

# The Backlash of Gender Quotas

Juan B. González      Alejandro Martínez-Marquina \*

April 7, 2026

## Abstract

Gender quotas are used to address gender disparities, but they may trigger backlash that undermines their effectiveness. In an online hiring experiment, requiring recruiters to hire a female candidate for one position leads them to penalize women in other, uncovered positions, offering lower salaries and making them twice as likely to receive no offer. No such effect arises when quotas mandate hiring men. Backlash emerges among supporters and opponents of equity policies, and renders the quota policy ineffective. Under the quota, women are evaluated jointly with quota beneficiaries, unintentionally raising the bar for the group the policy aims to promote.

**JEL classifications:** C91, D91, J16, J23, J71, M51

**Keywords:** gender; quotas; backlash; experiment

In recent years, gender quotas have become one of the most widely used policy tools to address gender disparities, spanning electoral lists, corporate boards, university admissions, and sports.<sup>1</sup> In practice, however, most quotas are partial: they mandate female representation in specific targeted positions—such as board seats or electoral lists—while leaving other positions uncovered and subject to discretion. Whether quotas achieve their core goal of improving aggregate outcomes for women, therefore, depends critically on whether they trigger backlash: negative spillovers to women in discretionary positions not covered by the policy. Backlash might explain why evidence on the effectiveness of quotas is mixed. While quotas reliably increase female representation in targeted positions, their broader impact varies

---

\*González: Department of Economics, University of Southern California, juanbgon@usc.edu. Martínez-Marquina: Marshall School of Business, University of Southern California, am04817@usc.edu. We thank Muriel Niederle, and the many Stanford SITE, and USC seminar participants for their invaluable feedback. We also thank all the participants from the LAX conference and Caltech workshop. This research was supported by an Institute for Outlier Research in Business (IORB) grant from the University of Southern California. The authors have no other financial relationships or conflicts of interest to disclose.

<sup>1</sup>Some examples include quotas in electoral lists in India (Beaman et al., 2009; Green et al., 2026), South Korea (Lee & Zanella, 2024), Italy (De Paola et al., 2010), or Spain (Casas-Arce & Saiz, 2015; Bagues & Campa, 2021); in corporate boards in Norway (Bertrand et al., 2019), France (Ferreira et al., 2020), or Italy (Ferrari et al., 2022); and in sports, such as chess in France (De Sousa & Niederle, 2022).

widely, sometimes leaving aggregate outcomes unchanged (Bertrand et al., 2019; Karekurve-Ramachandra & Sood, 2025), or even backfiring entirely (Bagues et al., 2017; Bian et al., 2024; Deschamps, 2024). Strikingly, negative spillovers have been documented across vastly different contexts, from traditional societies like South Korea and India (Beaman et al., 2009; Lee & Zanella, 2024) to progressive ones like California, France, and Spain (Gertsberg et al., 2021; Lippmann, 2021; Deschamps, 2024).

Despite decades of quota policies worldwide, the extent and source of backlash remain poorly understood. Quotas can have several confounding effects, making it difficult to identify backlash in observational studies. They may change the composition of the candidate pool (Arcidiacono & Lovenheim, 2016; Bleemer, 2022; De Sousa & Niederle, 2022), lower incentives for women to exert effort (Coate & Loury, 1993), or make evaluators update downwards their beliefs about female competence (Kogelnik & Strack, 2025). Therefore, although empirical studies have identified adverse impacts of gender quotas in a variety of settings<sup>2</sup>, it remains an open question whether they reflect a rational response to adverse selection or if they constitute a retaliatory response against women not covered by the policy. Therefore, identifying whether backlash exists and what drives it is crucial to understanding when quotas promote gender equity and when they backfire.

In this paper, we provide experimental evidence that quota backlash exists, identify its sources, and trace its consequences for policy effectiveness. To do so, we develop a new experimental design that introduces gender quotas while holding all other factors constant. Participants, acting as recruiters, send offers to available job candidates.<sup>3</sup> In the *Baseline*, recruiters allocate a fixed budget across two candidates—for example, a man and a woman. In the *Quota* treatment, recruiters face the same two discretionary candidates but must also hire an additional female candidate at a fixed cost, financed by an equivalent budget increase. Effectively, recruiters in both conditions choose how to allocate the same discretionary resources between the same two candidates, with the only difference being the mandatory quota hire. Because the quota leaves the discretionary candidate pool unchanged, any negative spillovers to non-targeted women reflect backlash rather than a response to adverse selection or compositional change.<sup>4</sup>

---

<sup>2</sup>Negative spillovers of gender quotas have been found in electoral lists (Beaman et al., 2009; Clayton & Tang, 2018; Bagues & Campa, 2021; Lippmann, 2021; Lee & Zanella, 2024; Karekurve-Ramachandra & Sood, 2025), corporate boards (Ahern & Dittmar, 2012; Gertsberg et al., 2021; Maida & Weber, 2022; Bian et al., 2024), and hiring committees (Bagues et al., 2017; Deschamps, 2024). In contrast, other studies have found positive spillovers (De Paola et al., 2010; Bagde et al., 2016; Ferreira et al., 2020; De Sousa & Niederle, 2022; Ferrari et al., 2022).

<sup>3</sup>See Bohren, Hull, and Imas (2025) and Exley and Nielsen (2024) for other examples in the literature where participants act as recruiters.

<sup>4</sup>Bleemer (2022) documents a decline of 250 Black and 900 Hispanic applicants per year to the UC system following the passage of Prop 209, which banned race-based affirmative action. Conversely, De Sousa and Niederle (2022) finds that the introduction of a gender quota in French chess teams led to a trickle-down effect, encouraging more women to join.

Figure 1: Main Experimental Design - First Decision

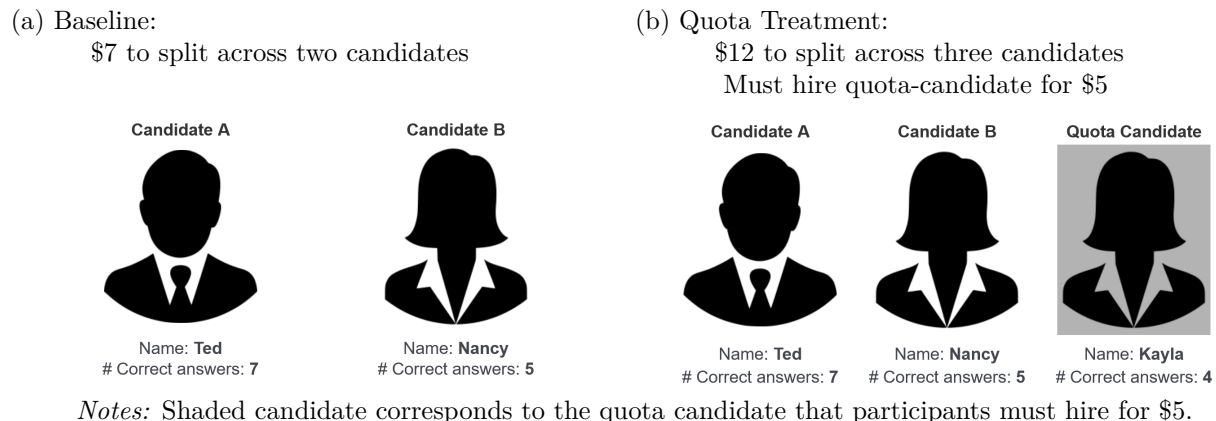


Figure 1 summarizes our main between-subject treatment variation. In the *Baseline*, participants split \$7 between two potential candidates, with the option to either split the money between the two or uniquely pay one of them. If a candidate is hired, the recruiter receives \$1 for each correct answer the candidate obtains in the last half of a 20-question quiz. In addition to the candidates' gender, we provide a measure of their proficiency at the task by showing how many of the first 10 quiz questions they answered correctly. Hence, we provide an objective measure of productivity that is highly predictive of future performance, and most importantly, does not depend on the quota. In the *Quota Treatment*, we increase the hiring budget to \$12, but \$5 of those funds must be used to hire the quota candidate. Hence, in both treatments, participants effectively have \$7 to be split among the two potential candidates. In the first part of the experiment, participants face 12 hiring decisions, similar to the one shown in Figure 1, in which we randomized candidates' gender and performance. This provides variation in gender composition and relative candidate performance, which will be central to identifying and dissecting quota backlash.

Our central hypothesis, referred to as *quota backlash*, posits that forcing the hiring of a female quota candidate (Kayla) creates a retaliatory response against other women not directly covered by the policy (Nancy), reducing their salaries and likelihood of being hired relative to the *Baseline*. Thereby, undermining the very goal of the quota in promoting the hiring of women. To rule out alternative explanations, such as a general aversion to quotas or a preference for hiring a person from each gender, we implement an additional treatment in which the quota candidate is male.

We find strong evidence of quota backlash. When recruiters are required to hire a female quota candidate, they penalize other women in discretionary positions. In the first hiring decision, where all participants face identical candidates and incentives (Figure 1), introducing a quota widens the gender wage gap by 24 percent—from \$2.49 to \$3.08—and nearly doubles the likelihood that the female candidate receives no offer, from 11.9 to 23.1 percent. Expanding the analysis to all 12 hiring decisions, backlash is performance-specific: it emerges only when female candidates underperform relative to their male counterparts, and

is absent when women perform as well or better. Among underperforming women, quotas increase by over 50 percent the likelihood of receiving no offer at all. This pattern suggests that quotas do not simply redistribute opportunities away from women, but raise the bar against which they are evaluated: under the quota, even a small performance disadvantage becomes enough to be passed over entirely. Importantly, no such effect arises in the *Quota Male* treatment, where the quota candidate is male. This rules out that backlash reflects a general aversion to quotas or preferences for diversity, pointing instead to a gender-specific retaliatory response. Two further alternative explanations can also be ruled out. First, backlash is stronger among participants who correctly answered all comprehension questions, ruling out task confusion. Second, we find no evidence that quotas shift recruiters’ beliefs about female candidates’ ability, ruling out statistical discrimination as a driver of backlash.

We next examine the sources of backlash. A natural hypothesis is that backlash would be strongest among participants who oppose gender equity interventions. We find the opposite. Backlash is concentrated among participants who support DEI policies: in the absence of quotas, they are more likely than others to hire women, but once a quota is introduced, they become less likely to do so. This pattern is consistent with moral crowding out: when helping becomes mandatory, those who would have helped voluntarily pull back. A second, distinct source of backlash emerges among participants with gender-biased beliefs who favor merit-based evaluation. These participants respond to quotas by actively raising the bar for non-targeted women, particularly when performance signals are noisy and leave greater room for discretion. Together, these two mechanisms reconcile a puzzling pattern in the empirical literature: negative spillovers from gender quotas have been documented in both progressive societies like California (Gertsberg et al., 2021; Bian et al., 2024), France (Lippmann, 2021; Deschamps, 2024), and Spain (Bagues et al., 2017), where moral crowding out is likely to operate, and traditional ones like South Korea (Lee & Zanella, 2024), and India (Beaman et al., 2009; Karekurve-Ramachandra & Sood, 2025), where backlash among those with gender-biased beliefs might be more prevalent. Backlash might be universal, not because everyone responds the same way, but because different groups exhibit it for different reasons.

Two factors determine the magnitude of quota backlash. First, backlash is stronger when the quota candidate performs poorly. When recruiters are forced to hire a candidate they would not have otherwise chosen, they respond by raising the bar for other women more aggressively. This suggests that female candidates are evaluated jointly: the treatment of women in discretionary positions depends not only on their own performance, but also on the performance of the quota beneficiary. Second, backlash is stronger when performance signals are less precise. In an additional treatment where recruiters observe noisier measures of candidate ability, backlash increases substantially, particularly among participants who believe men outperform women on average. Imprecise signals create more room for discretionary responses, allowing prior beliefs about gender differences to more freely shape hiring decisions. Together, these findings suggest that backlash is most severe precisely in the environments where quotas are most likely to be introduced: those in which candidate quality is hard to observe and evaluators hold strong prior beliefs about the ability of an underrepresented group.

Backlash severely undermines the effectiveness of quotas as a policy to promote female hiring. By design, quotas mechanically increase female hiring in targeted positions. But their aggregate effectiveness depends on how recruiters adjust their behavior in discretionary positions. To quantify this, we construct a mechanical benchmark: the increase in female hiring that would result if recruiters behaved exactly as in the Baseline, with the quota-mandated hire added on top. Backlash offsets 36 percent of these mechanical gains on average. This erosion is substantially larger where backlash is strongest: among lower-performing female candidates, backlash offsets 64 percent of the mechanical gains, and in the worst cases, it fully reverses the policy’s intended effect, leaving women less likely to be hired than in the absence of any quota.

Importantly, backlash is not merely a redistribution of opportunities across women. Recruiters who exhibit backlash forgo significant profits by passing over female candidates they could have hired, thereby reducing their earnings. Because backlash is costly, competitive pressure should discipline it. We test this in the second part of the experiment, where each recruiter competes with two others for the same candidates, so that a hire only occurs if their offer is the highest. Competition substantially attenuates backlash, both by moderating recruiters’ own behavior and by shifting hiring outcomes toward less discriminatory recruiters. As a result, competition restores the effectiveness of quotas: under competitive hiring, backlash offsets only 9 percent of the mechanical gains, compared to 36 percent without competition, suggesting that the degree of competition in a market is a key determinant of whether quota policies achieve their intended goals.

We contribute to a large literature documenting mixed effects of gender quotas on aggregate outcomes for women. We show that these mixed findings reflect the presence and magnitude of quota backlash: a retaliatory response that spills over to women in discretionary positions not covered by the policy. Consistently, [Bian et al. \(2024\)](#) find that firms in California reacted to corporate board quotas by reducing female hiring in other positions, and [Lee and Zanella \(2024\)](#) show that an electoral quota in South Korea led to fewer women being nominated to races not covered by the mandate.<sup>5</sup> The key mechanism is joint evaluation: under a quota, women are no longer assessed solely on their own merits, but also against the performance of quota beneficiaries, raising the bar for the entire group.

This mechanism explains the variation in the literature: backlash is stronger when quota candidates are perceived as underqualified, giving evaluators more reason to retaliate, and when performance signals are imprecise, giving them more room to do so. This is consistent with evidence from public hiring committees ([Bagues et al., 2017](#); [Deschamps, 2024](#); [Mocanu, 2024](#)), and from corporate boards where quotas displace a qualified male incumbent ([Gertsberg et al., 2021](#)). Conversely, when quota candidates are highly qualified or performance is easy to observe, backlash is limited and spillovers can be positive. This is the case of qualified board members ([Bertrand et al., 2019](#)), politicians elected in competitive elections ([Baltrunaite et al., 2014](#); [Lee & Zanella, 2024](#)), and of chess competitions where

---

<sup>5</sup>Notably, [Lee and Zanella \(2024\)](#) find that this pattern reverses with exposure to competent female politicians, consistent with our finding that backlash is strongest when quota beneficiaries are perceived as underqualified.

ELO ratings provide a transparent signal of ability (De Sousa & Niederle, 2022). Finally, we find that backlash arises simultaneously from two distinct groups: progressives who voluntarily support women but pull back once helping becomes mandatory, and evaluators with gender-biased beliefs who actively offset the policy. This heterogeneity explains why negative spillovers have been documented in both progressive and traditional societies alike.

Our design builds on a growing experimental literature on gender gaps that has proven particularly valuable for identifying causal mechanisms, and whose findings have been shown to predict outcomes outside the lab (Buser et al., 2014). One strand of this literature has focused on the beliefs and behavior of workers themselves, documenting robust gender differences in confidence, competitiveness, willingness to negotiate, and self-promotion that shape labor market outcomes (Niederle & Vesterlund, 2007; Bordalo et al., 2019; Exley et al., 2020; Exley & Kessler, 2022). A complementary strand has examined the beliefs of evaluators, showing that recruiters hold inaccurate beliefs about female performance and update them asymmetrically, thereby sustaining gender gaps even in the absence of overt bias (Coffman et al., 2021; Exley & Nielsen, 2024; Bohren, Haggag, et al., 2025; Bohren, Hull, & Imas, 2025). We build on these findings to explore how evaluators’ beliefs interact with equity policies to generate backlash. Unlike previous work, our focus is not on whether evaluators accurately assess female ability, but on how a hiring mandate alters their behavior toward women not covered by the policy, and how this backlash varies with their prior beliefs and policy attitudes.

A growing literature documents how well-intentioned equity policies can backfire due to endogenous responses by firms and evaluators. For example, Cullen and Pakzad-Hurson (2023) show that pay transparency laws reduce wages as firms preempt costly renegotiations. More focused on equity policies, Gentile Passaro et al. (2026) find that equal pay mandates in Chile increased occupational segregation, unintentionally widened the gender pay gap, and Leibbrandt and List (2025) show that equal employment opportunity statements can discourage minority workers from applying to jobs. In the context of education, Otero et al. (2023) show that affirmative action policies can negatively affect outcomes for non-targeted individuals. We add to this literature by showing that gender quotas can elicit backlash from employers and undermine the policy’s effectiveness.

We organize the paper as follows. Section II describes the experimental design. Section III presents our main results on the existence and sources of quota backlash. Section IV examines the effectiveness of quotas and the role of competition, showing that competitive pressure limits backlash and largely restores the policy’s intended gains. Section V rules out alternative explanations such as demand effects, statistical discrimination, and preferences for diversity and demonstrates the robustness of our findings. Section VI discusses how our results help explain the mixed findings in the empirical literature on gender quotas, clarifying when and why quotas succeed or backfire. Section VII concludes.

## 2 Experimental Design

### The Quota Backlash Hypothesis

We describe a simple framework to guide the empirical analysis. Consider a recruiter who evaluates two candidates, one male and one female, each with a signal of productivity  $s_g$ , for gender  $g \in m, f$ . After observing the signal, the recruiter makes a wage offer  $w_g$  to each candidate, and hires a candidate if the offer meets or exceeds their reservation wage  $r_g$ . If hired, the candidate produces output  $p_g$  and receives their reservation wage, with the recruiter keeping the remainder. The recruiter also observes whether a gender quota policy is in place. Recruiters maximize expected profit from hiring, which depends on productivity signals and reservation wages, but may also have preferences over which candidates to hire—for instance, exhibiting taste-based gender bias or valuing workplace diversity—that interact with the presence of a quota.

Let  $w_m^*(Q)$  and  $w_f^*(Q)$  denote optimal wage offers as a function of quota policy  $Q \in \{0, 1\}$ . The baseline wage gap is:

$$\Delta_{\text{Base}} = w_m^*(Q = 0) - w_f^*(Q = 0) \quad (1)$$

This gap may reflect differences in expected productivity, reservation wages, or taste-based preferences in the absence of a quota. Now consider a quota policy that requires the recruiter to hire an additional female candidate at a fixed cost  $w_q$ , with the budget increased by the same amount, leaving discretionary resources unchanged. The optimal offers under the quota are  $w_m^*(Q = 1)$  and  $w_f^*(Q = 1)$ , and the quota wage gap is:

$$\Delta_{\text{Quota}} = w_m^*(Q = 1) - w_f^*(Q = 1) \quad (2)$$

Our central hypothesis, *Quota Backlash*, posits that the quota exacerbates the wage gap:  $\Delta_{\text{Quota}} - \Delta_{\text{Base}} > 0$ . That is, quotas worsen outcomes for women not targeted by the policy. Importantly, this hypothesis does not require a pre-existing baseline gap. If  $\Delta_{\text{Base}} = 0$ , any positive gap under the quota constitutes backlash. Conversely, pre-existing gaps due to statistical or taste-based discrimination do not, on their own, imply backlash unless the gap widens in response to the policy.

**Mechanisms.** Quotas might generate backlash through several channels. First, quotas could change recruiters’ beliefs about the expected productivity or reservation wages of male versus female candidates. In our setting, however, this channel is unlikely to operate. The productivity signal is highly informative and orthogonal to the quota policy. Moreover, candidates are randomly selected from a common pool of workers and are unaware of the policy; therefore, the quota provides no basis for updating beliefs about productivity or reservation wages. Consistent with this, we find no evidence that quotas shift recruiter beliefs.

Second, backlash may arise from changes in recruiters’ preferences. Quotas could interact with existing taste-based motives: for example, recruiters who value workplace diversity may place less weight on hiring an additional woman once one has already been mandated, and recruiters who value autonomy may derive less utility from hiring candidates of the same type as those targeted by the policy.

Third, quotas could affect preferences in more subtle ways. While recruiters might derive utility from supporting under-performing candidates at baseline—a form of warm-glow utility—the mandatory nature of quota hiring could crowd out this intrinsic motivation, reducing their willingness to hire additional women in discretionary positions.

Our experimental design distinguishes among these channels by holding budgets constant and systematically varying candidate performance and the gender of the quota beneficiary. We show that quota backlash exists, is concentrated among decisions where female candidates underperform ( $s_f < s_m$ ), and is not driven by changes in beliefs.

## Design Overview

Our experimental design involves two main types of participants: “workers” and “recruiters.” Workers complete a general knowledge quiz with 20 questions (including verbal and quantitative questions, similar to GRE or SAT standardized tests), earning 10 cents per correct answer. Additionally, we also elicit their preference for being paid based on the quiz or a specific monetary amount, which we refer to as their *reservation wage*.

