctors apply to short-term posts and platform agents recommend doctors to institutions. The data identify the route initiating each contact and record observed acceptance or rejection decisions on both sides. I estimate an equilibrium model linking route-specific contact, bilateral acceptance, and the value of remaining unmatched, then compare recommendation policies as doctor search intensity adjusts. First, retaining doctor search makes short- and long-run welfare and the distribution of benefits across posts less sensitive to which recommendation objective the platform prioritizes. Second, with doctor search retained, the long-run rule raises both immediate expected surplus and long-run value relative to the platform's existing recommendations, although gains are uneven across posts. Third, doubling recommendation volume and equalizing totals across doctors raise value, but the subsequent redesign across posts adds more value than those two changes combined. That final gain is comparable to adding about 2 percent more doctor job slots or 6 percent more posts.

Section 2 describes the Japanese two-sided labor platform and the records used in the analysis. A contact occurs either when a doctor searches and applies for a post or when a platform agent recommends a doctor for a post, and the data record both the initiating route and the observed acceptance or rejection decisions on each side. The analysis focuses on 1,132 doctors and 2,412 posts scheduled in the Kanto region in December 2024, constructs predetermined doctor-facility histories and a doctor activity proxy from earlier platform use, and documents a sharp reversal across routes: following doctor-initiated search, doctor and post acceptance rates are 98.0 and 41.0 percent, whereas following agent-initiated screening they are 28.3 and 98.4 percent.

Section 3 develops a continuous-time two-sided search model that preserves these two contact routes. Each doctor operates one or more interchangeable model units, called doctor slots, that can each form a match; several slots allow a doctor to consider several spot jobs concurrently. Each doctor-post pair has route-specific Poisson contact rates, and a match forms only if both sides accept after contact. The doctor's continuation value affects doctor-initiated contact rates, while the post's continuation value affects agent-initiated contact rates; contact rates and bilateral acceptance jointly determine these continuation values. Stationary renewal holds the market's composition fixed by replacing exiting units with identical ones. This provides the

equilibrium system used in estimation and the policy exercises.

Section 4 parameterizes two route-specific contact indices and two side-specific acceptance indices using doctor, post, and pair characteristics, including five predetermined measures of prior doctor–facility relationships. The likelihood combines route-specific Poisson approach counts with doctor- and post-side acceptance decisions, and the parameters and continuation values are estimated jointly subject to the stationary-equilibrium restriction. The fitted model closely reproduces aggregate approaches and preserves the reversal in acceptance across routes. To interpret the estimated magnitudes, I report average percentage-point changes in contact and acceptance probabilities, holding continuation values and other covariates fixed. Distance is negatively associated with contact in both routes, while its conditional association with acceptance is imprecisely estimated. Prior contracting is positively associated with contact through both routes and post acceptance; the doctor-acceptance association is less precise. Analytic standard errors for the structural parameters account for dependence among observations sharing a doctor or a post and for the equilibrium response of continuation values.

Section 5 first compares retaining doctor search with assigning all contacts through platform recommendations. In each case, I compare a rule maximizing immediate expected bilateral surplus with a rule maximizing long-run value. Within each class, objective values are reported as percentages of their highest computed levels. Post-side access summarizes how evenly the benefits of recommendations are distributed across posts relative to evenly spread recommendations; the index equals 100 when post benefits are proportional to those under that reference allocation. On these scales, retaining doctor search makes objective values and post-side access at both horizons less sensitive to this choice. The differences are particularly small for long-run value and access, consistent with doctors adjusting their search intensity as recommendations change.

Conditional on retaining doctor search, the comparisons favor prioritizing long-run user value. The platform’s existing recommendations already score about 99 on long-run value and 98 on long-run access, whereas both short-run measures are around 80. Maximizing immediate surplus lowers both long-run value and access relative to the platform’s existing recommendations. By contrast, the best-found long-run rule raises both immediate expected surplus and long-run value, at the cost of lower post-side access at both horizons. These gains are uneven: post-side value rises, but fewer than half of posts gain, and doctor-side value falls slightly. Relative to current recommendations, the redesigned allocation favors shorter jobs and doctors’ preferred work schedules, while accepting longer distances and fewer specialty matches.

Section 5 then shows that improving allocation remains valuable even when the platform can make more recommendations. With doctor search retained, doubling the current recommendation profile raises aggregate continuation value by 0.65 percent. Holding the doubled volume fixed, equalizing recommendation totals across doctors adds 0.44 percent; holding those

equal totals fixed, redesigning their allocation across posts adds 1.80 percent. Each percentage is relative to the preceding policy, and together the three changes raise value by 2.91 percent relative to the platform’s existing recommendations. Pairwise redesign contributes more than the expansion and equalization steps combined. In this final step, recommendations generate more matches while fewer doctor applications sustain roughly the same number of doctor-initiated matches. The aggregate gains nevertheless conceal different effects on the two sides: relative to the platform’s existing recommendations, final doctor-side value remains lower while post-side value rises. To put the 1.80 percent redesign gain in perspective, I compare it with targeted market expansion among doctors or posts in the top decile of estimated continuation values. The interpolated equivalents are 1.98 percent more aggregate baseline doctor slots with post stock fixed, or 5.87 percent more aggregate baseline posts with doctor stock fixed.

Related literature. The paper first contributes to the frictional matching literature on sorting and screening (Eeckhout and Kircher, 2010; Moen, 1997; Burdett and Coles, 1997; Shimer and Smith, 2000; Chade, Eeckhout and Smith, 2017) by making the contrast between doctor-initiated search and platform-initiated screening empirically operational. The same platform contains both routes, the data identify which route generated each contact, and the model evaluates alternative channel architectures and allocation rules for a given number of platform-assigned contacts.

Second, the paper relates to empirical work on online matching platforms (Hitsch, Hortaçsu and Ariely, 2010; Fradkin, 2017; Gong, 2016; Agarwal, 2015; Galichon, Kominers and Weber, 2019). Its distinctive feature is identifying the initiating route and observing doctor- and post-side acceptance or rejection decisions, rather than only realized matches.

Third, the paper connects to work on recommendations and steering in labor and platform markets (Horton, 2017; Belot, Kircher and Muller, 2019; Le Barbanchon, Hensvik and Rathelot, 2023; Chen and Tsai, 2024), which estimates the effects of providing or changing recommendations. Here recommendations are one route inside a structural two-route system, so the counterfactuals compare alternative channel architectures and allocation rules on a common footing.

Finally, the paper contributes to algorithmic and operations work on search design, multichannel traffic, fairness, and two-sided recommendation (Immorlica et al., 2023; Manshadi et al., 2025; Shi, 2023, 2025; Chen, Hsieh and Lin, 2023; Sürer, Burke and Malthouse, 2018; Biega, Gummadi and Weikum, 2018): the recommendation rule is evaluated through equilibrium continuation values under finite doctor capacity rather than through immediate conversion or ranking metrics alone.

2 Setting and Data

2.1 Platform and Exposure Channels

The empirical setting is a Japanese two-sided labor platform operated by Medical Principle Co., Ltd. The platform matches licensed doctors with short-term “spot” positions at hospitals and clinics. Participating doctors typically hold regular positions elsewhere. Hospitals and clinics post temporary openings, primarily for health checkups, ward coverage, and overnight on-call shifts. The platform records the contacts and applications associated with each opening, as well as the subsequent decisions by both sides that lead to a match.

Contact between a doctor and a post can arise through either of two routes. A doctor may search the platform and apply to a post; I refer to the resulting recorded contact as *doctor-initiated search*, or the *sorting route*. Alternatively, a platform agent may identify a doctor and recommend that doctor for a post; I refer to this as *agent-initiated screening*, or the *screening route*. Medical institutions do not directly search the platform’s doctor pool. Under the platform’s current operation, doctor-initiated search therefore reflects information and initiative on the doctor side, whereas agent-initiated screening is the route that the platform can directly redesign.¹

I use *approach* to refer to either type of recorded contact. An approach is distinct from a listing view: it brings a particular doctor–post pair into contact and initiates an exchange of detailed information through the intermediary, often in the form of an interview. Both sides then decide whether to proceed, and a match forms only if both sides accept. A doctor who initiated an approach may still decline after learning the details, and either side may reject a recommendation initiated by an agent.

2.2 Market and Measurement

The online job-matching platform operates continuously, so the empirical analysis must isolate a specific market from the ongoing flow of doctors and posts. This subsection first explains how I define the estimation cohort. It then describes how I construct bilateral acceptance decisions and predetermined measures of prior platform activity within that cohort, before presenting its main descriptive patterns.

Estimation cohort. I define the estimation cohort as posts scheduled in the Kanto region in December 2024 and all doctors appearing in at least one approach to those posts. Because an application may be submitted before a post’s scheduled work date, I retain approaches

¹The platform’s user interface is itself a potential object of design. For example, the platform could remove the search function available to doctors and instead generate all approaches through algorithmic recommendations. Section 5 considers this more substantial change to platform functionality as a counterfactual benchmark.

within each selected post’s episode-specific half-open observation window rather than impose a December cutoff on approach dates. The final cohort contains 1,132 doctors and 2,412 posts. Here, for the purpose of defining the active market, I classify doctors who neither approach a selected post nor receive an agent-initiated approach to one as *inactive* in the December 2024 spot-job market. In other words, these doctors generate no recorded indication that they are actively considering a post in this cohort. This is an operational definition based on observed platform activity; it does not rule out unrecorded browsing or interest in other spot jobs.

The data also contain basic characteristics of both sides. Doctor characteristics include age, experience, specialty, location, and stated preferences. Post characteristics include location, task content, desired specialty, schedule, working hours, and compensation. The preferred-location indicator records whether the post lies in the doctor’s preferred prefecture or region. The work-schedule indicator, labeled “Preferred-day match” or “Day match” in the tables, records whether the post matches a preferred combination of weekday and daytime or overnight work.² The specialty indicator records whether the doctor has at least one specialty in the post’s specialty list.

Route and decision construction. The data record the initiating route and subsequent status of each approach, from which I construct separate doctor- and post-side acceptance indicators. Repeated in-window approaches remain distinct events in the Poisson exposure counts. For acceptance, a doctor–post pair whose panel decision can be attributed to exactly one current-episode route contributes at most one Bernoulli trial for each side with an observed decision; both-route and unassignable pairs do not enter the route-specific acceptance denominators. For each approached doctor–post pair, the platform records one of five statuses: *Contract*, *Cancelled After Contract*, *Not Hired*, *Approach*, or *Inquiry Handled*. The mapping from these statuses to binary acceptance indicators, constructed in consultation with the platform’s staff, is reported in Table 1.

Past platform activity. I construct five predetermined variables separately for each doctor–facility pair, using only activity recorded on the platform strictly before the order date of the current post: the number of approaches made by the doctor to the facility during the preceding 365 days; an indicator for any prior contract between the doctor and the facility; the number of days since the doctor’s most recent approach to the facility; the share of the pair’s approaches during the preceding 365 days that were initiated by an agent; and an indicator that no prior relationship between the doctor and the facility is observed. These variables place recorded traces of persistent familiarity in the four empirical indices. They do not eliminate unobserved

²The post’s shift category is constructed from job duration: posts longer than ten hours are classified as overnight and the remaining posts as daytime.

Table 1. Mapping from Approach Status to Acceptance Indicators

Approach status	Doctor accepts	Post accepts
Contract	1	1
Cancelled After Contract	1	1
Not Hired	1	0
Approach	0	1
Inquiry Handled	0	–

Notes: 1 denotes acceptance, 0 denotes rejection, and – denotes an undefined decision. An undefined decision is treated as missing rather than as a rejection.

doctor–facility heterogeneity or latent pair quality.

The formal model also allows doctors to differ in how many posts they can search simultaneously. I interpret this latent quantity as an activity or throughput measure, motivated by the substantial differences in recorded platform contacts among otherwise observably similar licensed doctors. Because this quantity is not directly observed, I proxy it with the maximum, over September–November 2024, of a doctor’s monthly number of recorded approaches across both routes, excluding approaches to the final December cohort and imposing a floor of one slot.

Descriptive patterns. A notable feature of the data is the sharp asymmetry according to who initiates an approach. Among observed decisions following doctor-initiated search, doctors accept 98.0 percent and posts accept 41.0 percent. Following agent-initiated screening, the corresponding acceptance rates are 28.3 percent for doctors and 98.4 percent for posts. Thus, doctor-initiated search reflects substantial doctor-side preselection, whereas agent-initiated screening reflects substantial post-side preselection.

Controlling for past doctor–facility activity is particularly important when estimating these acceptance patterns. Only 121,065 of the 2,730,384 possible doctor–post pairs, or 4.43 percent, have a nonzero prior doctor–facility history. Among the 3,281 route-attributable doctor–post cells eligible for at least one side’s acceptance likelihood, however, 1,685 cells, or 51.36 percent, have such history. These route-attributable pair cells are therefore much more likely than an average possible doctor–post pair to involve a doctor and facility with a pre-existing relationship. Without these recorded-history variables, differences associated with prior familiarity could be attributed to the initiating route or to contemporaneous pair characteristics. Appendix B.2 presents separate summary tables for past activity and other analysis variables.

3 Model

I develop a continuous-time variant of the two-sided search model of Adachi (2003). The two sides of the market are doctors and posts. On the doctor side, I call the unit that can form a match a *doctor slot*. A doctor may operate several interchangeable slots, which allows the model to represent differences in how actively doctors search and, in particular, how many spot jobs they consider concurrently. Observed characteristics and preferences are shared across slots belonging to the same doctor, but each slot is treated as an agent in the model. No analogous construction is needed on the post side: because an open post leaves the market once it matches, each post is itself one model agent.

The sequence of events in continuous time is as follows. At any instant, each unmatched doctor slot and each open post is in its own waiting state and has a continuation value that summarizes its future matching opportunities. For each doctor–post pair, the doctor chooses the intensity of doctor-initiated search, while a platform agent acting on behalf of the post chooses the intensity of agent-initiated screening. These pair- and route-specific intensities govern independent Poisson clocks. When one of these clocks rings, an approach occurs. The doctor and post then observe the information revealed by that contact and independently decide whether to proceed. If both accept, a match forms and the matched doctor slot and post exit. If either rejects, the doctor slot remains unmatched and the post remains open, so their contact processes continue.

The remainder of the section follows this order: I first describe approach-intensity choice and arrival, then post-approach acceptance, then continuation values and stationary renewal.

3.1 Agents, Timing, and Approach Decisions

Agents and exposure routes. Time is continuous. Let I be the set of doctors, J the set of posts, and $\mathcal{C} = \{\text{self}, \text{agency}\}$ the two exposure routes. The route *self* denotes a doctor-initiated approach generated by the doctor’s search over listings. The route *agency* denotes an agent-initiated recommendation generated by screening on behalf of the post. Doctor i has $K_i \geq 1$ interchangeable doctor slots. Each post j is one open position and therefore one model agent. Each doctor slot can receive approaches through both routes.

