

Interview hoarding

VIKRAM MANJUNATH

Department of Economics, University of Ottawa

THAYER MORRILL

Department of Economics, North Carolina State University

Many centralized matching markets are preceded by interviews between participants, including the residency matches between doctors and hospitals. Due to the COVID-19 pandemic, interviews in the National Resident Matching Program were switched to a virtual format, which resulted in a dramatic and asymmetric decrease in the cost of accepting interview invitations. We study the impact of an increase in the number of doctors' interviews on their final matches. We show analytically that if doctors can accept more interviews, but hospitals do not increase the number of interviews they offer, then no doctor who would have matched in the setting with more limited interviews is better off and many doctors are potentially harmed. This adverse effect is the result of what we call *interview hoarding*. We characterize optimal mitigation strategies for special cases and use simulations to extend these insights to more general settings.

KEYWORDS. NRMP, deferred acceptance, interviews, hoarding.

JEL CLASSIFICATION. C78, D47, J44.

1. INTRODUCTION

Perhaps the most well known application of matching theory is the entry-level labor market for physicians. In 2021, 37,470 positions were matched through the National Resident Matching Program (NRMP). The matching process consists of two steps. First, each physician interviews with a set of residency programs. Second, programs and physicians submit rank-ordered lists of those they interview to a centralized clearinghouse. This clearinghouse, run by the NRMP, matches physicians to residency programs using a version of Gale and Shapley's (1962) deferred acceptance (DA) algorithm (Roth and Peranson (1999)).

Vikram Manjunath: vikram@dosamobile.com

Thayer Morrill: thayer_morrill@ncsu.edu

We thank Anna Sorensen and Alkas Baybas for raising the question that sparked this paper. We also thank Alex Chan, Adrienne Quirouet, Assaf Romm, Al Roth, Erling Skancke, Colin Sullivan, William Thomson, and seminar audiences at North Carolina State University, Stanford, the University of Arizona, and the University of Lausanne for helpful comments and discussions. We are grateful for the feedback from two anonymous referees, who helped us improve the paper substantially. Manjunath gratefully acknowledges the support of the Social Sciences and Humanities Research Council of Canada through Grant 435-2019-0110.

Programs and applicants are both constrained in the number of interviews they can take part in. Prior to the COVID-19 pandemic, interviews were conducted in person. These interviews were particularly costly for physicians, who not only had to bear travel expenses, but also had to take days off from clinical rotations. The cost to programs was mainly in terms of time. For the 2020–2021 matching season, interviews were conducted virtually. While this dramatically decreased the cost of interviews for physicians, it did not substantially change the costs for programs. We are interested in the implications of this asymmetric change on the eventual match.

Intuitively, it seems possible that a doctor might receive a better match if she accepts more interviews. However, we show a surprising result: as long as she would have been matched with a program under the more constrained number of interviews, an increase in her interview capacity does not improve her match. Further, we show that even if she would not have been matched, she will only benefit from an increase in interviews if she would also have been part of a doctor–hospital pair that blocked the original matching.¹ The match rate for U.S. medical school graduates is typically around 94%.² Our results thus suggest that virtually all doctors would not benefit, and could in fact be harmed, by an increase in the number of interviews in which they participate.

Increasing the number of interviews a doctor accepts has a negative externality: these interviews can no longer be allocated to other doctors. As an illustration, consider a highly sought-after physician, one who is offered interviews at the leading programs and will end up matched with her favorite program. If interviews become cheaper, she accepts more interviews.³ However, the additional interviews she accepts are from less-preferred programs, and these interviews do not help her, as she ultimately matches with her favorite program, the same one she would have matched with given fewer interviews. Thus, her additional interviews are, in effect, wasted. We refer to this as *interview hoarding*. Interview hoarding has a cascading effect. The physicians who otherwise would have filled these wasted interview slots now must interview with programs they consider inferior. These physicians may have more interviews, but they do not have better interviews in a precise sense: they rate every new interview as worse than their previous match.

This implies a striking result. If there is an increase in doctors' interview capacities, but programs do not react, this increase causes the ultimate match to be (Pareto) worse from the matched physicians' perspective.⁴ These doctors fall into three categories: those who hoard interviews that are worse than their eventual match, those who receive more but worse interviews, and those who receive fewer and worse interviews. The first category is indifferent between the new outcome and the old. The latter two categories are harmed. Even among unmatched physicians, only those who fall through the

¹ While the NRMP match is stable with respect to submitted preferences over those one interviews, it may not be stable with respect to actual preferences over all possible matching partners.

² Specifically, the match rates in 2017 through 2020 were 94.3%, 94.3%, 93.9%, and 93.7%, respectively (see the NRMP's "Results and Data: 2021 Main Residency Match" (https://mk0nrmp3oyqui6wqfm.kinstacdn.com/wp-content/uploads/2021/05/MRM-Results_and-Data_2021.pdf)).

³ While we assume that agents have complete information, if there were an arbitrarily small amount of uncertainty about others' preferences, she would accept additional costless interviews.

⁴ There need not be a Pareto ranking from the programs' perspective; see the example in Section 1.2.

cracks—meaning that they are unmatched but would be welcomed by some programs—could potentially benefit from an increase in interviews. Given a typical match rate of 94%, this means fewer than 6% of doctors could possibly gain, while the overwhelming majority could be harmed.

We are not suggesting that there is no benefit to an increase in physicians' interviews, and we recognize that two of our assumptions are unrealistic: that doctors and programs know their preferences perfectly, and that agents, faced with a constraint on the number of interviews they accept, are not strategic.⁵ In reality, of course, neither doctors nor programs perfectly know their preferences. The point of interviews is to learn more about a candidate or program. Similarly, both doctors and programs are strategic in the interviews they accept. We expect a doctor to accept an interview with a “safety” program, one she is sure to match with, while lower-ranked programs likely do not invite the very best candidates for interviews so that they do not “waste” their slots. By having more interviews, a doctor learns about more programs and has more flexibility to strategize.⁶ However, our results show that these are the only two channels through which doctors may gain from an increase in the number of interviews they accept. The advantage of our modeling choices is that they allow us to identify a subtle bottleneck created by interviews that would likely be lost in the analysis of a more complex model.

Having shown that increases in doctors' interview capacities have adverse welfare consequences, we turn to mitigation. We consider policies that restrict the number of interviews that programs can offer and that candidates can accept. Though there are essentially no such policies that *always* (for every preference profile) yield a stable final matching (Proposition 1), we characterize such policies for “common preferences” (Proposition 2). These are salient preference profiles in which every doctor ranks the programs the same way and every program ranks the doctors the same way. The policies we characterize place a common cap on the number of interviews any program can offer and any candidate can accept. We also show that if the programs' interview capacities are fixed, say at l , then the number of blocking pairs increases and the match rate decreases as the doctors' interview cap moves further away from l in either direction (Proposition 3).

Our analytical results can inform policies for more general settings in which preferences are not quite common, but have a common component. In Section 7, we use simulations to show that the lessons from our analytical results hold up under weaker assumptions. Our results provide evidence that the optimal policy is for doctors and programs to have the same interview capacities.⁷

⁵We follow the approach of Echenique, Gonzalez, Wilson, and Yariv (2020) in assuming complete information about preferences and nonstrategic offers and acceptance of interviews.

⁶If a doctor was previously unmatched but formed a blocking pair with a hospital, then she must have declined that hospital's interview. Ex post, this was a strategic mistake. We interpret the doctor's benefit from increasing her interviews and being matched as a strategic benefit.

⁷This is true whether we define the optimality of a policy as maximizing the expected proportion of positions that are filled or minimizing the expected number of blocking pairs. In fact, these objectives are equivalent.

1.1 *Related literature*

While there is a large literature on the post-interview NRMP match,⁸ relatively few papers incorporate the pre-match interview process. One of the first to explicitly model interviews in the classic one-to-one matching model is that of [Lee and Schwarz \(2017\)](#). In their model, before participating in a centralized, two-sided match, firms learn their preferences over workers by engaging in costly interviews. They show that even if firms and workers interview with exactly the same numbers of agents, the extent of unemployment in the final match depends critically on the overlap between the sets of workers that firms interview.

Three other recent papers that incorporate pre-match interviews are those by [Kadam \(2021\)](#), [Beyhaghi \(2019\)](#), and [Echenique et al. \(2020\)](#). As in our paper, [Kadam \(2021\)](#) considers the implications of loosened interview constraints for doctors. However, the focus is on the strategic allocation of scarce interview slots. For the sake of tractability, his analysis features a stylized model of large markets. Under the assumption of common preferences over programs, Kadam shows that increasing doctors' capacities may increase total surplus, but not in a Pareto-improving way. Moreover, the match rate decreases. Kadam also highlights that when preferences are not necessarily common, the effect is ambiguous, since increased interview capacities dilute doctors' signaling ability.

[Beyhaghi \(2019\)](#) also performs a strategic analysis of a stylized large market model. However, she considers a slightly different setup with *application* caps for doctors and interview caps for programs. While similar, application caps are not the same as interview caps: they constrain the number of programs a doctor can express interest in at the outset of the interview matching phase, but not the number of interviews she can accept at the end. In Beyhaghi's model, inequity in the application caps decreases the expected total surplus. Moreover, when interview capacity is low, low application caps are socially desirable.