Recruiters choose which workers to hire across a series of decisions. In each decision, they receive a fixed budget and observe at least two candidates, for each of whom they see gender (proxied by name) and prior performance (the number of correct answers in the first 10 quiz questions). Recruiters then submit an offer to each candidate. A candidate is hired if the offer meets or exceeds their reservation wage, which is unknown to the recruiter; if hired, the recruiter pays the *reservation wage* and keeps the remainder. For each hired candidate, the recruiter earns \$1 per correct answer in the final 10 quiz questions. Offers to candidates who are not hired are simply kept by the recruiter. Since any unused budget is lost, they have no reason not to use all their funds to make offers. In fact, we later use this condition as an indicator of participants’ confusion.

In the *Baseline* treatment, recruiters receive \$7 and choose how to allocate it across two randomly selected candidates. Our main treatment variation introduces a *Quota Candidate*. In the *Quota* treatment, the budget is increased to \$12, but \$5 must be used to hire the quota candidate, leaving recruiters with the same \$7 of discretionary funds as in the Baseline. In this condition, all quota candidates are female, a requirement that we relax in other conditions. In both treatments, recruiters therefore face the same effective choice: how to split \$7 between two non-quota candidates. Our central hypothesis is that being required to hire a female quota candidate will lead recruiters to reduce offers to other female candidates, undermining the policy’s effectiveness.

Table 1: Summary of Experimental Design

|                                                               |                                                                             |
|---------------------------------------------------------------|-----------------------------------------------------------------------------|
| 12 hiring decisions                                           |                                                                             |
| First decision is the same non-quota candidates for everyone  |                                                                             |
| Part I                                                        | Candidates’ gender and performance randomized in the remaining 11 decisions |
| Drawn from a gender-balanced pool of 100 workers              |                                                                             |
| Part II                                                       | 12 hiring decisions, competing with two other participants                  |
| Candidate hired by the recruiter submitting the highest offer |                                                                             |
| End                                                           | Beliefs elicitation, strategy description, demographics, and policy views   |

The experiment consists of two parts. In part I, recruiters face 12 hiring decisions with two non-quota candidates each and without feedback between rounds. Candidates are randomly selected from the worker pool, introducing variation in gender and performance. The only decision not randomized is the first one, in which all participants face the same pair of candidates shown in Figure 1. Part II introduces competition: each recruiter is matched with two others bidding for the same candidates, and a candidate is hired only by the recruiter who submits the highest offer. This creates strategic incentives that could either amplify or dampen gender differences. A recruiter who anticipates that others are discriminating against women may find it profitable to outbid them, while a discriminating recruiter may bid more aggressively to secure their preferred candidate. To ensure response quality, participants must correctly answer comprehension questions before proceeding to hiring and spend at least 15 seconds on each decision. At the end of the study, participants report their strategies, estimate the average performance of male and female candidates, and provide demographic information and views on quota policies.

*Participants and Procedures.*— Participants were recruited through Prolific. We assembled a gender-balanced sample of 1,019 US-based participants with unique IP addresses, all with an approval rate above 90% and at least 100 prior submissions. The sample is also politically balanced, comprising equal shares of self-declared Democrats and Republicans.<sup>6</sup> A total of 408 participants were assigned to the two main treatments and 611 to the additional ones. During the instruction period, participants were required to correctly answer several comprehension questions before proceeding to the main decisions. We record the number of errors each subject makes in the comprehension questions and use this as a control in our analysis. Participants seem to understand the instructions: 47% make no mistakes on the comprehension questions, and 75% make no more than one. All main results are robust to restricting the sample to participants who answer all comprehension questions correctly. Subjects who finish the entire experiment are paid a \$4 participation fee and a bonus based on their performance. One decision in Part I and another in Part II are randomly selected for payment. In addition, subjects are paid 50 cents if they correctly guess the number of correct answers for each group, men, women, and quota candidates, if applicable. The average subject made \$12 and took 30 minutes to complete the survey.

<sup>6</sup>See Table A.3 for demographic characteristics and balance checks.

*Additional Treatments.*— We implement three additional treatments designed to shed light on when backlash arises and through which mechanisms.

In *Quota Male*, the designated quota candidate is male rather than female. This treatment allows us to test whether backlash reflects gender-specific bias or instead stems from a more general aversion to quotas or preferences for diversity.

To examine the role of signal precision and whether it exacerbates quota backlash, we conduct two additional treatments—collectively referred to as *Noisy Signal*—in which recruiters observe candidates’ answers to only two quiz questions rather than the full set of the first ten, yielding a markedly noisier signal of productivity. All other aspects of the experiment remain unchanged, with participants randomly assigned to the Baseline or Quota treatment.

## Why an Experiment

Despite its simplicity, our design enables us to isolate the causal effect of introducing gender quotas while ruling out several confounding factors common in observational settings.

A first challenge in empirical settings is that quotas often reduce opportunities for non-targeted groups unless capacity is simultaneously expanded. For instance, a college that reserves spots for a particular group may do so at the expense of others. Our design avoids this by increasing the budget alongside the quota mandate, leaving discretionary resources unchanged across treatments.

A second challenge is compositional: introducing a quota may alter the pool of candidates who apply, either by encouraging applications from the targeted group or by discouraging them in its absence. Bleemer (2022) documents a decline of 250 Black and 900 Hispanic applicants per year to the UC system following the passage of Prop 209, while De Sousa and Niederle (2022) finds that a gender quota in French chess teams encouraged more women to join. Such changes make it difficult to separate the response to the quota itself from changes in the applicant pool. Our design holds the candidate pool constant across treatments, isolating the causal effect of the mandate on recruiter behavior.

A third challenge is that who directly benefits from a quota is often unclear, since individuals from the targeted group may have qualified independently of the policy. We address this by explicitly identifying the quota beneficiary in each decision, even when all candidates belong to the targeted group. We also provide recruiters with an objective measure of each candidate’s productivity, allowing us to examine how outcomes vary with relative performance and minimizing the scope for statistical discrimination, enabling a cleaner identification of taste-based responses.

Beyond these identification advantages, a key strength of the lab setting is flexibility. Rather than examining a quota policy in a single specific context—as most empirical evidence

necessarily does—our design allows us to freely vary the decision environment in ways that field settings cannot. By systematically varying the level of competition recruiters face, candidates’ performance, and the precision of the performance signals they observe, we can not only document the existence of backlash but also identify what exacerbates it and what can limit its impact. This makes experimental evidence a useful complement to the growing body of empirical work on gender quotas, helping to organize its seemingly puzzling variety of findings under a common framework.

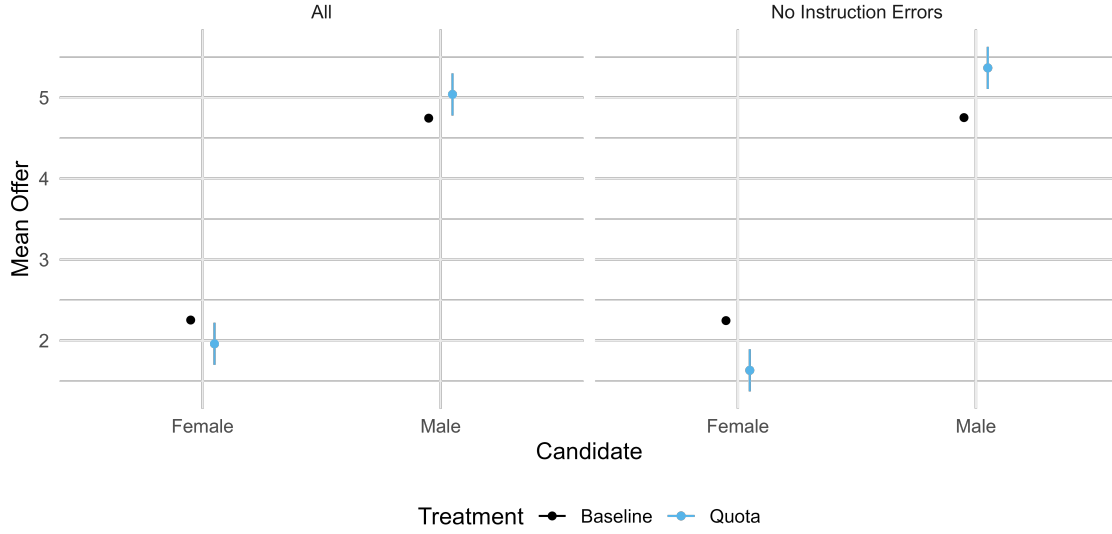
### 3 Main Result - The Backlash of Gender Quotas

This section presents evidence of backlash induced by gender quotas. When participants are required to hire a female quota candidate, they subsequently penalize other women not covered by the policy, and the backlash scales with the quota candidate’s performance. By contrast, we find no evidence of backlash when the quota candidate is male. Together, these findings suggest that female candidates are evaluated jointly, with the performance of one affecting the treatment of others.

We begin by examining the gender gap in the first hiring decision (Figure 2), in which the set of non-quota candidates is identical across participants. In this decision, the male candidate has a clear productivity advantage (seven versus five correct answers), so a gender gap is expected insofar as participants condition wages on performance. We find that introducing a quota candidate substantially amplifies this gap. Specifically, the difference between male and female wage offers increases from \$2.49 in the Baseline to \$3.08 in the Quota Treatment, a 24 percent widening ( $p = 0.028$ ). Restricting attention to participants who made no instruction errors yields a larger effect: the gender gap increases by \$1.23 across treatments ( $p = 0.001$ ).

In absolute terms, female candidates in the first decision of the Quota Treatment receive 86 cents for every dollar earned by women in the Baseline, a figure that falls to 73 cents when restricting the sample to participants without instruction errors. These patterns generalize across all twelve hiring decisions. In decisions where male candidates outperform female candidates, women earn, on average, 90.3 cents on the dollar relative to women in the Baseline ( $p = 0.042$ ), with the gap widening further—to 84.4 cents—among participants without instruction errors ( $p = 0.032$ ). By contrast, when female candidates outperform male candidates or receive equivalent performance signals, we find no evidence of quota-induced backlash, and women earn the same amount in both treatments ( $p = 0.425$ ). Notice that no quota effect for over-performing women implies no backlash against under-performing men, showing that backlash is not directed against all under-performing candidates but specifically against women. Thus, quota backlash exists, but it is not indiscriminate: it is performance-specific and affects only lower-performing female candidates, effectively raising the bar for hiring women.

Figure 2: Mean Offers in the First Decision

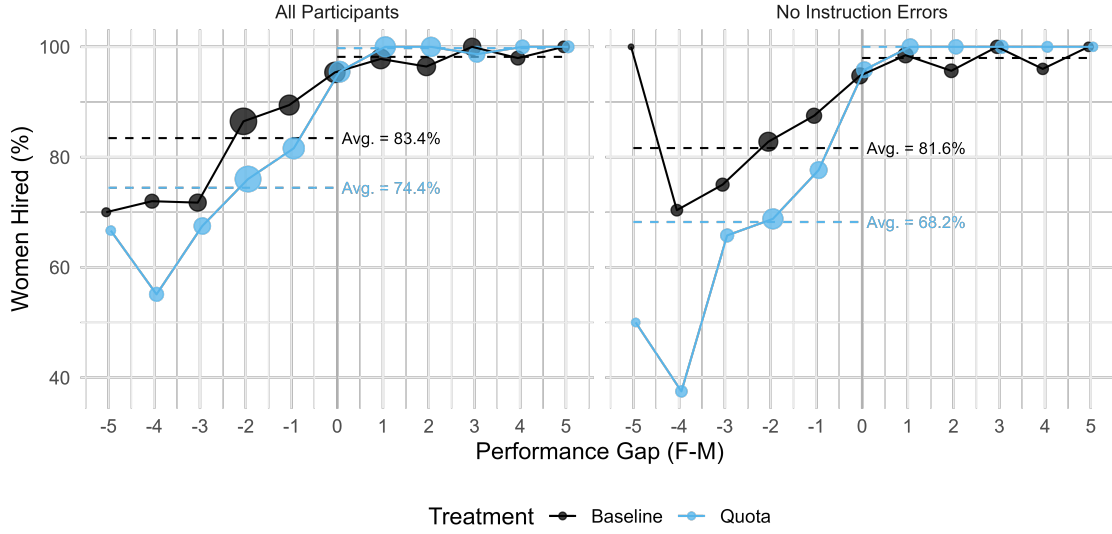


*Notes:* Error bars correspond to the 95% confidence interval of the difference in means after controlling for demographic characteristics and errors in the instructions. Controls include dummies for whether participants did not make any errors in the instructions, female gender, age above the median, white ethnicity, and no college degree. See Appendix Table A.5 for details.

After establishing an effect on the intensive margin—the average offer received by women—we next turn to the extensive margin: whether non-targeted female candidates receive an offer at all, and, conditional on receiving an offer, whether they are ultimately hired. In our setting, offers are automatically accepted if they weakly exceed the worker’s reservation wage, which we elicit independently. We therefore examine two complementary outcomes: the share of zero offers made to women, which guarantees non-hiring, and the share of women hired, which also depends on reservation wages. Reassuringly, both measures tell a consistent story.

As before, we begin with the first hiring decision, in which the male candidate outperforms the female candidate. In both treatments, the male candidate receives an offer in 100 percent of cases, indicating that participants attend to performance signals. In the Baseline, female candidates receive no offer in 11.9 percent of cases. This likelihood nearly doubles in the Quota Treatment, rising to 23.1 percent ( $p = 0.011$ ). Restricting the sample to participants without instruction errors further amplifies this gap. While the share of female candidates receiving no offer in the Baseline remains largely unchanged (13.2 percent), the corresponding share in the Quota Treatment increases to 30.1 percent ( $p = 0.007$ ). Because the reservation wage of the female candidate in this decision is relatively low (\$0.8), the same pattern emerges when we instead examine the share of hired women.

Figure 3: Likelihood of a Woman being Hired, by Performance Gap



*Notes:* Circle size represents sample size. Performance gap refers to the difference in performance between female and male candidates in decisions where candidates are of different genders. Horizontal lines depict the average likelihood of a female worker being hired, divided by whether she is outperformed by a male candidate. The left panel uses only decisions in which the entire budget was allocated and in which the candidates are of different genders. The right panel also restricts to participants who answer all understanding questions correctly. See Appendix Table A.8 for details.

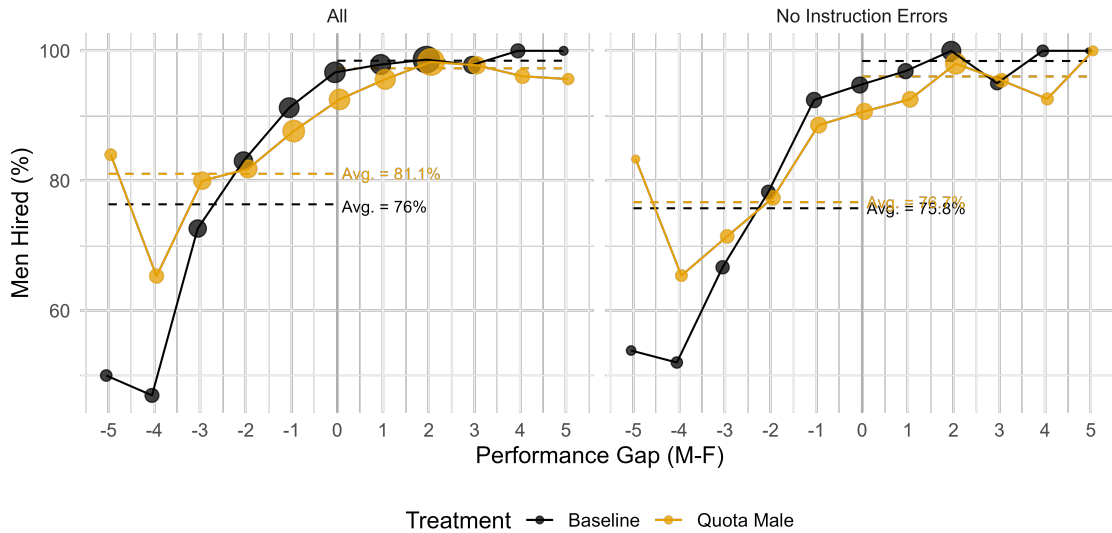
Figure 3 summarizes these effects across all decisions in which participants face two candidates of different genders. We further stratify decisions by the performance gap between candidates, shown on the horizontal axis, with positive values indicating cases in which women outperform men and negative values indicating the opposite. The exacerbated gap induced by the quota in the first decision generalizes across all decisions. While hiring rates are similar across treatments when women are the top-performing candidate, in cases where they are not, women are 9 percentage points less likely to be hired in the Quota Treatment ( $p = 0.009$ ). Consistent with our earlier findings, this backlash is even stronger among participants without instruction errors, with a 13.4 percentage-point reduction in hiring likelihood ( $p = 0.007$ ). Taken together, these extensive-margin effects confirm that quota backlash is performance-specific, affecting only lower-performing female candidates. We next examine how the performance of the quota candidate itself shapes backlash.

When the quota candidate is high performing, we find no evidence of backlash: hiring rates for women are indistinguishable across treatments. By contrast, the treatment gap widens as the quota candidate's performance declines. Figure A.7 plots the likelihood that a woman is hired in the Quota Treatment, stratified by the performance of the quota candidate. These patterns indicate that backlash arises precisely when the quota forces participants to hire a candidate they would not otherwise have hired. When quota candidates are strong, they would likely have been hired absent the policy, and consequently, participants do not adjust their behavior. However, when quota candidates are less exceptional—scoring below

seven out of ten—we observe clear evidence of retaliation against other female candidates. Thus, quota backlash depends not only on the performance of the female candidate under consideration but also on the performance of other female candidates. In contrast to male candidates, who are evaluated largely in isolation, female candidates appear to be judged jointly, with the performance of one affecting the treatment of others.

Same-gender decisions further show that women are penalized even in the absence of male comparison. Figure A.8 presents hiring outcomes as a function of the performance gap within same-gender pairs. We again find substantial backlash against under-performing female candidates: in the Quota Treatment, the likelihood that a low-performing woman receives no offer increases by 10.7 percentage points ( $p = 0.008$ ). By contrast, we find no corresponding effect for under-performing male candidates ( $p = 0.719$ ). That low-performing women are penalized even when evaluated only against other women indicates that quota backlash is not driven by changes in beliefs about female performance relative to men, but instead reflects a broader penalty applied to women following the imposition of a quota.

Figure 4: Likelihood of Hire by Performance Gap – *Quota Male* Treatment



*Notes:* Circle size represents sample size. Performance gap refers to the difference in performance between male candidates and female candidates, in decisions where candidates have different genders. Positive values reflect male candidates with better performance than the female candidates. Horizontal lines depict the average likelihood of receiving zero offers. The left panel uses only decisions in which the entire budget was allocated and in which the candidates are of different genders. The right panel also restricts to participants who answer all understanding questions correctly. See Appendix Table A.13 for details.

Finally, we test whether quota backlash reflects a general aversion to quotas or preferences for diversification, rather than a gender-specific bias, by examining the Quota Male treatment, in which the quota candidate is male. If backlash were driven by quota policies per se, we would also observe retaliation against non-targeted male candidates. We find no such evidence. Figure 4 shows that male quotas do not generate backlash against other male

candidates, in stark contrast to the effects observed when quota candidates are female. The estimated difference in hiring rates is only 5.1 percentage points, statistically insignificant ( $p = 0.163$ ), and opposite in sign to what a backlash interpretation would predict. This null result is even more pronounced among participants without instruction errors: 76.7 percent of low-performing men are hired in Quota Male and 75.6 percent in the Baseline ( $p = 0.862$ ). Taken together, these findings indicate that quota backlash is specific to women and is consistent with the presence of gender bias rather than a general opposition to quota policies.

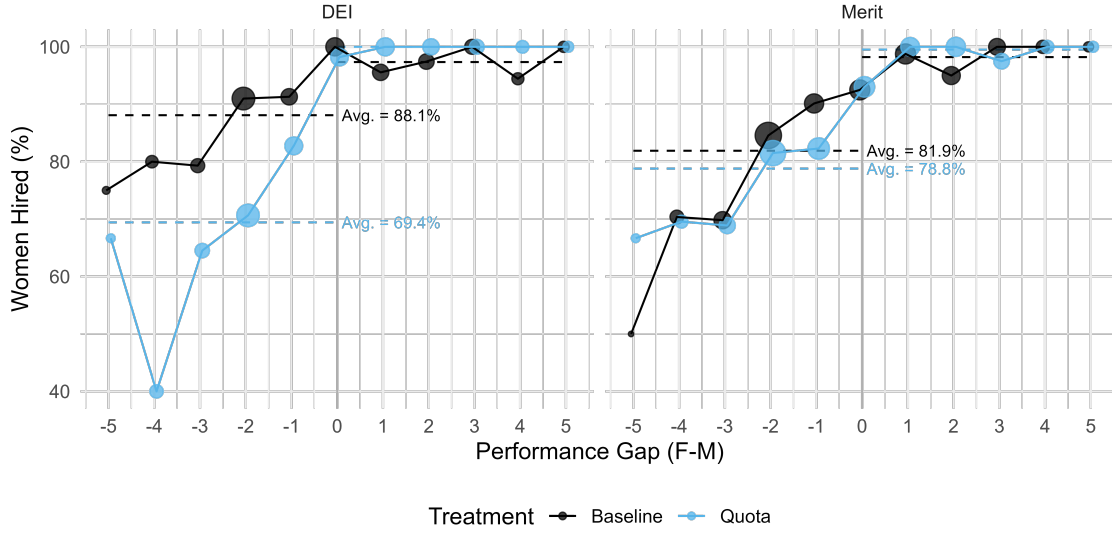
### 3.1 Who is Reacting Negatively to the Quotas

To understand why quotas generate backlash against female candidates, we next examine how our findings vary with the characteristics and beliefs of the participants making the offers. For this analysis, we combine demographic characteristics (such as gender) with self-reported measures of political ideology and support for equity-oriented policies, all elicited at the end of our study. The wording of all questions is provided in Appendix A.4.