Choosing approach intensities. For each doctor–post pair, the initiator of each route chooses the instantaneous rate at which that route generates an approach. On the self route, a higher rate represents more search attention by the doctor toward the post and therefore a higher arrival hazard of an application. On the agency route, it represents more screening attention directed toward the pair and therefore a higher arrival hazard of a recommendation. I represent this decision as a pair- and route-specific choice of a Poisson arrival rate. The rate

determines both the expected number of approaches per unit of time and the instantaneous hazard that an approach occurs. The initiator weighs the gross value of generating an approach against the continuation value of waiting, while moving the rate away from a route-specific baseline incurs a convex adjustment penalty.

Let s_{ij}^c denote the route-specific gross value that the initiator assigns to generating one additional approach between doctor i and post j . This is a pre-contact value index. It is distinct from the match payoff that the two sides observe after an approach occurs and from the total-user-value objective used for policy evaluation.

The relevant opportunity cost in the intensity choice is the continuation value of the initiating side:

$$\omega_{ij}^c(\alpha) = \begin{cases} \alpha_i^D, & c = \text{self}, \\ \alpha_j^P, & c = \text{agency}, \end{cases}$$

where α_i^D is the value of keeping one slot of doctor i available for future matching opportunities, and α_j^P is the value of keeping post j open. At this stage, these are candidate continuation values. The stationary equilibrium defined below requires them to coincide with the values generated by the resulting approach-arrival and acceptance processes.

For a candidate Poisson rate $\lambda > 0$ and a route-specific reference rate $\bar{\lambda}^c > 0$, define

$$D_{\text{KL}}(\lambda \| \bar{\lambda}^c) = \lambda \log\left(\frac{\lambda}{\bar{\lambda}^c}\right) - \lambda + \bar{\lambda}^c.$$

This expression is the relative-entropy rate of the Poisson process with intensity λ relative to the reference process with intensity $\bar{\lambda}^c$. It is nonnegative, equals zero at the reference rate, and is strictly convex in λ . I use it as an auxiliary convex penalty for moving the search or screening process away from its baseline, rather than interpreting it as a separately identified structural effort cost.³

I measure the two gross approach-value indices in units that normalize the intensity-adjustment scales to $\zeta^{\text{self}} = \zeta^{\text{agency}} = 1$. The initiator's auxiliary choice problem in the current market is therefore

$$\max_{\lambda > 0} \lambda(s_{ij}^c - \omega_{ij}^c(\alpha)) - D_{\text{KL}}(\lambda \| \bar{\lambda}^c). \quad (1)$$

The first term is the flow value of generating approaches, net of the initiator's waiting-value opportunity cost. The second term penalizes moving the approach rate away from its reference level. Its marginal cost is $\log(\lambda/\bar{\lambda}^c)$, so the objective is strictly concave and the unique conditional choice satisfies

$$\log \lambda_{ij}^c(\alpha) = \log \bar{\lambda}^c + s_{ij}^c - \omega_{ij}^c(\alpha). \quad (2)$$

³Relative-entropy penalties for moving a controlled stochastic law away from a passive reference law are used in KL-control formulations; see Todorov (2006). Wang et al. (2017) applies this idea directly to the control of point-process intensities. Here the device only rationalizes the functional form of the approach-rate equation: the penalty itself does not enter continuation values or the policy objective.

A positive net value raises the rate above that baseline, while a negative net value lowers it. On the self route, the doctor makes this choice. On the agency route, I use the same rule for the platform agent screening on behalf of the post.

Given the chosen rates, approaches between each slot of doctor i and post j arrive through two independent Poisson processes, with rates $\lambda_{ij}^{\text{self}}$ and $\lambda_{ij}^{\text{agency}}$. The two clocks almost surely do not ring simultaneously, so their rates add without an overlap correction. Across doctor i 's K_i slots, the aggregate route- c approach rate to post j is $K_i \lambda_{ij}^c$.

Acceptance after an approach. An approach is not a match. It is an event at which the two sides observe their match-specific information and then make bilateral acceptance decisions. Conditional on route c , their one-time payoffs are

$$U_{ij}^c = u_{ijc}^D + \varepsilon_{ijc}^D, \quad (3)$$

$$V_{ji}^c = u_{jic}^P + \varepsilon_{jic}^P. \quad (4)$$

The shocks are independent logistic random variables with scales $\zeta^D, \zeta^P > 0$, drawn independently across sides and exposure events. Utility is nontransferable, and the pair matches if and only if both sides accept. The systematic terms u_{ijc}^D and u_{jic}^P may depend on the route.

The acceptance rules are

$$a_{ijc}^D = \mathbf{1}\{U_{ij}^c \geq \alpha_i^D\}, \quad a_{jic}^P = \mathbf{1}\{V_{ji}^c \geq \alpha_j^P\}.$$

Write the resulting route-specific acceptance probabilities as

$$P_{ijc}^D(\alpha) = \frac{1}{1 + \exp\{(\alpha_i^D - u_{ijc}^D)/\zeta^D\}}, \quad P_{jic}^P(\alpha) = \frac{1}{1 + \exp\{(\alpha_j^P - u_{jic}^P)/\zeta^P\}},$$

where u_{ijc}^D and u_{jic}^P denote the systematic parts of (3) and (4). Bilateral acceptance on route c occurs with probability $p_{ijc}(\alpha) = P_{ijc}^D(\alpha)P_{jic}^P(\alpha)$.

3.2 Stationary Renewal and Continuation Values

Let $\delta_i^a \geq 0$ be doctor i 's exogenous abandonment hazard, $w_j \geq 0$ post j 's withdrawal hazard, and $r > 0$ the discount rate. The following benchmark turns the two acceptance thresholds into stationary continuation values.

Assumption 1 (Stationary renewal).

- (i) (Time-invariant primitives) *Gross approach values, baseline rates, payoff distributions, stocks, exit hazards, and discounting do not vary over time.*
- (ii) (Renewal) *When a doctor slot exits through a match or abandonment, or a post exits through a match or withdrawal, it is immediately renewed by a unit with the same modelled type. Hence the type composition and contact environment remain fixed.*

(iii) (Absorbing payoffs) *A matched unit receives its one-time payoff and exits; exogenous exit yields zero. The renewed unit starts a new search problem with independent future shocks.*

I model the market as stationary. In addition to leaving after a match, a doctor slot may become inactive and an open post may be withdrawn before matching. Either event removes that unit from the active market for reasons outside the model. Under stationary renewal, every exiting unit is immediately replaced by another unit of the same modeled type. Market stocks and the marginal type distributions on both sides therefore remain fixed. This is an analytical approximation, not a claim that the platform literally replaces each departing unit with an identical clone.⁴ Because the empirical market is calibrated to a single monthly cohort, I hold the marginal type distributions fixed when comparing stationary policies.

For route c , define the expected own-side gain above continued waiting,

$$q_{ijc}^D(\alpha) = \mathbb{E}[(U_{ij}^c - \alpha_i^D)_+], \quad q_{ijc}^P(\alpha) = \mathbb{E}[(V_{ji}^c - \alpha_j^P)_+],$$

where $x_+ = \max\{x, 0\}$. Because the two payoff shocks are independent, the expected gains from an exposure, taking the other side's acceptance into account, are

$$\psi_{ijc}^D(\alpha) = P_{jic}^P(\alpha)q_{ijc}^D(\alpha), \quad \psi_{ijc}^P(\alpha) = P_{ijc}^D(\alpha)q_{ijc}^P(\alpha).$$

The route matters both through its contact rate and through the systematic post-approach payoffs inside P , q , and ψ .

Balancing the exposure gains against discounting and exogenous exit gives the route-aware continuation-value system

$$\begin{aligned} \alpha_i^D &= \Phi_{i,K}^D(\alpha; \lambda) := \frac{\sum_{j \in J} \sum_{c \in \mathcal{C}} \lambda_{ij}^c \psi_{ijc}^D(\alpha)}{r + \delta_i^a}, \\ \alpha_j^P &= \Phi_{j,K}^P(\alpha; \lambda) := \frac{\sum_{i \in I} \sum_{c \in \mathcal{C}} K_i \lambda_{ij}^c \psi_{ijc}^P(\alpha)}{r + w_j}. \end{aligned} \tag{5}$$

The difference between the two equations reflects units: λ_{ij}^c is the contact rate for one doctor slot, whereas $K_i \lambda_{ij}^c$ is the aggregate flow from doctor i to post j .

Let $\Phi_K(\alpha; \lambda)$ stack the right-hand sides of (5), and define

$$\mathcal{Z}(\lambda^{\text{self}}, \lambda^{\text{agency}}, K) = \{\alpha \in \mathbb{R}_{\geq 0}^{|I|+|J|} : \alpha = \Phi_K(\alpha; \lambda)\}.$$

Definition 1 (Fixed-rate stationary-renewal equilibrium). *Fix $(\lambda^{\text{self}}, \lambda^{\text{agency}}, K)$. A fixed-rate stationary-renewal equilibrium is a pair (a, α) such that each side uses its route-specific threshold rule after every exposure and α equals the expected discounted value of waiting generated by those rules under Assumption 1.*

⁴Richer steady-state search models instead allow entry and matching jointly to determine the composition of the unmatched stocks; see, for example, Lauermaun and Nöldeke (2025). Allowing the type distributions to evolve as units enter and leave the market would be substantially less tractable at the scale of this application.

Proposition 1 (Fixed-rate characterisation and existence). *For nonnegative finite route rates and finite I, J , the fixed-rate stationary-renewal equilibria are exactly the threshold rules generated by $\alpha \in \mathcal{Z}(\lambda^{\text{self}}, \lambda^{\text{agency}}, K)$, and this set is nonempty.*

Proof. See Appendix A.3. □

Proposition 1 characterizes equilibrium conditional on fixed route-specific approach rates. In the current market, however, these rates are themselves chosen and therefore depend on the continuation values. Let $\Gamma(\alpha; s)$ stack the behavioral rate rules in equation (2). The current-market equilibrium combines rate optimality with value consistency:

$$\lambda = \Gamma(\alpha; s), \quad \alpha = \Phi_K(\alpha; \lambda). \quad (6)$$

A behavioral stationary-renewal equilibrium is a triple (λ, a, α) that satisfies equation (6) and the route-specific threshold rules. Because the behavioral rates are continuous, nonnegative, and bounded above on the nonnegative orthant by their values at zero, the same compact-box argument extends fixed-rate existence to this joint system; Appendix A.3 gives the argument.

4 Empirical Strategy and Results

4.1 Parameterization and Likelihood

The model leaves the gross approach values s_{ij}^c and the systematic post-approach payoffs u_{ijc}^D and u_{jic}^P unrestricted. To connect these primitives to the observed approach counts and bilateral decisions, I parameterize them using doctor characteristics, post characteristics, and variables defined for each doctor–post pair. I then impose the equilibrium consistency condition from Section 3.2 when evaluating the likelihood.

Let $c \in \{\text{self}, \text{agency}\}$ denote doctor-initiated search and agent-initiated screening. Let Z_{ij}^D and Z_{ij}^P denote the two design rows used by the implementation. Each collects three groups of observed covariates: doctor characteristics, post characteristics, and variables defined for the doctor–post pair. The third group includes distance, indicators for preference, schedule, specialty, and task compatibility, and the five predetermined doctor–facility history measures. The first column of each design row is identically one. Its coefficient is the only intercept in the corresponding index.

Gross approach values and contact rates. Each doctor–post pair has two route-specific contact rates: one for doctor-initiated search and one for agent-initiated screening. These rates determine the expected number of approaches observed through each route. Because the reference rate $\bar{\lambda}^c$ in equation (2) is constant within a route, I absorb $\log \bar{\lambda}^c$ into a route-specific

constant. I write the resulting empirical parameterization as

$$\begin{aligned}s_{ij}^{\text{self}} + \log \bar{\lambda}^{\text{self}} &= (Z_{ij}^D)' \beta_{\text{self}}, \\ s_{ij}^{\text{agency}} + \log \bar{\lambda}^{\text{agency}} &= (Z_{ij}^P)' \beta_{\text{agency}}.\end{aligned}$$

Substituting these expressions into the model’s intensity rule gives

$$\begin{aligned}\log \lambda_{ij}^{\text{self}} &= (Z_{ij}^D)' \beta_{\text{self}} - \alpha_i^D, \\ \log \lambda_{ij}^{\text{agency}} &= (Z_{ij}^P)' \beta_{\text{agency}} - \alpha_j^P.\end{aligned}$$

A doctor’s continuation value reduces doctor-initiated search, while a post’s continuation value reduces agent-initiated screening. We normalize $\zeta^{\text{self}} = \zeta^{\text{agency}} = 1$; there is no separately estimated exposure-scale parameter.

Post-approach payoffs and acceptance. The same three groups of covariates enter the two post-approach payoff indices, but their coefficients differ by side and from the coefficients in the contact-rate equations. I parameterize the systematic payoffs in equations (3)–(4) as

$$\begin{aligned}u_{ijc}^D &= (Z_{ij}^D)' \beta_D + \chi_D \mathbf{1}\{c = \text{self}\}, \\ u_{jic}^P &= (Z_{ij}^P)' \beta_P + \chi_P \mathbf{1}\{c = \text{agency}\}.\end{aligned}$$

Because payoff coefficients and logit scales are not separately identified, I normalize $\zeta^D = \zeta^P = 1$. The resulting acceptance probabilities are

$$\begin{aligned}\Pr(a_{ijc}^D = 1) &= \Lambda\left((Z_{ij}^D)' \beta_D - \alpha_i^D + \chi_D \mathbf{1}\{c = \text{self}\}\right), \\ \Pr(a_{jic}^P = 1) &= \Lambda\left((Z_{ij}^P)' \beta_P - \alpha_j^P + \chi_P \mathbf{1}\{c = \text{agency}\}\right).\end{aligned}$$

The two χ coefficients are route-specific differences in acceptance intercepts. For doctors, agent-initiated screening is the reference route, so χ_D is the change in the doctor acceptance index when the doctor initiates the contact. For posts, doctor-initiated search is the reference route, so χ_P is the change in the post acceptance index when a platform agent initiates the contact.

Observed likelihood and equilibrium restriction. The likelihood is defined over the doctor–post pairs constructed in Section 2.2. For each selected post episode in the monthly market, the observed route-specific approach count satisfies

$$n_{ij}^c \sim \text{Poisson}\left(K_i T_j \lambda_{ij}^c\right),$$

where λ_{ij}^c is the monthly contact rate for one doctor slot and T_j is the post episode’s observed duration in months. Doctor i has K_i such slots, so $K_i T_j \lambda_{ij}^c$ is the exposure-adjusted expected approach count for the doctor–post pair. The predetermined activity proxy fixes K_i before estimation and it enters with no estimated coefficient.