In our model, agents do not choose interviews strategically. Determining the optimal set of interviews is closely related to the portfolio choice problems in [Chade and Smith \(2006\)](#) and [Ali and Shorrer \(2021\)](#). Both of these studies solve for the optimal portfolio when an agent chooses among costly, stochastic options but only consumes one of the realizations. To apply the optimal solution to the interview scheduling problem, one would need to know precisely the probability of any given pair matching. However, this probability depends not only on the agents' preferences, but also on their strategies. Solving for equilibria when agents choose their interviews optimally is intractable without severe simplifying assumptions (such as those in the papers mentioned above).

Like the study by [Echenique et al. \(2020\)](#), which is methodologically closest to ours, we sidestep this issue. They explain a puzzling empirical pattern resulting from the NRMP match: 46.3% of the physicians were matched to their top-ranked residency programs, and 71.1% were matched to a program they ranked in their top three. These statistics seem to contradict surveys indicating that many doctors have similar preferences over residency programs. In explaining this phenomenon, Echenique et al.

⁸See the multitude of papers following [Roth and Peranson \(1999\)](#).

highlight the importance of the interviewing process that precedes the match. Roughly speaking, the pre-match interviewing process restricts the preferences that the physicians actually submit to the NRMP. Therefore, a proper interpretation is not that the physicians matched with their most-preferred programs, but rather that they matched with the programs they most-preferred *among those with which they interviewed*.

Our work complements these papers by providing further evidence of the importance of understanding the pre-match interviews to properly evaluate the NRMP match itself.

1.2 Motivating example

We present the intuition behind the welfare loss from doctors' increased interview capacity with a simple example. Consider a market with five doctors $\{d_1, \dots, d_5\}$ and four hospitals $\{h_1, \dots, h_4\}$. Each doctor and hospital seeks one successful match. Their preferences are

d_1	d_2	d_3	d_4	d_5	h_1	h_2	h_3	h_4
h_1	h_1	h_1	h_3	h_4	d_1	d_1	d_1	d_1
h_4	h_2	h_2	h_4	h_3	d_2	d_2	d_4	d_5
h_3	h_4	h_3	h_2	h_1	d_3	d_4	d_5	d_3
h_2	h_3	h_4	h_1	h_2	d_5	d_3	d_2	d_2
					d_4	d_5	d_3	d_4

Suppose that the interview capacities of the doctors and hospitals are

d_1	d_2	d_3	d_4	d_5	h_1	h_2	h_3
2	1	1	1	1	2	1	1

Interviews are initially offered by hospitals. In the first round, h_1 invites d_1 and d_2 , and h_2, h_3 , and h_4 all invite d_1 . Because d_1 can accept only two invitations, she declines those from h_2 and h_3 . Hospital h_2 then offers an interview to d_2 and h_3 invites d_4 . Since d_2 can accept only one interview, she declines h_2 's invitation. Then h_2 invites and is turned down by d_4 . Finally, h_2 invites d_3 , who accepts the invitation. The final interviews are

d_1	d_2	d_3	d_4	d_5
$\{h_1, h_4\}$	$\{h_1\}$	$\{h_2\}$	$\{h_3\}$	$\emptyset$

The final matching is computed by applying the doctor-proposing DA algorithm to the agents' preferences restricted to agents they interview with. The outcome is

d_1	d_2	d_3	d_4	d_5
h_1		h_2	h_3	

Now suppose each doctor can accept one more invitation, but the hospitals' interview capacities remain the same. In this case, d_1 does not reject h_3 's invitation in the first

round of interview invitations. Similarly, d_2 does not reject h_2 's invitation in the second round. The interview schedule is

d_1	d_2	d_3	d_4	d_5
$\{h_1, h_3, h_4\}$	$\{h_1, h_2\}$	$\emptyset$	$\emptyset$	$\emptyset$

This leads to the final matching:

d_1	d_2	d_3	d_4	d_5
h_1	h_2			

Despite being able to accept more interviews, doctors d_3 and d_4 are worse off and d_5 is no better. The only doctor who gains is d_2 .⁹ We make the following observations about the winners and losers from the increased interview capacity.

- Doctor d_2 , the only doctor who gained, was originally unmatched.
- The original matching was not stable. Each of d_2 and d_5 was part of a blocking pair (d_2 blocked with h_2 and d_5 blocked with h_4). Specifically, the only doctor who gained from an increase in interview capacity was among the doctors who blocked the original match.
- Despite being part of a pair that blocked the original matching, d_5 did not regret turning down any interview invitations, while d_2 did. By “regret” we mean the doctor turned down an interview invitation from a hospital that she preferred to her final matching. In particular, the only doctor who gained *did* previously reject an invitation that she ended up regretting.

These observations are not specific to this example. In Theorem 1, we show that in general a doctor can only gain from an increase in interview capacities if she was originally unmatched, she was part of a blocking pair, *and* she regretted rejecting an interview invitation.

2. THE MODEL

A *market* consists of a triple (D, H, P) , where D is a finite set of doctors, H is a finite set of hospitals, and P is a profile of strict preferences for the doctors and hospitals. For each $h \in H$, $\mathcal{P}_h$ is the set of strict preferences over $D \cup \{h\}$, and for each $d \in D$, $\mathcal{P}_d$ is the set of strict preferences over $H \cup \{d\}$. The set of preference profiles is $\mathcal{P} \equiv \times_{i \in H \cup D} \mathcal{P}_i$.

There are two phases to the matching process. The first is a decentralized interview phase; the second is the centralized matching phase. The former involves many-to-many matching, while the latter is a standard one-to-one matching problem (Roth and Sotomayor (1990)).

A *many-to-many* matching is a function $\nu : H \cup D \rightarrow 2^{H \cup D}$ such that, for each $d \in D$ and $h \in H$, $\nu(d) \subseteq H$, $\nu(h) \subseteq D$, and $h \in \nu(d)$ if and only if $d \in \nu(h)$.

⁹The programs, however, are not uniformly better or worse off: h_2 is better off while h_3 is worse off.

For each $h \in H$, let $\iota_h \in \mathbb{N}$ be h 's interview capacity. Similarly, for each $d \in D$, let $\kappa_d \in \mathbb{N}$ be d 's interview capacity. We call the profile $(\iota, \kappa) = ((\iota_h)_{h \in H}(\kappa_d)_{d \in D})$ the *interview capacity profile*. An *interview matching* is a many-to-many matching ν such that for every doctor d , $|\nu(d)| \leq \kappa_d$ and for every hospital h , $|\nu(h)| \leq \iota_h$.

An interview matching ν is *pairwise stable* if there is no doctor–hospital pair (d, h) such that $h \notin \nu(d)$ but

- either $|\nu(h)| < \iota_h$ and $d P_h h$ or there exists a $d' \in \nu(h)$ such that $d P_h d'$
- either $|\nu(d)| < \kappa_d$ and $h P_d d$ or there exists an $h' \in \nu(d)$ such that $h P_d h'$.

A *matching* is a function $\mu : H \cup D \rightarrow H \cup D$ such that for each $d \in D$ and $h \in H$, $\mu(h) \in D \cup \{h\}$, $\mu(d) \in H \cup \{d\}$, and $\mu(d) = h$ if and only if $\mu(h) = d$. We say that the matching μ is *individually rational* if for each $i \in D \cup H$, $\mu(i) R_i i$. The pair (d, h) *blocks* the matching μ if $h P_d \mu(d)$ and $d P_h \mu(h)$. A matching is *stable* if it is individually rational and not blocked by any pair.

To describe how the market works, we follow the approach of Echenique et al. (2020) by assuming complete information and nonstrategic behavior.¹⁰ This means that hospitals naively make offers to their most-preferred doctors, and these offers, if rejected, trickle down to less-preferred doctors. Thus, given (ι, κ) and $P \in \mathcal{P}$, the final matching, which we call the (ι, κ) *matching*, is the outcome of the following two-phase process.

Phase 1. The interview matching ν is the hospital-optimal pairwise stable many-to-many matching where the capacities of the hospitals and doctors are given by ι and κ , respectively. This can be computed by applying the hospital-proposing DA algorithm: each $h \in H$ is matched with up to ι_h doctors and each $d \in D$ is matched with up to κ_d hospitals. Since we abstract away the informational aspect of the problem, the input to DA is a choice function for each agent that is responsive to her preference relation and constrained by her interview capacity.¹¹ The hospital-proposing DA algorithm is an approximation of the decentralized process by which hospitals invite doctors, extending invitations to further doctors when invitations are declined.

Phase 2. The (ι, κ) matching is computed by applying the doctor-proposing DA algorithm. The input to DA is the true preference profile restricted to the interview match, $(P_i|_{\nu(i)})_{i \in D \cup H}$.

The DA algorithm is used twice. To avoid confusion, we refer to the two instances as interview-DA and match-DA. Our approach differs from Echenique et al. (2020) in that we set the interview matching to be the *hospital-optimal* many-to-many stable matching, while they set it to be the *doctor-optimal* one. Their choice is appropriate for the question they ask, while for our question, the hospital-proposing DA is more appropriate, as it is the hospitals that typically make interview invitations. Our results hold

¹⁰In the working-paper version of this paper (<https://arxiv.org/abs/2102.06440>), we discuss a version of our model and some of our results when preferences are formed during the interviews.