A natural hypothesis is that backlash would be strongest among participants who oppose gender-equity interventions. Instead, we find the opposite. Figure 5 presents results on the extensive margin, distinguishing participants who support DEI initiatives in college admissions from those who favor a purely merit-based system. The negative impact of quotas on hiring is concentrated entirely among DEI supporters; participants who prefer merit-only admissions exhibit no backlash. In the absence of quotas, DEI supporters are 6.2 percentage points more likely than others to hire women ( $p = 0.119$ ), suggesting an active willingness to promote female candidates. Once quotas are introduced, however, DEI supporters become 9.4 percentage points less likely to hire women ( $p = 0.077$ ). Thus, DEI supporters switch from favoring female candidates to penalizing them under quotas, whereas merit-based proponents behave similarly across treatments.

These results are robust across alternative survey measures and self-declared political alignment. Figure A.11 in the Appendix examines heterogeneity by demographic characteristics and political views. We find no systematic role for the manager’s gender, and the effects are strongest among participants who answered the instruction-comprehension questions correctly, making confusion an unlikely explanation. Instead, backlash is consistently concentrated among participants who identify as liberal or who express support for equity-oriented policies across multiple survey questions. In other words, regardless of how we classify political attitudes, supporters of DEI-type policies exhibit larger treatment effects.

Figure 5: Heterogeneity by Policy Views - College Admissions



*Notes:* Circle size represents sample size. Performance gap refers to the difference in performance between female and male candidates in decisions where candidates are of different genders. Positive values reflect female candidates that outperform the male candidate. Horizontal lines depict the average likelihood of a female being hired when the male candidate outperforms. The panels divide the sample according to their views on diversity in college admissions: the left panel shows decisions by participants who believe “applicants’ racial and ethnic background should be considered to help promote diversity” (41.5% of the sample), while the right panel shows those who believe “applicants should be admitted solely on the basis of merit” (58.5%). We only use decisions where the entire budget was allocated and when the candidates are of different genders. See Appendix Table A.14 for details.

We next study the role of expectations about candidate performance as a possible channel. One possibility is that the introduction of quotas changes beliefs about women’s productivity—for example, that quotas signal female candidates must otherwise be weaker. However, we find no evidence that quotas shift beliefs: expected performance of male and female candidates is statistically indistinguishable across treatments (Appendix Table A.24). If anything, beliefs move weakly in the opposite direction, though the effects are small and insignificant.

Beliefs still matter, however, for understanding who displays backlash. We re-estimate our DEI support results, conditioning on whether participants believe that women candidates, on average, perform at least as well as male candidates. Among participants who believe that women are at least as good as men<sup>7</sup>, DEI supporters again show the same pattern: higher hiring of women in the Baseline, but substantially lower hiring in the Quota Treatment (Appendix Figure A.10). DEI supporters also display some backlash even when they believe men perform better, though the effect is attenuated. Interestingly, the only circumstance in which merit-based participants exhibit backlash is when they believe that

<sup>7</sup>About 67 percent of the sample believe women to be at least as good candidates as men. This fraction is the identical across DEI supporters and detractors.

women perform worse than men. Taken together, these patterns suggest that backlash arises from two distinct groups. The first consists of participants who support gender-equity policies but reduce their support once the policy becomes mandatory, consistent with a form of moral crowding out. The second consists of merit-oriented participants who also believe women perform worse and respond to quotas by actively offsetting them. Quantitatively, however, the first group accounts for most of the aggregate effect, as favoring gender policies consistently predicts larger treatment effects (Appendix Figure A.11).

### 3.1.1 Noisy Measures of Performance

That performance beliefs play only a limited role is not surprising in our setting, given that performance signals are highly informative and allow participants to easily distinguish high- and low-performing candidates. However, many real-world contexts in which quotas are implemented lack such precise ability signals. For example, while De Sousa and Niederle (2022) examine the effect of quotas in chess, where each player has a precise ELO rating, evaluators in other settings—such as scientific hiring committees (Bagues et al., 2017; Deschamps, 2024)—must rely on more subjective assessments.

To examine the role of signal precision and whether it exacerbates quota backlash, we conduct two additional treatments—collectively referred to as Noisy Signal—in which recruiters observe a substantially less informative measure of candidate performance. In these treatments, participants observe candidates’ answers to only two quiz questions rather than the full set of the first ten, yielding a markedly noisier signal of productivity. All other aspects of the experiment remain unchanged, with participants randomly assigned to the Baseline or Quota treatment.

Under less informative signals, quota backlash is not only robust but also stronger in magnitude. Appendix Figure A.12 shows that quotas reduce the hiring of under-performing women by 12.9 percentage points ( $p = 0.001$ ) and Appendix Table A.22 replicates our earlier patterns: backlash is directed exclusively against lower-performing women, with no corresponding effects for other women or for men. Importantly, Appendix Figure A.13 reveals that when performance signals are noisier, backlash becomes substantially more pronounced among merit-oriented participants. For these participants, quotas interact with prior beliefs about gender differences in performance. Among merit-based policy supporters who also believe that men outperform women on average, quotas generate a 26.9 percentage-point decline in the likelihood of hiring under-performing female candidates ( $p < 0.001$ ). By comparison, the corresponding effect is 15.1 percentage points when signals are precise. These findings suggest that precise performance signals constrain backlash, whereas noisier informational environments create greater scope for discretionary responses.

Overall, our heterogeneity analyses indicate that quota backlash is not primarily driven by overt hostility toward women or by changes in performance beliefs. Instead, it reflects how quotas reshape the behavior of participants who would otherwise help female candidates. Participants who express support for equity-oriented policies exhibit the strongest

behavioral responses, reducing female hiring once the policy becomes mandatory. This pattern is consistent with a crowding-out mechanism in which externally imposed constraints weaken previously voluntary behavior. By contrast, backlash among merit-oriented participants emerges only when quotas interact with prior beliefs that women perform worse, and intensifies when performance signals become noisier. Taken together, these findings highlight the interaction between backlash, discretion, and prior beliefs, and show that policy-induced backlash can arise even among individuals who are otherwise positively predisposed toward the targeted group.

## 4 The Effectiveness of Gender Quotas

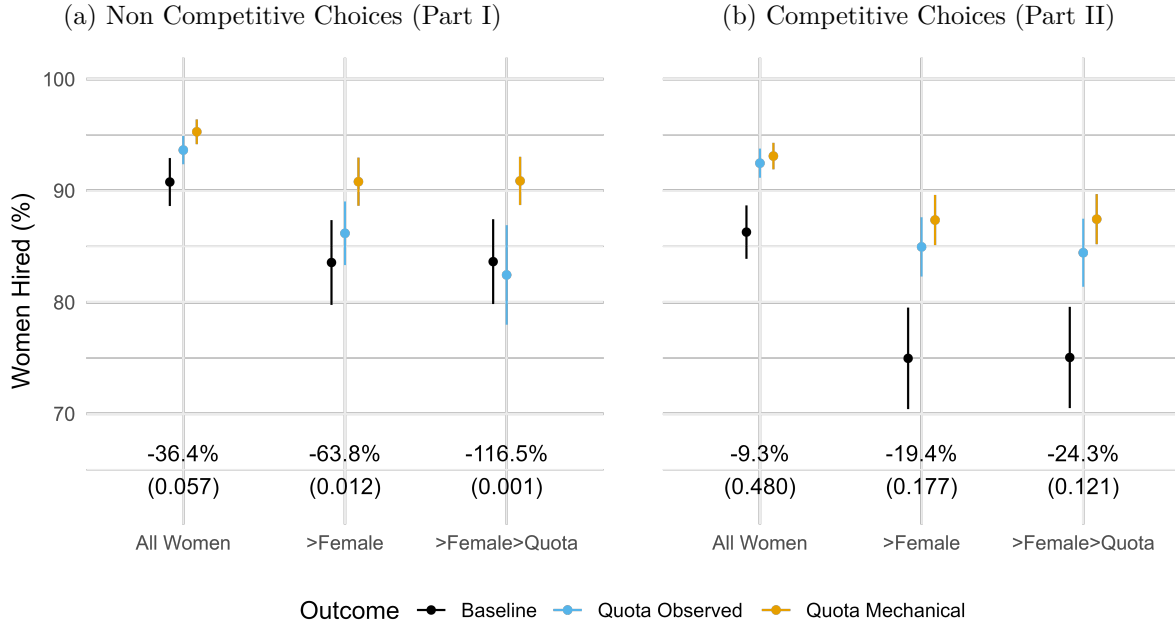
Partial quotas —policies mandating the hiring of some women while leaving discretion over others— mechanically increase the hiring of targeted women. Their overall effectiveness, however, depends on how recruiters adjust their behavior in the remaining discretionary positions. We have shown that many recruiters exhibit backlash, reducing the hiring of women in positions not directly covered by the quota. In this section, we examine the implications of this response for the policy’s overall effectiveness. We find that backlash against non-targeted women offsets a substantial share of the policy’s mechanical gains, leading to inefficient hiring decisions. Importantly, this erosion is not inevitable: because backlash entails real costs for recruiters, increased competition disciplines these responses and largely restores the effectiveness of quotas. We begin by quantifying how backlash erodes the mechanical gains of quotas before turning to its efficiency consequences and the role of competition.

To assess the effectiveness of quotas, we evaluate how backlash against non-targeted women offsets the gains generated by quota-mandated hires. Specifically, we compare three scenarios. First, we use our Baseline as a benchmark for recruiter behavior in the absence of any quota policy. Second, we construct a mechanical benchmark—referred to as Quota Mechanical—that captures how female hiring would change if quotas were imposed, but recruiters continued to behave as in the Baseline.<sup>8</sup> Finally, we measure the aggregate effect of the policy—the Quota Observed outcome—by computing the overall probability that a woman is hired under the Quota treatment, including both quota-mandated hires and discretionary decisions. The difference between Quota Mechanical and Quota Observed reveals the extent to which backlash undermines the policy’s effectiveness.

---

<sup>8</sup>We obtain this benchmark by adding a quota-mandated female hire to the average likelihood of hiring a female candidate in the Baseline treatment.

Figure 6: Likelihood of Female Candidates Being Hired



*Notes:* Error bars correspond to 95% confidence interval of the mean, with standard errors clustered at the participant level and controlling for whether participants did not make any errors in the instructions, female gender, age above the median, white ethnicity, and no college degree. The figures in the graph represent the effectiveness loss due to backlash —how much of the mechanical gains is lost, along with the p-value of a t-test of equality between the Quota Mechanical and the Quota Observed outcomes in parentheses. The left panel shows the share of women hired when hiring is not competitive (Part I), and the right panel when hiring is competitive (Part II). The first column shows the effects among all women, the second column restricts the sample to under-performing female candidates, and the third column further restricts it to decisions with under-performing quota candidates. The share of women hired in the Quota Observed and Quota Mechanical outcomes is calculated by adding the quota-mandated \$5 offer to an additional female candidate for each decision in the *Quota* and *Baseline* treatment, respectively. The Quota Mechanical should be interpreted as the likelihood of receiving an offer if the introduction of quotas triggered no backlash on non-targeted candidates. Both panels only use decisions where the entire budget was allocated. See Appendix Table A.17 for details.

Figure 6a summarizes the effectiveness of partial quotas. At Baseline, the average female candidate has a 90.8% probability of being hired. Under the Quota Mechanical benchmark—holding discretionary behavior fixed—female hiring would increase by 4.5 percentage points. In Quota Observed, however, recruiters reduce female hiring in discretionary positions, so overall female hiring rises by only 2.8 percentage points. Backlash, therefore, offsets 36.4% of the policy’s mechanical gains ( $p = 0.057$ ). This erosion is substantially larger where backlash is strongest. Among under-performing female candidates, precisely the cases where quotas have the greatest potential to increase hiring, the mechanical benchmark implies a 7.2 percentage-point increase, yet backlash offsets 63.8% of these gains ( $p = 0.012$ ). Restricting attention to decisions in which the mandated quota candidate is also low-performing, backlash fully offsets the mechanical effect and even reverses the sign of the policy’s impact: women become *less* likely to be hired relative to Baseline. Taken together, these results show

that quotas induce endogenous responses that significantly erode their intended effects.

This erosion reflects more than a redistribution of hiring opportunities across women. Backlash leads recruiters to overlook qualified female candidates, resulting in inefficient outcomes and reduced recruiter profits. We quantify these costs by measuring forgone profits relative to profit-maximizing offers, considering only discretionary candidates. Appendix Table A.18 shows that, in decisions with at least one female candidate, recruiters in the Quota treatment leave 24% more money on the table than in Baseline ( $p = 0.123$ ). These profit losses are driven by backlash rather than by the increased decision complexity of adding the quota candidate. First, forgone profits are largest precisely where backlash is strongest: recruiters leave 36% more money on the table when the non-targeted female candidate underperforms ( $p = 0.056$ ), and 52% more when the mandated quota candidate has a particularly low signal ( $p = 0.021$ ). Second, quotas do not affect forgone profits in decisions involving only male candidates, that is, decisions with no scope for backlash ( $p = 0.919$ ). Together, these results indicate that backlash represents a costly response for recruiters, rather than an optimal adjustment to the policy.

Because backlash is costly, competitive pressure should limit its scope. When recruiters face tighter competition, the margin for sustaining inefficient hiring decisions shrinks. To test this prediction, we introduce competition in Part II by randomly matching each recruiter with two others. A recruiter hires a worker only if their offer is the highest within the matched group. In this environment, competition may mitigate backlash through two complementary channels. First, it disciplines recruiters' own behavior. In our initial design, recruiters have monopsonistic power. Under competition, the cost of favoring a preferred male candidate increases, as doing so may require a higher—and more costly—offer to outbid competitors. Second, competition alters how individual offers translate into outcomes. When multiple recruiters bid for the same candidates, hiring outcomes increasingly reflect the behavior of the recruiter making the highest offer, reducing the influence of those who exhibit stronger backlash.<sup>9</sup> Together, these channels imply that greater competition should mitigate backlash and restore the effectiveness of quota policies.

Competition substantially attenuates quota backlash, with both recruiters' own behavior and competitive pressures playing important roles. We first examine recruiters' behavior by focusing on offers made in Part II rather than hiring decisions, which mechanically depend on matching. Backlash against underperforming women is markedly reduced. In Part I, quotas decrease the likelihood that underperforming women receive an offer by 9 percentage points. In Part II, this gap nearly halves, falling to 5.6 percentage points ( $p = 0.150$ ). Learning over the course of the experiment does not explain this difference, as backlash does not vary significantly across rounds in Part I (Appendix Table A.21).

We next examine hiring outcomes under competition. For each decision, we randomly construct 2,000 recruiter triads and assign each worker the highest offer received within the

---

<sup>9</sup>This mechanism corresponds to the Becker hypothesis (Becker, 2010), which predicts that discrimination creates profit opportunities for non-discriminating recruiters; as competition intensifies, these opportunities are exploited, improving outcomes for the discriminated group.

triad, so that hiring outcomes reflect the behavior of the marginal recruiter. Under this allocation rule, the consequences of backlash are substantially reduced: quotas lower the probability that underperforming women are hired by only 1.1 percentage points. These results indicate that competition disciplines backlash by both moderating recruiters’ behavior and shifting outcomes toward recruiters exhibiting weaker backlash. The two channels contribute approximately equally.

Table 2: Decomposition of the effect of competition

|                         | Hired                 |                    |                               |
|-------------------------|-----------------------|--------------------|-------------------------------|
|                         | No Competition<br>(1) | Competition<br>(2) | Competition + Matching<br>(3) |
| Constant                | 0.834<br>(0.000)      | 0.743<br>(0.000)   | 0.969<br>(0.000)              |
| Quota                   | -0.090<br>(0.009)     | -0.056<br>(0.150)  | -0.011<br>(0.092)             |
| Observations            | 1,171                 | 1,146              | 6                             |
| Dependent variable mean | 0.790                 | 0.715              | 0.963                         |

*Notes:* Results from a linear regression with  $p$ -values in parentheses. Standard errors are clustered at the individual level in Columns (1) and (2), and at the candidate pair level in Column (3). The dependent variable is a dummy that equals 1 if the female candidate received an offer higher or equal than her reservation wage. Quota is a dummy that equals 1 if the participant was in the *Quota Treatment*. Column (1) reports the results for Part I, where recruiters don’t compete in offers. Column (2) reports the non-matched results for Part II, this is, whether the recruiters’ offers are greater or equal than the reservation wage without matching recruiters with each other. Column (3) reports the matched results for Part II, where each recruiter is randomly matched with two others. We simulate 2,000 random matches of 3 recruiters for each candidate pair, calculate the average hire probability across simulations for each pair and treatment, and use this average as the dependent variable. Note the small number (3) of clusters in Column (3): we present the result only to illustrate the magnitude of the Quota treatment.

Because competition limits backlash, it also restores the effectiveness of quotas. Figure 6b repeats the effectiveness analysis under competition, focusing on the average recruiter’s behavior and abstracting from recruiter matching. At Baseline, the average woman is hired in 86.3% of decisions. In the absence of backlash, the quota’s mechanical effect would increase this likelihood by 6.8 percentage points. Under competition, quotas increase female hiring by 6.2 percentage points, implying that backlash erodes only 9.3% of the policy’s intended gains ( $p = 0.480$ ). Even in decisions that exhibit the strongest backlash in Part I, competition substantially attenuates its effects. This pattern also emerges in the matched sample of recruiters (Figure A.14), where competition similarly reduces efficiency losses.

Taken together, our findings highlight that the success or failure of quotas depends critically on market structure. In environments with substantial discretion, quotas may trigger costly behavioral responses that undermine their intended effects. In competitive

settings, these distortions are constrained by market discipline: competition limits backlash and preserves the policy’s ability to reshape outcomes.

## 5 Robustness and alternative explanations

In this section, we rule out alternative explanations for our results and present a robustness analysis using a treatment with male quota candidates.

### Are our results driven by confusion?

When interpreting our findings, we rely on the assumption that participants understand the experimental decision problem equally well in both treatments. To minimize differential confusion across treatments, we limited differences between instructions and included a comprehensive set of understanding questions. Almost 47% of participants answered all questions correctly, and 75% make no more than one mistake, with understanding errors in the instructions balanced across treatments<sup>10</sup>. Notably, all our results remain or become even more pronounced when restricting the sample to those who answer all questions correctly.

### Are our results driven by statistical discrimination?

Recruiters might have different productivity expectations for men and women. However, as discussed in Section 2, statistical discrimination can only generate quota backlash if quotas lower beliefs about women’s performance relative to men, this is, if quotas amplify statistical discrimination. For example, recruiters might take quotas as a negative signal of female productivity and revise downwards their beliefs. Because participants observe precise performance signals and know that the candidate pool is unaffected by the quotas, it’s unlikely that the quota policy lowers beliefs about women in our experiment. Indeed, we find no evidence of statistical discrimination driving quota backlash. We elicit recruiters’ beliefs about performance across genders at the end of the experiment, and find that quotas improve belief accuracy and close the small gender gap in expected performance that arises at baseline<sup>11</sup>. Thus, if anything, quotas in our experimental labor market *reduce* statistical discrimination against women, ruling out this mechanism as the cause of the quota backlash.

---

<sup>10</sup>In our main treatments, the average number of mistakes in the instructions is 1.12 for the *Baseline* and 1.08 for the *Quota Treatment* (p-value = 0.826).

<sup>11</sup>Predicted performance of female candidates is 5.55 in the *Baseline* and 5.89 in the *Quota Treatment*. In comparison, average predictions for males are 5.66 and 5.84, respectively. The ground truth for both genders is 6. See Appendix Table A.24.

## Are our results driven by demand effects?

The quota policy in the *Quota treatment* is intentionally salient. A potential explanation for our results is that quotas make participants realize this is an experiment about gender, and they respond to experimenter demands by following what they perceive as social norms. We believe this mechanism to be unlikely to drive our results for several reasons. First, if participants believe the study is designed to show that quotas are beneficial or harmful, we would expect effects to emerge regardless of whether the quota candidate is male or female; instead, backlash arises specifically when the quota candidate is female, inconsistent with a demand-driven account. Second, participants driving the results should be those who oppose quotas and gender diversity, yet we observe the opposite: backlash is concentrated among participants who support affirmative action in hiring and college admissions. Finally, all decisions in our experiment are incentivized and, as shown in Section 4, exhibiting backlash is costly for recruiters. This further mitigates demand effects concerns, as conforming to perceived experimenter expectations comes at a real financial cost to participants.

## Are our results driven by preferences for diversity or autonomy?

Quota backlash can arise from preferences for workplace diversity, which place less value on an additional female worker if one has already been hired through the quota mandate. Preferences for autonomy can also explain backlash: recruiters might derive less utility from candidates of the same type as those they’re mandated to hire. These mechanisms have been hypothesized by previous research (Meyerinck et al., 2019), but have not been tested directly: quotas targeting men are rare, so responses to the policy can’t usually be separated from gender bias. We directly test these hypotheses by implementing an additional treatment (*Quota Male*) where recruiters are mandated to hire a male candidate. If preferences for diversity or autonomy explain backlash, recruiters in this treatment should exhibit backlash against male candidates in discretionary positions. Instead, we find no backlash against men in the *Quota Male* treatment. Figure 4 shows that men become, if anything, slightly more likely to be hired for discretionary positions when they are the target of the quota policy. These results suggest preferences for diversity or autonomy are not the drivers of quota backlash.