Write $\mathcal{C} = \{\text{self}, \text{agency}\}$ and let $\mathcal{M}$ denote the set of doctor–post pairs in the estimation market. For route c , let $\mathcal{A}_c^D$ and $\mathcal{A}_c^P$ denote the subsets with an observed doctor- and post-side acceptance decision, respectively. Set $m_{ij}^c = K_i T_j \lambda_{ij}^c$, $P_{ijc}^D = \Pr(a_{ijc}^D = 1)$, and $P_{jic}^P = \Pr(a_{jic}^P = 1)$. Suppressing the dependence of the rates and acceptance probabilities on (θ, α) , the likelihood is

$$\begin{aligned} L(\theta, \alpha) = & \prod_{c \in \mathcal{C}} \prod_{(i,j) \in \mathcal{M}} \frac{\exp(-m_{ij}^c) (m_{ij}^c)^{n_{ij}^c}}{(n_{ij}^c)!} \\ & \times \prod_{c \in \mathcal{C}} \prod_{(i,j) \in \mathcal{A}_c^D} (P_{ijc}^D)^{a_{ijc}^D} (1 - P_{ijc}^D)^{1-a_{ijc}^D} \\ & \times \prod_{c \in \mathcal{C}} \prod_{(i,j) \in \mathcal{A}_c^P} (P_{jic}^P)^{a_{jic}^P} (1 - P_{jic}^P)^{1-a_{jic}^P}. \end{aligned} \quad (7)$$

Because $\mathcal{C}$ contains two routes, equation (7) has two Poisson count components and four route-by-side Bernoulli components. The route-specific constants inside λ_{ij}^c remain elements of θ and are estimated jointly rather than profiled out.

4.2 Identification and Estimation Procedure

The contact records and the decisions made after contact discipline different parts of the model jointly with the continuation values. Variation in doctor-initiated and agent-initiated approach counts across pairs informs the two contact-rate indices. Conditional on a route-attributable approach, variation in doctor and post decisions across pairs and routes informs the two acceptance indices and the two initiator-specific acceptance intercept differences.

A remaining identification concern is pair-specific information that doctors or platform agents observe but the researcher does not. For example, a doctor may apply to a hospital, or an agent may recommend that doctor, because the doctor has worked successfully with the facility before. The same familiarity can also make both sides more likely to accept after the approach. If it is omitted, the estimated coefficients on distance, recorded match indicators, or the identity of the initiating side may partly capture relationship-driven contact and acceptance. I therefore include the five predetermined doctor–facility history measures in every contact-rate and acceptance index. These measures record prior approaches, prior contracting, recency, the route used in earlier contacts, and the absence of any recorded prior relationship. They place the most direct observable traces of persistent doctor–facility familiarity in both stages of the model.

The monthly doctor-abandonment hazard, post-withdrawal hazard, and discount rate are calibrated to 0.05, 0.10, and $-\log(0.99)$, respectively. I formulate the numerical target as joint maximization of $\ell(\theta, \alpha) = \log L(\theta, \alpha)$ over the parameters and continuation values, subject to equation (6). Appendix Section B.3 describes the MPEC problem, solver implementation, and convergence diagnostics.

I report standard errors clustered by doctor and by post, allowing dependence among doctor–post pairs that share a doctor or a post. The calculation also accounts for the dependence of continuation values on the structural parameters through the equilibrium conditions. When differentiating the likelihood, I therefore include both the direct effects of parameter changes and their indirect effects through changes in continuation values. Appendix Section B.4 describes the analytic covariance calculation.

Remark 1. *Another limitation arises because existence of the stationary-renewal fixed point does not imply that its solution is unique. If the same structural parameters support several continuation-value solutions, the mapping from parameters to observed outcomes is not single valued, and the data alone need not reveal which equilibrium generated the market. As in many dynamic structural models, the present analysis does not resolve this general multiple-equilibrium identification problem. The estimates therefore condition on the equilibrium branch implied by the documented initialization.*

4.3 Estimates and Channel Fit

The fitted model preserves the sharp reversal in acceptance across routes and closely reproduces doctor acceptance after sorting. It predicts 2,521.52 doctor-initiated approaches, compared with 2,503 observed, and 955.70 agent-initiated approaches, compared with 1,013 observed. Doctor acceptance is 97.98 percent in the data and 97.97 percent in the model after sorting, and 28.35 versus 24.85 percent after screening. Post acceptance is 40.98 versus 41.97 percent after sorting and 98.43 versus 94.76 percent after screening. Thus, the model reproduces the route-side ordering while leaving some differences in acceptance levels, especially after screening.

Table 2 shows how the estimated associations with observed characteristics differ between contact and acceptance. Entries are average probability changes, in percentage points, for a one-standard-deviation increase in a continuous variable or a change from zero to one in an indicator. Other covariates and continuation values are held fixed. Contact entries average changes in the probability of at least one approach during the observed post episode over doctor–post pairs; acceptance entries average over observed decisions for the indicated route and side. The six columns distinguish the two contact routes and each side’s acceptance after either route. Shared acceptance coefficients can imply different probability changes because the starting probabilities and decision samples differ across routes. Parentheses contain approximate transformed standard errors; Appendix Section B.4 explains the calculation, and Table B.4 reports the underlying coefficients and clustered standard errors.

Greater distance is associated primarily with a lower probability of contact. A one-standard-deviation increase in log distance lowers this probability by 0.030 percentage points through doctor-initiated search and 0.006 points through agent-initiated screening, relative to fitted

Table 2. Selected Changes in Contact and Acceptance Probabilities

	Contact		Doctor acceptance		Post acceptance	
	Self	Agency	Self	Agency	Self	Agency
Log distance	-0.030 (0.003)	-0.006 (0.002)	0.176 (0.215)	1.494 (1.934)	0.129 (1.350)	0.026 (0.275)
Hours	0.003 (0.006)	0.018 (0.014)	-0.831 (0.488)	-5.186 (2.397)	1.854 (7.329)	0.366 (1.397)
Log posted salary	-0.010 (0.004)	-0.029 (0.002)	-0.167 (0.353)	-1.262 (2.518)	-31.159 (4.394)	-16.465 (6.414)
Preference match	0.022 (0.009)	0.006 (0.005)	0.271 (0.453)	2.186 (4.001)	-0.899 (2.909)	-0.185 (0.608)
Day match	0.097 (0.020)	0.010 (0.005)	1.504 (0.302)	12.438 (4.202)	4.440 (2.722)	0.888 (0.500)
Specialty match	0.068 (0.011)	0.039 (0.011)	-0.765 (0.672)	-5.476 (3.787)	7.006 (3.862)	1.370 (0.664)
Prior approaches [†] (past 365 days)	0.059 (0.013)	0.037 (0.006)	0.089 (0.036)	0.782 (0.324)	-0.347 (0.250)	-0.078 (0.057)
Prior contract [†]	1.125 (0.198)	1.598 (0.422)	-1.379 (1.246)	-17.629 (8.565)	28.799 (6.216)	9.128 (1.318)

Notes: Entries are probability changes in percentage points. Continuous covariates increase by one full-estimation-sample standard deviation; indicators change from zero to one. Self and Agency denote doctor-initiated and agent-initiated contacts. Contact entries average changes in the probability of at least one contact over the post episode across all retained doctor–post pairs, using their observed stock and duration offsets. Acceptance entries average over the original observed decision trials for the indicated side and route. Continuation values and all other covariates and offsets remain fixed. [†]History comparisons restrict each averaging population to pairs with a positive prior approach count. Approximate standard errors in parentheses transform each coefficient’s doctor–post clustered standard error by the delta method, holding the starting prediction indices fixed; they omit uncertainty in those starting predictions. Appendix B.4 gives the calculation.

averages of 0.090 and 0.033 percent. The acceptance changes are imprecisely estimated on both sides. Geographic proximity therefore predicts which pairs come into contact more clearly than it predicts acceptance once contact occurs.

Job characteristics have different associations with the two sides' acceptance decisions. A one-standard-deviation increase in hours lowers doctor acceptance by 0.83 percentage points after search and 5.19 points after screening. Higher posted salary is associated mainly with lower screening and post acceptance: a one-standard-deviation increase in log salary lowers the screening contact probability by 0.029 points and post acceptance by 31.16 and 16.47 points after search and screening, respectively. Its associations with doctor acceptance are small and imprecisely estimated.

Compatibility with a doctor's preferred work schedule predicts both contact and doctor acceptance. A schedule match raises the two contact probabilities by 0.097 and 0.010 percentage points, and doctor acceptance by 1.50 points after search and 12.44 points after screening. The difference between the doctor-side acceptance changes partly reflects the already high acceptance probability after doctor-initiated search. The post-side point estimates are also positive, at 4.44 and 0.89 points. Specialty matches likewise predict more contact, but their acceptance associations differ by side: the point estimates are positive for posts and negative for doctors. Preferred-location matches predict more contact, with less precisely estimated associations with acceptance. Compatibility measures thus have more consistent positive associations with contact than with both sides' subsequent acceptance.

The relationship-history rows focus on pairs with at least one approach in the preceding 365 days. Within this group, more prior approaches predict more contact through both routes and slightly higher doctor acceptance, with little change in post acceptance. A prior contract has a stronger association with post acceptance: the probability changes are 28.80 points after search and 9.13 points after screening. It also raises the two contact probabilities by 1.125 and 1.598 points. Its doctor-side point estimates are negative, although the underlying coefficient's 95 percent interval includes zero. Prior contact frequency and prior contracting therefore capture different aspects of the relationship history associated with current contact and acceptance.

5 Policy Exercises

The policy exercises change how the platform recommends jobs to doctors. They are motivated by the platform's consideration of an algorithmic recommendation system and ask three broad questions. First, should algorithmic recommendations operate alongside doctors' own search and applications, or should they replace doctor search entirely? Second, what should the algorithm prioritize: generating greater immediate expected surplus or creating greater long-run value after doctors and posts adjust? Third, should the platform increase the number of

recommendations rather than only change who receives them?

Section 5.1 first explains how counterfactual recommendation policies are constructed. The next three subsections answer the three questions in the same order. Section 5.2 compares keeping doctor search with replacing it and shows why the model favors retaining decentralized search. Section 5.3 then conditions on doctor search remaining available, compares what the algorithm should prioritize, and shows how the long-run rule changes who receives recommendations. Finally, Section 5.4 keeps doctor search available and asks what the platform gains from making more recommendations, separates that change from reallocating recommendations across doctors and posts, and compares the gains from redesign with what could be achieved by expanding the market itself.

5.1 Policy Classes and Objectives

All counterfactual policies choose the same object. Let

$$A = (A_{ij})_{i \in I, j \in J}, \quad A_{ij} \geq 0,$$

where A_{ij} is the expected number of monthly contacts that the platform assigns to doctor–post pair (i, j) through recommendations. Equivalently, A_{ij}/K_i is the recommendation rate for one slot of doctor i . A policy class determines which matrices A are feasible and whether doctor search remains available. A policy objective determines how the platform chooses A within that class. I introduce these two ingredients in turn and then describe how the resulting policies are computed and compared.

5.1.1 Policy Class

For each policy class, the comparisons at current contact budgets keep the expected number of contacts that the platform assigns to each doctor at the level estimated for the current market and change only which posts receive those contacts. This is because a policy that sharply changes this number for an individual doctor would be difficult to justify in practice. Instead, I do not impose the same requirement on posts and the number of platform-assigned contacts received by an individual post may therefore change substantially, subject to the pair-specific limits below. The later recommendation-volume exercise changes the doctor-level budgets.

To implement these comparisons, I use fitted current contacts to define pair-specific upper bounds on the policy matrix. Let

$$\hat{S}_{ij} = K_i \lambda_{ij}^{\text{self}}(\hat{\alpha}_i^D) \quad \text{and} \quad \hat{A}_{ij} = K_i \lambda_{ij}^{\text{agency}}(\hat{\alpha}_j^P)$$

denote the fitted monthly numbers of self-initiated and agency-initiated contacts under the current policy. For the one-month horizon, define the pair-specific reference level

$$\overline{M}_{ij} = \max\{K_i, \hat{S}_{ij} + \hat{A}_{ij}\}.$$

This reference allows one contact per doctor slot for pair (i, j) unless the fitted current total is larger, so it never excludes the fitted current allocation. I define the portion that remains available for platform recommendations after accounting for fitted self-initiated contacts as

$$\bar{A}_{ij} = \max\{\bar{M}_{ij} - \hat{S}_{ij}, 0\}.$$

The two policy classes use these quantities differently: $\bar{A}_{ij}$ caps recommendations when doctor search remains available, whereas $\bar{M}_{ij}$ caps recommendations when all contacts are assigned by the platform.

Search-retained class. The *search-retained* class changes platform recommendations while leaving doctor-initiated search governed by the estimated behavioral rule. Let b_i denote the expected monthly number of contacts that the platform assigns to doctor i through recommendations. For a feasible vector $b = (b_i)_{i \in I}$, the class has feasible set

$$\mathcal{B}^{\text{retain}}(b) = \left\{ A \geq 0 : \sum_j A_{ij} = b_i \ \forall i, \quad A_{ij} \leq \bar{A}_{ij} \right\}. \quad (8)$$

The baseline exercise uses each doctor's fitted current recommendation budget, $b_i^{\text{agency}} = \sum_j \hat{A}_{ij}$. For a chosen A , the platform replaces the behavioral agency-intensity equation with the forced rate A_{ij}/K_i . Doctor search continues to satisfy

$$M_{ij}^{\text{self}}(\alpha; K) = K_i \lambda_{ij}^{\text{self}}(\alpha_i^D).$$

Continuation values and acceptance therefore solve

$$\begin{aligned} (r + \delta_i^a) \alpha_i^D &= \sum_j \left[\lambda_{ij}^{\text{self}}(\alpha_i^D) \psi_{ij, \text{self}}^D(\alpha) + \frac{A_{ij}}{K_i} \psi_{ij, \text{agency}}^D(\alpha) \right], \\ (r + w_j) \alpha_j^P &= \sum_i \left[K_i \lambda_{ij}^{\text{self}}(\alpha_i^D) \psi_{ij, \text{self}}^P(\alpha) + A_{ij} \psi_{ij, \text{agency}}^P(\alpha) \right]. \end{aligned}$$

Write $\alpha^{\text{retain}}(A; K)$ for the solution. Thus A is the only object assigned directly by the platform; doctor search, acceptance, and continuation values respond in equilibrium. At fixed market stocks and estimated parameters, a change in α_i^D scales doctor i 's self-search rates across all posts by the same factor. The response changes that doctor's total search intensity, not the relative allocation of that search across posts.

Full-replacement class. The *full-replacement* class removes doctor-initiated search and assigns all contacts through platform recommendations. Its feasible set is

$$\mathcal{B}^{\text{replace}}(b) = \left\{ A \geq 0 : \sum_j A_{ij} = b_i \ \forall i, \quad A_{ij} \leq \bar{M}_{ij} \right\}.$$

The baseline exercise uses each doctor's fitted current total-contact budget, $b_i^{\text{total}} = \sum_j (\hat{S}_{ij} + \hat{A}_{ij})$. A candidate matrix A sets $M^{\text{self}} = 0$ and $M^{\text{agency}} = A$, so continuation values solve

$$\begin{aligned} (r + \delta_i^a) \alpha_i^D &= \sum_j \frac{A_{ij}}{K_i} \psi_{ij, \text{agency}}^D(\alpha), \\ (r + w_j) \alpha_j^P &= \sum_i A_{ij} \psi_{ij, \text{agency}}^P(\alpha). \end{aligned}$$

Write the solution as $\alpha^{\text{replace}}(A; K)$.