¹¹Specifically, the choice function for each hospital selects its most-preferred doctors from each set.

whether we assume interview-DA to be doctor-proposing or hospital-proposing, as we show in the [Appendix](#).

3. WELFARE IMPACT OF INCREASED INTERVIEWS

Our aim is to study how a change in interview costs impacts a market. It is intuitive that when doctors interview with more hospitals, the interview market becomes more competitive. However, the impact of this on a doctor's ultimate match is not clear. We expect a doctor to benefit from an increase in her interviews but to be harmed when other doctors also increase their interviews; therefore, the ultimate impact is ambiguous.

We show that in fact only certain doctors who were previously unmatched can possibly benefit from additional interviews.¹² Specifically, to gain from additional interviews, a doctor must have been previously unmatched, been part of a blocking pair for the original matching, and regretted turning down an interview invitation.¹³

In 2020, 93.7% of doctors graduating from U.S. medical schools were matched to a program by the NRMP.¹⁴ Therefore, few doctors could benefit from an increase in interviews, while potentially many could be harmed.

THEOREM 1. *Suppose that for each $d \in D$, $\kappa'_d \geq \kappa_d$, and let μ and μ' be the (ι, κ) matching and (ι, κ') matching, respectively. If $\mu'(d) P_d \mu(d)$, then $\mu(d) = d$, $(d, \mu'(d))$ blocks μ , and, for each $h \in \nu(d)$, $h P_d \mu'(d)$. That is, a doctor can only gain from increased interview capacities if she (i) was unmatched, (ii) was part of a blocking pair, and (iii) regretted turning down an interview from a hospital she blocked with.*

Stability of the original matching is a natural notion of equilibrium in a well functioning market. An implication of Theorem 1 is that if such an equilibrium were shocked with increased interview capacities, the consequence would be a Pareto worsening of the match from the doctors' perspective.

COROLLARY 1. *Suppose that for each $d \in D$, $\kappa'_d \geq \kappa_d$, and let μ and μ' be the (ι, κ) matching and the (ι, κ') matching, respectively. If μ is stable, then for every $d \in D$, $\mu_d R_d \mu'_d$.*

We show a series of lemmas to prove Theorem 1. In what follows, let ν and μ be the interview and the final matchings, respectively, under (ι, κ) . Similarly, let ν' and μ' be the interview and the final matchings under (ι, κ') . We frame the temporal language below in reference to a hypothetical change in doctors' interview capacities from

¹²As a reminder, we are assuming that doctors have perfect information and are nonstrategic. While doctors can benefit from the increased information and easing of strategic constraints that additional interviews provide, our result indicate that these are the only ways a doctor can benefit.

¹³As a reminder, by "regretted" we mean that this unmatched doctor declined an interview invitation in the first phase from a hospital that she would have matched with if she had accepted it.

¹⁴This number is from the NRMP's "Results and Data: 2021 Main Residency Match" (https://mk0nrmp3oyqui6wqfm.kinstacdn.com/wp-content/uploads/2021/05/MRM-Results_and-Data_2021.pdf). The 2020 match rate was in fact slightly lower than in previous years. In 2016–2020, the match rates were 93.8%, 94.3%, 94.3%, 93.9%, and 93.7%, respectively.

κ (“before”) to κ' (“after”). As a reminder, we use interview-DA and match-DA to refer to the instances of deferred acceptance for computing the interview matchings and the final matchings, respectively.

We first establish some properties of the interview matchings. The intuition for these results comes from one of the classical results in two-sided matching theory: when the set of men increases, no man benefits from this increased competition while no woman is harmed.¹⁵ In our setting, an increase in the number of interviews a doctor accepts plays the role of additional men participating in the market. This means that the hospitals are able to interview better doctors. However, there is a tension between interviewing better doctors and interviewing the right doctors. Thus, improving the set of candidates that a hospital interviews does not necessarily translate to an improvement in its eventual match. Lemma 1 is the key to understanding the effect of additional interview capacities on the interview matching.

LEMMA 1. *Suppose that for each $d \in D$, $\kappa'_d \geq \kappa_d$, and let ν be the interview matching under κ . If $h \in \nu(d)$, then d does not reject h when interview-DA is run with capacities κ' . That is, no doctor rejects a hospital that previously interviewed her.*

PROOF. Suppose not. When interview-DA is run (with capacities κ'), let d be one of the first doctor to reject a hospital h that interviewed her under capacities κ . As $\kappa'_d \geq \kappa_d$, d must have received a new interview proposal from some hospital h' . As h' did not propose to d when capacities were κ , it must have been rejected by some doctor $d' \in \nu(h)$, a doctor it previously interviewed. But this contradicts d being among the first doctor to reject a hospital with which she previously interviewed. $\square$

We cannot say whether a doctor “prefers” her interviews under κ versus κ' , as we only have doctor’s preferences over individual hospitals and not sets of hospitals. However, we show—in a specific sense—that while a doctor may get new interviews, she does not get better interviews.

LEMMA 2. *Suppose that for each $d \in D$, $\kappa'_d \geq \kappa_d$, and let ν and ν' be the interview matchings under κ and κ' , respectively. For every $d \in D$, if $h' \in \nu'(d) \setminus \nu(d)$, then for each $h \in \nu(d)$, $h P_d h'$. That is, any new interview a doctor receives is worse than all her prior interviews.*

PROOF. Suppose d receives an interview proposal from some $h \notin \nu(d)$. If h did not propose an interview to d under κ , then h must have been rejected by a doctor who it previously interviewed. However, this contradicts Lemma 1. If h did propose an interview to d under κ , then since $h \notin \nu(d)$, d rejected h ’s interview proposal. By revealed preference, for each $h' \in \nu(d)$, $h' P_d h$. $\square$

By Lemmas 1 and 2, if a doctor was previously matched to a hospital, then every new interview she receives is worse than her previous assignment.

¹⁵See Theorem 2.25 of Roth and Sotomayor (1990).

In our framework, the impact on the hospitals is analogous to the classical result in which no man benefits from the increased competition due to additional men and also no woman is harmed. A hospital either has the same set of interviews, has additional interviews, or interviews new doctors it prefers to its previous interviews. In each of these scenarios, the hospital's set of interviews (weakly) improves. The next lemma shows that whenever a program interviews a new doctor, the program “keeps” all the interviews with doctors it prefers.

LEMMA 3. *Suppose that for each $d \in D$, $\kappa'_d \geq \kappa_d$, and let ν and ν' be the interview matchings under κ and κ' , respectively. For every $h \in H$, if $d \in \nu'(h)$, $d' \in \nu(h)$, and $d' P_h d$, then $d' \in \nu'(h)$. That is, if a hospital interviews a doctor d , it interviews every doctor it used to interview among those that it prefers to d .*

PROOF. Since $d' P_h d$, h proposes an interview to d' before it proposes an interview to d under κ' . By Lemma 1, h is not rejected by any doctor it previously interviewed. As h proposes to d under κ' , it must have already proposed to, but not have been rejected by, d' . Therefore, h continues to interview d' . $\square$

Having established the above properties of the interview matching, we now consider the matching phase. We start by showing that rejections in match-DA are monotonic with regard to doctors' interview capacities.

LEMMA 4. *Suppose that for each $d \in D$, $\kappa'_d \geq \kappa_d$, and let ν and ν' be the (ι, κ) matching and (ι, κ') matching, respectively. Suppose hospital h rejected doctor d when match-DA was run with interviews ν . If $h \in \nu'(d)$, then h rejects d when match-DA is run with interviews ν' . That is, a hospital continues to reject any doctor who it previously rejected.*

PROOF. We prove a stronger statement. We prove that if h rejected d in round n when match-DA was run under ν and if $h \in \nu'(d)$, then h rejects d in round n or earlier when match-DA is run under ν' . We proceed by induction on n . For the base step, consider a doctor d who hospital h rejected in the first round of match-DA under ν . Let d' be the doctor who h tentatively accepted when it rejected d . By assumption, $d \in \nu'(h)$. By Lemma 3, since h prefers d' to d and h interviews d under κ' , h also interviews d' under κ' ($d' \in \nu'(h)$). Moreover, by Lemma 2, any new interview d' receives is worse for her than h . Therefore, d' still proposes to h in the first round of match-DA under the new capacities, and h still rejects d in favor of a doctor it finds at least as good as d' .

Our inductive hypothesis is that for any $d \in \nu'(h)$, any $n > 1$, and any $m < n$, if h rejected d in round m of match-DA under ν , then h rejects d in round m or earlier of match-DA under ν' .

First we show that for any doctor d and any hospital h that interviews d under both κ and κ' ($h \in \nu(d) \cap \nu'(d)$), if d proposed to hospital h in round n or earlier of match-DA under ν , then d proposes to h in round n or earlier of match-DA under ν' . Consider any $h' \in \nu'(d)$ such that $h' P_d h$. By Lemma 2, h' is not a new interview ($h' \in \nu(d)$). Since when match-DA was run under ν , d proposed to h in round n or earlier, $h' \in \nu(d)$, and $h' P_d h$,

d proposes to and is rejected by h' prior to round n . By the inductive hypothesis, when match-DA is run under ν' , h' rejects d' prior to round n . Therefore, d' proposes to h by round n .