## Are our results driven by expectations on reservation wages?

A priori, quotas could change recruiters’ beliefs about reservation wages of women, depressing offers and generating backlash <sup>12</sup>. Although unlikely in our setting, where the candidate

---

<sup>12</sup>It’s theoretically unclear whether a change in beliefs about reservation wages would generate backlash. Lower reservation wages have two opposing effects on optimal offers: candidates with lower wages are at the same time more likely to be hired at a given offer (which depresses offers) and less costly at the margin (which incentivizes higher offers).

pool is unaffected by the quotas, recruiters might adapt their offers to their expectations of reservation wages. However, we find no evidence that participants even consider reservation wages when making offers. When asked what strategy they followed to make offers, virtually all participants report adjusting their offers to performance, while only one participant out of 1,019 mentions reservation wages at all. This is a sensible strategy, as reservation wages are not observable to participants. Decisions in the *Baseline Treatment* back the strategies reported by participants. Table A.4 shows that recruiters adjust their offers by candidates' performance. Importantly, although recruiters could (accurately) expect a gender gap in reservation wages, they make equal offers to men and women, suggesting lack of sophistication about wages. Therefore, it's unlikely that changes in expected reservation wages explain quota backlash.

## **Are our results robust to recruiters gaining more experience?**

Competitive hiring is only introduced after twelve decisions, so the muted backlash in Part II could be due to recruiters gaining experience with the task rather competition. We find no evidence to support this alternative explanation. Table A.20 shows that quota backlash in the last three rounds of Part I is of similar magnitude (8.6 percentage points) than across all of Part I (8.7 percentage points with that same specification). Furthermore, Table A.21 finds no significant change in quota backlash over the twelve decisions of Part I, neither when focusing on offers nor on hiring outcomes. In other words, we find no pre-trends before the introduction of competitive hiring, so the results in Part II can be attributed to competition.

## **Are our results robust to noisier performance signals?**

When making their offers, recruiters see each candidate's performance in the first ten questions of the quiz. This signal is highly informative of the productivity of the candidate, which is their performance in the last ten questions. Arguably, quality signals in labor markets or politics are not as predictive of future performance. To test whether backlash extends to settings with less informative signals, we implement an additional set of treatments called Noisy Signal. In these, recruiters only see the candidates' performance in two of the first ten questions, which is much less predictive of productivity. The rest of the experiment remains the same, and participants are randomly assigned to the Baseline or the Quota treatment. Under noisier signals, Table A.22 finds that quotas reduce the likelihood of an under-performing woman being hired for discretionary position by 12.9 percentage points, with no effects for any male candidates. Quota backlash is thus not only robust to noisier performance signals, but of greater magnitude once when recruiters have more discretion to interpret signals.

## 6 Discussion

Gender quotas, by design, increase female representation among positions directly targeted by the policy. In practice, however, quotas are typically partial rather than exhaustive, leaving scope for endogenous responses in positions not covered by the mandate. Electoral quotas often cover only specific races, while corporate quotas typically target boards but not lower-ranked positions. In these settings, the central puzzle is not whether quotas raise representation among targeted candidates, but why they sometimes fail to generate broader gains. In this section, we use our experimental results as a lens to organize the empirical literature, highlighting how backlash against non-targeted women, candidate quality, information environments, and competitive pressure shape the effectiveness of quota policies.

**Backlash as a Negative Spillover.** Our results indicate that partial quotas can trigger backlash, leading to negative spillovers for women in positions not covered by the policy. Importantly, this backlash is not uniform: quotas raise the bar for evaluating other women, disproportionately affecting under-performing candidates. This pattern helps organize a range of empirical findings documenting limited or negative spillovers from gender quotas. For instance, Lee and Zanella (2024) show that political parties in South Korea respond to electoral quotas by nominating fewer women to races not covered by the policy, offsetting aggregate gains. Similarly, Lippmann (2021) document that French parties counteract gender quotas by placing mandated women in non-winnable districts. Beyond politics, Bian et al. (2024) shows that firms in California reduced demand for female labor in lower-ranked positions following board quota mandates. Backlash in empirical settings appears to operate by raising standards for non-targeted women. For instance, gender quotas in scientific hiring committees in France, Italy, and Spain raised standards for female candidates, resulting in fewer women being hired (Bagues et al., 2017; Deschamps, 2024). More broadly, our results contribute to a growing body of evidence showing that well-intentioned gender equity policies can generate unintended and sometimes counterproductive responses (Cullen & Pakzad-Hurson, 2023; Gentile Passaro et al., 2026).

**When Does Backlash Arise?** While backlash often takes the form of raising standards for non-targeted women, our results also clarify when such responses are most likely to arise. Backlash arises primarily when the quota significantly changes the evaluator’s hiring decision—namely, when recruiters are mandated to hire lower-performing quota candidates they would not have selected otherwise. In contrast, when quotas target highly qualified women who would have been hired even in the absence of the policy, the quota is effectively non-binding and generates little backlash. This distinction helps reconcile contrasting findings in the empirical literature. In settings where quotas apply to highly qualified candidates—such as competitive electoral lists—quotas have increased female representation and often generated positive spillovers (De Paola et al., 2010; Baltrunaite et al., 2014; Besley et al., 2017; Ferrari et al., 2022). By contrast, backlash tends to arise when quotas displace higher-experience incumbents or mandate the selection of candidates perceived as under-performing. Consistent with this pattern, Gertsberg et al. (2021) show that backlash against board quotas in California is concentrated among firms where a more experienced male direc-

tor is displaced, Lee and Zanella (2024) document stronger backlash when quota politicians have lower observable qualifications, and Karekurve-Ramachandra and Sood (2025) show low-“skill” quota politicians in India had null effects for female representation. Overall, these findings suggest that quotas are most effective when targeted candidates have higher qualifications.

**The Role of Information and Discretion.** Beyond candidate quality, the information environment in which quotas are implemented plays a central role in determining the scope for backlash. Even when quotas target high-quality candidates, they are more likely to trigger backlash when evaluators face noisy performance signals and retain greater discretion. Consistent with this logic, we find that backlash is stronger when signals of performance are less informative, giving evaluators more room to exercise retaliatory responses. This pattern helps interpret evidence from scientific hiring committees, where candidate quality is difficult to observe and quotas have raised standards for female candidates, reducing their hiring probabilities (Bagues et al., 2017). Similarly, Mocanu (2024) show that discrimination against women is more pronounced when public hiring evaluations rely on less precise metrics. In contrast, when performance signals are very precise, evaluators have less discretion to adjust their behavior in response to quotas, and backlash is correspondingly weaker. In line with this prediction, gender quotas in French competitive chess—a setting in which ability is very precisely measured by the ELO rating—generated positive spillovers for women (De Sousa & Niederle, 2022). Taken together, these findings suggest that quotas are most effective when the ability of candidates can be measured with sufficient precision to limit discretionary backlash.

**Who Exhibits Backlash and Why?** Our experiment highlights heterogeneity in who responds to quotas with backlash and through which channels. Unlike observational settings, our design allows us to measure recruiters’ beliefs and policy attitudes, revealing two distinct responses. One group of recruiters tends to believe that women are at least as capable as men and supports diversity policies. Among them, backlash appears to operate through moral crowding out: while they actively support lower-performing women in the absence of quotas, this discretionary support reverses once assistance becomes mandated. A second group of recruiters is more likely to believe that men outperform women on average and to emphasize meritocratic principles over equity considerations. These recruiters respond to quotas by raising standards for non-targeted women, particularly in environments with noisy performance signals that allow greater discretion. This heterogeneity helps interpret why quota backlash has been documented across very different institutional and cultural contexts. Moral crowding out among pro-equity evaluators offers a lens to understand backlash observed in relatively progressive settings such as California, Norway, France, and Spain (Ahern & Dittmar, 2012; Bagues & Campa, 2021; Gertsberg et al., 2021; Bian et al., 2024). At the same time, merit-based backlash among evaluators with gender-biased beliefs aligns with evidence from more traditional contexts, including South Korea, Lesotho, and rural India (Beaman et al., 2009; Clayton & Tang, 2018; Lee & Zanella, 2024; Green et al., 2026). Together, these patterns suggest that backlash can arise from distinct motivations, helping explain why gender quotas have encountered resistance across a wide range of social and institutional environments.

**Competition as a Market-Driven Solution.** Finally, our results point to competition as a key institutional feature that limits quota backlash and restores policy effectiveness. Because backlash entails real costs for evaluators, competitive pressure reduces the scope for sustaining such responses. This insight helps interpret why quotas tend to perform better in settings where decision makers face strong incentives to select high-performing candidates. In scientific hiring committees, where evaluators often face little direct competition and bear few costs from biased decisions, quotas have generated substantial backlash (Bagues et al., 2017). By contrast, in highly competitive environments, such as professional chess or closely contested electoral races, quotas have increased female representation with limited adverse spillovers (De Sousa & Niederle, 2022; Lee & Zanella, 2024). More broadly, these findings highlight that quota policies are most likely to succeed when embedded in environments that limit discretionary backlash through competitive pressure.

## 7 Conclusion

This paper provides experimental evidence for the existence of backlash against gender quotas and traces its consequences for policy effectiveness. We find that mandating the hire of a female candidate significantly reduces the likelihood that other women, those not targeted by the quota, receive offers. This backlash is concentrated among lower-performing candidates, indicating that quotas raise the bar for evaluating women rather than simply redistributing opportunities.

Our results offer a cautionary tale about equity-promoting policies. Endogenous responses can undermine well-intended mandates and ultimately produce the opposite of their intended effects. As our heterogeneity results show, backlash is not confined to those who hold negative views of women. It arises among conservatives with gender-biased beliefs, but also among progressives who voluntarily support women in the absence of a mandate. Although the mechanisms differ, moral crowding out vs. active offsetting, the eventual outcome is the same: women not covered by the policy end up worse off. These results align with a growing body of evidence showing that well-intentioned gender equity policies can generate unintended and sometimes counterproductive responses (Cullen & Pakzad-Hurson, 2023; Gentile Passaro et al., 2026).

Although the effects we document are sizable—quotas increase the likelihood that non-beneficiary women receive no offer by 50 percent—the extent to which they generalize will likely vary depending on context. Numerous empirical studies already point in the same direction, documenting quota backlash across countries and institutional settings (Bagues et al., 2017; Gertsberg et al., 2021; Lippmann, 2021; Bian et al., 2024; Lee & Zanella, 2024). We believe the value of our experimental evidence lies not in its magnitude alone but in what it reveals about the underlying mechanisms. One key finding is that backlash interacts with performance expectations and prior beliefs about gender gaps: quotas do not affect all women equally, but disproportionately harm those who are already perceived as lower-performing. A distinct advantage of the lab setting is that, rather than examining a quota

policy in a specific context—as most empirical evidence necessarily does—our design allows us to freely vary the decision environment in ways that field settings cannot. For example, by controlling the level of competition or the candidates’ performance. This allows us not only to document the existence of backlash but also to identify what exacerbates it and what can limit its impact. This makes experimental evidence a useful complement to the growing body of empirical work on gender quotas, helping to organize its seemingly puzzling variety of findings under a common framework.

Our results show that once gender quotas are introduced, women face heightened scrutiny: they are evaluated not only on their own performance, but also on the performance of those who directly benefit from the policy. The reaction to gender quotas therefore hinges critically on the perceived competence of quota beneficiaries.<sup>13</sup> This challenge is compounded in real-world settings where performance is noisy or difficult to observe, creating greater scope for distorted beliefs. As shown in Bohren, Haggag, et al. (2025), such misperceptions can play a central role in discriminatory behavior, and our results suggest that backlash is most severe precisely when women’s abilities are hardest to assess—the same environments where quotas are most likely to be introduced. Taken together, these findings point to two complementary strategies for mitigating backlash. First, quota policies should pay careful attention to how those directly targeted are perceived, and where possible pair representation goals with a clear emphasis on the qualifications of selected candidates. One example is the Texas top 10% rule, which guarantees acceptance to students who graduate at the top of their high school class (Black et al., 2023).<sup>14</sup> Because high schools vary in their demographic composition, this policy promotes diversity while foregrounding the strong performance of those admitted. Second, correcting biased beliefs about women’s abilities—even without changing the policy itself—may be an effective complementary strategy for reducing backlash.

Finally, our results on competition identify an important dimension for predicting where quotas will succeed and where they will backfire. In settings where evaluators face little direct competition and bear few costs from biased decisions, quotas are more likely to generate adverse responses, consistent with evidence from scientific hiring committees and public hiring (Bagues et al., 2017; Deschamps, 2024; Mocanu, 2024). By contrast, highly competitive environments appear to limit backlash and preserve the policy’s intended effects (De Sousa & Niederle, 2022; Lee & Zanella, 2024). The degree of competition in a given setting is therefore a crucial dimension to consider when deciding where gender quotas should be implemented. More broadly, our findings underscore that to design better policies for promoting equity, it is not enough to study their intended effects; it is equally important to anticipate and understand their unintended ones.

---

<sup>13</sup>Bohren et al. (2019) provide an example where female discrimination reverses after a sequence of positive evaluations in a large online field experiment.

<sup>14</sup>Students must graduate in the top 10% of their class for automatic admission to the UT system, except at UT-Austin, where the threshold is the top 5%.

## References

- Coate, S., & Loury, G. C. (1993). Will Affirmative-Action Policies Eliminate Negative Stereotypes? *American Economic Review*, 83(5), 1220–1240.
- Niederle, M., & Vesterlund, L. (2007). Do Women Shy Away From Competition? Do Men Compete Too Much? *The Quarterly Journal of Economics*, 122(3), 1067–1101.
- Beaman, L., Chattopadhyay, R., Duflo, E., Pande, R., & Topalova, P. (2009). Powerful Women: Does Exposure Reduce Bias? *The Quarterly Journal of Economics*, 124(4), 1497–1540.
- Becker, G. S. (2010). *The economics of discrimination*. University of Chicago press.
- De Paola, M., Scoppa, V., & Lombardo, R. (2010). Can gender quotas break down negative stereotypes? Evidence from changes in electoral rules. *Journal of Public Economics*, 94(5), 344–353.
- Ahern, K. R., & Dittmar, A. K. (2012). The Changing of the Boards: The Impact on Firm Valuation of Mandated Female Board Representation. *The Quarterly Journal of Economics*, 127(1), 137–197.
- Baltrunaite, A., Bello, P., Casarico, A., & Profeta, P. (2014). Gender quotas and the quality of politicians. *Journal of Public Economics*, 118, 62–74.
- Buser, T., Niederle, M., & Oosterbeek, H. (2014). Gender, competitiveness, and career choices. *The Quarterly Journal of Economics*, 129(3), 1409–1447.
- Casas-Arce, P., & Saiz, A. (2015). Women and Power: Unpopular, Unwilling, or Held Back? *Journal of Political Economy*, 123(3), 641–669.
- Arcidiacono, P., & Lovenheim, M. (2016). Affirmative Action and the Quality-Fit Trade-Off. *Journal of Economic Literature*, 54(1), 3–51.
- Bagde, S., Epplé, D., & Taylor, L. (2016). Does Affirmative Action Work? Caste, Gender, College Quality, and Academic Success in India. *American Economic Review*, 106(6), 1495–1521.
- Bagues, M., Sylos-Labini, M., & Zinovyeva, N. (2017). Does the Gender Composition of Scientific Committees Matter? *American Economic Review*, 107(4), 1207–1238.
- Besley, T., Folke, O., Persson, T., & Rickne, J. (2017). Gender Quotas and the Crisis of the Mediocre Man: Theory and Evidence from Sweden. *American Economic Review*, 107(8), 2204–2242.
- Clayton, A., & Tang, B. (2018). How women’s incumbency affects future elections: Evidence from a policy experiment in Lesotho. *World Development*, 110, 385–393.
- Bertrand, M., Black, S. E., Jensen, S., & Lleras-Muney, A. (2019). Breaking the Glass Ceiling? The Effect of Board Quotas on Female Labour Market Outcomes in Norway. *The Review of Economic Studies*, 86(1), 191–239.
- Bohren, J. A., Imas, A., & Rosenberg, M. (2019). The dynamics of discrimination: Theory and evidence. *American Economic Review*, 109(10), 3395–3436.
- Bordalo, P., Coffman, K., Gennaioli, N., & Shleifer, A. (2019). Beliefs about Gender. *American Economic Review*, 109(3), 739–773.
- Meyerinck, F. v., Niessen-Ruenzi, A., Schmid, M., & Solomon, S. D. (2019). As California goes, so goes the nation? Board gender quotas and the legislation of non-economic values. *Unpublished Manuscript*.

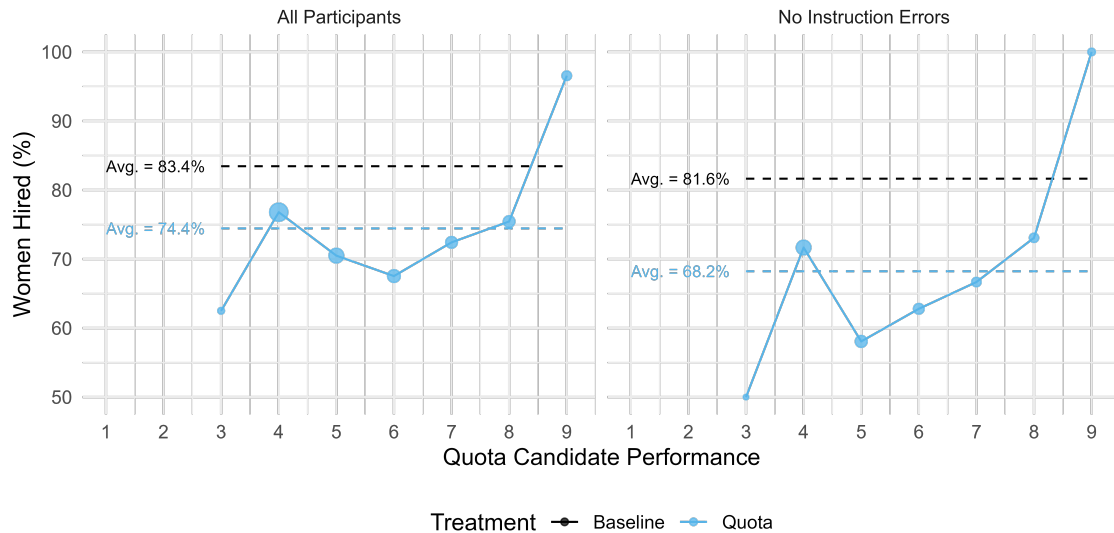
- Exley, C. L., Niederle, M., & Vesterlund, L. (2020). Knowing when to ask: The cost of leaning in. *Journal of Political Economy*, 128(3), 816–854.
- Ferreira, D., Ginglinger, E., Laguna, M.-A., & Skalli, Y. (2020). Closing the Gap: Board Gender Quotas and Hiring Practices. (2992213).
- Bagues, M., & Campa, P. (2021). Can gender quotas in candidate lists empower women? Evidence from a regression discontinuity design. *Journal of Public Economics*, 194, 104315.
- Coffman, K. B., Exley, C. L., & Niederle, M. (2021). The Role of Beliefs in Driving Gender Discrimination. *Management Science*, 67(6), 3551–3569.
- Gertsberg, M., Mollerstrom, J., & Pagel, M. (2021). *Gender quotas and support for women in board elections* (tech. rep.). National Bureau of Economic Research Cambridge, MA.
- Lippmann, Q. (2021). Are gender quotas on candidates bound to be ineffective? *Journal of Economic Behavior & Organization*, 191, 661–678.
- Bleemer, Z. (2022). Affirmative Action, Mismatch, and Economic Mobility after California’s Proposition 209. *The Quarterly Journal of Economics*, 137(1), 115–160.
- De Sousa, J., & Niederle, M. (2022). Trickle-Down Effects of Affirmative Action: A Case Study in France. (30367).
- Exley, C. L., & Kessler, J. B. (2022). The Gender Gap in Self-Promotion. *The Quarterly Journal of Economics*, 137(3), 1345–1381.
- Ferrari, G., Ferraro, V., Profeta, P., & Pronzato, C. (2022). Do Board Gender Quotas Matter? Selection, Performance, and Stock Market Effects. *Management Science*, 68(8), 5618–5643.
- Maida, A., & Weber, A. (2022). Female Leadership and Gender Gap within Firms: Evidence from an Italian Board Reform. *ILR Review*, 75(2), 488–515.
- Black, S. E., Denning, J. T., & Rothstein, J. (2023). Winners and losers? The effect of gaining and losing access to selective colleges on education and labor market outcomes. *American Economic Journal: Applied Economics*, 15(1), 26–67.
- Cullen, Z. B., & Pakzad-Hurson, B. (2023). Equilibrium effects of pay transparency. *Econometrica*, 91(3), 765–802.
- Otero, S., Barahona, N., & Dobbin, C. (2023). Affirmative action in centralized college admission systems. *Unpublished manuscript*. Retrieved June 19, 2025, from [https://sebotero.github.io/papers/cotas\\_submission.pdf](https://sebotero.github.io/papers/cotas_submission.pdf)
- Bian, B., Li, J., & Li, K. (2024, May). Does Mandating Women on Corporate Boards Backfire?
- Deschamps, P. (2024). Gender Quotas in Hiring Committees: A Boon or a Bane for Women? *Management Science*, 70(11), 7486–7505.
- Exley, C. L., & Nielsen, K. (2024). The gender gap in confidence: Expected but not accounted for. *American Economic Review*, 114(3), 851–885.
- Lee, J. E., & Zanella, M. (2024). Learning about women’s competence: The dynamic response of political parties to gender quotas in South Korea.
- Mocanu, T. (2024). Designing Gender Equity: Evidence from Hiring Practices.
- Bohren, J. A., Haggag, K., Imas, A., & Pope, D. G. (2025). Inaccurate statistical discrimination: An identification problem. *Review of Economics and Statistics*, 1–16.