5.1.2 Policy Objective

For a fixed policy class $k \in \{\text{retain}, \text{replace}\}$ and feasible b , I compare two objectives for choosing $A \in \mathcal{B}^k(b)$: immediate expected surplus and long-run user value. Both hold fixed the policy class and the number of recommendations assigned to each doctor. I also define an even-allocation benchmark for measuring post-side access.

Benchmark for access. The benchmark spreads each doctor's recommendations as evenly as the pair-specific limits allow.

Let $\bar{C}^{\text{retain}} = \bar{A}$ and $\bar{C}^{\text{replace}} = \bar{M}$. The benchmark allocation is

$$R_{ij}^k(b) = \min\{\bar{C}_{ij}^k, \tau_i^k(b)\}, \quad \sum_j R_{ij}^k(b) = b_i,$$

where $\tau_i^k(b)$ is chosen so that doctor i receives exactly b_i recommendations. It is a benchmark for an even allocation, not the solution to a separate fairness-maximization problem. For the baseline comparison, write $R^{\text{retain}} = R^{\text{retain}}(b^{\text{agency}})$ and $R^{\text{replace}} = R^{\text{replace}}(b^{\text{total}})$.

Short run: immediate expected surplus. The short-run problem allocates recommendations to maximize immediate expected bilateral surplus at the fitted continuation values. To state this problem, let the route-specific contact matrices realized at continuation vector α be

$$M_{ij}^{c,k}(A; \alpha) = \begin{cases} K_i \lambda_{ij}^{\text{self}}(\alpha_i^D), & k = \text{retain}, c = \text{self}, \\ A_{ij}, & k = \text{retain}, c = \text{agency}, \\ 0, & k = \text{replace}, c = \text{self}, \\ A_{ij}, & k = \text{replace}, c = \text{agency}. \end{cases}$$

The objective is

$$S^k(A) = \sum_{i,j,c} M_{ij}^{c,k}(A; \hat{\alpha}) \left\{ \psi_{ijc}^D(\hat{\alpha}) + \psi_{ijc}^P(\hat{\alpha}) \right\}.$$

The short-run rule maximizes $S^k(A)$ over $A \in \mathcal{B}^k(b)$. It holds opportunity costs and acceptance fixed at $\hat{\alpha}$. In the search-retained class it also holds doctor contacts at $\hat{S}$, so only the recommendation matrix changes immediately.

Long run: value after equilibrium adjustment. The long-run rule accounts for how a recommendation allocation changes the equilibrium. For each candidate A , I re-solve acceptance and continuation values. In the search-retained class, this calculation also updates the self-search contact rates $\lambda^{\text{self}}(\alpha^D)$. I then add the resulting continuation values across all doctor slots and posts. I refer to this sum as total user value in the model:

$$W^k(A; K) = \sum_i K_i \alpha_i^{D,k}(A; K) + \sum_j \alpha_j^{P,k}(A; K).$$

The first term counts one doctor-side continuation value for each of doctor i 's K_i slots, and the second counts one post-side continuation value for each post. The long-run problem maximizes $W^k(A; K)$ over $A \in \mathcal{B}^k(b)$. I report the best allocation found by a numerical search that satisfies a first-order stopping criterion.⁵

Comparison across policies. I compare the policies along two dimensions: how evenly they distribute post-side outcomes and how much value they generate according to the two objective functions. Because posts are the institutions that pay for the platform's service, the platform must also avoid recommendation policies that distribute the benefits of that service too unevenly across them; I use the post-side access index to capture this practical equity concern.

To compare the distribution of outcomes across posts, I define a post-specific outcome for each objective. In the short-run comparison, $y_j^{SR,k}(A)$ is post j 's immediate expected surplus flow at the fitted continuation values. In the long-run comparison, $y_j^{LR,k}(A)$ is post j 's equilibrium continuation value:

$$y_j^{SR,k}(A) = \sum_{i,c} M_{ij}^{c,k}(A; \hat{\alpha}) \psi_{ijc}^P(\hat{\alpha}), \quad y_j^{LR,k}(A) = \alpha_j^{P,k}(A; K).$$

For each post, I compare its outcome under A with its outcome under the corresponding even-allocation benchmark. For $q \in \{SR, LR\}$, the resulting post-side access index is

$$\mathcal{A}^{q,k}(A) = 100 \left[1 - \text{Gini}_{j \in J} \left\{ \frac{y_j^{q,k}(A)}{y_j^{q,k}(R^k)} \right\} \right],$$

provided $y_j^{q,k}(R^k) > 0$. The benchmark scores 100 because, when $A = R^k$, every ratio inside the Gini coefficient equals one and the Gini coefficient of this constant vector is zero. A lower score

⁵The search uses the gradient of total user value obtained by differentiating the equilibrium system, including the response of doctor search when it remains available. Each accepted update must raise the value computed after re-solving the equilibrium and satisfy the doctor-level budgets and pair-specific bounds. Each reported long-run allocation attains a relative first-order policy gap no greater than 10^{-6} and a maximum equilibrium residual no greater than 10^{-10} . Appendix C defines the gap and reports independent checks of the final allocations. These checks establish approximate first-order stationarity, not global optimality or equilibrium uniqueness. I therefore retain the term "best-found" long-run rule. All reported values are conditional on the maintained equilibrium branch.

means that posts' outcomes change by more unequal proportions relative to the benchmark. The index is neither the share of posts receiving recommendations nor a measure of inequality in the raw levels of post outcomes.

I also evaluate each displayed allocation under both objective functions. Within each policy class, the short-run objective is reported as a percentage of the value attained by the short-run maximizing rule, so that rule equals 100. The long-run objective is reported as a percentage of the value attained by the best-found long-run rule, so that rule also equals 100.

Within each policy class, I therefore report the short-run maximizing rule and the best-found long-run rule, together with the current hybrid as a comparison. In the full-replacement panel, the current hybrid is an external two-route comparator rather than an element of $\mathcal{B}^{\text{replace}}(b)$. I evaluate its fitted outcomes using the same full-replacement benchmarks. Its normalized objective values can therefore exceed 100: in that panel, 100 is the value attained within the full-replacement policy class, not an upper bound for the external current hybrid.

5.2 Should Doctor Search Remain Available?

This subsection answers the first question: should doctor search remain available alongside algorithmic recommendations? Our simulations show that retaining doctor search makes recommendation design less sensitive to which criterion the platform prioritizes. As shown in Table 3, across the two counterfactual rules, the short- and long-run objective values and access indices all vary less when doctor search is retained, using the within-class benchmarks defined in Section 5.1. The platform can therefore prioritize immediate surplus or long-run value without the extreme deterioration in other dimensions in contrast to the case of full replacement. For example, when all contacts are assigned by the platform, the short-run rule attains only 42.20 percent of the best-found long-run value and a long-run access score of 12.03.

The differences are particularly small in the long run, consistent with doctor search providing an adjustment margin as recommendations change. With doctor search retained, the two rules produce long-run values between 98.64 and 100 percent of the best-found value and long-run access scores between 92.53 and 94.48. Doctors adjust how intensively they search as continuation values change: although the platform assigns the same 1,138.2 expected monthly recommendations under each rule, equilibrium doctor-initiated contacts range from 2,242.2 to 2,876.8. These responses are absent under full replacement, where all contacts are assigned by the platform.

5.3 What Should the Algorithm Prioritize?

Conditional on retaining doctor search, I argue that the platform should prioritize long-run user value when designing its recommendations. Three considerations support this choice.

Table 3. Access, Value, and Realized Exposure under Alternative Policies

Policy rule	Agency	Realized self	Realized total	Short access	Long access	Short objective	Long value
<i>Panel A: Doctor search retained; agency screening redesigned</i>							
Current hybrid	1,138.2	2,773.4	3,911.6	82.11	97.83	80.19	98.89
Short-run best found	1,138.2	2,242.2	3,380.5	43.37	94.48	100.00	98.64
Long-run best found	1,138.2	2,876.8	4,015.1	64.60	92.53	93.32	100.00
<i>Panel B: Doctor search replaced; all contacts assigned by platform</i>							
Current hybrid	1,138.2	2,773.4	3,911.6	72.17	92.95	63.23	100.88
Short-run best found	3,911.6	0.0	3,911.6	4.80	12.03	100.00	42.20
Long-run best found	3,911.6	0.0	3,911.6	76.24	97.14	71.05	100.00

Notes: Agency is the platform-controlled recommendation count. Realized self is the equilibrium doctor-initiated count, and realized total is their sum. The displayed Realized self and Realized total columns are the long-run equilibrium counts for each rule. In Panel A, doctor search is recomputed from continuation values under every screening matrix for the long-run evaluation; it is not held fixed. The short-run objective instead holds baseline self-search contacts and baseline continuation and surplus terms fixed. In Panel B, realized self contacts are zero for full-replacement candidates; the current hybrid is a comparison benchmark outside that candidate class. Access is the post-side index normalized to the even-allocation benchmark. The short objective and long value are normalized within policy class. Best-found is under maintained equilibrium selection and does not claim a global optimum. Agency budget and agency decision-pair caps are feasibility gates. Any realized-total overage created by endogenous doctor search is reported as a non-gating diagnostic.

First, the current hybrid already performs close to the long-run benchmarks, whereas its short-run performance leaves more room for improvement. Panel A of Table 3 reports a long-run access score of 97.83 and long-run value equal to 98.89 percent of the best-found value. In contrast, its short-run access score is 82.11 and its short-run surplus is 80.19 percent of the maximum. Thus, both long-run measures are close to their benchmarks even though the short-run measures remain around 80. This pattern suggests that the platform already places implicit weight on long-run outcomes. That interpretation is consistent with my communications with the platform, in which its representatives emphasized that agents’ work extends beyond raising short-term match success rates.

Second, maximizing short-run surplus can make both long-run outcomes worse than under the current hybrid. The short-run rule reduces long-run access from 97.83 to 94.48 and normalized long-run value from 98.89 to 98.64, so it is dominated by the current hybrid on these two measures. This comparison qualifies the robustness discussed in Section 5.2: retaining doctor search limits the sensitivity of outcomes to the objective, but it does not prevent a short-run rule from sacrificing both long-run value and access. If recommendation design should work well beyond the criterion it directly maximizes, this is a reason to avoid focusing solely on immediate surplus.

Third, optimizing long-run value can improve welfare at both horizons relative to the current hybrid, while accepting a reduction in access. The long-run rule raises the normalized short-run

surplus objective from 80.19 to 93.32, a gain of 13.13 percentage points, while raising normalized long-run value from 98.89 to 100. Its short-run access score falls from 82.11 to 64.60, a reduction of about one fifth; long-run access also falls, from 97.83 to 92.53. The rule therefore improves welfare under either horizon’s criterion, although it requires a tradeoff with post-side access. Together with the current hybrid’s strong long-run performance, these comparisons support prioritizing long-run user value in recommendation design.

Uneven gains from redesign. Although the redesign raises total user value only modestly relative to the current hybrid, it produces uneven gains and losses across posts. Doctor-side value falls by 0.20 percent and post-side value rises by 6.78 percent, but only 45.85 percent of posts gain value. Grouping posts by their continuation values under the current hybrid, the lowest quintile gains 16.18 percent in aggregate, whereas the highest quintile loses 0.96 percent. These fixed baseline groups show that the increase in post-side value includes gains among initially lower-value posts, while leaving many other posts worse off. It is not a uniform improvement in post-side outcomes.

The search-retained redesign raises the total number of matches, although it changes the routes through which they occur. Relative to the current hybrid, recommendation-initiated matches rise from 203.4 to 355.6 per month, while doctor-initiated matches fall from 945.8 to 894.4. Total matches therefore rise from 1,149.2 to 1,250.0; they also exceed the 1,168.3 matches under the short-run rule. The gain in continuation value reflects this larger number of matches after allowing for the different gains obtained from each doctor–post pair and route. Appendix Table C.2 separates the match number, match composition, and conditional-gain contributions. The allocation of recommendations is therefore important both for the matches it creates directly and for the matches that continue to arise through doctor search.

How does the redesign change recommendations? Table 4 describes which recommendations change by comparing the current and best-found allocations with doctor search retained. Each doctor’s recommendation total is fixed, so the differences show which posts receive the existing recommendations.

The best-found long-run rule shifts recommendations toward shorter jobs and more matches to doctors’ preferred work schedules, while assigning more weight to distant posts and fewer recommendations to specialty matches. Recommendation-weighted job duration falls from 11.19 to 6.43 hours. Recommendation-weighted geometric-mean distance rises from 28.7 to 37.2 kilometers. The share of recommendations matching the doctor’s preferred combination of weekday and shift type rises by 18.62 percentage points. Preferred-location matches rise slightly, by 0.73 percentage points, whereas specialty matches fall by 24.03 percentage points. Thus, the higher-value allocation does not improve every familiar measure of pair compatibility: it favors shorter

Table 4. How the Best-Found Long-Run Rule Retargets Recommendations

Recommendation-weighted attribute	Current	Best found	Change
Job duration (hours)	11.19	6.43	−4.76
Mean log distance	3.356	3.617	+0.261
Preference match (%)	36.32	37.05	+0.73 pp
Day match (%)	35.73	54.35	+18.62 pp
Specialty match (%)	31.61	7.58	−24.03 pp

Notes: Entries weight observed doctor–post attributes by platform-controlled recommendations. Current refers to fitted agency recommendations, and Best found refers to the best-found long-run recommendation matrix when doctor search remains available. Each doctor’s total number of recommendations is held at its fitted current level. Equilibrium doctor search is recomputed when evaluating the policy but is not used as a weight in this table.

jobs and preferred work schedules while accepting longer distances and fewer specialty matches.

This exercise changes which posts receive the existing recommendations. The remaining question is whether the platform can create larger gains by increasing the number of recommendations rather than only changing their allocation.

5.4 How Should the Platform Make More Recommendations?

I organize this exercise around three sequential policy changes: recommendation expansion, equalization across doctors, and pairwise redesign. These changes correspond to three operational decisions: how many recommendations the platform makes, whether doctors receive equal recommendation opportunities, and which doctor–post pairs receive those recommendations. First, the platform doubles the total number of recommendations. Second, holding that larger total fixed, it assigns the same expected number of recommendations to every doctor. Third, holding both the aggregate volume and each doctor’s recommendation total fixed, it redesigns which posts receive those recommendations.