To complete the induction step, suppose that hospital h rejected doctor d in favor of doctor d' in round n of match-DA under ν . By assumption, $d \in \nu'(h)$. By Lemma 3, since h continues to interview d but prefers d' , h also continues to interview d' . Since $h \in \nu(d) \cap \nu'(d)$ and $h \in \nu(d') \cap \nu'(d')$, and both d and d' proposed to h in round n or earlier of match-DA under ν , we have shown that both d and d' propose to h by round n of match-DA under ν' . Thus, by round n , under ν' , h receives a proposal it prefers to d . Therefore, h rejects d under ν' in round n or earlier. $\square$

We are now ready to prove Theorem 1.

PROOF OF THEOREM 1. Consider any doctor $d \in D$. First, suppose that $\mu(d) = h \in H$. Consider any h' such that $h' P_d h$. We show that $\mu'(d) \neq h'$. If $h' \notin \nu'(d)$, then we are done. If $h' \in \nu'(d)$, by Lemma 2, $h' \in \nu(d)$ (all new interviews are worse than h). Since $h' P_d h$, $h' \in \nu(d)$, and $\mu(d) = h$, when match-DA is run under ν , d proposed to h' and was rejected. By Lemma 4, when match-DA is run under ν' , h' rejects d . Therefore, $\mu'(d) \neq h'$. We conclude that no doctor who was matched by μ is matched to a better hospital by μ' .

Now suppose $\mu(d) = d$ but $\mu'(d) = h \in H$. We show that d and h block μ . If $\mu(h) = h$, then since μ' is individually rational, d and h block μ . If $\mu(h) = d' \in D$, since no doctor who was matched by μ is matched to a better hospital by μ' and since $\mu'(d') \neq h = \mu(d')$, we deduce that $h P_{d'} \mu'(d')$. If $d' P_h d$, then by Lemma 3, $d' \in \nu'(h)$ (h interviewed d' before, so since h now interviews doctor d who it likes less, it continues to interview d'). But this contradicts the stability of μ' at the restricted preferences, since d' and h interview under ν' and prefer each other to their respective assignments. Therefore, it must be that $d P_h d'$, and, indeed, d and h block μ . Finally, we show that d regretted rejecting h . Since h and d block μ and d was unmatched, $d \notin \nu(h)$. Moreover, since $d P_h \mu(h)$, h proposed to d when interview-DA was run under capacities κ and d rejected h : for each $h' \in \nu(d)$, $h' P_d h$. $\square$

Theorem 1 tells us that when the number of interviews doctors accept increases, there is little scope for improving doctor welfare, but great potential for harm. The example in Section 1.2 illustrates that only certain unmatched doctors can gain from increased capacities. This example is not pathological. Lemmas 1 and 2 highlight the root cause of the inferior match, which is interview hoarding. The set of doctors who escape the adverse effects of an increase in capacities is a subset of unmatched doctors. When the match rate is high, this set is small.

Under our simplifying assumptions that agents are nonstrategic and have complete information, Theorem 1 implies that the shift to virtual interviews for the 2020–2021 season of the NRMP ought to have led to an inferior matching.¹⁶

¹⁶In the working-paper version of this paper (<https://arxiv.org/abs/2102.06440>), we contrast simulation results for this naive behavior with the common heuristic of including a “safety” candidate when choosing a set. While the NRMP has touted the high match rate for 2021, this may be driven by hospitals being matched to safety candidates (under heuristic behavior) rather than being unmatched (under naive behavior). Other than the match rate, heuristic choice does not qualitatively affect the results.

4. CAPACITY PROFILES THAT ENSURE STABILITY

As stated in Corollary 1, when the initial capacity profile leads the two-phase process to a stable matching, no doctor benefits from increased interview capacities. Understanding the relationship between capacity profiles and stability is crucial in designing policies related to interview capacities. Of course, this depends on specifics of the market, such as the ratio of doctors to hospitals and the degree to which preferences are correlated. However, we are able to provide tight characterizations for certain endpoint cases that provide intuition for more general markets.

4.1 *Stability for all preferences*

In studying the stability of the two-phase process, we first discuss worst-case performance: what capacity profiles yield stable matchings *for every* preference profile? It turns out that only very extreme capacity profiles satisfy this property. We characterize these capacity profiles in our next result.

PROPOSITION 1. *Capacity profile (ι, κ) yields a stable matching for all preferences if and only if either of the following scenarios occurs:*

- (i) *Every doctor and every hospital has only unit interview capacity (that is, for each $d \in D$, $\kappa_d = 1$, and for each $h \in H$, $\iota_h = 1$).*
- (ii) *Every doctor and every hospital has high interview capacity (that is, for each $d \in D$, $\kappa_d \geq \min\{|D|, |H|\}$, and for each $h \in H$, $\iota_h \geq \min\{|D|, |H|\}$).*

PROOF. This result is trivial if $|D| = 1$ or $|H| = 1$, so suppose that $|D| > 1$ and $|H| > 1$. We first prove necessity. Suppose that (ι, κ) yields a stable matching for all preferences.

We start by establishing that if one doctor or hospital has greater than unit interview capacity, then every doctor and hospital has interview capacity of at least 2. Stated differently, if any doctor or hospital has unit capacity, then all doctors and hospitals have unit capacity. We denote by ν the interview matching and by μ the (ι, κ) matching.

CLAIM 1. (i) *If there is $d \in D$ such that $\kappa_d > 1$, then for each $d' \in D$, $\kappa_{d'} > 1$, and for each $h \in H$, $\iota_h > 1$.*

(ii) *If there is $h \in H$ such that $\iota_h > 1$, then for each $h' \in H$, $\iota_{h'} > 1$, and for each $d \in D$, $\kappa_d > 1$.*

PROOF. We prove only the first statement, as the proof of the second statement is analogous and requires only a reversal of the roles of doctors and hospitals.

Suppose, for the sake of contradiction, that (ι, κ) always yields a stable matching and there are $d_1 \in D$ such that $\kappa_{d_1} > 1$ and $h_2 \in H$ such that $\iota_{h_2} = 1$. Let $h_1 \in H \setminus \{h_2\}$ and $d_2 \in D \setminus \{d_1\}$. Consider $P \in \mathcal{P}$ where each doctor ranks h_1 first and h_2 second, and each hospital ranks d_1 first and d_2 second. All hospitals offer an interview to d_1 , and as $\kappa_{d_1} > 1$, d_1 accepts interviews from at least h_1 and h_2 . Since $\iota_{h_2} = 1$, h_2 only interviews d_1 . Let μ be the (ι, κ) matching. Since (ι, κ) always yields a stable matching, μ

is stable, so $\mu(d_1) = h_1$, as h_1 and d_1 are mutual favorites. Therefore, $\mu(h_2) = h_2$, as h_2 only interviews d_1 . Note that (d_2, h_2) forms a blocking pair of μ as $h_2 P_{d_2} \mu(d_2)$, since $\mu(d_2) \notin \{h_1, h_2\}$ and $d_2 P_{h_2} h_2$. This contradicts the stability of μ and, thus, the assumption that (ι, κ) always yields a stable matching. We have, therefore, established that if there is $d \in D$ such that $\kappa_d > 1$, then for each $h \in H$, $\iota_h > 1$.

We now prove that if there is a $d_1 \in D$ such that $\kappa_{d_1} > 1$, then for each $d \in D$, $\kappa_d > 1$. Suppose for the sake of contradiction that there is $d_2 \in D$ such that $\kappa_{d_2} = 1$. Let $h_1, h_2 \in H$. Consider $P \in \mathcal{P}$ such that each doctor ranks h_1 first and h_2 second, and each hospital ranks d_1 first and d_2 second. As we have shown above, $\iota_{h_1}, \iota_{h_2} > 1$, so both h_1 and h_2 offer interviews to both d_1 and d_2 . Since h_1 is her favorite hospital, d_2 accepts its offer. Thus, $\nu(d_2) = \{h_1\}$. However, $\mu(d_1) = h_1$ since d_1 and h_1 are mutual favorites, so $\mu(d_2) = d_2$. This means that (d_2, h_2) form a blocking pair of μ , as the only hospital that d_2 prefers to h_2 is h_1 . This contradicts the stability of μ and, thus, the assumption that (ι, κ) always yields a stable matching. $\square$

We complete the proof of necessity by showing that neither a doctor nor a hospital can have an intermediate capacity.

CLAIM 2. *There is no $d \in D$ such that $1 < \kappa_d < \min\{|D|, |H|\}$, and there is no hospital h such that $1 < \iota_h < \min\{|D|, |H|\}$.*

PROOF. We prove this statement for the case where $|D| \leq |H|$. The proof when $|H| < |D|$ is symmetric.

Suppose for the sake of contradiction that $d_1 \in D$ is such that $\kappa_{d_1} = k$, where $1 < k < |D|$. Let $P \in \mathcal{P}$ be such that for i from 1 through $k + 1$,

$$P_{d_1} : h_2, h_3, \dots, h_{k+1}, h_1, \dots$$

$$P_{h_i} : h_i, h_1, \dots, h_{i-1}, h_{i+1}, \dots$$

$$P_{h_1} : d_1, d_2, \dots$$

$$P_{h_i} : d_i, d_1, \dots, d_{i-1}, d_{i+1}, \dots$$

We have constructed the preference profile P such that the following conditions hold:

- For each i from 1 through $k + 1$, d_i and h_i are matched in every stable matching.
- Each of the $k + 1$ hospitals $h_1, \dots, h_{k+1}$ offers d_1 an interview.
- Doctor d_1 accepts interview offers from hospitals $h_2, \dots, h_{k+1}$, but not from h_1 .