- Bohren, J. A., Hull, P., & Imas, A. (2025). Systemic Discrimination: Theory and Measurement. *The Quarterly Journal of Economics*, 140(3), 1743–1799.
- Karekurve-Ramachandra, V., & Sood, G. (2025, November). The Limits of Electoral Gender Quotas in Rural Local Bodies.
- Kogelnik, M., & Strack, P. (2025). (Mis-)Understanding Quotas.
- Leibbrandt, A., & List, J. A. (2025). Do equal employment opportunity statements encourage racial minorities? evidence from a large natural field experiment. *European Economic Review*, 174, 104987.
- Gentile Passaro, D., Kojima, F., & Pakzad-Hurson, B. (2026). Equal Pay for Similar Work. *American Economic Review*, 116(3), 977–1013.
- Green, D. P., Singh, M., & Sood, G. (2026). New Evidence on the Effects of Randomly Assigned Reservations for Women Leaders in Indian Local Government.

# A Appendix

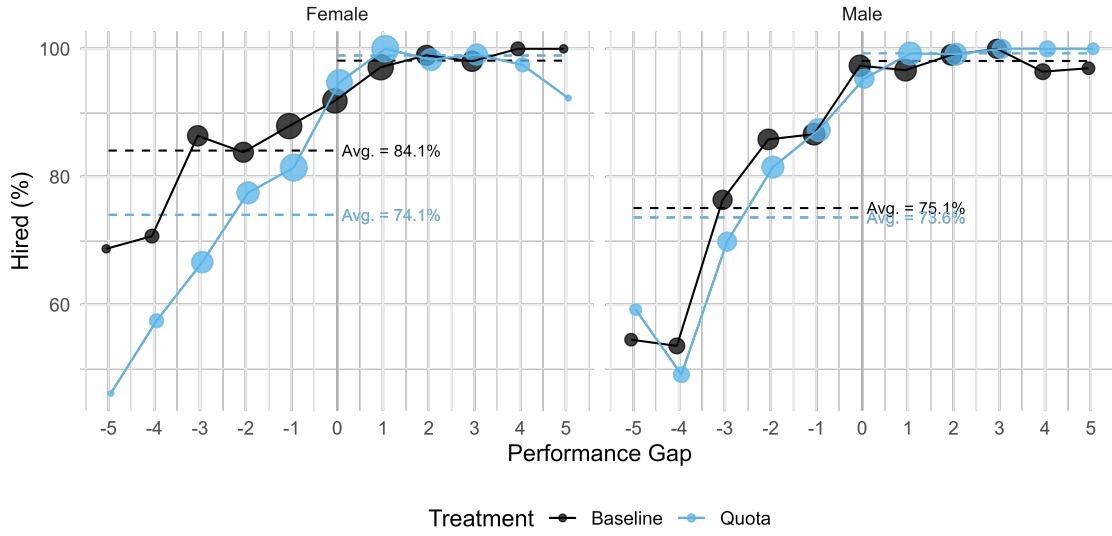
## A.1 Additional Figures and Tables

Figure A.7: Likelihood of a Woman being Hired, by Quota Candidate Performance



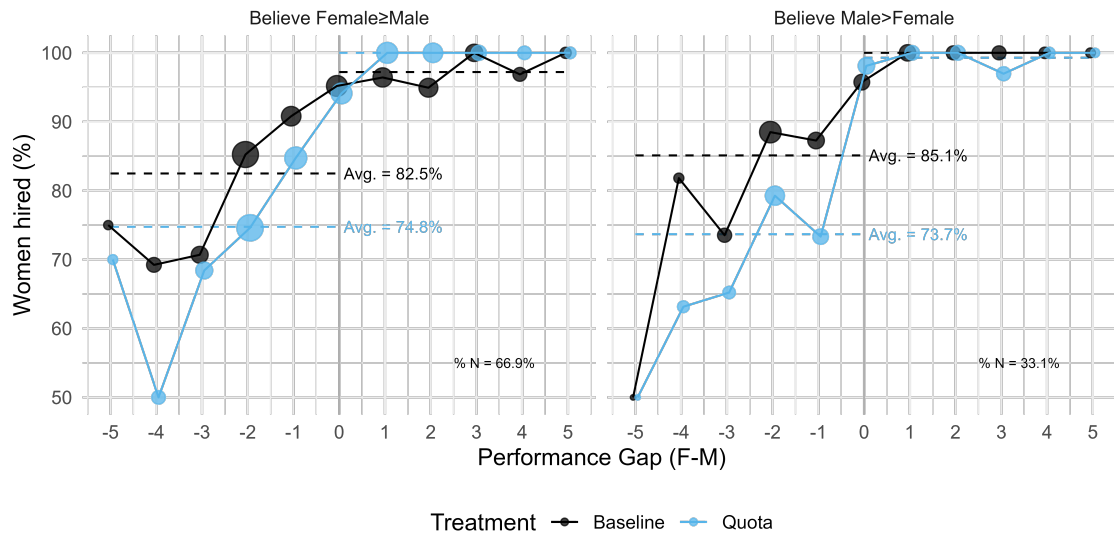
*Notes:* Circle size represents sample size. Quota Candidate Performance refers to the performance signal of the mandated quota candidate. We only include decisions where the female candidate is outperformed by the male one. Horizontal lines depict the average likelihood of a female worker being hired when outperformed by a male candidate. The left panel only uses decisions where the entire budget was allocated and when the candidates are of different genders. The right panel also restricts to participants who answer all understanding questions correctly.

Figure A.8: Likelihood of Hire by Performance Gap – Same Sex Pairs



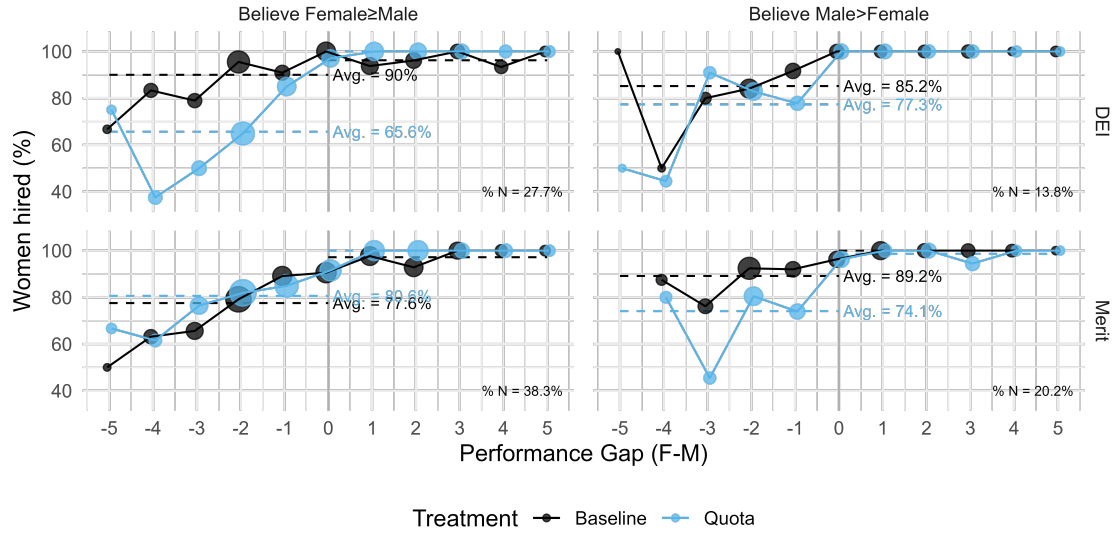
*Notes:* Circle size represents sample size. Performance gap refers to the difference in performance between candidates in decisions where both have the same gender. Positive values reflect candidates with better performance than the other candidate in the pair. Horizontal lines depict the average likelihood of being hired. Both panels use only decisions where the entire budget was allocated. See Appendix Table A.12 for details.

Figure A.9: Likelihood of Hire by Performance Gap – Heterogeneity by Performance Beliefs



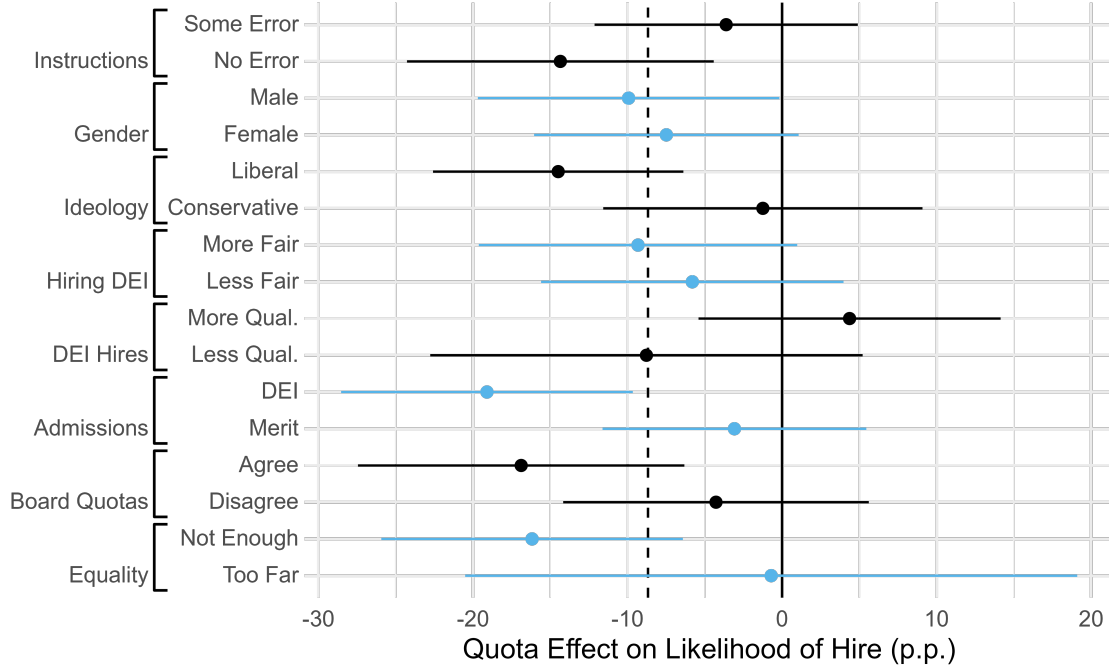
*Notes:* Circle size represents sample size. Performance gap refers to the difference in performance between candidates in decisions where both have the same gender. Positive values reflect candidates with better performance than the other candidate in the pair. Horizontal lines depict the average likelihood of being hired. The left panel shows participants who believe male candidates to have a better average performance, while the right panel shows those who believe female candidates to be better on average. Both panels use only decisions where the entire budget was allocated. See Appendix Table A.15 for details.

Figure A.10: Likelihood of Hire – Heterogeneity by Beliefs and Policy Views



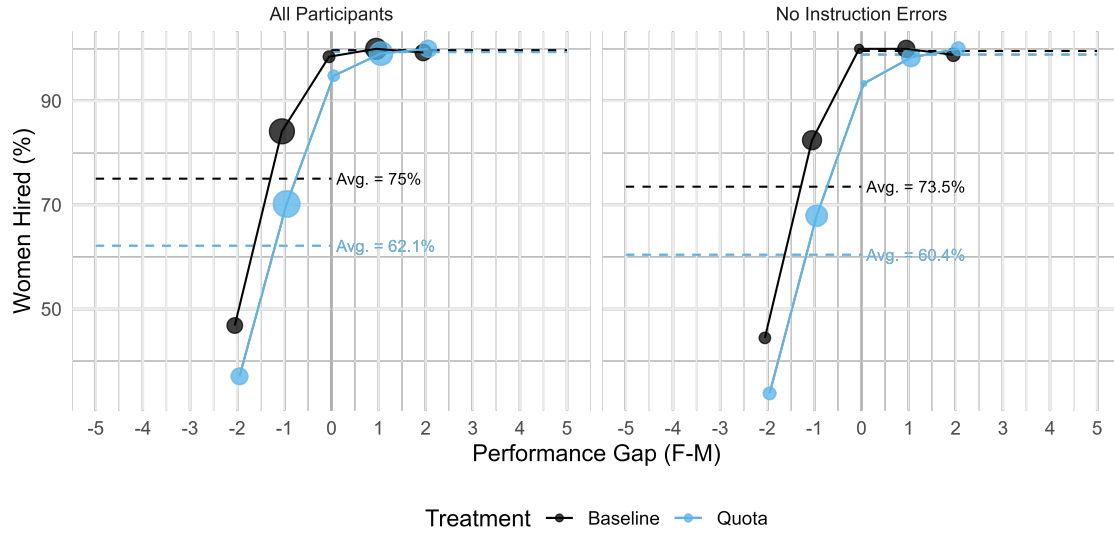
*Notes:* Circle size represents sample size. Performance gap refers to the difference in performance between candidates in decisions where both have the same gender. Positive values reflect candidates with better performance than the other candidate in the pair. Horizontal lines depict the average likelihood of being hired. The top left panel shows participants who believe female candidates to have a better average performance and support DEI in college admissions (11.5% of the sample), the top right panel shows participants who believe male candidates to have a better performance and support DEI in college admissions (13.8%), the bottom left panel shows participants who believe female candidates to have a better average performance and don't support DEI in college admissions (18.4%), and the bottom right panel shows participants who believe male candidates to have a better average performance and don't support DEI in college admissions (20.2%). All panels use only decisions where the entire budget was allocated. See Appendix Table A.15 for details.

Figure A.11: Heterogeneity of *Quota Treatment* Effect on Hire Likelihood Offers



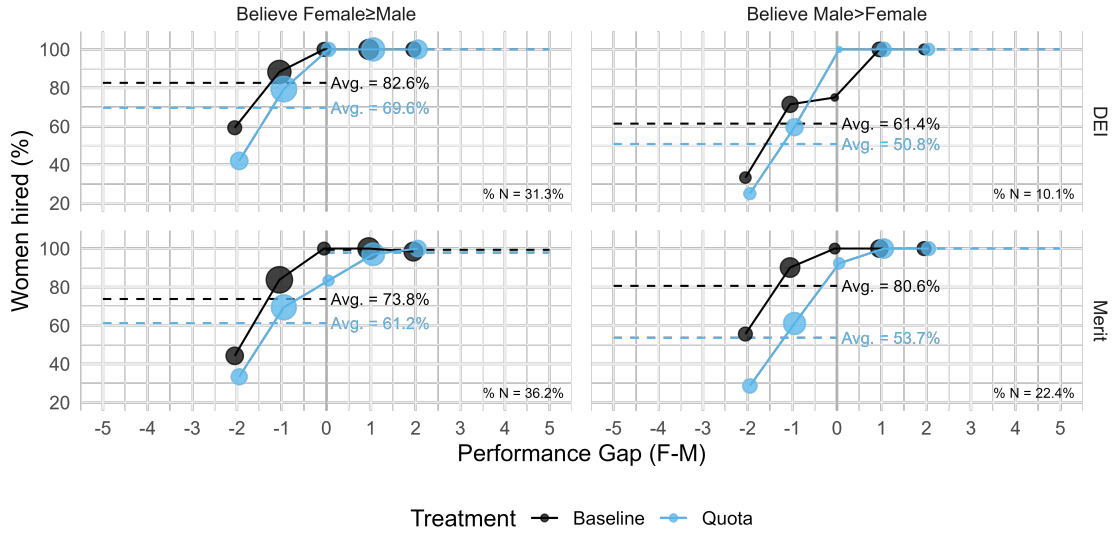
*Notes:* The vertical dashed line shows the Average Treatment Effect in Part I. Each point corresponds to the coefficient of *Quota Treatment* when regressed on whether a female candidate receives a zero offer, including the usual controls: whether the participant committed any mistake in the instructions, female gender, age above the median, white ethnicity, and no college degree. Error bars correspond to 95% confidence interval of the coefficient. Each regression is estimated by interacting the *Quota Treatment* with a different variable. We first explore different decisions: those with *Competitive Hiring*, while *Quota* interacts the treatment effect with whether the quota candidate has a higher or lower performance than the non-quota female one. Then we focus on heterogeneity by participant characteristics: errors in the *Instructions*, and participant's *Gender* and *Ideology*. Finally, we show heterogeneity by policy views: whether gender diversity makes hiring processes more or less fair (*Hiring DEI*) and the candidates more or less qualified (*DEI Hires*); whether *Admissions* should only consider merit or also DEI, whether there should be gender quotas on the board of directors of firms (*Board Quotas*); and whether the country has gone too far or not far enough in giving women equal opportunities (*Equality*).

Figure A.12: Likelihood of Hire – *Noisy Signal* Treatment



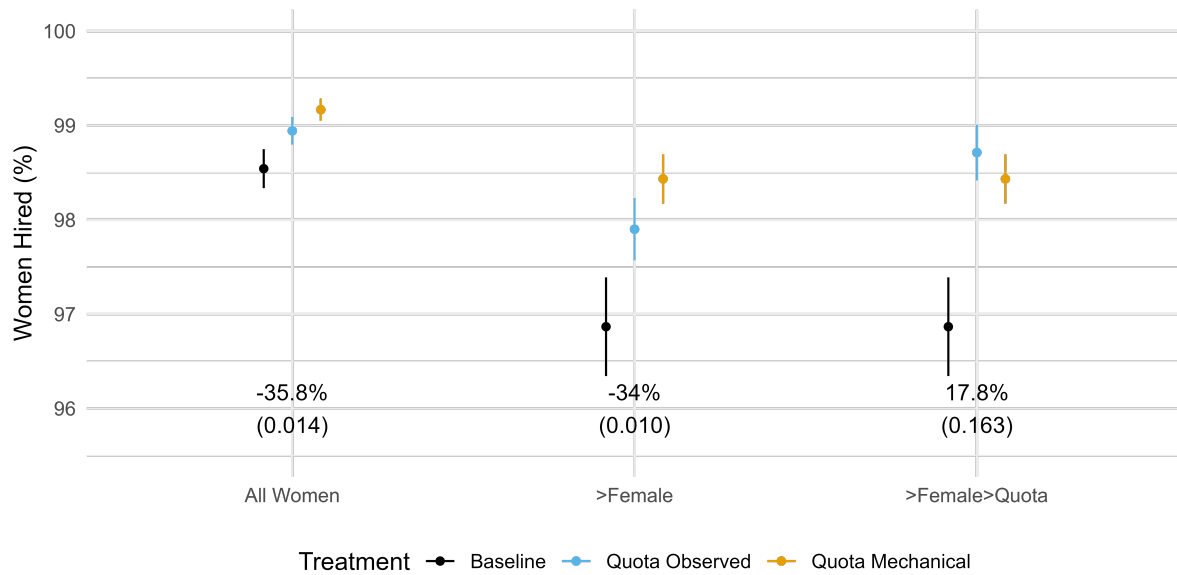
*Notes:* Circle size represents sample size. Performance gap refers to the difference in performance between candidates in decisions where both have the same gender. Positive values reflect candidates with better performance than the other candidate in the pair. Horizontal lines depict the average likelihood of being hired. All panels restrict the sample to participants in the *Noisy Signal* treatment. The left panel shows all participants, while the right panel restricts the sample to those participants who made no error in the instruction comprehension questions. All panels use only decisions where the entire budget was allocated. See Appendix Table A.16 for details.

Figure A.13: Likelihood of Hire in *Noisy Signal* – Heterogeneity



*Notes:* Circle size represents sample size. Performance gap refers to the difference in performance between candidates in decisions where both have the same gender. Positive values reflect candidates with better performance than the other candidate in the pair. Horizontal lines depict the average likelihood of being hired. All panels restrict the sample to participants in the *Noisy Signal* treatment. The top left panel shows participants who believe female candidates to have a better average performance and support DEI in college admissions (13.5% of the sample), the top right panel shows participants who believe male candidates to have a better performance and support DEI in college admissions (10.1%), the bottom left panel shows participants who believe female candidates to have a better average performance and don't support DEI in college admissions (14.4%), and the bottom right panel shows participants who believe male candidates to have a better average performance and don't support DEI in college admissions (22.4%). All panels use only decisions where the entire budget was allocated. See Appendix Table A.15 for details.

Figure A.14: Likelihood of Female Candidates Being Hired - Matched Recruiters



*Notes:* Error bars correspond to 95% confidence interval of the mean, with standard errors clustered at the participant level and controlling for whether participants did not make any errors in the instructions, female gender, age above the median, white ethnicity, and no college degree. The figures in the graph represent the effectiveness loss due to backlash —how much of the mechanical gains is lost, along with the p-value of a t-test of equality between the Quota Mechanical and the Quota Observed outcomes in parentheses. The first column shows the effects among all women, the second column restricts the sample to under-performing female candidates, and the third column further restricts it to decisions with under-performing quota candidates. The share of women hired in the Quota Observed and Quota Mechanical outcomes is calculated by adding the quota-mandated \$5 offer to an additional female candidate for each decision in the *Quota* and *Baseline* treatment, respectively. The Quota Mechanical should be interpreted as the likelihood of receiving an offer if the introduction of quotas triggered no backlash on non-targeted candidates. Both panels only use decisions where the entire budget was allocated.