To construct the first step, I start from doctor i ’s fitted current recommendation total,

$$b_i^{\text{agency}} = \sum_j \hat{A}_{ij},$$

in the feasible set in equation (8). The doubled profile assigns doctor i a recommendation total of $2b_i^{\text{agency}}$. For the second step, I keep the same doubled aggregate volume but assign every doctor the same expected number of recommendations:

$$b_i^{\text{eq}} = \frac{2 \sum_{\ell} b_{\ell}^{\text{agency}}}{|I|}.$$

The equalization step applies only to recommendations assigned by the platform; endogenous doctor search can still make total contacts differ across doctors. For the third step, I hold b^{eq}

Table 5. Recommendation Expansion, Equalization across Doctors, and Pairwise Redesign

Panel A: Value changes and who gains

Policy step	Value change (%)			Share with higher value (%)	
	Total	Doctors	Posts	Doctor slots	Posts
Double current recommendations	+0.65	−1.77	+10.97	2.17	100.00
Equalize recommendation totals	+0.44	−0.83	+5.22	16.25	84.04
Redesign pair allocation	+1.80	+1.60	+2.50	97.18	57.88
Overall: current to final	+2.91	−1.03	+19.68	17.47	70.44

Panel B: Contacts and matches in the final redesign step

Route	Change (%)		Conversion rate (%)	
	Contacts	Matches	Before	After
Doctor search	−25.40	+0.09	16.26	21.81
Recommendations	0.00	+29.94	24.42	31.73

Notes: Doctor search remains available throughout. The first three rows of Panel A compare each step with the preceding policy; Overall compares the final redesigned allocation with the current hybrid. Each percentage change uses the corresponding total or side-specific value under the baseline for that row. Share with higher value is the percentage of doctor slots or posts whose continuation value exceeds that baseline by more than 10^{-4} . Doctors are weighted by their numbers of slots; posts receive equal weight. Overall shares are computed directly from current and final values for each unit. Panel B compares doubled recommendations with equal doctor totals against the final redesigned allocation. Conversion is accepted matches divided by contacts on the indicated route. All entries condition on the estimated model and re-solved equilibria.

fixed and search over

$$A \in \mathcal{B}^{\text{retain}}(b^{\text{eq}})$$

for the best-found long-run allocation of each doctor’s recommendations across posts.

Panel A of Table 5 reports the percentage gains in aggregate continuation value from these three changes. Doubling the current recommendation profile raises value by 0.65 percent. Equalizing recommendation totals across doctors adds a further 0.44 percent, and redesigning the allocation across posts at those equal totals adds 1.80 percent. Each step’s percentage gain is measured relative to the policy immediately before it. The overall comparison in Panel A shows that the three changes together raise value by 2.91 percent relative to the current hybrid. Pairwise redesign contributes more than the first two steps combined, giving the platform a reason to improve the allocation of recommendations even when it can increase their number.

Uneven gains across the three steps. Panel A also shows that the three steps distribute their gains differently between doctors and posts. Relative to the current hybrid, the final allocation lowers doctor-side value by 1.03 percent while raising post-side value by 19.68 percent. Continuation values increase for 17.47 percent of doctor slots and 70.44 percent of posts. Expansion and equalization raise post-side value while reducing doctor-side value. Pairwise redesign then recovers part of those doctor-side losses: it raises doctor-side value by 1.60 percent and post-side value by 2.50 percent. Doctors receive 69.52 percent of this last increment. In this final step, continuation values increase for 97.18 percent of doctor slots and 57.88 percent of posts. Among posts, the lowest baseline-value quintile gains 11.17 percent, while the highest quintile loses 4.94 percent. These comparisons show why the aggregate increments should be assessed together with who receives them.

To understand the 1.80 percent value gain from the final pairwise redesign, I examine how contacts and accepted matches change through doctor search and recommendations. Panel B of Table 5 compares the doubled, equal-budget profile with the redesigned allocation. Doctor-initiated contacts fall by 25.40 percent, while matches through that route rise by just 0.09 percent. The share of these contacts that both sides accept rises from 16.26 to 21.81 percent. At the same time, the fixed number of recommendations generates 29.94 percent more matches, and their conversion rate rises from 24.42 to 31.73 percent. Thus, the model predicts more successful recommendations while fewer doctor applications sustain roughly the same number of doctor-initiated matches.

Market expansion versus recommendation redesign. To put the 1.80 percent value gain from pairwise redesign in perspective, I compare it with the gains from attracting additional doctors or posts. From the platform’s perspective, recommendation redesign and market expansion are both ways to raise user value. I therefore ask how much targeted market expansion would generate the same increase in aggregate continuation value.

The market-expansion benchmark starts from the doubled recommendation profile with equal totals across doctors, before pairwise redesign. It therefore compares market expansion with a 1.80 percent increase in value relative to that profile, matching the final step in Table 5. Doctor-side expansion increases the number of slots available from selected doctor types. On the post side, let L_j denote the number of identical copies of observed post j , with $L_j = 1$ in the observed market. Aggregate contacts directed to post type j are divided across its copies, so each copy receives route- c contact flow M_{ij}^c/L_j . The corresponding long-run value is

$$W^{\text{retain}}(A; K, L) = \sum_i K_i \alpha_i^{D, \text{retain}}(A; K, L) + \sum_j L_j \alpha_j^{P, \text{retain}}(A; K, L).$$

I rank doctors by $\hat{\alpha}_i^D$ and observed posts by $\hat{\alpha}_j^P$, and add stock to the top decile on each side. Each addition is distributed across the selected types in proportion to their baseline stocks.

During this exercise, I hold the platform recommendation matrix fixed at that pre-redesign profile and re-solve the equilibrium after each change in market stocks, allowing doctor search and continuation values to adjust.⁶

With post stock fixed, linear interpolation between the evaluated grid points matches the redesign gain with additional slots from top-decile doctor types equal to 1.98 percent of aggregate baseline doctor-slot stock. With doctor stock fixed, the corresponding top-decile post addition is 5.87 percent of aggregate baseline post stock. Appendix Figure C.5 shows the interpolated value contours, with the contour matching the redesign gain highlighted. These results put the scale of pairwise redesign in perspective. Improving the allocation of existing recommendations can substitute for some market expansion, although the preferred investment also depends on the costs of redesign and of attracting additional doctors or posts.

References

- Adachi, Hiroyuki.** 2003. “A Search Model of Two-Sided Matching under Nontransferable Utility.” *Journal of Economic Theory*, 113(2): 182–198.
- Agarwal, Nikhil.** 2015. “An Empirical Model of the Medical Match.” *American Economic Review*, 105(7): 1939–1978.
- Belot, Michèle, Philipp Kircher, and Paul Muller.** 2019. “Providing Advice to Job-seekers at Low Cost: An Experimental Study on Online Advice.” *Review of Economic Studies*, 86(4): 1411–1447.
- Biega, Asia J., Krishna P. Gummadi, and Gerhard Weikum.** 2018. “Equity of Attention: Amortizing Individual Fairness in Rankings.” 405–414.
- Burdett, Ken, and Melvyn G. Coles.** 1997. “Marriage and Class.” *Quarterly Journal of Economics*, 112(1): 141–168.
- Cameron, A. Colin, Jonah B. Gelbach, and Douglas L. Miller.** 2011. “Robust Inference With Multiway Clustering.” *Journal of Business & Economic Statistics*, 29(2): 238–249.
- Chade, Hector, Jan Eeckhout, and Lones Smith.** 2017. “Sorting through Search and Matching Models in Economics.” *Journal of Economic Literature*, 55(2): 493–544.

⁶I consider seven levels of doctor-slot additions and seven levels of post additions, and solve the equilibrium for all 49 combinations. Each addition is expressed as a percentage of the corresponding side’s total stock in the baseline market. For example, a 1 percent doctor-slot addition allocates new slots totaling 1 percent of all baseline doctor slots to the selected top-decile doctor types; it does not increase each selected type’s slots by 1 percent. Post additions are defined analogously.

- Chen, Kuan-Ming, Yu-Wei Hsieh, and Ming-Jen Lin.** 2023. “Reducing Recommendation Inequality via Two-Sided Matching: A Field Experiment of Online Dating.” *International Economic Review*, 64(3): 1201–1221.
- Chen, Nan, and Hsin-Tien Tsai.** 2024. “Steering via Algorithmic Recommendations.” *RAND Journal of Economics*, 55(4): 501–518.
- Eeckhout, Jan, and Philipp Kircher.** 2010. “Sorting versus Screening: Search Frictions and Competing Mechanisms.” *Journal of Economic Theory*, 145(4): 1354–1385.
- Fradkin, Andrey.** 2017. “Search, Matching, and the Role of Digital Marketplace Design in Enabling Trade: Evidence from Airbnb.” SSRN Working Paper.
- Galichon, Alfred, Scott Duke Kominers, and Simon Weber.** 2019. “Costly Concessions: An Empirical Framework for Matching with Imperfectly Transferable Utility.” *Journal of Political Economy*, 127(6): 2875–2925.
- Gong, Jing.** 2016. “A Structural Two-Sided Matching Model of Online Labor Markets.” SSRN Working Paper.
- Hitsch, Gunter J., Ali Hortaçsu, and Dan Ariely.** 2010. “Matching and Sorting in Online Dating.” *American Economic Review*, 100(1): 130–163.
- Horton, John J.** 2017. “The Effects of Algorithmic Labor Market Recommendations: Evidence from a Field Experiment.” *Journal of Labor Economics*, 35(2): 345–385.
- Immorlica, Nicole, Brendan Lucier, Vahideh Manshadi, and Alexander Wei.** 2023. “Designing Approximately Optimal Search on Matching Platforms.” *Management Science*, 69(8): 4609–4626.
- Lauermann, Stephan, and Georg Nöldeke.** 2025. “Matching with Frictions.” In *Handbook of the Economics of Matching*. Vol. 2, Chapter 7, 579–642. Elsevier B.V.
- Le Barbanchon, Thomas, Lena Hensvik, and Roland Rathelot.** 2023. “How Can AI Improve Search and Matching? Evidence from 59 Million Personalized Job Recommendations.” SSRN Working Paper.
- Manshadi, Vahideh, Scott Rodilitz, Daniela Saban, and Akshaya Suresh.** 2025. “Online Algorithms for Matching Platforms with Multichannel Traffic.” *Management Science*, 71(9): 7674–7691.
- Moen, Espen R.** 1997. “Competitive Search Equilibrium.” *Journal of Political Economy*, 105(2): 385–411.

- Shimer, Robert, and Lones Smith.** 2000. “Assortative Matching and Search.” *Econometrica*, 68(2): 343–369.
- Shi, Peng.** 2023. “Optimal Matchmaking Strategy in Two-Sided Marketplaces.” *Management Science*, 69(3): 1323–1340.
- Shi, Peng.** 2025. “Optimal Match Recommendations in Two-Sided Marketplaces with Endogenous Prices.” *Management Science*, 71(9): 7431–7448.
- Sürer, Özge, Robin Burke, and Edward C. Malthouse.** 2018. “Multistakeholder Recommendation with Provider Constraints.” 54–62.
- Todorov, Emanuel.** 2006. “Linearly Solvable Markov Decision Problems.” Vol. 19. MIT Press.
- Wang, Yichen, Grady Williams, Evangelos Theodorou, and Le Song.** 2017. “Variational Policy for Guiding Point Processes.” Vol. 70 of *Proceedings of Machine Learning Research*, 3684–3693. PMLR.

A Model Details and Proofs

A.1 Route-Specific Closed Forms

For route c , write $\mu_{ijc}^D := u_{ijc}^D$ and $\mu_{jic}^P := u_{jic}^P$ for the systematic payoffs parameterized in Section 4.1. If ε is logistic with scale ζ and location zero, then

$$\mathbb{E}[(\mu + \varepsilon - \alpha)_+] = \zeta \log(1 + \exp\{(\mu - \alpha)/\zeta\}).$$

Consequently,

$$\begin{aligned} q_{ijc}^D(\alpha) &= \zeta^D \log\left(1 + \exp\{(\mu_{ijc}^D - \alpha_i^D)/\zeta^D\}\right), \\ q_{ijc}^P(\alpha) &= \zeta^P \log\left(1 + \exp\{(\mu_{jic}^P - \alpha_j^P)/\zeta^P\}\right), \\ \psi_{ijc}^D(\alpha) &= \frac{q_{ijc}^D(\alpha)}{1 + \exp\{(\alpha_j^P - \mu_{jic}^P)/\zeta^P\}}, \quad \psi_{ijc}^P(\alpha) = \frac{q_{ijc}^P(\alpha)}{1 + \exp\{(\alpha_i^D - \mu_{ijc}^D)/\zeta^D\}}. \end{aligned} \tag{9}$$

These expressions make clear why route rates cannot be collapsed once either initiator-specific acceptance intercept difference is nonzero. For example,

$$\lambda_{ij}^{\text{self}} \psi_{ij,\text{self}}^D + \lambda_{ij}^{\text{agency}} \psi_{ij,\text{agency}}^D \neq (\lambda_{ij}^{\text{self}} + \lambda_{ij}^{\text{agency}}) \psi_{ij}^D$$

for any route-common ψ_{ij}^D in general. Setting both χ terms to zero recovers the nested route-common acceptance specification.

A.2 Auxiliary Representation of Behavioral Intensities

The behavioral exposure equations admit the relative-entropy representation in equation (1). For a Poisson intensity λ and baseline $\bar{\lambda} > 0$, the KL divergence rate defined in the main text is

$$D_{\text{KL}}(\lambda \parallel \bar{\lambda}) = \lambda \log(\lambda/\bar{\lambda}) - \lambda + \bar{\lambda}.$$

Its derivative with respect to λ is $\log(\lambda/\bar{\lambda})$. The auxiliary objective

$$\lambda(s - \omega) - D_{\text{KL}}(\lambda \parallel \bar{\lambda})$$

is strictly concave for $\lambda > 0$, and its unique first-order solution is

$$\log \lambda = \log \bar{\lambda} + s - \omega.$$

These expressions use the maintained normalization $\zeta^{\text{self}} = \zeta^{\text{agency}} = 1$ for the two exposure routes. Hence the relevant continuation value enters each behavioral log-intensity with coefficient minus one, not with an estimated route-specific loading. For the doctor-search route, $\omega = \alpha_i^D$; for behavioral agent screening at the estimated current-market equilibrium, $\omega = \alpha_j^P$. A constant $\log \bar{\lambda}$ is observationally part of the route's exposure intercept, so the implementation does not estimate a separate baseline-rate parameter.

This representation supplies intuition for the log-intensity specification, not a second payoff term in the stationary-renewal problem. In particular, the KL expression is absent from the continuation-value equations and from the counterfactual value objective. In the search-retained counterfactual the platform also replaces the behavioral agency first-order condition with a fixed agency contact matrix, while the doctor-search first-order condition continues to determine $\lambda^{\text{self}}(\alpha^D)$.