The first and third points are immediate consequences of the preferences. The second point is a consequence of the first part of Claim 1: Since $\kappa_{d_1} > 1$, every hospital has an interview capacity of at least 2 and ranks d_1 in its top two. However, this contradicts the definition of μ as the (ι, κ) matching, since $h_1 \notin \nu(d_1)$ yet, by stability, $h_1 = \mu(d_1)$.

A similar construction shows that there is no $h \in H$ such that $1 < \iota_h < |D|$. Suppose for the sake of contradiction that $h_1 \in H$ is such that $\iota_{h_1} = l$, where $1 < l < |D|$. Let $P \in \mathcal{P}$

be such that for i from 1 through $l + 1$,

$$\begin{aligned} P_{d_1} &: h_1, h_2, \dots \\ P_{d_i} &: h_i, h_1, \dots, h_{i-1}, h_{i+1}, \dots \\ P_{h_1} &: d_2, d_3, \dots, d_{l+1}, d_1 \\ P_{h_i} &: d_i, d_1, \dots, d_{i-1}, d_{i+1}, \dots \end{aligned}$$

By the second part of Claim 1, since $\iota_{h_1} > 1$, every doctor has a capacity of at least 2. Therefore, the following conditions hold:

- For each i from 1 through $l + 1$, d_i and h_i are matched in every stable matching.
- Each of the l doctors $d_2, \dots, d_{l+1}$ accepts an interview from h_1 .
- Hospital h_1 does not offer d_1 an interview.

Thus, $h_1 \notin \nu(d_1)$, so $h_1 \neq \mu(d_1)$. This contradicts the stability of μ , the (ι, κ) matching, and, in turn, the assumption that (ι, κ) always yields a stable matching. $\square$

We now turn to sufficiency. If every agent has an interview capacity of 1, then the interview matching is actually a matching. Moreover, it is a stable matching. So suppose that each agent has an interview capacity of at least $\min\{|D|, |H|\}$. If $|D| = |H|$, then the interview matching involves an interview between every mutually acceptable doctor–hospital pair. This means that the (ι, κ) matching is the doctor-optimal stable matching under unrestricted preferences, which is stable. We now show that even if $|D| < |H|$ or $|D| > |H|$, the (ι, κ) matching, μ , is stable. Suppose the doctor–hospital pair (d, h) blocks μ . By definition of μ as the (ι, κ) matching, if $h P_d \mu(d)$ and $d P_h \mu(h)$, then $h \notin \nu(d)$.

Suppose $|D| < |H|$. Since $\iota_h \geq |D|$, h would have offered an interview to d and would have been rejected when interview-DA is run, so $\nu(d)$ contains κ_d hospitals that d prefers to h . Since $h P_d \mu(d)$, and $\mu(d) \in \nu(d) \cup \{d\}$, this means $\mu(d) = d$. Then d is rejected by every hospital in $\nu(d)$ when match-DA is run. However, $|\nu(d)| = \kappa_d \geq |D|$ and since d is acceptable to every hospital in $\nu(d)$, she is only rejected when another doctor applies. However, this implies that when match-DA terminates, every hospital in $\nu(d)$ has tentatively accepted some doctor other than d , which is a contradiction as there are not enough such doctors.

Suppose $|H| < |D|$. Since $\kappa_d \geq |H|$, d does not reject any interviews she is offered. Since $h \notin \nu(d)$, h offers interviews to and has them accepted by $\iota_h \geq |H|$ doctors who it prefers to d . Since $d P_h \mu(h)$, h does not receive a proposal from any $d' \in \nu(h)$ when match-DA is run, since it finds all such d' better than d . This implies that each $d' \in \nu(h)$ is tentatively accepted by some hospital other than h when DA terminates, which is a contradiction, as there are not enough such hospitals. $\square$

Proposition 1 highlights a previously overlooked role that the interview phase plays in determining whether or not the ultimate NRMP match is stable. While interviews are

necessary for agents to gain information, we learn from Proposition 1 that interviews can also act as a bottleneck. Even with complete information, once any agent is capable of participating in more than one interview, all agents must interview with essentially the entire market to be certain that the ultimate match is stable.

4.2 Homogeneous capacity profiles

One potential intervention that has been suggested to deal with interview hoarding is a cap on the number of interviews each doctor can accept.¹⁷ Here we consider *homogeneous capacity profiles*, where all doctors face the same cap and all hospitals face the same cap. Thus, the intervention would be described by two numbers: an interview capacity $l \in \mathbb{N}$ for hospitals and an interview capacity $k \in \mathbb{N}$ for doctors. The pair (l, k) corresponds to the capacity profile (ι, κ) , where for each $h \in H$, $\iota_h = l$, and for each $d \in D$, $\kappa_d = k$.

By Proposition 1, a homogeneous capacity profile (l, k) always yields a stable matching only if $l = k = 1$ or $l, k \geq \min\{|D|, |H|\}$. Nonetheless, (l, k) may yield stable matchings for specific profiles of preferences. One might ask whether, starting at a profile $P \in \mathcal{P}$ and capacity profile (l, k) that yields a stable matching at P , the comparative statics with respect to l and k are consistent. The following example demonstrates that this is not so. It may be that, depending on P , increasing k renders a previously stable matching unstable or the opposite. In other words, the effect of an increase in k is specific to P and l .

EXAMPLE 1. Incrementing or decrementing l or k can either create or eliminate instability.

Suppose $|D| = |H| = 3$ and consider $P \in \mathcal{P}$ such that for each $i = 1, 2, 3$,¹⁸

$$\begin{array}{cc} \frac{P_{h_i}}{d_1} & \frac{P_{d_i}}{h_1} \\ d_2 & h_2 \\ d_3 & h_3 \\ h_i & d_i \end{array}$$

For P , $(2, 2)$ yields a stable matching: The interview matching is ν such that $\nu(h_1) = \nu(h_2) = \{d_1, d_2\}$ and $\nu(h_3) = \{d_3\}$. So the (l, k) matching is μ such that for each $i = 1, 2, 3$, $\mu(h_i) = d_i$, which is the unique stable matching.

We now observe that if we increment or decrement either l or k by 1, the matching is no longer stable. In other words, none of $(1, 2)$, $(3, 2)$, $(2, 1)$, or $(2, 3)$ yields a stable matching. We summarize the interview matching and the (l, k) matching for each of these in Table 1. All four of the (l, k) matchings are unstable. $\diamond$

The mechanics of Example 1 are robust, and it is not by accident that $(2, 2)$ yields a stable outcome to start with. The preferences in the example have a particularly salient

¹⁷For instance, see Morgan et al. (2020).

¹⁸This can be embedded into a larger problem.

TABLE 1. Interviews and final matchings in Example 1 for different values of k and l .

(l, k)	Interview Matching	(l, k) Matching
(1, 2)	$\nu(h_1) = \nu(h_2) = \{d_1\}, \nu(h_3) = \{d_2\}$	$\mu(h_1) = d_1, \mu(h_3) = d_2, \mu(h_2) = h_2, \mu(d_3) = d_3$
(3, 2)	$\nu(h_1) = \nu(h_2) = D, \nu(h_3) = \{\}$	$\mu(h_1) = d_1, \mu(h_2) = d_2, \mu(h_3) = h_3, \mu(d_3) = d_3$
(2, 1)	$\nu(h_1) = \{d_1, d_2\}, \nu(h_2) = \{d_3\}, \nu(h_3) = \{\}$	$\mu(h_1) = d_1, \mu(h_2) = d_3, \mu(h_3) = h_3, \mu(d_2) = d_2$
(2, 3)	$\nu(h_1) = \nu(h_2) = \nu(h_3) = \{d_1, d_2\}$	$\mu(h_1) = d_1, \mu(h_2) = d_2, \mu(h_3) = h_3, \mu(d_3) = d_3$

configuration, on which we focus here. A profile $P \in \mathcal{P}$ has *common preferences* if all doctors rank the hospitals in the same way and all hospitals rank the doctors in the same way. To further restrict the definition, we also require that each doctor finds each hospital acceptable and each hospital finds each doctor acceptable. That is, for each pair $d, d' \in D$ and each pair $h, h' \in H$, $P_d|_H = P_{d'}|_H$, $P_h|_D = P_{h'}|_D$, $d P_h h$, and $h P_d d$.¹⁹

As we see from Example 1, a result like Proposition 1 does not hold if we restrict ourselves to common preferences. Our next result is a characterization of homogeneous capacity profiles that yield stable matchings for common preferences.²⁰

PROPOSITION 2. *Let $P \in \mathcal{P}$ be such that there are common preferences. A homogeneous capacity profile (l, k) yields a stable matching at P if and only if $l = k$ or $l, k \geq \min\{|D|, |H|\}$.*

PROOF. Let $\{d_t\}_{t=1}^{|D|}$ and $\{h_t\}_{t=1}^{|H|}$ be enumerations of D and H , respectively, such that every hospital prefers d_t to d_{t+1} and every doctor prefers h_t to h_{t+1} . Let $m = \min\{|D|, |H|\}$. There is a unique stable matching μ^* , such that for each $t = 1, \dots, m$, $\mu^*(h_t) = d_t$.