Table A.3: Balance Table

| Variable                  | Baseline | Quota           | Quota Male      | Base Noisy | Quota Noisy     |
|---------------------------|----------|-----------------|-----------------|------------|-----------------|
| Instructions Correct (P1) | 0.44     | 0.50<br>[0.20]  | 0.47<br>[0.55]  | 0.55       | 0.56<br>[0.82]  |
| Instructions Correct (P2) | 0.51     | 0.44<br>[0.20]  | 0.45<br>[0.24]  | 0.53       | 0.48<br>[0.37]  |
| All Budget Allocated      | 0.90     | 0.92<br>[0.31]  | 0.88<br>[0.63]  | 0.79       | 0.88<br>[0.00]  |
| Age                       | 39.01    | 40.76<br>[0.15] | 39.31<br>[0.81] | 42.55      | 42.10<br>[0.72] |
| Male                      | 0.52     | 0.48<br>[0.43]  | 0.48<br>[0.49]  | 0.51       | 0.40<br>[0.03]  |
| White                     | 0.75     | 0.79<br>[0.38]  | 0.62<br>[0.01]  | 0.68       | 0.73<br>[0.28]  |
| No College Education      | 0.28     | 0.34<br>[0.18]  | 0.11<br>[0.00]  | 0.20       | 0.27<br>[0.10]  |
| Ideology (Conservative)   | 3.99     | 3.62<br>[0.07]  | 4.01<br>[0.91]  | 4.12       | 3.62<br>[0.02]  |
| Participants              | 205      | 203             | 205             | 205        | 201             |

*Notes:* Each cell reports the mean of the variable for the corresponding treatment. Brackets report the p-value of a t-test of difference in means against the Baseline (for Quota, Quota Male) or against Baseline Noisy (for Quota Noisy). *No Instruction Errors, P1* and *No Instruction Errors, P2* show the share of participants who made no understanding errors in the instructions for Part 1 and Part 2, respectively. *All Budget Allocated* shows the share of decisions that used the entire available budget to make offers. *Male* and *Age* reflect the demographic composition of the sample, *White* and *No College Education* display the share of participants who identify as racially White and have no college education, respectively. *Ideology* shows the average political alignment of the participants, where 1 reflects very liberal and 7 very conservative views.

Table A.4: Baseline Gap in Offers and Hire Likelihood in Part I

|                         | Offer            |                  |                   | Hired             |                   |                   |
|-------------------------|------------------|------------------|-------------------|-------------------|-------------------|-------------------|
|                         | (1)              | (2)              | (3)               | (4)               | (5)               | (6)               |
| Constant                | 3.479<br>(0.000) | 3.488<br>(0.000) | 3.488<br>(0.000)  | 1.046<br>(0)      | 1.075<br>(0.000)  | 1.108<br>(0.000)  |
| Male                    | 0.034<br>(0.386) | 0.021<br>(0.603) | 0.017<br>(0.770)  | -0.026<br>(0.002) | -0.029<br>(0.001) | -0.026<br>(0.051) |
| Performance             | 0.544<br>(0)     | 0.545<br>(0)     | 0.600<br>(0.000)  | 0.029<br>(0.000)  | 0.030<br>(0.000)  | 0.030<br>(0.000)  |
| Performance Other       | -0.543<br>(0)    | -0.545<br>(0)    | -0.599<br>(0.000) | -0.051<br>(0.000) | -0.051<br>(0.000) | -0.059<br>(0.000) |
| Individual Controls     |                  | Y                | Y                 |                   | Y                 | Y                 |
| No Instruction Errors   |                  |                  | Y                 |                   |                   | Y                 |
| Observations            | 4,356            | 4,262            | 1,848             | 4,356             | 4,262             | 1,848             |
| Dependent variable mean | 3.500            | 3.500            | 3.500             | 0.899             | 0.899             | 0.891             |

*Notes:* Results from a linear regression with  $p$ -values in parentheses and standard errors clustered at the individual level. The dependent variable Offer is the offer made to each candidate, in dollars (Columns (1)-(3)). The dependent variable Hired is a dummy that equals 1 if the candidate received an offer higher or equal than her reservation wage (Columns (4)-(6)). Male is a dummy that equals 1 if the candidate is male. Performance is the performance signal of the candidate. Performance Other is the performance signal of the other available candidate. All columns use the 12 hiring decisions in Part I. Columns (2), (3), (5) and (6) also include participant-level dummies as controls: whether participants did not make any errors in the instructions, female gender, age above the median, white ethnicity, and no college degree. Columns (3) and (6) restrict the sample to participants who made no error in the instruction comprehension questions.

Table A.5: Offers in the first decision of Part I

|                         | Offer             |                   |                   |
|-------------------------|-------------------|-------------------|-------------------|
|                         | (1)               | (2)               | (3)               |
| Constant                | 2.254<br>(0.000)  | 2.254<br>(0.000)  | 2.247<br>(0.000)  |
| Male                    | 2.492<br>(0.000)  | 2.492<br>(0.000)  | 2.506<br>(0.000)  |
| Quota                   | -0.303<br>(0.020) | -0.293<br>(0.028) | -0.614<br>(0.001) |
| Male $\times$ Quota     | 0.607<br>(0.020)  | 0.586<br>(0.028)  | 1.227<br>(0.001)  |
| Individual Controls     |                   | Y                 | Y                 |
| No Instruction Errors   |                   |                   | Y                 |
| Observations            | 750               | 736               | 342               |
| Dependent variable mean | 3.500             | 3.500             | 3.500             |

*Notes:* Results from a linear regression with  $p$ -values in parentheses and standard errors clustered at the individual level. The dependent variable is the offer made to each candidate, in dollars, in the first decision of Part I. Male is a dummy that equals 1 if the candidate is male. Quota is a dummy that equals 1 if the participant was in the *Quota Treatment*. Columns (2) and (3) also include participant-level dummies as controls: whether participants did not make any errors in the instructions, female gender, age above the median, white ethnicity, and no college degree. Column (3) restricts the sample to participants who made no error in the instruction comprehension questions.

Table A.6: Offers to women across all decisions of Part I

|                         | Offer             |                   |                   |                  |                   |                   |
|-------------------------|-------------------|-------------------|-------------------|------------------|-------------------|-------------------|
|                         | Under-performing  |                   |                   | Over-performing  |                   |                   |
|                         | (1)               | (2)               | (3)               | (4)              | (5)               | (6)               |
| Constant                | 2.202<br>(0.000)  | 3.968<br>(0.000)  | 4.197<br>(0.000)  | 4.593<br>(0.000) | 2.747<br>(0.000)  | 2.621<br>(0.000)  |
| Quota                   | -0.240<br>(0.028) | -0.217<br>(0.042) | -0.326<br>(0.032) | 0.093<br>(0.408) | 0.080<br>(0.425)  | 0.156<br>(0.258)  |
| Performance             |                   | 0.323<br>(0.000)  | 0.257<br>(0.000)  |                  | 0.525<br>(0.000)  | 0.603<br>(0.000)  |
| Performance Other       |                   | -0.408<br>(0.000) | -0.437<br>(0.000) |                  | -0.390<br>(0.000) | -0.421<br>(0.000) |
| Individual Controls     |                   | Y                 | Y                 |                  | Y                 | Y                 |
| No Instruction Errors   |                   |                   | Y                 |                  |                   | Y                 |
| Observations            | 1,171             | 1,147             | 538               | 1,181            | 1,154             | 532               |
| Dependent variable mean | 2.083             | 2.088             | 1.907             | 4.639            | 4.652             | 4.766             |

*Notes:* Results from a linear regression with  $p$ -values in parentheses and standard errors clustered at the individual level. The dependent variable is the offer made to a female candidate, in dollars. Quota is a dummy that equals 1 if the participant was in the *Quota Treatment*. Performance is the performance signal of the candidate. Performance Other is the performance signal of the other available candidate. Columns (1)-(3) restrict the sample to under-performing female candidates, this is, female candidates with a lower performance signal than their male counterparts. Columns (4)-(6) repeat the analysis for the rest of female candidates in different-gender pairs. Columns (2)-(3) and (5)-(6) also include participant-level dummies as controls: whether participants did not make any errors in the instructions, female gender, age above the median, white ethnicity, and no college degree. Columns (3) and (6) restrict the sample to participants who made no error in the instruction comprehension questions.

Table A.7: Zero offers to and hire likelihood of female candidates in the 1 decision of Part I

|                         | Zero Offer       |                  |                  | Hired             |                   |                   |
|-------------------------|------------------|------------------|------------------|-------------------|-------------------|-------------------|
|                         | (1)              | (2)              | (3)              | (4)               | (5)               | (6)               |
| Constant                | 0.119<br>(0.000) | 0.057<br>(0.281) | 0.148<br>(0.088) | 0.881<br>(0.000)  | 0.943<br>(0.000)  | 0.852<br>(0.000)  |
| Quota                   | 0.112<br>(0.005) | 0.101<br>(0.011) | 0.171<br>(0.007) | -0.112<br>(0.005) | -0.101<br>(0.011) | -0.171<br>(0.007) |
| Individual Controls     |                  | Y                | Y                |                   | Y                 | Y                 |
| No Instruction Errors   |                  |                  | Y                |                   |                   | Y                 |
| Observations            | 375              | 368              | 171              | 375               | 368               | 171               |
| Dependent variable mean | 0.173            | 0.177            | 0.228            | 0.827             | 0.823             | 0.772             |

*Notes:* Results from a linear regression with  $p$ -values in parentheses and standard errors clustered at the individual level. The dependent variable Zero Offer is a dummy that equals 1 if the female candidate received no offer (Columns (1)-(3)). The dependent variable Hired is a dummy that equals 1 if the female candidate received an offer higher or equal than her reservation wage (Columns (4)-(6)). Quota is a dummy that equals 1 if the participant was in the *Quota Treatment*. Columns (2)-(3) and (5)-(6) also include participant-level dummies as controls: whether participants did not make any errors in the instructions, female gender, age above the median, white ethnicity, and no college degree. Columns (3) and (6) restrict the sample to participants who made no error in the instruction comprehension questions.

Table A.8: Hire likelihood of female candidates across Part I

|                         | Hired             |                   |                   |                  |                   |                   |
|-------------------------|-------------------|-------------------|-------------------|------------------|-------------------|-------------------|
|                         | Under-performing  |                   |                   | Over-performing  |                   |                   |
|                         | (1)               | (2)               | (3)               | (4)              | (5)               | (6)               |
| Constant                | 0.834<br>(0.000)  | 1.507<br>(0.000)  | 1.620<br>(0.000)  | 0.975<br>(0.000) | 0.994<br>(0.000)  | 0.997<br>(0.000)  |
| Quota                   | -0.090<br>(0.009) | -0.087<br>(0.009) | -0.137<br>(0.007) | 0.012<br>(0.270) | 0.014<br>(0.238)  | 0.021<br>(0.185)  |
| Performance             |                   | 0.023<br>(0.109)  | 0.014<br>(0.484)  |                  | 0.007<br>(0.031)  | 0.006<br>(0.163)  |
| Performance Other       |                   | -0.094<br>(0.000) | -0.106<br>(0.000) |                  | -0.009<br>(0.001) | -0.012<br>(0.015) |
| Individual Controls     |                   | Y                 | Y                 |                  | Y                 | Y                 |
| No Instruction Errors   |                   |                   | Y                 |                  |                   | Y                 |
| Observations            | 1,171             | 1,147             | 538               | 1,181            | 1,154             | 532               |
| Dependent variable mean | 0.790             | 0.789             | 0.742             | 0.981            | 0.981             | 0.981             |

*Notes:* Results from a linear regression with  $p$ -values in parentheses and standard errors clustered at the individual level. The dependent variable is a dummy that equals 1 if the female candidate received an offer higher or equal than her reservation wage. Quota is a dummy that equals 1 if the participant was in the *Quota Treatment*. Performance is the performance signal of the candidate. Performance Other is the performance signal of the other available candidate. Columns (1)-(3) restrict the sample to underperforming female candidates, this is, female candidates with a lower performance signal than their male counterparts. Columns (4)-(6) repeat the analysis for the rest of female candidates in different-gender pairs. Columns (2)-(3) and (5)-(6) also include participant-level dummies as controls: whether participants did not make any errors in the instructions, female gender, age above the median, white ethnicity, and no college degree. Columns (3) and (6) restrict the sample to participants who made no error in the instruction comprehension questions.

Table A.9: Zero offers to female candidates across Part I

|                         | Zero Offer       |                   |                   |                   |                   |                   |
|-------------------------|------------------|-------------------|-------------------|-------------------|-------------------|-------------------|
|                         | Under-performing |                   |                   | Over-performing   |                   |                   |
|                         | (1)              | (2)               | (3)               | (4)               | (5)               | (6)               |
| Constant                | 0.145<br>(0.000) | -0.356<br>(0.000) | -0.471<br>(0.003) | 0.025<br>(0.006)  | 0.014<br>(0.438)  | 0.003<br>(0.911)  |
| Quota                   | 0.076<br>(0.025) | 0.070<br>(0.035)  | 0.125<br>(0.014)  | -0.015<br>(0.146) | -0.017<br>(0.134) | -0.021<br>(0.185) |
| Performance             |                  | -0.033<br>(0.016) | -0.025<br>(0.241) |                   | -0.007<br>(0.020) | -0.006<br>(0.163) |
| Performance Other       |                  | 0.078<br>(0.000)  | 0.090<br>(0.000)  |                   | 0.009<br>(0.001)  | 0.012<br>(0.015)  |
| Individual Controls     |                  | Y                 | Y                 |                   | Y                 | Y                 |
| No Instruction Errors   |                  |                   | Y                 |                   |                   | Y                 |
| Observations            | 1,171            | 1,147             | 538               | 1,181             | 1,154             | 532               |
| Dependent variable mean | 0.183            | 0.185             | 0.234             | 0.018             | 0.017             | 0.019             |

*Notes:* Results from a linear regression with  $p$ -values in parentheses and standard errors clustered at the individual level. The dependent variable is a dummy that equals 1 if the female candidate received no offer. Quota is a dummy that equals 1 if the participant was in the *Quota Treatment*. Performance is the performance signal of the candidate. Performance Other is the performance signal of the other available candidate. Columns (1)-(3) restrict the sample to underperforming female candidates, this is, female candidates with a lower performance signal than their male counterparts. Columns (4)-(6) repeat the analysis for the rest of female candidates in different-gender pairs. Columns (2)-(3) and (5)-(6) also include participant-level dummies as controls: whether participants did not make any errors in the instructions, female gender, age above the median, white ethnicity, and no college degree. Columns (3) and (6) restrict the sample to participants who made no error in the instruction comprehension questions.

Table A.10: Offers to and hire likelihood of under-performing women in Part I

|                         | Offer            |                   |                   | Hired            |                   |                   |
|-------------------------|------------------|-------------------|-------------------|------------------|-------------------|-------------------|
|                         | (1)              | (2)               | (3)               | (4)              | (5)               | (6)               |
| Constant                | 1.790<br>(0.000) | 3.713<br>(0.000)  | 3.349<br>(0.000)  | 0.693<br>(0.000) | 1.429<br>(0.000)  | 1.474<br>(0.000)  |
| Quota Performance       | 0.032<br>(0.304) | 0.059<br>(0.043)  | 0.101<br>(0.014)  | 0.009<br>(0.321) | 0.018<br>(0.056)  | 0.027<br>(0.055)  |
| Performance             |                  | 0.314<br>(0.000)  | 0.258<br>(0.003)  |                  | 0.024<br>(0.254)  | 0.029<br>(0.303)  |
| Performance Other       |                  | -0.413<br>(0.000) | -0.400<br>(0.000) |                  | -0.100<br>(0.000) | -0.123<br>(0.000) |
| Individual Controls     |                  | Y                 | Y                 |                  | Y                 | Y                 |
| No Instruction Errors   |                  |                   | Y                 |                  |                   | Y                 |
| Observations            | 579              | 567               | 288               | 579              | 567               | 288               |
| Dependent variable mean | 1.962            | 1.963             | 1.753             | 0.744            | 0.743             | 0.677             |

*Notes:* Results from a linear regression with  $p$ -values in parentheses and standard errors clustered at the individual level. The dependent variable Offer is the offer made to a female candidate, in dollars (Columns (1)-(3)). The dependent variable Hired is a dummy that equals 1 if the female candidate received an offer higher or equal than her reservation wage (Columns (4)-(6)). Quota Performance is the performance signal of the quota-mandated candidate. Performance is the performance signal of the candidate. Performance Other is the performance signal of the other available candidate. All regressions are estimated from decisions within the Quota treatment. Columns (2)-(3) and (5)-(6) also include participant-level dummies as controls: whether participants did not make any errors in the instructions, female gender, age above the median, white ethnicity, and no college degree. Columns (3) and (6) restrict the sample to participants who made no error in the instruction comprehension questions.

Table A.11: Hire likelihood of male candidates across Part I

|                         | Hired             |                   |                   |                  |                   |                   |
|-------------------------|-------------------|-------------------|-------------------|------------------|-------------------|-------------------|
|                         | Under-performing  |                   |                   | Over-performing  |                   |                   |
|                         | (1)               | (2)               | (3)               | (4)              | (5)               | (6)               |
| Constant                | 0.813<br>(0.000)  | 1.310<br>(0.000)  | 1.341<br>(0.000)  | 0.985<br>(0)     | 1.003<br>(0.000)  | 0.985<br>(0.000)  |
| Quota                   | -0.037<br>(0.261) | -0.020<br>(0.516) | -0.037<br>(0.426) | 0.008<br>(0.167) | 0.008<br>(0.211)  | 0.013<br>(0.127)  |
| Performance             |                   | 0.063<br>(0.000)  | 0.059<br>(0.000)  |                  | 0.003<br>(0.327)  | 0.003<br>(0.531)  |
| Performance Other       |                   | -0.105<br>(0.000) | -0.115<br>(0.000) |                  | -0.008<br>(0.040) | -0.004<br>(0.267) |
| Individual Controls     |                   | Y                 | Y                 |                  | Y                 | Y                 |
| No Instruction Errors   |                   |                   | Y                 |                  |                   | Y                 |
| Observations            | 1,181             | 1,154             | 532               | 1,171            | 1,147             | 538               |
| Dependent variable mean | 0.794             | 0.791             | 0.765             | 0.989            | 0.989             | 0.991             |

*Notes:* Results from a linear regression with  $p$ -values in parentheses and standard errors clustered at the individual level. The dependent variable is a dummy that equals 1 if the male candidate received an offer higher or equal than their reservation wage. Quota is a dummy that equals 1 if the participant was in the *Quota Treatment*. Performance is the performance signal of the candidate. Performance Other is the performance signal of the other available candidate. Columns (1)-(3) restrict the sample to under-performing male candidates, this is, male candidates with a lower performance signal than their female counterparts. Columns (4)-(6) repeat the analysis for the rest of male candidates in different-gender pairs. Columns (2)-(3) and (5)-(6) also include participant-level dummies as controls: whether participants did not make any errors in the instructions, female gender, age above the median, white ethnicity, and no college degree. Columns (3) and (6) restrict the sample to participants who made no error in the instruction comprehension questions.

Table A.12: Hire likelihood of under-performing candidates in same-gender pairs - Part I

|                         | Hired             |                   |                   |                   |                   |                   |
|-------------------------|-------------------|-------------------|-------------------|-------------------|-------------------|-------------------|
|                         | Female Candidate  |                   |                   | Male Candidate    |                   |                   |
|                         | (1)               | (2)               | (3)               | (4)               | (5)               | (6)               |
| Constant                | 0.841<br>(0.000)  | 1.285<br>(0.000)  | 1.396<br>(0.000)  | 0.751<br>(0.000)  | 1.354<br>(0.000)  | 1.395<br>(0.000)  |
| Quota                   | -0.100<br>(0.009) | -0.106<br>(0.007) | -0.081<br>(0.165) | -0.014<br>(0.714) | -0.012<br>(0.738) | -0.032<br>(0.556) |
| Performance             |                   | 0.031<br>(0.037)  | 0.023<br>(0.271)  |                   | 0.081<br>(0.000)  | 0.098<br>(0.000)  |
| Performance Other       |                   | -0.075<br>(0.000) | -0.093<br>(0.000) |                   | -0.115<br>(0.000) | -0.135<br>(0.000) |
| Individual Controls     |                   | Y                 | Y                 |                   | Y                 | Y                 |
| No Instruction Errors   |                   |                   | Y                 |                   |                   | Y                 |
| Observations            | 912               | 898               | 424               | 863               | 845               | 389               |
| Dependent variable mean | 0.788             | 0.791             | 0.738             | 0.744             | 0.744             | 0.707             |

*Notes:* Results from a linear regression with  $p$ -values in parentheses and standard errors clustered at the individual level. The dependent variable is a dummy that equals 1 if the candidate received an offer higher or equal than their reservation wage. Quota is a dummy that equals 1 if the participant was in the *Quota Treatment*. Performance is the performance signal of the candidate. Performance Other is the performance signal of the other available candidate. Columns (1)-(3) restrict the sample to female candidates in same-gender pairs. Columns (4)-(6) repeat the analysis for male candidates in same-gender pairs. Columns (2)-(3) and (5)-(6) also include participant-level dummies as controls: whether participants did not make any errors in the instructions, female gender, age above the median, white ethnicity, and no college degree. Columns (3) and (6) restrict the sample to participants who made no error in the instruction comprehension questions.