A.3 Proof of Proposition 1

Proof. Characterisation. Let (a, α) be a fixed-rate stationary-renewal equilibrium. Because the payoff shocks are continuous, ties occur with probability zero, and optimality requires the threshold rules in Section 3.1. Consider one waiting slot of doctor i over an interval of length dt . A route- c exposure to post j occurs with probability $\lambda_{ij}^c dt + o(dt)$. Relative to continued waiting, its expected value is

$$\mathbb{E}[a_{jic}^P(U_{ij}^c - \alpha_i^D)_+] = \psi_{ijc}^D(\alpha).$$

Discounting and abandonment reduce the waiting value at rate $r + \delta_i^a$. Stationary renewal leaves the environment after every event unchanged. Thus

$$\alpha_i^D = (1 - (r + \delta_i^a)dt)\alpha_i^D + \sum_{j,c} \lambda_{ij}^c dt \psi_{ijc}^D(\alpha) + o(dt),$$

which gives the doctor equation in (5). The post equation follows in the same way, except that the K_i independent slots generate a total route- c flow $K_i \lambda_{ij}^c$ from doctor i to post j . Hence every equilibrium continuation vector is a fixed point. Conversely, a fixed point together with the associated threshold rules satisfies optimal acceptance and value consistency, and therefore is a stationary-renewal equilibrium.

Existence. For fixed nonnegative finite route rates, every function ψ_{ijc}^D and ψ_{ijc}^P in (9) is continuous and nonnegative. It is also nonincreasing in every continuation-value argument on which it depends. Therefore $\Phi_K(\cdot; \lambda)$ is continuous and, for all $\alpha \geq 0$,

$$0 \leq \Phi_K(\alpha; \lambda) \leq \Phi_K(0; \lambda)$$

componentwise. Define the box

$$B = \prod_k [0, \Phi_{K,k}(0; \lambda)].$$

It is nonempty, compact, and convex, and $\Phi_K(B; \lambda) \subseteq B$. Brouwer's fixed-point theorem gives at least one fixed point in B . $\square$

Under the behavioral rule (2), route rates depend on the relevant continuation value. Substituting $\Gamma(\alpha; s)$ into Φ_K preserves continuity and nonnegativity. Both Γ and the expected-gain terms are componentwise nonincreasing on the nonnegative orthant, so

$$0 \leq \Phi_K(\alpha; \Gamma(\alpha; s)) \leq \Phi_K(0; \Gamma(0; s))$$

componentwise. The box whose upper corner is the last expression is invariant, so Brouwer's theorem establishes existence for the behavioral system (6). The search-retained counterfactual is a hybrid instance of the same argument: the agency matrix is fixed, while $\lambda^{\text{self}}(\alpha^D)$ remains continuous, nonnegative, and bounded above by its value at zero. Only the full-replacement benchmark holds both route matrices fixed while continuation values and acceptance are resolved.

A.4 A Sufficient Condition for Fixed-Exposure Uniqueness

Define the zero-continuation upper bounds

$$\bar{q}_{ijc}^D = \zeta^D \log \left(1 + e^{\mu_{ijc}^D / \zeta^D} \right), \quad \bar{q}_{ijc}^P = \zeta^P \log \left(1 + e^{\mu_{ijc}^P / \zeta^P} \right).$$

Lemma 1 (Fixed-exposure contraction). *Suppose that both λ^{self} and λ^{agency} are fixed matrices*

that do not vary with α or with the fixed-point iterate. Let

$$\kappa_\Phi = \max \left\{ \max_i \frac{\sum_{j,c} \lambda_{ij}^c \left(1 + \frac{\bar{q}_{ijc}^D}{4\zeta^P} \right)}{r + \delta_i^a}, \max_j \frac{\sum_{i,c} K_i \lambda_{ij}^c \left(1 + \frac{\bar{q}_{ijc}^P}{4\zeta^D} \right)}{r + w_j} \right\}.$$

If $\kappa_\Phi < 1$, then Φ_K is a contraction under the sup norm and the stationary-renewal equilibrium is unique.

Proof. From the closed forms,

$$\left| \frac{\partial \psi_{ijc}^D}{\partial \alpha_i^D} \right| = P_{jic}^P P_{ijc}^D \leq 1.$$

The derivative of a logistic acceptance probability is bounded in absolute value by $1/(4\zeta)$, while $q_{ijc}^D(\alpha) \leq \bar{q}_{ijc}^D$ for $\alpha \geq 0$. Hence

$$\left| \frac{\partial \psi_{ijc}^D}{\partial \alpha_j^P} \right| \leq \frac{\bar{q}_{ijc}^D}{4\zeta^P}.$$

The symmetric bounds hold for ψ_{ijc}^P . It follows that every absolute row sum of the Jacobian of Φ_K is no greater than κ_Φ . The mean-value inequality therefore gives $\|\Phi_K(\alpha) - \Phi_K(\alpha')\|_\infty \leq \kappa_\Phi \|\alpha - \alpha'\|_\infty$, and Banach's theorem applies. $\square$

This condition requires discounting and exit to dominate arrival-weighted surplus. It is sufficient, not necessary, and applies when the route-rate matrices are fixed. It therefore applies to the fixed-route full-replacement map, but it does not establish uniqueness of the behavioral exposure map used in estimation or of the search-retained hybrid map. Reported long-run results instead condition on the maintained numerical equilibrium selection and are labelled best found rather than optimal.

A.5 Continuous Time and Route Accounting

Independent continuous-time Poisson clocks imply that the probability of a self- and an agency-initiated approach arriving at exactly the same instant is zero. This justifies adding the two arrival rates. It does not justify pooling their post-exposure outcomes. When acceptance depends on the initiating route, the fixed point must retain

$$\sum_c \lambda_{ij}^c p_{ijc}, \quad \sum_c \lambda_{ij}^c \psi_{ijc}^D, \quad \text{and} \quad \sum_c K_i \lambda_{ij}^c \psi_{ijc}^P$$

as route-weighted sums. The estimator and policy solver use this accounting throughout.

The stationary-renewal assumption enters the structural continuation values and all long-run policy levels. The raw route-specific acceptance patterns do not require it. The short-run

Table B.1. Sample, Route, and History Construction

Object	Count	Positive cells	Share (%)
<i>A. Sample flow</i>			
Raw December doctor–post cells	2,768,872		
Timeline-invalid cells excluded	38,488		
Final doctor–post cells	2,730,384		
Active doctors	1,132		
Final posts	2,412		
Timeline-invalid posts excluded	34		
<i>B. Route-specific approaches entering the MPEC likelihood</i>			
Sorting (doctor initiated)	2,503	2,481	0.091
Screening (agent initiated)	1,013	1,009	0.037
<i>C. Prior-relationship support</i>			
All final cells	2,730,384	121,065	4.43
Cells with any panel-recorded approach	3,518	1,758	49.97
Cells with a recorded bilateral-status decision	3,381	1,716	50.75
Route-attributable acceptance-likelihood cells	3,281	1,685	51.36

Notes: The raw data contain 2,504 doctor-initiated events. The estimation sample excludes one audited out-of-window Web event, recorded on November 18, 2024 after the post’s half-open observation window ended on November 15. Consequently, 2,503 doctor-initiated events enter the MPEC likelihood. The same exclusion removes one positive sorting cell and one doctor- and post-side sorting acceptance success/trial. Panel C’s first three rows describe the raw final panel; its last row and Panel B describe the post-policy likelihood sample. Repeated in-window approaches remain distinct exposure events. Histories exclude the current post episode and use only information available before its order date.

objective does not re-solve the stationary system after a policy change, but it still inherits the baseline continuation values from the baseline MPEC estimate. Accordingly, neither the short-run nor the long-run exercise should be read as a literal transition forecast under nonstationary entry, seasonal composition, or permanent depletion of matched types.

B Additional Data and Estimation Material

B.1 Episode-Aware Route and Decision Construction

This subsection reports the post-episode selection rule, treatment of reused post identifiers, raw-media route mapping, bilateral status mapping, and the acceptance exclusions described in Section 2.2.

B.2 Descriptive Statistics

Table B.2 summarizes doctor, post, and doctor–post characteristics, while Table B.3 reports the five measures of past doctor–facility activity. Contract history is dated by the recorded contract date; approach counts, route shares, and recency use approach timestamps. No outcome resolved after the current post’s order date enters the main feature set.

Table B.2. Market Descriptive Statistics

Variable	<i>N</i>	Mean	SD	P25	Median	P75	Nonzero (%)
<i>A. Doctor characteristics</i>							
Approach events per doctor	1,132	3.11	3.91	—	2.00	—	—
Age	1,132	42.74	12.06	33.00	40.00	50.25	100.00
Experience	1,132	16.49	11.77	7.00	13.00	24.00	100.00
Lagged activity/throughput proxy K_i	1,132	3.70	4.91	1.00	2.00	4.00	100.00
<i>B. Post characteristics</i>							
Approach events per post	2,412	1.46	1.32	—	1.00	—	—
Hours	2,412	10.12	8.68	4.00	8.50	12.50	100.00
Advertised pay	2,412	72.55	43.42	44.00	60.50	91.25	100.00
Advertised wage per hour	2,412	9.13	3.91	5.83	10.00	11.43	100.00
<i>C. Doctor–post characteristics</i>							
Distance	2,730,384	61.08	104.81	22.76	34.40	58.32	100.00
Preferred-location match (%)	2,730,384	25.63	43.66	0.00	0.00	100.00	25.63
Preferred-day match (%)	2,730,384	31.49	46.45	0.00	0.00	100.00	31.49
Specialty match (%)	2,730,384	15.84	36.51	0.00	0.00	0.00	15.84

Notes: Entries are raw sample statistics. Indicator variables are shown in percentage points. Age and experience are in years; hours are hours; distance is in kilometers; advertised pay is in thousand yen; and hourly wage is in thousand yen per hour. Dashes mark unavailable quantiles or nonzero shares. The lagged activity measure is a throughput proxy, not observed physical capacity.

Table B.3. Predetermined Doctor–Facility Relationship Variables

Variable	<i>N</i>	Mean	SD	P25	Median	P75	Nonzero (%)
Prior approaches (365 days)	2,730,384	0.11	1.92	0.00	0.00	0.00	2.38
Any prior contract (%)	2,730,384	2.00	13.98	0.00	0.00	0.00	2.00
Days since last approach	2,730,384	19.58	127.11	0.00	0.00	0.00	4.43
Prior agency share (365 days) (%)	2,730,384	0.49	6.55	0.00	0.00	0.00	0.66
No prior relationship (%)	2,730,384	95.57	20.58	100.00	100.00	100.00	95.57

Notes: History is measured strictly before each post’s order date and shown before estimation standardization. Days since last approach and prior agency share are coded zero when no prior relationship exists; the separate no-prior-relationship indicator records that boundary.

B.3 MPEC Implementation and Convergence Diagnostics

Let $\ell(\theta, \alpha)$ denote the sum of the two route-specific Poisson approach-count components and the side-specific Bernoulli trials supplied by route-attributable doctor–post pairs, and let $c(\theta, \alpha) = \alpha - \Phi(\theta, \alpha)$. For numerical evaluation, the implementation uses the algebraically equivalent update

$$\begin{aligned}\tilde{\Phi}_i^D &= \frac{\sum_{j,c} \lambda_{ij}^c \{\psi_{ijc}^D + \alpha_i^D p_{ijc}\}}{r + \delta_i^a + h_i^D}, & h_i^D &= \sum_{j,c} \lambda_{ij}^c p_{ijc}, \\ \tilde{\Phi}_j^P &= \frac{\sum_{i,c} K_i \lambda_{ij}^c \{\psi_{ijc}^P + \alpha_j^P p_{ijc}\}}{r + w_j + h_j^P}, & h_j^P &= \sum_{i,c} K_i \lambda_{ij}^c p_{ijc}.\end{aligned}$$

For either side k , with the corresponding exit rate e_k and match hazard h_k , the residuals obey the pointwise identity

$$(r + e_k + h_k) \{\alpha_k - \tilde{\Phi}_k(\theta, \alpha)\} = (r + e_k) \{\alpha_k - \Phi_k(\theta, \alpha)\}.$$

The two residual maps therefore have exactly the same zero set. Writing $\tilde{c}(\theta, \alpha) = \alpha - \tilde{\Phi}(\theta, \alpha)$, the implemented MPEC problem is

$$\max_{\theta, \alpha} \ell(\theta, \alpha) \quad \text{subject to} \quad \tilde{c}(\theta, \alpha) = 0,$$

which is equivalent to imposing $c(\theta, \alpha) = 0$ in the model. The active parameter space contains 92 structural coordinates, including both route contact-rate constants, and 3,544 continuation values (1,132 doctor values and 2,412 post values). Fourteen direct coordinates with exact separation or no support are fixed at zero. The calculation does not profile the route constants or impose the fixed- α likelihood score as an additional restriction. The exposure scales are normalized to $\zeta^{\text{self}} = \zeta^{\text{agency}} = 1$; neither an inverse exposure scale nor its logarithm is an estimated coordinate. Consequently the same continuation values used in acceptance enter the two exposure equations with coefficient minus one.

The parameter search solves the equilibrium conditions at each proposed structural parameter vector and uses derivatives that include the resulting change in continuation values. Proposals are accepted only when they retain equilibrium feasibility and improve the likelihood. The full reduced Hessian is then used to check whether an additional local correction would materially change the parameter vector or likelihood.

The reported estimate has a six-block unprofiled log likelihood of $-28,664.909007$, with a maximum absolute equilibrium residual of 2.486900×10^{-13} . At this point the reduced Hessian of the negative log likelihood is positive definite in all 92 active coordinates. The largest normalized Newton correction, computed without a step cap or Hessian regularization, is 2.234477×10^{-5} , and its local quadratic prediction of the increase in total log likelihood is 1.553244×10^{-10} .⁷

⁷For the numerical step diagnostic, write $\theta_k = s_k u_k$, where $s_k = 1/\max\{1, \sqrt{\text{mean}(z_k^2)}\}$ and the mean is taken over doctor–post pairs for the corresponding design column. The two initiator-specific acceptance intercept differences have $s_k = 1$. Reported coefficients and their covariance retain the original parameter coordinates.

Together with two successive small updates accepted without truncation, these values satisfy the practical numerical criteria: an equilibrium residual at most 10^{-8} , a remaining normalized correction below 0.001, and a predicted likelihood gain below 0.01. The final curvature also passes directional finite-difference checks. Table B.10 reports the numerical diagnostics. This is a local numerical convergence assessment; it does not establish a global optimum or uniqueness of the equilibrium.

B.4 Analytic Standard Errors

The standard errors allow the continuation values to respond to changes in the structural parameters. Let $f = -\ell$ be the raw total negative log likelihood. On the maintained regular equilibrium branch, define

$$D = \frac{\partial \alpha}{\partial \theta'} = -\tilde{c}_\alpha^{-1} \tilde{c}_\theta, \quad \nu = -\tilde{c}_\alpha^{-T} f_\alpha, \quad R = \begin{bmatrix} I \\ D \end{bmatrix}.$$

All quantities are evaluated at the reported estimate. The observed curvature of the likelihood restricted to the equilibrium branch is

$$H = R' \nabla_{(\theta, \alpha)}^2 \{f(\theta, \alpha) + \nu' \tilde{c}(\theta, \alpha)\} R,$$

where ν is held fixed when taking the Hessian. Including the curvature of the constraints is necessary; holding continuation values fixed would give a different information matrix.