Let ν be the interview matching under (l, k) and let μ be the (l, k) matching.

First we show that (l, k) yields a stable matching at P only if $l = k$ or $l, k \geq \min\{|D|, |H|\}$. Suppose $l \neq k$. If $l < k$ and $l < \min\{|D|, |H|\}$, then for each $t = 1, \dots, k$, $\nu(h_t) = \{d_1, \dots, d_l\}$. In particular, $d_k \notin \nu(h_k)$, so $\mu(h_k) \neq d_k$. On the other hand, if $l > k$ and $k < \min\{|D|, |H|\}$, then for each $t = 1, \dots, l$, $\nu(d_t) = \{h_1, \dots, h_k\}$. In particular, $h_l \notin \nu(d_l)$, so $\mu(d_l) \neq h_l$. In either case, the (l, k) matching is not stable.

Now we show that if $l = k \leq m$, then (l, k) yields a stable matching. For each $t = 1, \dots, m$, let $\underline{t} = \lfloor \frac{t-1}{l} \rfloor$. Then, for each $t = 1, \dots, l$, $\nu(h_t) = \{d_{\underline{t}l+1}, \dots, d_{(\underline{t}+1)l}\}$ and $\nu(d_t) = \{h_{\underline{t}l+1}, \dots, h_{(\underline{t}+1)l}\}$. Thus, for each $t = 1, \dots, m$, $\mu(h_t) = d_t$ and so μ is stable.

Finally, if $l, k \geq m$, then for each $t = 1, \dots, m$, $\nu(d_t) \supseteq \{h_1, \dots, h_m\}$. Since $h_t \in \nu(d_t)$, $h_t = \mu(d_t)$ and so μ is stable. $\square$

If the hospitals' interview capacity is fixed at some specific l , the question of where to set the doctors' interview cap, k , is an important policy decision. Proposition 2 says that the optimal value for k is exactly at l , whether the objective is to minimize the number of

¹⁹Under common preferences, there is a unique stable matching.

²⁰The characterization of Proposition 2 does not hold for capacity profiles that are not homogeneous. For a counterexample, suppose $|D| = 4$, $|H| = 3$, there is $d \in D$ such that $\kappa_d = 3$ for each $d' \in D \setminus \{d\}$, $\kappa_{d'} = 2$, there is $h \in H$ such that $\iota_h = 4$, and for each $h' \in H \setminus \{h\}$, $\iota_{h'} = 2$. For any common preferences, (ι, κ) yields a stable matching.

blocking pairs or to maximize the match rate (the proportion of positions that are filled). Our next result sheds more light on this.

PROPOSITION 3. *Fix the hospitals' interview capacity at l and consider k and k' such that either $k' < k \leq l$ or $l \leq k < k'$. Suppose $P \in \mathcal{P}$ has common preferences. The (l, k') matching has more blocking pairs and a weakly lower match rate than the (l, k) matching.*

PROOF. Let $P \in \mathcal{P}$ be such that there are common preferences. Let $\{d_t\}_{t=1}^{|D|}$ and $\{h_t\}_{t=1}^{|H|}$ be enumerations of D and H , respectively, such that every hospital prefers d_t to d_{t+1} and every doctor prefers h_t to h_{t+1} .

Let $m = \min\{\lfloor \frac{|H|}{k} \rfloor, \lfloor \frac{|D|}{l} \rfloor\}$. The interview matching is such that for each d_t , if $t \leq ml$,

$$\nu(d_t) = \{h_{(n-1)k+1}, \dots, h_{nk}\}, \quad \text{where } n \text{ is such that } (n-1)l < t \leq nl,$$

if $ml < t \leq (m+1)l$,

$$\nu(d_t) = \begin{cases} \{h_{mk+1}, \dots, h_n\} & \text{if } |H| \geq mk + 1, \\ \emptyset & \text{otherwise,} \end{cases} \quad \text{where } n = \min\{|H|, (m+1)k\}$$

and if $(m+1)n < t$, $\nu(d_t) = \emptyset$.

We first consider the case in which $k > l$ and show that the number of matched hospitals is decreasing in k and that the number of blocking pairs is increasing in k .

Given P and its restriction to ν , the (l, k) matching, μ , at P is such that for each d_t , if $t \leq ml$,

$$\mu(d_t) = h_{(n-1)k+(t \bmod l)}, \quad \text{where } n \text{ is such that } (n-1)l < t \leq nl,$$

if $ml < t \leq (m+1)l$,

$$\mu(d_t) = \begin{cases} h_{mk+(t \bmod l)} & \text{if } |H| \geq mk + (t \bmod l), \\ d_t & \text{otherwise,} \end{cases}$$

and if $(m+1)l < t$, $\mu(d_t) = d_t$.

Let $n = \min\{|H| - mk, |D| - ml\}$. Given the (l, k) matching above, the set of matched hospitals is

$$\{h_{ik+s} : i = 0, \dots, m-1, s = 1, \dots, l\} \cup \{h_t : t = mk + 1, \dots, mk + n\}.$$

Therefore, the number of matched hospitals is $ml + n$. Holding l fixed, this is decreasing in k .

The (l, k) matching is blocked by all pairs consisting of an unmatched hospital and any doctor with a higher index, that is, $(h_t, d_{t'})$ such that $t \leq mk$, $t - 1 \bmod k \geq l$, and $t' > t$. These are the only pairs that block it. Thus, the number of blocking pairs is

$$\sum_{n=0}^{m-1} \sum_{i=l+1}^k |D| - (nk + i).$$

Holding l fixed, this is increasing in k .

Now we consider the case in which $k < l$ and show that the number of matched hospitals is increasing in k and the number of blocking pairs is decreasing in k .

Given P and its restriction to ν , the (l, k) matching at P is such that for each h_t , if $t \leq mk$,

$$\mu(h_t) = d_{(n-1)l+(t \bmod k)}, \quad \text{where } n \text{ is such that } (n-1)l < t \leq nl,$$

if $mk < t \leq (m+1)k$,

$$\mu(h_t) = \begin{cases} h_{ml+(t \bmod k)} & \text{if } |D| \geq ml + (t \bmod k), \\ h_t & \text{otherwise,} \end{cases}$$

and if $(m+1)k < t$, $\mu(h_t) = h_t$.

Let $n = \min\{|H| - mk, |D| - ml\}$. Given the (l, k) matching above, the set of matched hospitals is $\{h_t : t \leq mk + n\}$. Therefore, the number of matched hospitals is $mk + n$. Since $k < l$, this is weakly increasing in k .

The (l, k) matching is blocked by all pairs consisting of an unmatched doctor and any hospital with a higher index, that is, $(d_t, h_{t'})$ such that $t \leq ml$, $t - 1 \bmod l \geq k$, and $t' > t$. These are the only pairs that block it. Thus, the number of blocking pairs is

$$\sum_{n=0}^{m-1} \sum_{i=k+1}^l |D| - (nl + i).$$

Holding l fixed, this is decreasing in k . □

5. SIMULATIONS

Our analytical results are of two sorts. On one hand, Theorem 1 applies without restrictions on preferences. However, it only helps identify the problem caused by increases in doctors' interview capacities, without suggesting a remedy. On the other hand, when we focus on common preferences, Propositions 2 and 3 deliver a clear-cut policy prescription. In this section, we use simulations to bridge the gap. This allows us to consider how changes in the doctors' interview capacities affect hospitals' welfare, match rates, stability, and so on in a more general setting.

While there is evidence that preferences do indeed have a common component (Agarwal (2015), Rees-Jones (2018)), agents care about fit as well. Moreover, an idiosyncratic component is to be expected. We adopt the random utility model of Ashlagi, Kano-ria, and Leshno (2017).²¹ Each hospital $h \in H$ has a common component to its quality, x_h^C , and a fit component, x_h^F . Similarly, each doctor $d \in D$ has a common component to her quality, x_d^C , and a fit component, x_d^F . The utilities that h and d enjoy from being matched to one another are

$$u_h(d) = \beta x_d^C - \gamma (x_h^F - x_d^F)^2 + \varepsilon_{hd},$$

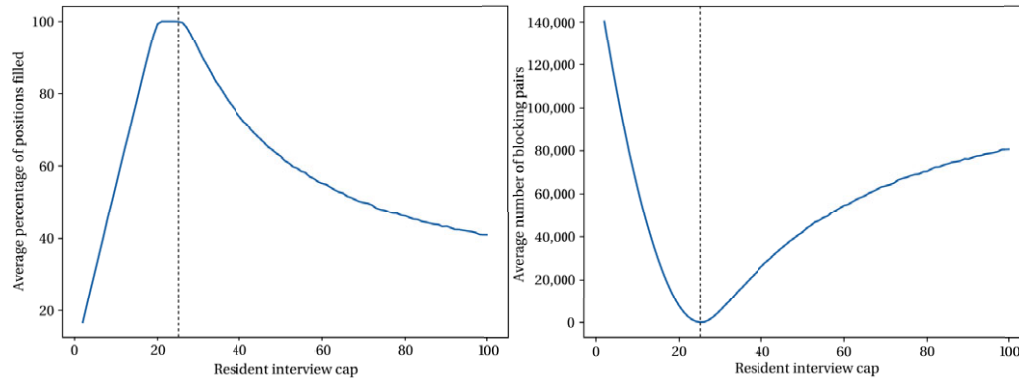
$$u_d(h) = \beta x_h^C - \gamma (x_h^F - x_d^F)^2 + \varepsilon_{dh},$$

²¹This, in turn, is adapted from Hitsch, Hortaçsu, and Ariely (2010).

respectively, where ε_{hd} and ε_{dh} are drawn independently from the standard logistic distribution. Each x_h^C , x_h^F , x_d^C , and x_d^F is drawn independently from the uniform distribution over $[0, 1]$. The coefficients β and γ weight the common and fit components, respectively. When β and γ are both zero, preferences are drawn uniformly at random. As $\beta \rightarrow \infty$, these approach common preferences. As γ increases, preferences become more aligned (the fit, which is orthogonal to the common component, becomes more important).