Table A.13: Hire likelihood of male candidates, *Quota Male* Treatment

|                         | Hired            |                   |                   |                   |                   |                   |
|-------------------------|------------------|-------------------|-------------------|-------------------|-------------------|-------------------|
|                         | Under-performing |                   |                   | Over-performing   |                   |                   |
|                         | (1)              | (2)               | (3)               | (4)               | (5)               | (6)               |
| Constant                | 0.760<br>(0.000) | 1.339<br>(0.000)  | 1.413<br>(0.000)  | 0.981<br>(0)      | 1.044<br>(0.000)  | 1.028<br>(0.000)  |
| Quota Male              | 0.051<br>(0.163) | 0.025<br>(0.520)  | -0.022<br>(0.691) | -0.017<br>(0.104) | -0.020<br>(0.092) | -0.034<br>(0.129) |
| Performance             |                  | 0.054<br>(0.000)  | 0.062<br>(0.000)  |                   | 0.007<br>(0.095)  | 0.011<br>(0.140)  |
| Performance Other       |                  | -0.103<br>(0.000) | -0.127<br>(0.000) |                   | -0.015<br>(0.003) | -0.018<br>(0.025) |
| Individual Controls     |                  | Y                 | Y                 |                   | Y                 | Y                 |
| No Instruction Errors   |                  |                   | Y                 |                   |                   | Y                 |
| Observations            | 897              | 871               | 394               | 1,486             | 1,452             | 698               |
| Dependent variable mean | 0.786            | 0.780             | 0.756             | 0.972             | 0.972             | 0.961             |

*Notes:* Results from a linear regression with  $p$ -values in parentheses and standard errors clustered at the individual level. The dependent variable is a dummy that equals 1 if the male candidate received an offer higher or equal than their reservation wage. Quota is a dummy that equals 1 if the participant was in the *Quota Male* Treatment. Performance is the performance signal of the candidate. Performance Other is the performance signal of the other available candidate. Columns (1)-(3) restrict the sample to underperforming male candidates, this is, male candidates with a lower performance signal than their female counterparts. Columns (4)-(6) repeat the analysis for the rest of male candidates in different-gender pairs. Columns (2)-(3) and (5)-(6) also include participant-level dummies as controls: whether participants did not make any errors in the instructions, female gender, age above the median, white ethnicity, and no college degree. Columns (3) and (6) restrict the sample to participants who made no error in the instruction comprehension questions.

Table A.14: Hire likelihood of women in Part I — Heterogeneity by policy views

|                         | Hired             |                   |                   |                   |                   |                   |
|-------------------------|-------------------|-------------------|-------------------|-------------------|-------------------|-------------------|
|                         | Under-performing  |                   |                   | Over-performing   |                   |                   |
|                         | (1)               | (2)               | (3)               | (4)               | (5)               | (6)               |
| Constant                | 0.882<br>(0.000)  | 1.600<br>(0.000)  | 1.698<br>(0.000)  | 0.982<br>(0.000)  | 1.003<br>(0.000)  | 1.003<br>(0.000)  |
| Quota                   | -0.188<br>(0.000) | -0.191<br>(0.000) | -0.207<br>(0.004) | 0.014<br>(0.239)  | 0.010<br>(0.397)  | 0.015<br>(0.524)  |
| Merit                   | -0.062<br>(0.118) | -0.084<br>(0.034) | -0.103<br>(0.127) | -0.013<br>(0.471) | -0.023<br>(0.225) | -0.023<br>(0.339) |
| Quota $\times$ Merit    | 0.157<br>(0.019)  | 0.159<br>(0.015)  | 0.174<br>(0.082)  | -0.004<br>(0.837) | 0.006<br>(0.782)  | 0.013<br>(0.678)  |
| Performance             |                   | 0.020<br>(0.175)  | 0.018<br>(0.400)  |                   | 0.007<br>(0.042)  | 0.007<br>(0.176)  |
| Performance Other       |                   | -0.098<br>(0.000) | -0.115<br>(0.000) |                   | -0.009<br>(0.001) | -0.012<br>(0.013) |
| Individual Controls     |                   | Y                 | Y                 |                   | Y                 | Y                 |
| No Instruction Errors   |                   |                   | Y                 |                   |                   | Y                 |
| Observations            | 1,076             | 1,055             | 510               | 1,088             | 1,065             | 503               |
| Dependent variable mean | 0.796             | 0.794             | 0.747             | 0.980             | 0.980             | 0.980             |

*Notes:* Results from a linear regression with  $p$ -values in parentheses and standard errors clustered at the individual level. The dependent variable is a dummy that equals 1 if the female candidate received an offer higher or equal than her reservation wage. Quota is a dummy that equals 1 if the participant was in the *Quota Treatment*. Merit is a dummy that equals zero if the participant supports DEI in college admissions, and 1 if the participant believes admissions should only be based on merit. Performance is the performance signal of the candidate. Performance Other is the performance signal of the other available candidate. Columns (1)-(3) restrict the sample to underperforming female candidates, this is, female candidates with a lower performance signal than their male counterparts. Columns (4)-(6) repeat the analysis for the rest of female candidates in different-gender pairs. Columns (2)-(3) and (5)-(6) also include participant-level dummies as controls: whether participants did not make any errors in the instructions, female gender, age above the median, white ethnicity, and no college degree. Columns (3) and (6) restrict the sample to participants who made no error in the instruction comprehension questions.

Table A.15: Hire likelihood of under-performing women in Part I — Heterogeneity by policy views and beliefs

|                                           | Hired             |                   |                   |                   |                   |                   |
|-------------------------------------------|-------------------|-------------------|-------------------|-------------------|-------------------|-------------------|
|                                           | (1)               | (2)               | (3)               | (4)               | (5)               | (6)               |
| Constant                                  | 0.825<br>(0.000)  | 1.501<br>(0.000)  | 1.603<br>(0.000)  | 0.900<br>(0.000)  | 1.615<br>(0.000)  | 1.746<br>(0.000)  |
| Quota                                     | -0.077<br>(0.061) | -0.074<br>(0.063) | -0.131<br>(0.041) | -0.244<br>(0.000) | -0.250<br>(0.000) | -0.300<br>(0.002) |
| Believe M>F                               | 0.026<br>(0.578)  | 0.024<br>(0.620)  | 0.035<br>(0.590)  | -0.048<br>(0.378) | -0.060<br>(0.275) | -0.098<br>(0.283) |
| Quota $\times$ Believe M>F                | -0.037<br>(0.614) | -0.039<br>(0.598) | -0.012<br>(0.912) | 0.165<br>(0.080)  | 0.164<br>(0.086)  | 0.249<br>(0.095)  |
| Performance                               |                   | 0.022<br>(0.119)  | 0.013<br>(0.512)  |                   | 0.022<br>(0.150)  | 0.021<br>(0.339)  |
| Performance Other                         |                   | -0.094<br>(0.000) | -0.105<br>(0.000) |                   | -0.097<br>(0.000) | -0.118<br>(0.000) |
| Merit                                     |                   |                   |                   | -0.124<br>(0.016) | -0.149<br>(0.004) | -0.199<br>(0.024) |
| Quota $\times$ Merit                      |                   |                   |                   | 0.274<br>(0.001)  | 0.280<br>(0.001)  | 0.301<br>(0.020)  |
| Believe M>F $\times$ Merit                |                   |                   |                   | 0.164<br>(0.036)  | 0.173<br>(0.031)  | 0.255<br>(0.063)  |
| Quota $\times$ Believe M>F $\times$ Merit |                   |                   |                   | -0.346<br>(0.011) | -0.350<br>(0.011) | -0.346<br>(0.104) |
| Individual Controls                       |                   | Y                 | Y                 |                   | Y                 | Y                 |
| No Instruction Errors                     |                   |                   | Y                 |                   |                   | Y                 |
| Observations                              | 1,171             | 1,147             | 538               | 1,076             | 1,055             | 510               |
| Dependent variable mean                   | 0.796             | 0.797             | 0.751             | 0.796             | 0.794             | 0.747             |

*Notes:* Results from a linear regression with  $p$ -values in parentheses and standard errors clustered at the individual level. The dependent variable is a dummy that equals 1 if the female candidate received an offer higher or equal than her reservation wage. Quota is a dummy that equals 1 if the participant was in the *Quota Treatment*. Believe M>F is a dummy that equals 1 if the participant believes the average male candidate to have strictly higher performance than the average woman. Merit is a dummy that equals zero if the participant supports DEI in college admissions, and 1 if the participant believes admissions should only be based on merit. Performance is the performance signal of the candidate. Performance Other is the performance signal of the other available candidate. All models only use decisions from Part I with underperforming female candidates, this is, female candidates with a lower performance signal than their male counterparts. Columns (2)-(3) and (5)-(6) also include participant-level dummies as controls: whether participants did not make any errors in the instructions, female gender, age above the median, white ethnicity, and no college degree. Columns (3) and (6) restrict the sample to participants who made no error in the instruction comprehension questions.

Table A.16: Hire likelihood of under-performing women in Part I — Heterogeneity by policy views and beliefs, Noisy Treatment

|                                           | Hired             |                   |                   |                   |                   |                   |
|-------------------------------------------|-------------------|-------------------|-------------------|-------------------|-------------------|-------------------|
|                                           | (1)               | (2)               | (3)               | (4)               | (5)               | (6)               |
| Constant                                  | 0.773<br>(0.000)  | 0.987<br>(0.000)  | 0.908<br>(0.000)  | 0.826<br>(0.000)  | 1.027<br>(0.000)  | 0.955<br>(0.000)  |
| Quota                                     | -0.117<br>(0.018) | -0.099<br>(0.037) | -0.126<br>(0.065) | -0.130<br>(0.037) | -0.083<br>(0.178) | -0.087<br>(0.353) |
| Believe M>F                               | -0.038<br>(0.478) | -0.026<br>(0.611) | 0.001<br>(0.994)  | -0.212<br>(0.014) | -0.148<br>(0.071) | -0.178<br>(0.123) |
| Quota $\times$ Believe M>F                | -0.092<br>(0.271) | -0.106<br>(0.192) | -0.069<br>(0.523) | 0.024<br>(0.862)  | -0.035<br>(0.787) | 0.071<br>(0.655)  |
| Performance                               |                   | 0.402<br>(0.000)  | 0.400<br>(0.000)  |                   | 0.399<br>(0.000)  | 0.395<br>(0.000)  |
| Performance Other                         |                   | -0.234<br>(0.000) | -0.238<br>(0.000) |                   | -0.241<br>(0.000) | -0.248<br>(0.000) |
| Merit                                     |                   |                   |                   | -0.088<br>(0.142) | -0.050<br>(0.395) | -0.026<br>(0.779) |
| Quota $\times$ Merit                      |                   |                   |                   | 0.004<br>(0.970)  | -0.050<br>(0.588) | -0.080<br>(0.544) |
| Believe M>F $\times$ Merit                |                   |                   |                   | 0.280<br>(0.009)  | 0.197<br>(0.059)  | 0.290<br>(0.054)  |
| Quota $\times$ Believe M>F $\times$ Merit |                   |                   |                   | -0.167<br>(0.334) | -0.092<br>(0.578) | -0.213<br>(0.326) |
| Individual Controls                       |                   | Y                 | Y                 |                   | Y                 | Y                 |
| No Instruction Errors                     |                   |                   | Y                 |                   |                   | Y                 |
| Observations                              | 1,046             | 1,025             | 566               | 1,046             | 1,025             | 566               |
| Dependent variable mean                   | 0.684             | 0.682             | 0.657             | 0.684             | 0.682             | 0.657             |

*Notes:* Results from a linear regression with  $p$ -values in parentheses and standard errors clustered at the individual level. The dependent variable is a dummy that equals 1 if the female candidate received an offer higher or equal than her reservation wage. Quota is a dummy that equals 1 if the participant was in the *Quota Treatment*. Believe M>F is a dummy that equals 1 if the participant believes the average male candidate to have strictly higher performance than the average woman. Merit is a dummy that equals zero if the participant supports DEI in college admissions, and 1 if the participant believes admissions should only be based on merit. Performance is the performance signal of the candidate. Performance Other is the performance signal of the other available candidate. All models only use decisions from Part I in the Noisy Treatment with underperforming female candidates, this is, female candidates with a lower performance signal than their male counterparts. Columns (2)-(3) and (5)-(6) also include participant-level dummies as controls: whether participants did not make any errors in the instructions, female gender, age above the median, white ethnicity, and no college degree. Columns (3) and (6) restrict the sample to participants who made no error in the instruction comprehension questions.

Table A.17: Net effect of quotas on hire likelihood of female candidates

|                                 | Hired             |                   |                   |                   |                   |                   |
|---------------------------------|-------------------|-------------------|-------------------|-------------------|-------------------|-------------------|
|                                 | All Women         |                   | > Women           |                   | > Women > Quota   |                   |
|                                 | (1)               | (2)               | (3)               | (4)               | (5)               | (6)               |
| Constant                        | 0.952<br>(0.000)  | 0.936<br>(0.000)  | 0.919<br>(0.000)  | 0.871<br>(0.000)  | 0.923<br>(0.000)  | 0.874<br>(0.000)  |
| Quota                           | 0.029<br>(0.025)  | 0.028<br>(0.134)  | 0.026<br>(0.281)  | 0.034<br>(0.364)  | -0.012<br>(0.691) | -0.018<br>(0.693) |
| Mechanical                      | 0.045<br>(0.000)  | 0.050<br>(0.000)  | 0.072<br>(0.000)  | 0.087<br>(0.000)  | 0.072<br>(0.000)  | 0.087<br>(0.000)  |
| Competitive                     | -0.045<br>(0.000) | -0.073<br>(0.000) | -0.086<br>(0.000) | -0.139<br>(0.000) | -0.086<br>(0.000) | -0.139<br>(0.000) |
| Quota $\times$ Competitive      | 0.033<br>(0.010)  | 0.055<br>(0.008)  | 0.074<br>(0.002)  | 0.126<br>(0.002)  | 0.106<br>(0.000)  | 0.167<br>(0.000)  |
| Mechanical $\times$ Competitive | 0.023<br>(0.000)  | 0.038<br>(0.000)  | 0.052<br>(0.000)  | 0.082<br>(0.000)  | 0.052<br>(0.000)  | 0.082<br>(0.000)  |
| Individual Controls             | Y                 | Y                 | Y                 | Y                 | Y                 | Y                 |
| No Instruction Errors           |                   | Y                 |                   | Y                 |                   | Y                 |
| Observations                    | 21,818            | 10,123            | 9,367             | 4,317             | 7,495             | 3,378             |
| Dependent variable mean         | 0.926             | 0.910             | 0.857             | 0.819             | 0.854             | 0.810             |

*Notes:* Results from a linear regression with  $p$ -values in parentheses and standard errors clustered at the individual level. The dependent variable is a dummy that equals 1 if the female candidate received an offer higher or equal than their reservation wage. Quota is a dummy that equals 1 if the participant was in the *Quota* Treatment. Mechanical is a dummy that equals 1 for decisions in the counterfactual Quota Mechanical case, where quotas increase hiring among mandated candidates but has no spillovers on non-mandated ones. Competitive is a dummy that equals 1 for decisions in Part II, where recruiters compete among each other to hire candidates. Columns (1)-(2) use all decisions with a female candidate, including same-gender pairs. Columns (3)-(4) restrict the sample to under-performing female candidates, this is, women with a lower performance signal than their counterparts. Columns (5)-(6) further restrict the sample to decisions where the under-performing female candidate has a higher signal than the quota candidate. All models include participant-level dummies as controls: whether participants did not make any errors in the instructions, female gender, age above the median, white ethnicity, and no college degree. Even-numbered columns restrict the sample to participants who made no error in the instruction comprehension questions.

Table A.18: Profit forgone by recruiters, Part I

|                 | Forgone Profit (dollars) |                  |                  |                  |                  |                  |                  |                  |
|-----------------|--------------------------|------------------|------------------|------------------|------------------|------------------|------------------|------------------|
|                 | All Women                |                  | > Women          |                  | > Women > Quota  |                  | Male-Male        |                  |
|                 | (1)                      | (2)              | (3)              | (4)              | (5)              | (6)              | (7)              | (8)              |
| Constant        | 0.452<br>(0.015)         | 0.592<br>(0.012) | 0.406<br>(0.076) | 0.506<br>(0.100) | 0.508<br>(0.041) | 0.562<br>(0.082) | 0.562<br>(0.011) | 0.429<br>(0.098) |
| Quota           | 0.213<br>(0.123)         | 0.299<br>(0.131) | 0.331<br>(0.056) | 0.519<br>(0.045) | 0.479<br>(0.021) | 0.738<br>(0.021) | 0.015<br>(0.919) | 0.072<br>(0.734) |
| Controls        | Y                        | Y                | Y                | Y                | Y                | Y                | Y                | Y                |
| No Inst. Errors |                          | Y                |                  | Y                |                  | Y                |                  | Y                |
| Observations    | 3,360                    | 1,568            | 1,147            | 538              | 854              | 396              | 944              | 439              |
| Dep. var. mean  | 0.996                    | 1.170            | 1.085            | 1.314            | 1.072            | 1.302            | 0.949            | 0.996            |

*Notes:* Results from a linear regression with  $p$ -values in parentheses and standard errors clustered at the individual level. The dependent variable is the profit forgone by recruiters: the difference between the maximum profit in each round and the profit actually accrued by their decision. Quota is a dummy that equals 1 if the participant was in the *Quota* Treatment. Only decisions in Part I are included. Columns (1)-(2) use all decisions with a female candidate, including same-gender pairs. Columns (3)-(4) restrict the sample to underperforming female candidates, this is, women with a lower performance signal than their counterparts. Columns (5)-(6) further restrict the sample to decisions where the under-performing female candidate has a higher signal than the quota candidate. Columns (7)-(8) only use decisions with two male candidates, for comparison. All models include participant-level dummies as controls: whether participants did not make any errors in the instructions, female gender, age above the median, white ethnicity, and no college degree. Even-numbered columns restrict the sample to participants who made no error in the instruction comprehension questions.

Table A.19: Hire likelihood of female candidates across Part II

|                         | Hired             |                   |                   |                  |                  |                   |
|-------------------------|-------------------|-------------------|-------------------|------------------|------------------|-------------------|
|                         | Under-performing  |                   |                   | Over-performing  |                  |                   |
|                         | (1)               | (2)               | (3)               | (4)              | (5)              | (6)               |
| Constant                | 0.743<br>(0.000)  | 2.273<br>(0.000)  | 2.468<br>(0.000)  | 0.974<br>(0.000) | 0.987<br>(0.000) | 0.965<br>(0.000)  |
| Quota                   | -0.056<br>(0.150) | -0.038<br>(0.318) | -0.015<br>(0.804) | 0.006<br>(0.617) | 0.006<br>(0.622) | -0.002<br>(0.931) |
| Performance             |                   | -0.098<br>(0.000) | -0.114<br>(0.001) |                  |                  |                   |
| Performance Other       |                   | -0.104<br>(0.000) | -0.133<br>(0.000) |                  | 0.002<br>(0.719) | 0.003<br>(0.734)  |
| Individual Controls     |                   | Y                 | Y                 |                  | Y                | Y                 |
| No Instruction Errors   |                   |                   | Y                 |                  |                  | Y                 |
| Observations            | 1,146             | 1,122             | 529               | 1,101            | 1,077            | 515               |
| Dependent variable mean | 0.715             | 0.716             | 0.635             | 0.977            | 0.978            | 0.967             |

*Notes:* Results from a linear regression with  $p$ -values in parentheses and standard errors clustered at the individual level. The dependent variable is a dummy that equals 1 if the female candidate received an offer higher or equal than her reservation wage. Quota is a dummy that equals 1 if the participant was in the *Quota Treatment*. Performance is the performance signal of the candidate. Performance Other is the performance signal of the other available candidate. Columns (1)-(3) restrict the sample to underperforming female candidates, this is, female candidates with a lower performance signal than their male counterparts. Columns (4)-(6) repeat the analysis for the rest of female candidates in different-gender pairs. Columns (2)-(3) and (5)-(6) also include participant-level dummies as controls: whether participants did not make any errors in the instructions, female gender, age above the median, white ethnicity, and no college degree. Columns (3) and (6) restrict the sample to participants who made no error in the instruction comprehension questions.