For doctor–post pair (i, j) , let ℓ_{ij} collect its two contact-count contributions and all observed acceptance-decision contributions. Its score on the equilibrium branch is

$$s_{ij} = \nabla_\theta \ell_{ij} + D' \nabla_\alpha \ell_{ij}.$$

Following Cameron, Gelbach and Miller (2011), I cluster these scores by doctor and by post, allowing dependence among pairs sharing either index. Writing $S_i^D = \sum_j s_{ij}$ and $S_j^P = \sum_i s_{ij}$, the two-way cluster covariance uses

$$B = \sum_i S_i^D (S_i^D)' + \sum_j S_j^P (S_j^P)' - \sum_{i,j} s_{ij} s_{ij}', \quad \hat{V} = H^{-1} B H^{-1}.$$

The subtraction removes double counting of the doctor–post intersections. The main calculation uses no finite-cluster correction (CR0). Its covariance is positive semidefinite and all 92 marginal variances are positive, so no covariance adjustment is used. Model-based observed-information standard errors and a finite-cluster correction (CR1) are retained as supplementary diagnostics.⁸ Inference conditions on the calibrated hazards and discount rate, the activity

⁸The model-based covariance is H^{-1} . The CR1 sensitivity multiplies each of the three cluster sums by $[G/(G-1)][(N-1)/(N-92)]$, using its own number of groups G . Here $N = 2,730,384$ doctor–post pairs, with 1,132 doctor groups, 2,412 post groups, and N intersection groups. The scores are not demeaned.

proxies, observed covariates and market composition, and the maintained equilibrium branch. It does not incorporate uncertainty in those calibrations or allow arbitrary dependence across pairs sharing neither a doctor nor a post.

Probability changes and approximate standard errors. Table 2 translates selected coefficients into average probability changes with continuation values and other covariates held fixed. For acceptance, let $F(v) = \Lambda(v)$ and let $\hat{\eta}_{ij}$ be the fitted acceptance index for the relevant route and side. For contact, let $F(v) = 1 - \exp\{-\exp(v)\}$ and $\hat{\eta}_{ij} = \log m_{ij}^c$, including the observed $K_i T_j$ offset. For a standardized continuous covariate k , set $v_{ij,k} = \hat{\eta}_{ij}$; for an indicator, set $v_{ij,k} = \hat{\eta}_{ij} - \hat{\beta}_k z_{ij,k}$ to obtain the index when that indicator is zero. With averaging weights w_{ij} summing to one, define

$$g_k(b) = 100 \sum_{i,j} w_{ij} \{F(v_{ij,k} + b) - F(v_{ij,k})\}.$$

The reported probability change is $g_k(\hat{\beta}_k)$. Contact comparisons weight all retained doctor–post pairs equally; acceptance comparisons weight the observed decisions in the relevant route and side. The two relationship-history comparisons restrict these samples to pairs with a positive approach count in the preceding 365 days, keeping the original estimation-sample standardization.

For a simple uncertainty summary, I apply the delta method to each coefficient separately, treating the starting indices $v_{ij,k}$ and the weights as fixed:

$$\widehat{\text{se}}\{g_k(\hat{\beta}_k)\} \simeq 100 \left| \sum_{i,j} w_{ij} F'(v_{ij,k} + \hat{\beta}_k) \right| \widehat{\text{se}}(\hat{\beta}_k).$$

Here $F'(v) = \Lambda(v)\{1 - \Lambda(v)\}$ for acceptance and $F'(v) = \exp(v)\exp\{-\exp(v)\}$ for contact. The coefficient standard errors on the right-hand side are the doctor–post clustered standard errors derived above. This approximation propagates uncertainty in the coefficient being translated, but omits uncertainty in the starting predictions and covariance with other estimated parameters.

B.5 Covariate Coefficient Estimates

Table B.4 reports the selected index coefficients underlying Table 2.

The full tables report the 90 active coefficients in the four contact and acceptance indices, including their intercepts, and a separate panel reports the two initiator-specific acceptance intercept differences. Standard errors appear in parentheses beneath the estimates. Continuous covariates are standardized, while indicators retain their zero-to-one scale. Fourteen unsupported coordinates are held at zero and shown as dashes, rather than as estimated zeros. The exposure loadings on continuation values are fixed normalizations rather than estimated coefficients.

Table B.4. Selected Parameter Estimates

	Sorting	Screening	Doctor acc.	Post acc.
Log distance	-0.4190 (0.0442)	-0.1945 (0.0761)	0.1002 (0.1275)	0.0063 (0.0655)
Hours	0.0311 (0.0674)	0.4596 (0.3018)	-0.3822 (0.1918)	0.0896 (0.3526)
Log posted salary	-0.1151 (0.0525)	-2.1926 (0.5228)	-0.0877 (0.1780)	-2.0443 (0.4999)
Preference match	0.2417 (0.0851)	0.1929 (0.1519)	0.1489 (0.2653)	-0.0437 (0.1418)
Day match	0.9223 (0.1291)	0.3076 (0.1351)	0.8349 (0.2472)	0.2153 (0.1305)
Specialty match	0.6213 (0.0731)	0.9296 (0.1823)	-0.3856 (0.2882)	0.3381 (0.1841)
Prior approaches (past 365 days)	0.0448 (0.0100)	0.0485 (0.0076)	0.0552 (0.0226)	-0.0170 (0.0122)
Prior contract	0.8902 (0.1075)	2.9306 (0.2756)	-1.0797 (0.6210)	1.3925 (0.2941)

Notes: Entries are raw index coefficients. Continuous covariates are standardized; indicator coefficients correspond to a change from zero to one. Analytic standard errors in parentheses are clustered by doctor and post (two-way CR0). Scores include the response of continuation values through the equilibrium conditions. Inference conditions on the maintained equilibrium branch, calibrated hazards and discount rate, the observed market composition, and activity proxies. Calibration uncertainty is not included.

Table B.5. Parameter Estimates: Doctor acceptance

Covariate	Estimate (SE)
Intercept	17.4105 (1.9456)
Log distance	0.1002 (0.1275)
Doctor age	4.0249 (4.0325)
Experience	-2.3576 (3.7282)
Hours	-0.3822 (0.1918)
Log posted salary	-0.0877 (0.1780)
Preference match	0.1489 (0.2653)
Day match	0.8349 (0.2472)
Specialty match	-0.3856 (0.2882)
Outpatient match	0.1075 (0.2991)
Emergency match	—
Home-visit match	0.2504 (0.7098)
Ward match	-0.8459 (0.3551)
Dialysis match	-0.3295 (0.6658)
Health-check match	0.3603 (0.5175)
Endoscopy match	0.2761 (0.6401)
Surgery match	—
House-call match	-0.5780 (0.6325)
Work-with-peer match	—
Flexible-work match	-0.3642 (0.8010)
Occupational-health match	—
Prior approaches (past 365 days)	0.0552 (0.0226)
Prior contract	-1.0797 (0.6210)
Days since last approach	-0.1643 (0.0606)
Prior agency share (past 365 days)	0.0021 (0.0288)
No prior relationship	-0.9899 (0.6278)

Notes: Entries are raw index coefficients; continuous covariates are standardized and indicators are not. Dashes denote fixed unsupported coordinates, not estimated zeros. Exposure loadings on continuation values are normalized to one and are not estimated rows. Analytic standard errors in parentheses are clustered by doctor and post (two-way CR0). Scores include the response of continuation values through the equilibrium conditions. Inference conditions on the maintained equilibrium branch, calibrated hazards and discount rate, the observed market composition, and activity proxies. Calibration uncertainty is not included.

Table B.6. Parameter Estimates: Post acceptance

Covariate	Estimate (SE)
Intercept	6.8671 (0.6760)
Log distance	0.0063 (0.0655)
Doctor age	-0.1571 (0.1912)
Experience	-0.1059 (0.1870)
Hours	0.0896 (0.3526)
Log posted salary	-2.0443 (0.4999)
Preference match	-0.0437 (0.1418)
Day match	0.2153 (0.1305)
Specialty match	0.3381 (0.1841)
Outpatient match	0.0064 (0.1832)
Emergency match	—
Home-visit match	-1.0348 (0.4614)
Ward match	-0.0805 (0.2175)
Dialysis match	0.2633 (0.6531)
Health-check match	-0.0768 (0.3434)
Endoscopy match	0.7435 (0.4932)
Surgery match	—
House-call match	-1.1422 (0.4684)
Work-with-peer match	0.0096 (0.5326)
Flexible-work match	-2.7982 (0.4949)
Occupational-health match	—
Prior approaches (past 365 days)	-0.0170 (0.0122)
Prior contract	1.3925 (0.2941)
Days since last approach	-0.0879 (0.0460)
Prior agency share (past 365 days)	-0.0417 (0.0242)
No prior relationship	-0.2559 (0.2999)

Notes: Entries are raw index coefficients; continuous covariates are standardized and indicators are not. Dashes denote fixed unsupported coordinates, not estimated zeros. Exposure loadings on continuation values are normalized to one and are not estimated rows. Analytic standard errors in parentheses are clustered by doctor and post (two-way CR0). Scores include the response of continuation values through the equilibrium conditions. Inference conditions on the maintained equilibrium branch, calibrated hazards and discount rate, the observed market composition, and activity proxies. Calibration uncertainty is not included.

Table B.7. Parameter Estimates: Sorting exposure

Covariate	Estimate (SE)
Intercept	10.8022 (1.7046)
Log distance	-0.4190 (0.0442)
Doctor age	2.5331 (3.8711)
Experience	-0.9331 (3.5753)
Hours	0.0311 (0.0674)
Log posted salary	-0.1151 (0.0525)
Preference match	0.2417 (0.0851)
Day match	0.9223 (0.1291)
Specialty match	0.6213 (0.0731)
Outpatient match	0.3790 (0.0940)
Emergency match	—
Home-visit match	1.1057 (0.2269)
Ward match	0.3420 (0.1190)
Dialysis match	2.3813 (0.2438)
Health-check match	0.5623 (0.1108)
Endoscopy match	1.5732 (0.3521)
Surgery match	—
House-call match	0.2291 (0.4071)
Work-with-peer match	1.5811 (0.5591)
Flexible-work match	1.8620 (0.2457)
Occupational-health match	—
Prior approaches (past 365 days)	0.0448 (0.0100)
Prior contract	0.8902 (0.1075)
Days since last approach	-0.2692 (0.0293)
Prior agency share (past 365 days)	-0.0475 (0.0119)
No prior relationship	-2.5000 (0.1253)

Notes: Entries are raw index coefficients; continuous covariates are standardized and indicators are not. Dashes denote fixed unsupported coordinates, not estimated zeros. Exposure loadings on continuation values are normalized to one and are not estimated rows. Analytic standard errors in parentheses are clustered by doctor and post (two-way CR0). Scores include the response of continuation values through the equilibrium conditions. Inference conditions on the maintained equilibrium branch, calibrated hazards and discount rate, the observed market composition, and activity proxies. Calibration uncertainty is not included.

Table B.8. Parameter Estimates: Screening exposure

Covariate	Estimate (SE)
Intercept	-2.2979 (0.7053)
Log distance	-0.1945 (0.0761)
Doctor age	-0.5681 (0.2376)
Experience	0.4081 (0.2374)
Hours	0.4596 (0.3018)
Log posted salary	-2.1926 (0.5228)
Preference match	0.1929 (0.1519)
Day match	0.3076 (0.1351)
Specialty match	0.9296 (0.1823)
Outpatient match	-0.0060 (0.1919)
Emergency match	—
Home-visit match	-0.3495 (0.5036)
Ward match	-0.0708 (0.1668)
Dialysis match	2.6438 (0.5978)
Health-check match	0.3370 (0.3002)
Endoscopy match	1.3536 (0.4886)
Surgery match	—
House-call match	0.4815 (0.4785)
Work-with-peer match	—
Flexible-work match	1.7356 (0.3958)
Occupational-health match	—
Prior approaches (past 365 days)	0.0485 (0.0076)
Prior contract	2.9306 (0.2756)
Days since last approach	-0.2565 (0.0755)
Prior agency share (past 365 days)	0.1412 (0.0124)
No prior relationship	-1.3779 (0.3390)

Notes: Entries are raw index coefficients; continuous covariates are standardized and indicators are not. Dashes denote fixed unsupported coordinates, not estimated zeros. Exposure loadings on continuation values are normalized to one and are not estimated rows. Analytic standard errors in parentheses are clustered by doctor and post (two-way CR0). Scores include the response of continuation values through the equilibrium conditions. Inference conditions on the maintained equilibrium branch, calibrated hazards and discount rate, the observed market composition, and activity proxies. Calibration uncertainty is not included.

Table B.9. Initiator-Specific Acceptance Intercept Differences

Route-specific parameter	Estimate (SE)
Doctor-initiated acceptance intercept difference	5.8706 (0.3758)
Agent-initiated acceptance intercept difference	3.8459 (0.2206)

Notes: These two parameters are active components of the 92-coordinate structural vector. Analytic standard errors in parentheses are clustered by doctor and post (two-way CR0). Scores include the response of continuation values through the equilibrium conditions. Inference conditions on the maintained equilibrium branch, calibrated hazards and discount rate, the observed market composition, and activity proxies. Calibration uncertainty is not included.

B.6 Numerical Diagnostics

Table B.10. Estimation and Analytic Inference Diagnostics

Diagnostic	Value
Active / fixed coordinates	92 / 14
Total log likelihood	-28,664.909007
Equilibrium residual (maximum absolute value)	2.487e-13
Remaining normalized Newton correction (maximum)	2.234e-05
Local quadratic predicted log-likelihood gain	1.553e-10
Practical local numerical criteria	Satisfied
Doctor / post clusters	1,132 / 2,412
Doctor-post intersections	2,730,384
Primary analytic covariance	Two-way CR0
Primary covariance positive semidefinite	Yes

Notes: The remaining correction uses the full equilibrium-reduced Hessian without regularization. The likelihood gain is a local quadratic prediction, not a bound on global improvement. Numerical convergence does not establish a global optimum or equilibrium uniqueness. Analytic standard errors in parentheses are clustered by doctor and post (two-way CR0). Scores include the response of continuation values through the equilibrium conditions. Inference conditions on the maintained equilibrium branch, calibrated hazards and discount rate, the observed market composition, and activity proxies. Calibration uncertainty is not included.

C Additional Policy Results

Figure C.1 displays the access–value comparison with doctor search retained, and Figure C.5 reports the market-expansion comparison. In both exercises, every long-run candidate recomputes doctor search from the solved continuation values.

C.1 Who Gains from the Policy Changes?