Our simulated market has 400 hospitals.²² We have set the number of doctors at 470.²³ The parameters for the random utility model are $\beta = 40$ and $\gamma = 20$. Since our interest is in the effects of changes in doctors' interview capacities, we fix hospital interview capacities at $l = 25$.²⁴

In our first simulation, we vary k from 2 to 100.²⁵ As k increases, the match rate increases and then decreases (Figure 1a). On the other hand, as k increases, the number of blocking pairs decreases and then increases (Figure 1b). These results are consistent



(a) Average match rate relative to interview cap. (b) Average number of blocking pairs relative to interview cap.

FIGURE 1. Average match rate and average number of blocking pairs relative to doctors' interview cap. In this figure, k varies from 2 to 100 with l fixed at 25 (indicated by a vertical line). When $k = l$, the average match rate is 99.9525% and the average number of blocking pairs is 136.08.

²²The NRMP match is broken down into smaller matches by specialty. In 2020, among 50 specialties for PGY-1 programs, the largest had 8697 positions, the 10th largest had 849 positions, the 25th largest had 38 positions, the 49th largest had one position, and the smallest had no positions (these data are available from the NRMP (https://mk0nrmp3oyqui6wqfm.kinstacdn.com/wp-content/uploads/2020/06/MM_Results_and-Data_2020-1.pdf)). Our chosen number of hospitals is comparable to the 70th percentile among specialties (that is, 70% of specialties are smaller than this).

²³There were, on average, 0.85 PGY-1 positions per applicant in the 2020. Our chosen number of doctors reflects this ratio.

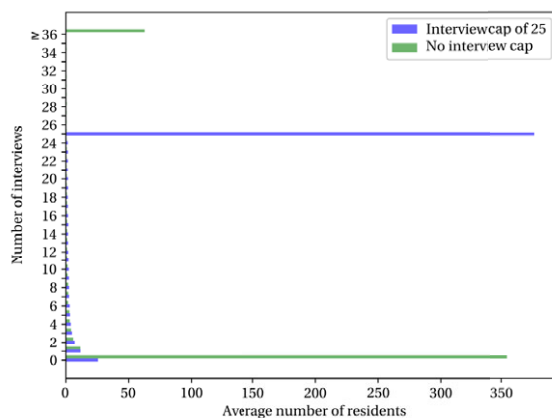
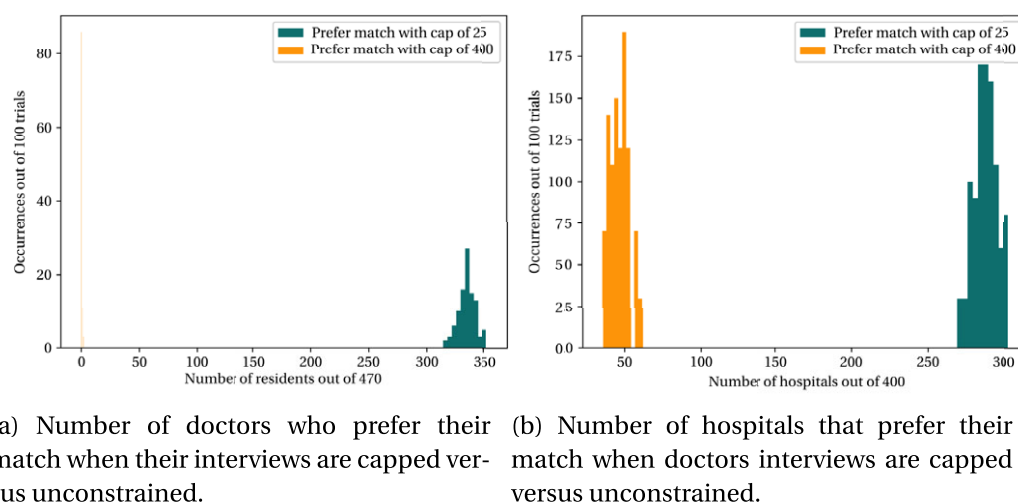
²⁴We have chosen values of β and γ so that there is a significant weight on both the common and fit components of preferences. However, the magnitude of these components is somewhat arbitrary. In the working-paper version of this paper (<https://arxiv.org/abs/2102.06440>), we discuss the robustness of our findings with regard to our choices of model and parameter values.

²⁵We have chosen this upper bound to be high enough that further increases have little effect. We interpret this as doctors being essentially unconstrained in the number of interviews they can accept.

with Proposition 3. Although the preferences are not common, both the match rate and stability (measured by the number of blocking pairs) are optimized at $k = l$.

Our next set of results evaluate a hypothetical policy of restricting doctors to a maximum of $k = l (= 25)$ interviews. We compare this policy with the benchmark of no intervention, where doctors are completely unconstrained.

Figure 2a shows the distribution of the number of doctors who prefer their match under the optimal k over the benchmark, as well as the distribution of those with the opposite preference. We see that the former is considerably higher than the latter. Theorem 1 allows for certain unmatched doctors to gain from relaxing this policy. However, Figure 2a shows that such doctors are rare: on average, only 0.17 doctors gain, while



(c) Distribution of interviews when doctors' interviews are capped versus unconstrained.

FIGURE 2. Comparison of capped versus unconstrained doctor interviews. This figure shows comparisons of the effects of the intervention of capping doctors' interview capacities at $k = l$ versus leaving the capacities unconstrained. In (c), since the tail of the distribution without caps is very thin, the total number of doctors with 36 or more interviews is consolidated into a single bar.

334.45 are harmed. Even comparing the supports of the distributions, no more than two doctors gain in any of the trials, while at least 314 are harmed in every trial. Figure 2b shows the analogous distributions for hospitals. The result is not as sharp as for the doctors, but, on average, 240.39 more hospitals are harmed than gain. In fact, in no trial do more than 62 hospitals gain or are fewer than 269 hospitals harmed. Although Theorem 1 does not address hospitals' welfare, this suggests that more hospitals prefer the optimal cap of $k = l$ over leaving the doctors unconstrained. The optimal cap also has the benefit of bringing stability to the final matching. There are, on average, 105,075.56 blocking pairs eliminated by this policy.²⁶

Finally, in Figure 2c we compare the distribution of interviews among the doctors between the two capacities. The results point to interview hoarding as the cause for the poor matching when there is no interview cap. A small number of doctors (on average 68.08) hoard so many (at least 36) interviews that a large number of doctors number (on average 354.58) end up with *none*. The constraint on doctors' interviews to $k = l$ binds for many doctors (376.13 on average). This leads to a better distribution of interviews and improves the final match, as we see above. Moreover, very few doctors (25.57 on average) end up with no interviews.

In our last simulations, we consider the possibility that the NRMP could not only set a cap on interviews that doctors can accept, but also control the number of interviews that hospitals offer. From Proposition 2, we know that if preferences are common, the match rate would be maximized at $k = l$. Figure 3 shows that even when preferences are not exactly common, this is still the optimal policy: along the diagonal, where k and l are equal, the match rate is close to 100% and the match is almost stable.

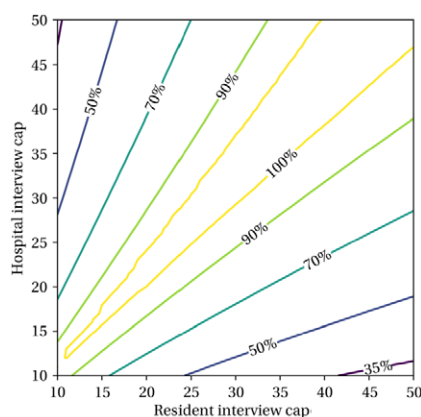


FIGURE 3. Average match rate for different values of k and l . The average match rate is highest around $k = l$. The asymmetry is driven by the imbalance of our simulated markets (there are 470 residents but only 400 hospitals).

²⁶To put this number in context, there are a total of 188,000 doctor–hospital pairs that could possibly block.

6. CONCLUSION

The COVID-19 pandemic has had a significant impact on the way interviews, including residency interviews between doctors and hospitals, are conducted. We have examined the impact of an increase in doctors' interviews on the final doctor–hospital residency matches. Our theoretical results show that an increase in doctors' interview capacities leads to an overall worsening of doctors' welfare due to a phenomenon we call interview hoarding, which worsens the interview bottleneck. We have identified an optimal mitigation strategy for this bottleneck in which the interview capacities of doctors and hospitals are held equal, a strategy that also leads to greater stability. We have extended these results through simulations.

Several implications of our study may be helpful in informing future policy. The negative effects of an increase in doctors' interviews demonstrated in our results suggest that the 2021 NRMP match was likely inferior to that of previous years. In the future, it could be beneficial for the NRMP to consider policies to mitigate interview hoarding. Our analysis and simulations provide evidence that interview caps could be effective in doing so. Such caps could be implemented with very limited centralization, for example, by the use of a ticketing system. Even if such an intervention is not possible in the short term, we recommend that residency programs be advised to increase the number of candidates they interview relative to previous years, so as to equalize, if possible, the interview capacities of doctors and hospitals.

The design of a fully centralized clearinghouse for interviews is an area that remains open. As earlier research on the interview pre-markets has shown, strategic analysis is only tractable under very stringent assumptions (Kadam (2021), Lee and Schwarz (2017), Beyhaghi (2019)). Nonetheless, our paper adds to the evidence (along with Echenique et al. (2020)) that a more holistic approach that includes the interview stage is critical.

The interview-driven bottleneck may be a factor in other matching contexts as well, such as fully decentralized labor markets. The economics job market is one such example, which (like the residency market) formerly involved in-person interviews and on-site visits, but now has transitioned to a virtual process. Further research focusing on the impact of virtual interviews in the matching process could provide valuable theoretical insights and policy direction in improving welfare and stability of such markets.

APPENDIX: DOCTOR-OPTIMAL INTERVIEW MATCHING

The only difference between our model and that of Echenique et al. (2020) is that we suppose that the interview matching is hospital-optimal rather than doctor-optimal. Their modeling choice is natural for the question they ask, as it gives each doctor her *best* stable set of interviews. Thus, their result that most doctors match with hospitals they rank highly can only be stronger for other interview matchings. For our analysis, the doctor-optimal interview matching does not have this natural appeal. To the contrary, hospital-proposing DA is a reasonable approximation of the interview matching process. Nonetheless, our results are not driven by this choice. The only proof that relies on this choice is that of Theorem 1. In this appendix, we show that the result holds

even for the doctor-optimal interview matching followed by the doctor-optimal final matching. In what follows, we use the same terminology and notation as before, with the understanding that the interview matching is doctor-optimal.

As in the statement of the theorem, suppose that for each $d \in D$, $\kappa_d \leq \kappa'_d$. We show below that Lemmas 2 and 3 hold even with the change from the hospital-optimal interview matching to the doctor-optimal interview matching. The key is to establish that, in the interview phase, if a doctor d is rejected by a hospital h under capacities κ , then h rejects her under κ' as well. Given capacities $\bar{\kappa}$, let $A_d(m; \bar{\kappa})$ be the hospitals that d proposes to and let $R_d(m; \bar{\kappa})$ be the hospitals that reject doctor d by the end of round m of the interview phase. We show that these sets are monotonic in $\bar{\kappa}$.

CLAIM 3. *For any positive integer m ,*

$$R_d(m; \kappa) \subseteq R_d(m; \kappa'),$$

$$A_d(m; \kappa) \subseteq A_d(m; \kappa').$$

PROOF. We proceed by induction on m , the base case being $m = 1$. In the first round of the interview phase, each $d \in D$ proposes to her favorite hospitals up to her interview capacity. Since every doctor proposes to at least as many hospitals under κ' as under κ , every hospital receives at least as many proposals under κ' as under κ . Therefore, if a doctor d is rejected by a hospital h in the first round of the interview phase under κ , she is also rejected by h in the first round under κ' . Now consider a round $m > 1$ of the interview phase and suppose for each doctor d that $R_d(m-1; \kappa) \subseteq R_d(m-1; \kappa')$ and $A_d(m-1; \kappa) \subseteq A_d(m-1; \kappa')$. In round m , each $d \in D$ proposes to her favorite hospitals that have not yet rejected her up to her interview capacity. Under κ , d proposes to her κ_d favorite hospitals in $H \setminus R_d(m-1; \kappa)$. Under κ' , she proposes to her κ'_d favorite hospitals in $H \setminus R_d(m-1; \kappa')$. By the inductive hypothesis, $H \setminus R_d(m-1; \kappa') \subseteq H \setminus R_d(m-1; \kappa)$. Therefore, if d proposes to h under κ , either she proposes to h under κ' as well (she is choosing more hospitals from a smaller set of options) or she has already proposed to and has been rejected by h under κ' . In either case, if $h \in A_d(m; \kappa)$, then $h \in A_d(m; \kappa')$. Since each hospital h receives more proposals but its capacity does not change, if h rejects doctor d under κ , it also rejects doctor d when choosing from a larger set of doctors who have proposed to it under κ' . Therefore, if $h \in R_d(m; \kappa)$, then $h \in R_d(m; \kappa')$. $\square$

We now explain how Claim 3 implies that Lemmas 2 and 3 hold even when we switch to the doctor-optimal interview matching. Let ν and μ be the interview and final matchings, respectively, under (ι, κ) . Similarly, let ν' and μ' be the interview and final matchings under (ι, κ') .

Lemma 2 says that for each $d \in D$, if $h \in \nu'(d) \setminus \nu(d)$, then $\mu(d) P_d h$. Given $\bar{\kappa}$, let $R_d(\bar{\kappa})$ be the set of hospitals that reject d in any round of the interview phase under capacities $\bar{\kappa}$. By Claim 3, $R_d(\kappa) \subseteq R_d(\kappa')$. Under κ , $\nu(d)$ consists of d 's κ_d most-preferred hospitals in $H \setminus R_d(\kappa)$. That is, the κ_d highest-ranked hospitals that did not reject her. Under κ' , $\nu'(d)$ comprises d 's κ'_d most-preferred hospitals in $H \setminus R_d(\kappa')$. As $H \setminus R_d(\kappa') \subseteq H \setminus R_d(\kappa)$, if $h' \in \nu'(d) \setminus \nu(d)$, then for every $h \in \nu(d)$, $h P_d h'$. In words,

since d is interviewed by h under κ , she was not rejected by h under κ . As more hospitals rejected d under κ' than under κ , h does not reject d under κ . Therefore, d could have proposed to h under κ , but she chose not to. Therefore, by revealed preference, she prefers all hospitals in $\nu(d)$ to any of her “new” interviews under κ' (those in $\nu'(h) \setminus \nu(h)$).

Lemma 3 said that if $d \in \nu'(h)$, $d' \in \nu(h)$, and $d' P_h d$, then $d' \in \nu'(h)$. By Claim 3, d' proposes to at least as many hospitals in the interview phase under κ' as under κ . Since d' proposes to h under κ , she also proposes to h under κ' . Each hospital h accepts its ι_h favorite applicants, so if it accepts d , it must also accept d' .

Since Lemmas 2 and 3 hold, the remainder of the proof follows exactly as in Section 3.

REFERENCES

- Agarwal, Nikhil (2015), “An empirical model of the medical match.” *American Economic Review*, 105, 1939–1978. [520]
- Ali, S. Nageeb and Ran I. Shorrer (2021), “The college portfolio problem.” Working paper, Penn State. [506]
- Ashlagi, Itai, Yash Kanoria, and Jacob D. Leshno (2017), “Unbalanced random matching markets: The stark effect of competition.” *Journal of Political Economy*, 125, 69–98. [520]
- Beyhaghi, Hedyeh (2019), “Approximately-optimal mechanisms in auction design, search theory, and matching markets.” PhD thesis, Cornell University. Chapter 5: Two-Sided Matching with Limited Number of Interviews. [506, 524]
- Chade, Hector and Lones Smith (2006), “Simultaneous search.” *Econometrica*, 74, 1293–1307. [506]
- Echenique, Federico, Ruy Gonzalez, Alistair Wilson, and Leeat Yariv (2020), *Top of the Batch: Interviews and the Match*. [505, 506, 509, 524]
- Gale, David and Lloyd Shapley (1962), “College admissions and the stability of marriage.” *American Mathematical Monthly*, 69, 9–15. [503]
- Hitsch, Gunter J., Ali Hortaçsu, and Dan Ariely (2010), “Matching and sorting in online dating.” *American Economic Review*, 100, 130–163. [520]
- Kadam, Sangram V. (2021), “Interviewing in matching markets with virtual interviews.” [506, 524]
- Lee, Robin S. and Michael Schwarz (2017), “Interviewing in two-sided matching markets.” *The RAND Journal of Economics*, 48, 835–855. [506, 524]
- Morgan, Helen Kang, Abigail F. Winkel, Taylor Standiford, Rodrigo Muñoz, Eric A. Strand, David A. Marzano, Tony Ogburn, Carol A. Major, Susan Cox, and Maya M. Hammoud (2020), “The case for capping residency interviews.” *Journal of Surgical Education*, 78, 755–762. [517]
- Rees-Jones, Alex (2018), “Suboptimal behavior in strategy-proof mechanisms: Evidence from the residency match.” *Games and Economic Behavior*, 108, 317–330. [520]

Roth, Alvin E. and Marilda A. Oliveira Sotomayor (1990), *Two-Sided Matching: A Study in Game-Theoretic Modeling and Analysis*. Econometric Society Monographs. Cambridge University Press. [508, 511]

Roth, Alvin E. and Elliott Peranson (1999), “The redesign of the matching market for American physicians: Some engineering aspects of economic design.” *American Economic Review*, 89, 748–780. [503, 506]

Co-editor Federico Echenique handled this manuscript.

Manuscript received 12 May, 2021; final version accepted 25 May, 2022; available online 13 June, 2022.