Table C.1 separates the value changes on the two sides of the market for the current-budget comparisons and reports how broadly each side’s gains are shared. Panel A of Table 5 reports the corresponding results for recommendation expansion, equalization across doctors, and pairwise redesign. In both tables, doctor-side value is $\sum_i K_i \alpha_i^D$, so the reported winning share weights doctors by their numbers of slots. Post-side value is $\sum_j \alpha_j^P$, and each post receives equal weight. A unit is classified as gaining or losing when its continuation value changes by more than 10^{-4} in the corresponding direction; smaller changes are classified as unchanged. These shares describe the distribution of the model’s predicted value changes and are not sampling probabilities.

For the group comparisons in the text, I rank posts by their continuation values under the current hybrid and divide them into five approximately equal-sized groups using their cumulative stock weights. Group membership remains fixed across every policy comparison. Each reported group percentage is the change in that group’s total value relative to its value under the preceding policy in the comparison. Thus, the expanded-budget redesign comparison uses the doubled, equal-budget profile as its denominator while retaining the groups defined under the current hybrid. This distinguishes initially low-value posts from posts selected after observing the policy gain.

Table C.1. Who Gains from Current-Budget Recommendation Policies?

Policy comparison	Change in continuation value			Share with higher value (%)	
	Doctors	Posts	Total	Doctor slots	Posts
Current to search-retained redesign	−155.04	+1,218.64	+1,063.60	31.36	45.85
Current to full replacement	−8,823.99	+8,003.08	−820.91	1.08	100.00

Notes: Each row compares the final allocation with the current hybrid. Doctor-side value weights each doctor by the number of doctor slots; post-side value sums over posts. Share with higher value is the percentage of doctor slots or posts whose continuation value exceeds its current-hybrid value by more than 10^{-4} . Doctors are weighted by their numbers of slots; posts receive equal weight. These are accounting comparisons conditional on the estimated model. They do not imply that every unit on a gaining side benefits.

C.2 Continuation Value and Accepted Matches

The number of matches alone does not determine the value of a recommendation policy. An accepted match generates a gain above continued waiting for each side, and the continuation-value equations aggregate these gains after accounting for discounting and exogenous exit. The following accounting separates changes in the number of matches from changes in who matches and the gains obtained in those matches. All quantities are evaluated at each policy's equilibrium, rather than at the fixed continuation values used in the short-run objective.

Let $m_{ijc} = M_{ij}^c P_{ijc}^D P_{jic}^P$ be the expected number of accepted matches for doctor–post pair (i, j) through route c . Independence of the two payoff shocks implies that each side's expected gain above continued waiting, conditional on a match, is

$$g_{ijc}^D = \frac{q_{ijc}^D}{P_{ijc}^D} = \mathbb{E}[U_{ij}^c - \alpha_i^D \mid U_{ij}^c \geq \alpha_i^D],$$

$$g_{ijc}^P = \frac{q_{ijc}^P}{P_{jic}^P} = \mathbb{E}[V_{ji}^c - \alpha_j^P \mid V_{ji}^c \geq \alpha_j^P].$$

Multiplying the doctor equation in (5) by K_i and summing both sides of the market gives

$$W = \sum_{i,j,c} m_{ijc} \left(\frac{g_{ijc}^D}{r + \delta_i^a} + \frac{g_{ijc}^P}{r + w_j} \right).$$

Thus, a match contributes more to this accounting when its conditional gain above continued waiting is larger. The gains are measured in the model's utility units; they are neither wages nor gross match payoffs. The calibration sets $r = -\log(0.99)$, $\delta_i^a = 0.05$ for every doctor, and $w_j = 0.10$ for every post. These rates are fixed across policies and constant within each side, so changes in match composition do not operate through differences in exit rates within a side. The different rates across sides do matter: an equal surplus flow receives a larger continuation-value weight on the doctor side.

To distinguish more matches from different matches, let $Q = \sum_{i,j,c} m_{ijc}$, $s_{ijc} = m_{ijc}/Q$, and $v_{ijc} = g_{ijc}^D/(r + \delta_i^a) + g_{ijc}^P/(r + w_j)$. Then $W = Q\mu$, where $\mu = \sum_{i,j,c} s_{ijc} v_{ijc}$ is the average contribution per accepted match. For any two policies, write Δ for the after-minus-before difference and an overbar for the mean of the two endpoint values. Applying the symmetric two-factor identity first to $Q\mu$ and then to each $s_{ijc} v_{ijc}$ gives

$$\Delta W = \underbrace{\Delta Q \bar{\mu}}_{\text{match number}} + \underbrace{\bar{Q} \sum_{i,j,c} \Delta s_{ijc} \bar{v}_{ijc}}_{\text{match composition}} + \underbrace{\bar{Q} \sum_{i,j,c} \bar{s}_{ijc} \Delta v_{ijc}}_{\text{conditional gains}}. \quad (10)$$

The second term changes which doctor–post pairs and routes generate the matches, evaluating their contributions at the average of the two endpoint values. The third changes the conditional gains at a fixed average match composition. This is an exact accounting identity at equilibrium, not a sequence of independently implementable policy interventions. Conditional gains

Table C.2. Match Counts, Match Composition, and Continuation Value

Comparison	Match change	Value change	Contributions to value change		
			Match number	Composition	Conditional gains
Current to search-retained redesign	+100.77	+1,063.60	+7,987.97	−5,123.50	−1,800.87
Short-run to long-run rule	+81.70	+1,295.54	+6,413.12	−5,122.38	+4.80
Current to full replacement	+692.02	−820.91	+46,014.76	−60,438.67	+13,603.00
Redesign at doubled equal budgets	+167.08	+1,717.40	+11,628.29	−3,783.03	−6,127.85

Notes: Values are in model utility units. The last three columns use the symmetric accounting in equation (10): match number, the distribution of accepted matches across doctor–post pairs and routes, and conditional gains above continued waiting at each pair and route. Their sum reproduces the change in the equilibrium flow representation of value. Before rounding, the largest difference from the saved continuation-value change is below 0.003, due to the equilibrium tolerance used for the fixed comparison policies. Components are accounting terms, not effects of independently changing one behavioral channel.

themselves subtract policy-specific continuation values, so the last term is not an estimate of a change in gross match quality.

Table C.2 reports these terms. Components are rounded independently, so the displayed numbers need not add exactly; the reported numerical residuals refer to the unrounded calculations.

The search-retained long-run rules increase both matches and value relative to their respective current, short-run, or equal-budget comparison policies. In each comparison, the positive match-number contribution exceeds the combined change from match composition and conditional gains. Full replacement, by contrast, generates more matches but lower total value than the current hybrid: its negative match-composition contribution more than offsets the positive match-number and conditional-gain contributions. Thus, the number of matches alone does not determine the value ranking; which doctor–post pairs match and through which routes also matters.

The contact and acceptance margins explain how the expanded-budget redesign can reduce doctor search while preserving doctor-initiated matches. For each pair and route, the match count is the product of contacts, doctor acceptance, and post acceptance. I allocate its change to these three factors by replacing them in every possible order and averaging the marginal contributions over the six orders. Within doctor search, changed contact counts contribute −211.86 matches, changes in post acceptance contribute 215.58, and changes in doctor acceptance contribute −3.09. The near cancellation is consistent with the almost unchanged doctor-initiated match count in the text. These terms compare equilibrium endpoints; holding only one factor fixed need not yield another equilibrium and does not identify the causal effect of separately changing that factor.

Table C.3. Numerical Checks for the Long-Run Policies

Policy class	Doctor budgets	Total value	Relative gap	Residual	Audit
Search retained	Current	95,478.79	9.07e-07	1.36e-14	Pass
Search retained	Doubled, equal	97,166.10	9.86e-07	2.67e-14	Pass
Full replacement	Current	93,594.29	8.90e-07	2.92e-15	Pass
Full replacement	Doubled, equal	94,188.29	9.97e-07	4.23e-14	Pass

Notes: Each policy satisfies a relative Frank–Wolfe gap of at most 10^{-6} on its fixed feasible set. The gap is the maximum directional derivative toward another feasible policy, divided by total value. Equilibrium residuals are at most 10^{-10} and relative adjoint residuals at most 10^{-9} . Independent audits check equilibrium replay, a separately assembled gradient, and feasible-direction finite differences. These checks certify approximate first-order stationarity, not global optimality or equilibrium uniqueness. Fixed comparison policies and market-expansion grid equilibria use an equilibrium tolerance of 10^{-5} .

C.3 Numerical Policy Search

The long-run searches stop when no feasible change in the recommendation matrix offers a sufficiently large first-order improvement in total user value. For policy class k and its fixed budget vector b , let $g^k(A) = \nabla_A W^k(A; K)$ include the response of the equilibrium to A . Define

$$G^k(A) = \max_{B \in \mathcal{B}^k(b)} \sum_{i,j} g_{ij}^k(A)(B_{ij} - A_{ij}).$$

The maximization ranges over the complete feasible policy set, including pairs that currently receive no recommendations. I stop when $G^k(A)/\max\{1, |W^k(A; K)|\} \leq 10^{-6}$ and require a maximum equilibrium residual no greater than 10^{-10} . All four reported long-run allocations satisfy these criteria, as shown in Table C.3. The gradient is obtained from the equilibrium equations; its adjoint solve must have relative residual no greater than 10^{-9} .

Independent checks reproduce each allocation’s value, feasibility, and first-order gap, and compare the gradient with finite differences along feasible directions. Initializing the equilibrium solver at either the reported continuation values or those of a different policy reproduces the same values. These checks support approximate first-order stationarity on the maintained equilibrium branch. They do not establish a global policy optimum, a bound on the remaining global value gain, or equilibrium uniqueness. The fixed comparison policies and the market-expansion grid are solved to an equilibrium tolerance of 10^{-5} ; the tighter residual criterion applies to the four optimized long-run allocations.

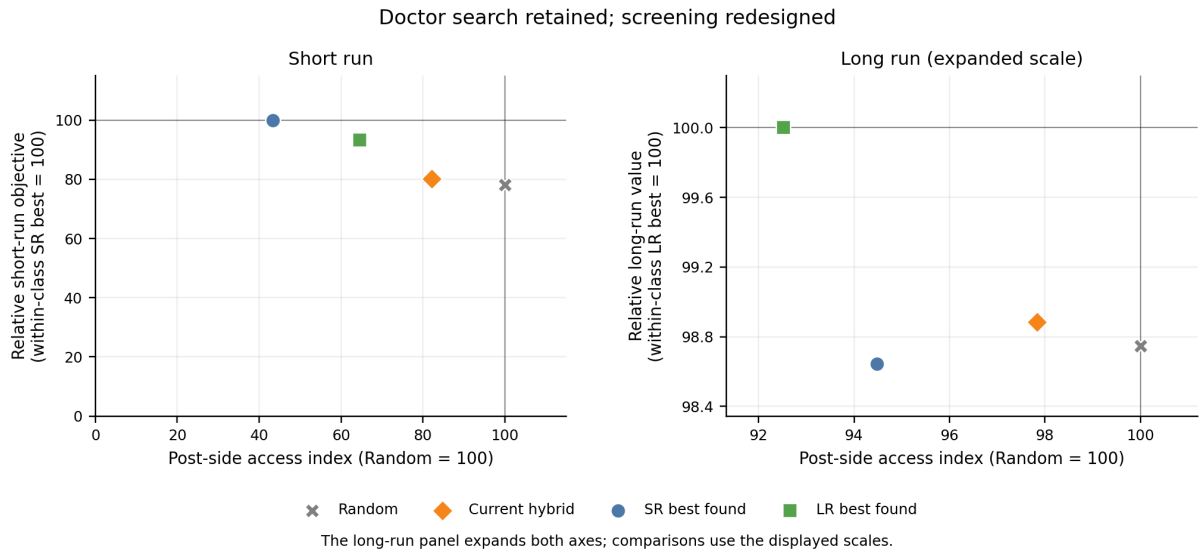


Figure C.1. Access and value when agent screening is redesigned and doctor search responds in equilibrium.

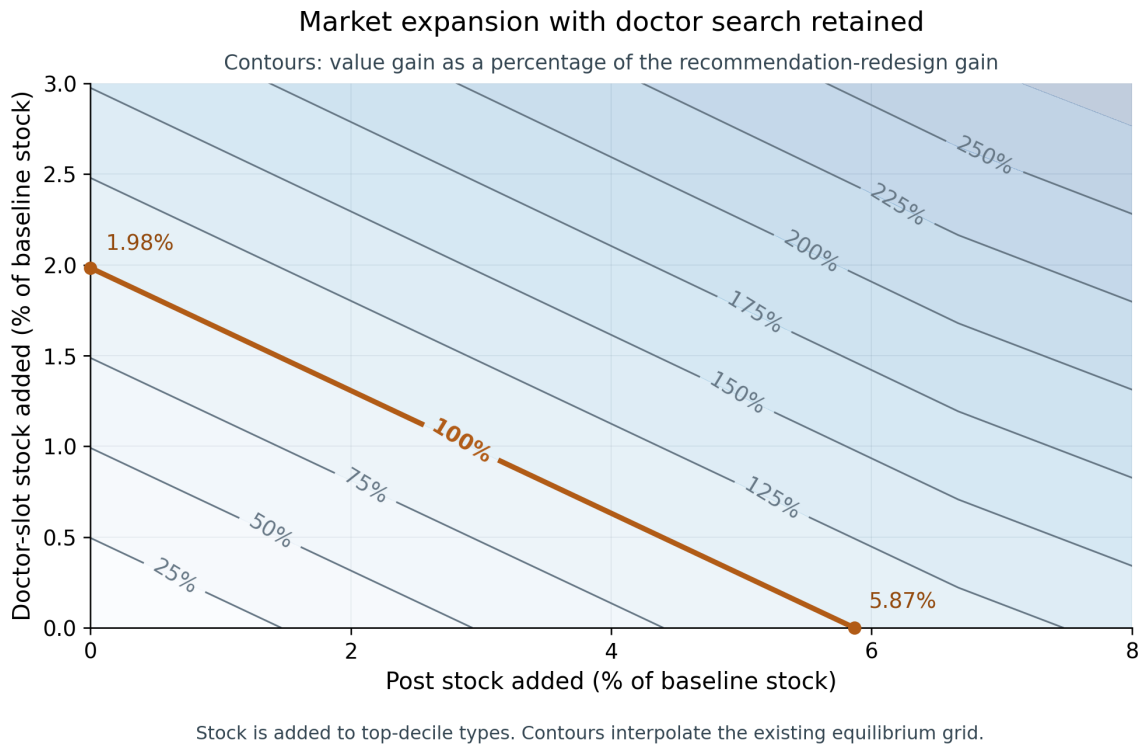


Figure C.5. Market-composition equivalents of the equal-budget redesign increment with doctor search retained. Contours show the value gain as a percentage of the redesign increment; the highlighted 100 percent contour matches that increment. All contours interpolate the evaluated equilibrium grid.