Table A.20: Backlash against under-performing women over groups of rounds of Part I

|                         | Hired             |                   |                   |                   |
|-------------------------|-------------------|-------------------|-------------------|-------------------|
|                         | Rounds 1-3        | Rounds 4-6        | Rounds 7-9        | Rounds 10-12      |
|                         | (1)               | (2)               | (3)               | (4)               |
| Constant                | 1.620<br>(0.000)  | 1.643<br>(0.000)  | 0.988<br>(0.000)  | 1.773<br>(0.000)  |
| Quota                   | -0.107<br>(0.007) | -0.106<br>(0.070) | -0.050<br>(0.352) | -0.086<br>(0.206) |
| Performance             | 0.015<br>(0.714)  | 0.019<br>(0.525)  | 0.053<br>(0.036)  | -0.008<br>(0.801) |
| Performance Other       | -0.106<br>(0.004) | -0.107<br>(0.000) | -0.056<br>(0.027) | -0.099<br>(0.000) |
| Individual Controls     | Y                 | Y                 | Y                 | Y                 |
| No Instruction Errors   |                   |                   |                   |                   |
| Observations            | 501               | 203               | 239               | 204               |
| Dependent variable mean | 0.796             | 0.773             | 0.824             | 0.745             |

*Notes:* Results from a linear regression with  $p$ -values in parentheses and standard errors clustered at the individual level. The dependent variable *Hired* is a dummy that equals 1 if the female candidate received an offer higher or equal than her reservation wage. *Quota* is a dummy that equals 1 if the participant was in the *Quota Treatment*. *Performance* is the performance signal of the candidate. *Performance Other* is the performance signal of the other available candidate. All regressions are estimated for under-performing female candidates and include participant-level dummies as controls: whether participants did not make any errors in the instructions, female gender, age above the median, white ethnicity, and no college degree.

Table A.21: Quota backlash against under-performing women over rounds of Part I

|                         | Offer             |                   |                   | Hired             |                   |                   |
|-------------------------|-------------------|-------------------|-------------------|-------------------|-------------------|-------------------|
|                         | (1)               | (2)               | (3)               | (4)               | (5)               | (6)               |
| Constant                | 2.317<br>(0.000)  | 4.151<br>(0.000)  | 4.502<br>(0.000)  | 0.824<br>(0.000)  | 1.503<br>(0.000)  | 1.692<br>(0.000)  |
| Round                   | -0.020<br>(0.341) | -0.023<br>(0.265) | -0.040<br>(0.178) | -0.002<br>(0.793) | 0.000<br>(0.962)  | -0.006<br>(0.498) |
| Quota                   | -0.487<br>(0.056) | -0.381<br>(0.119) | -0.645<br>(0.036) | -0.167<br>(0.033) | -0.121<br>(0.097) | -0.224<br>(0.033) |
| Round $\times$ Quota    | 0.039<br>(0.235)  | 0.025<br>(0.434)  | 0.063<br>(0.112)  | 0.012<br>(0.222)  | 0.006<br>(0.544)  | 0.014<br>(0.275)  |
| Performance             |                   | 0.326<br>(0.000)  | 0.264<br>(0.000)  |                   | 0.022<br>(0.123)  | 0.015<br>(0.456)  |
| Performance Other       |                   | -0.412<br>(0.000) | -0.442<br>(0.000) |                   | -0.092<br>(0.000) | -0.105<br>(0.000) |
| Individual Controls     |                   | Y                 | Y                 |                   | Y                 | Y                 |
| No Instruction Errors   |                   |                   | Y                 |                   |                   | Y                 |
| Observations            | 796               | 779               | 367               | 796               | 779               | 367               |
| Dependent variable mean | 2.072             | 2.076             | 1.899             | 0.773             | 0.773             | 0.728             |

*Notes:* Results from a linear regression with  $p$ -values in parentheses and standard errors clustered at the individual level. The dependent variable Offer equals the offer in dollars made to the under-performing female candidate (Columns (1)-(3)). The dependent variable Hired is a dummy that equals 1 if the female candidate received an offer higher or equal than her reservation wage (Columns (4)-(6)). Quota is a dummy that equals 1 if the participant was in the *Quota Treatment*. Performance is the performance signal of the candidate. Performance Other is the performance signal of the other available candidate. All regressions are estimated for under-performing female candidates. Columns (2)-(3) and (5)-(6) also include participant-level dummies as controls: whether participants did not make any errors in the instructions, female gender, age above the median, white ethnicity, and no college degree. Columns (3) and (6) restrict the sample to participants who made no error in the instruction comprehension questions.

Table A.22: Hire likelihood of female candidates — *Noisy Signal* Treatment

|                         | Hired             |                   |                   |                   |                   |                   |
|-------------------------|-------------------|-------------------|-------------------|-------------------|-------------------|-------------------|
|                         | Under-performing  |                   |                   | Over-performing   |                   |                   |
|                         | (1)               | (2)               | (3)               | (4)               | (5)               | (6)               |
| Constant                | 0.750<br>(0.000)  | 0.969<br>(0.000)  | 0.906<br>(0.000)  | 0.996<br>(0)      | 0.998<br>(0.000)  | 0.986<br>(0.000)  |
| Quota                   | -0.129<br>(0.001) | -0.117<br>(0.002) | -0.140<br>(0.008) | -0.007<br>(0.263) | -0.009<br>(0.186) | -0.018<br>(0.174) |
| Performance             |                   | 0.392<br>(0.000)  | 0.399<br>(0.000)  |                   | 0.003<br>(0.424)  | 0.004<br>(0.553)  |
| Performance Other       |                   | -0.236<br>(0.000) | -0.244<br>(0.000) |                   | -0.016<br>(0.075) | -0.012<br>(0.278) |
| Individual Controls     |                   | Y                 | Y                 |                   | Y                 | Y                 |
| No Instruction Errors   |                   |                   | Y                 |                   |                   | Y                 |
| Observations            | 1,106             | 1,084             | 591               | 1,012             | 978               | 570               |
| Dependent variable mean | 0.682             | 0.680             | 0.660             | 0.992             | 0.992             | 0.989             |

*Notes:* Results from a linear regression with  $p$ -values in parentheses and standard errors clustered at the individual level. The dependent variable is a dummy that equals 1 if the female candidate received an offer higher or equal than her reservation wage. Quota is a dummy that equals 1 if the participant was in the *Quota Treatment*. Performance is the performance signal of the candidate. Performance Other is the performance signal of the other available candidate. Columns (1)-(3) restrict the sample to underperforming female candidates, this is, female candidates with a lower performance signal than their male counterparts. Columns (4)-(6) repeat the analysis for the rest of female candidates in different-gender pairs. Columns (2)-(3) and (5)-(6) also include participant-level dummies as controls: whether participants did not make any errors in the instructions, female gender, age above the median, white ethnicity, and no college degree. Columns (3) and (6) restrict the sample to participants who made no error in the instruction comprehension questions.

Table A.23: Hire likelihood of male candidates — *Noisy Signal* Treatment

|                         | Hired             |                   |                   |                  |                   |                   |
|-------------------------|-------------------|-------------------|-------------------|------------------|-------------------|-------------------|
|                         | Under-performing  |                   |                   | Over-performing  |                   |                   |
|                         | (1)               | (2)               | (3)               | (4)              | (5)               | (6)               |
| Constant                | 0.656<br>(0.000)  | 0.880<br>(0.000)  | 0.806<br>(0.000)  | 0.994<br>(0)     | 1.012<br>(0.000)  | 1.017<br>(0.000)  |
| Quota                   | -0.042<br>(0.325) | -0.017<br>(0.676) | -0.024<br>(0.663) | 0.002<br>(0.614) | 0.002<br>(0.709)  | 0.009<br>(0.158)  |
| Performance             |                   | 0.363<br>(0.000)  | 0.383<br>(0.000)  |                  | -0.012<br>(0.086) | -0.015<br>(0.159) |
| Performance Other       |                   | -0.228<br>(0.000) | -0.208<br>(0.000) |                  | 0.008<br>(0.181)  | 0.015<br>(0.160)  |
| Individual Controls     |                   | Y                 | Y                 |                  | Y                 | Y                 |
| No Instruction Errors   |                   |                   | Y                 |                  |                   | Y                 |
| Observations            | 1,012             | 978               | 570               | 1,106            | 1,084             | 591               |
| Dependent variable mean | 0.633             | 0.634             | 0.625             | 0.995            | 0.995             | 0.997             |

*Notes:* Results from a linear regression with  $p$ -values in parentheses and standard errors clustered at the individual level. The dependent variable is a dummy that equals 1 if the male candidate received an offer higher or equal than her reservation wage. Quota is a dummy that equals 1 if the participant was in the *Quota Treatment*. Performance is the performance signal of the candidate. Performance Other is the performance signal of the other available candidate. Columns (1)-(3) restrict the sample to underperforming male candidates, this is, male candidates with a lower performance signal than their female counterparts. Columns (4)-(6) repeat the analysis for the rest of male candidates in different-gender pairs. Columns (2)-(3) and (5)-(6) also include participant-level dummies as controls: whether participants did not make any errors in the instructions, female gender, age above the median, white ethnicity, and no college degree. Columns (3) and (6) restrict the sample to participants who made no error in the instruction comprehension questions.

Table A.24: Expected Gender Performance by Treatment

| Sample                | Treatment           | Male<br>Candidates | Female<br>Candidates | Quota<br>Candidate | M-F<br>Paired t-test |
|-----------------------|---------------------|--------------------|----------------------|--------------------|----------------------|
| All Participants      | Baseline<br>(205)   | 5.69               | 5.62                 | -                  | 0.4774               |
|                       | Quota<br>(203)      | 5.82               | 5.85                 | 5.69               | 0.7661               |
|                       | Quota Male<br>(205) | 5.88               | 5.64                 | 6.17               | 0.0345**             |
| No Instruction Errors | Baseline<br>(90)    | 5.54               | 5.38                 | -                  | 0.3567               |
|                       | Quota<br>(102)      | 5.92               | 5.98                 | 5.42               | 0.6417               |
|                       | Quota Male<br>(98)  | 6.09               | 5.97                 | 6.31               | 0.4395               |

*Notes:* Expected Performance corresponds to participants' guess of the average number of correct answers for each group. In Baseline, these are male and female candidates. In Quota and Quota Male, the groups consist of male and female candidates, in addition to quota candidates that are female and male, respectively. Sample is divided between All Participants and those who answered all instructions questions correctly in Part I. Parentheses depict the number of observations. The last column reports the results of a paired t-test between the expected performance difference between male and female candidates for each sample and treatment.

## A.2 Survey Instructions — Baseline [Quota]

### Part 1

We begin with Part 1 of the study, which consists of **twelve hiring decisions**. One of these decisions will be randomly chosen to determine your dollar payoff from Part 1.

All hiring decisions proceed as follows:

- At the start of each hiring decision, you will be presented with a series of candidates that you can hire to increase your earnings.
- You have \$7 [\$12] that you can use to make offers to the candidates.
- You lose any money you don't use to make offers.

### How do offers work?

- Previously, each candidate declared a minimum wage, corresponding to the minimum amount they need to receive so they accept your offer.
- You do not know the candidate's minimum wage.
- If your offer is **lower** than their minimum wage, they will automatically decline your offer and you will just keep the money to yourself.
- If your offer is **greater or equal** than their minimum wage, they will automatically accept your offer and you will **pay them their minimum wage**, keeping the rest to yourself. Hence, your offer reflects your maximum willingness to pay to hire someone.

For example, imagine you offer \$7 to a candidate whose minimum wage is \$5. Then you will hire that candidate, paying them \$5 and keep the remaining \$2 to yourself. Now, imagine you offer \$4 to a candidate whose minimum wage is \$6. Then you will not hire that candidate and keep the \$4 to yourself.

### Why do you want to hire a worker?

The reason to hire a candidate is to potentially increase your earnings.

- Each candidate has previously answered 20 quiz questions.
- If you hired a candidate, you will earn \$1 for each question they got right in the **second half** of the quiz.
- To facilitate your decision, you will know how many questions they got right out of 10 questions of the first half of the quiz, which we call productivity.

The idea being that if a candidate did well in the first half of the quiz, they might be more likely to do well in the second half.

For example, if you hired a candidate for \$3 and later you learned they solved 5 questions right. Then you will earn \$5. That is \$2 more than if you hadn't hired them. If you hadn't hired that candidate, you would have lost the \$3 and missed on the \$5 payment.

The quiz questions are similar to those in standardized tests, and include quantitative and verbal questions. An example of a quantitative question is:

*If  $x^2 - 5x + 6 = 0$ , what are the possible values of  $x$ ?*

- A. 2, 3
- B. 1, 6
- C.  $-2, -3$
- D. 0, 5
- E.  $3, -2$

An example of a verbal question is:

*Select the word that is closest in meaning to 'ephemeral'.*

- A. Permanent
- B. Fleeting
- C. Tangible
- D. Significant
- E. Complex

### **Who are the candidates?**

Candidates are selected from a pool of 100 workers that have already completed the quiz and stated their minimum wage.

You will make 12 different hiring decisions.

- In each decision, we will randomly select **two candidates** from the pool of workers.
- Candidates are all unique and you will not see the same candidate more than once.

- Your decision is to decide how to split your initial \$7 [\$12] across the candidates.

**Quota Message** (only shown in Quota Female [Male]): “*To promote equality of opportunity, diversity, and inclusiveness in hiring decisions, you’ll be required to follow a policy of **gender quotas**.*”

- *In addition to the two candidates, we will randomly select a **female [male] worker** as the quota candidate. **You must offer the quota candidate \$5.***
- *You can’t opt out or choose what to do with those \$5. They must be used to hire the female [male] quota candidate."*

You are free to choose how many offers you make and how much to offer to each candidate as long as they do not add up to more than \$7 [\$12].

This is the end of the instructions. We provide a summary of the instructions before you proceed to make the hiring decisions.

You can access this summary when making your decisions by clicking on the button *Summary Instructions*.

### Instructions Summary

- You have \$7 to make offers to two randomly selected candidates from a pool of workers. You lose any money you don’t use to make offers.
- Each candidate declared a **minimum wage**, corresponding to the minimum amount they need to receive so they accept your offer.
- If your offer is lower than their minimum wage, they will automatically decline.
- If your offer is greater or equal than their minimum wage, they will automatically accept your offer and you will pay them their minimum wage, keeping the rest to yourself. Hence, **your offer reflects your maximum willingness to pay to hire someone.**
- Each candidate has previously answered 20 quiz questions.
- If you hire a candidate, you will earn \$1 for each question they answered correctly in the **second half** of the quiz.
- To facilitate your decision, you will know how many questions they got right out of the first half of the quiz, which we call productivity.

- *To promote gender diversity, you'll be required to follow a policy of **gender quotas**: in addition to the two candidates, we will randomly select a female [male] worker as the quota candidate. You **must offer \$5** to the quota candidate.*
- Offers cannot add up to more than \$7 [\$12, and you must offer \$5 to the quota candidate]. **You lose any money you don't use to make offers.**

When you are ready to continue, click below to start the first hiring decision.

## Part 2

Part 2 of the study consists of **twelve hiring decisions**. One of these decisions will be randomly chosen to determine your dollar payoff from Part 2.

All hiring decisions proceed in the same way as in Part I:

- In each decision, we will randomly select 2 candidates from the pool of workers.
- Candidates are all unique and you will not see the same candidate more than once.
- Your decision is to decide how to split your initial \$7 across the candidates [remaining \$7 across the non-quota candidates. You have to offer \$5 to the quota candidate].
- You lose any money you don't offer.
- You are free to choose how many offers you make and how much to offer to each candidate as long as they do not add up to more than \$7 [\$12].

## What's new?

Now, in Part 2, you will be **competing with other employers** to hire the candidates.

## Who are your competitors?

- In each decision, we will pair you with **two other participants** that will be your competitor employers.
- Your competitor employers are randomly selected, so you **will not be matched with the same competitors** every round.
- Competitors are other real participants like you, who make their decisions independently.

## How does competition work?

- You are free to make any offers to each candidate, and so are your competitors.
- You only hire a candidate **if your offer is the highest among the competitors**, and, as in Part I, if your offer is higher than the candidate's minimum wage. [As every competitor has to offer \$5 to the quota candidate, we will randomize who hires the quota candidate.]
- You don't know the offers of your competitors and they don't know your offers.

This is the end of the instructions. We provide a summary of the instructions before you proceed to make the hiring decisions.

You can access this summary when making your decisions by clicking on the button *Summary Instructions*.

### Instructions Summary

- You have \$7 [\$12] to make offers to two randomly selected candidates from a pool of workers. You lose any money you don't use to make offers.
- Each candidate declared a **minimum wage**, corresponding to the minimum amount they need to receive so they accept an offer.
- If your offer is lower than their minimum wage, they will automatically decline.
- In Part 2 you **compete against two other employers** at each hiring decision. The competitors are randomly assigned and are not the same across decisions.
- To hire a candidate, your offer has to be greater than your competitors' offers and equal or greater than the candidate's minimum wage.
- If your offer is accepted, you will pay the candidate their minimum wage, keeping the rest to yourself. Hence, **your offer reflects your maximum willingness to pay to hire someone**.
- Each candidate has previously answered 20 quiz questions.
- If you hire a candidate, you will earn \$1 for each question they answered correctly in the **second half** of the quiz.
- To facilitate your decision, you will know how many questions they got right in the first 10 questions of the quiz.
- *To promote gender diversity, you'll be required to follow a policy of gender quotas: in addition to the two candidates, we will randomly select a female [male] worker as the quota candidate. You must offer \$5 to the quota candidate. Offers cannot add up to more than \$12*

- You lose any money you don't use to make offers.

When you are ready to continue, click below to start the first hiring decision.

## A.3 Comprehension Questions

Questions for the Baseline Treatment are marked as B, questions for the Quota Treatment are marked as Q.

### Part 1

**B1/Q1.** How many decisions will you make in Part I?

- 12 decisions
- 5 decisions
- 20 decisions

**B2.** How much money do you have to make offers to the candidates?

- \$7
- As much as I want
- \$5

**Q2.** How much money do you have to make offers to the candidates?

- \$12
- As much as I want
- \$5

**B3/Q3.** Assume you offer \$6 to a candidate. You later learn that their minimum wage was \$4. Is your offer accepted? And if so, how much do you pay them?

- They decline your offer and you do not hire the candidate. You keep the \$6 to yourself.
- Your offer is accepted and you hire the candidate. You pay them \$6.
- Your offer is accepted and you hire the candidate. You pay them \$4.
- They decline your offer and you do not hire the candidate. You keep \$4 to yourself.

**B4/Q4.** Assume you offer \$7 to a candidate. You later learn that their minimum wage was \$4. Is your offer accepted? And if so, how much do you pay them?

- They decline your offer and you do not hire the candidate. You keep the \$7 to yourself.
- Your offer is accepted and you hire the candidate. You pay them \$7.
- Your offer is accepted and you hire the candidate. You pay them \$4.
- They decline your offer and you do not hire the candidate. You keep \$4 to yourself.

**B5/Q5.** Assume you offer \$3 to a candidate. You later learn that their minimum wage was \$5. Is your offer accepted? And if so, how much do you pay them?

- They decline your offer and you do not hire the candidate. You keep all \$3 to yourself.
- Your offer is accepted and you hire the candidate. You pay them \$3.
- Your offer is accepted and you hire the candidate. You pay them \$5.
- They decline your offer and you do not hire the candidate. You keep \$5 to yourself.

**B6/Q6.** Assume you hired a candidate for \$4. You later learn that they solved 7 questions right in the second half of the quiz. What is your payoff from hiring this candidate?

- Your payoff is \$4 because you paid them a \$4 wage.
- Your payoff is \$7 because the candidate answered 7 questions right.
- Your payoff is \$7. Your initial amount.
- Your payoff is \$0 because you did not hire the candidate.

**B7/Q7.** Assume you hired a candidate for \$7. You later learn that they solved 4 questions right in the second half of the quiz. What is your payoff from hiring this candidate?

- Your payoff is \$7 because you paid them a \$7 wage.
- Your payoff is \$4 because the candidate answered 4 questions right.
- Your payoff is \$7. Your initial amount.
- Your payoff is \$0 because you did not hire the candidate.

**Q8.** Can you make any offer to the female quota candidate?

- Yes, you can offer the female quota candidate any amount.

- No, there is a mandated offer of \$5 for the female quota candidate.

**B8** Assume that you have the following candidates to choose from: Candidate A and Candidate B. Candidate A — # Correct answers in first 10 Questions: 6. Candidate B — # Correct answers in first 10 Questions: 5. Can you offer \$6 to Candidate A and \$5 to Candidate B?

- Yes, you can.
- No, you can't since the offers add up to more than your initial amount of \$7.

**Q9.** Assume that you have the following candidates to choose from: Candidate A, Candidate B, and a Quota Candidate. Candidate A — # Correct answers in first 10 Questions: 6. Candidate B — # Correct answers in first 10 Questions: 5. Quota Candidate — # Correct answers in first 10 Questions: 5. Can you offer \$6 to Candidate A and \$7 to Candidate B?

- Yes, you can.
- No, you can't since after paying \$5 to the quota candidate the offers add up to more than your remaining funds.

**B9.** Can you offer \$7 to Candidate A and \$0 to Candidate B?

- Yes, you can since the offers do not add up to more than your initial amount of \$7.
- No, you can't.

**Q10.** Can you offer \$7 to Candidate A and \$0 to Candidate B?

- Yes, you can since after paying \$5 to the quota candidate the offers do not add up to more than your remaining funds.
- No, you can't.

**B10.** Can you offer \$3 to Candidate A and \$4 to Candidate B?

- Yes, you can since the offers do not add up to more than your initial amount of \$7.
- No, you can't.

**Q11.** Can you offer \$4 to Candidate A and \$3 to Candidate B?

- Yes, you can since after paying \$5 to the quota candidate the offers do not add up to more than your remaining funds.
- No, you can't.

## Part 2

**B11/Q12.** How many competitors will you have in each decision?

- Two
- Three

**B12/Q13.** You offer \$3 to a [non-quota] candidate with a \$2 reservation wage. If a competitor offers \$4 to that same candidate, do you hire the candidate?

- Yes, because your offer is higher than the candidate's reservation wage.
- No, because the competitor's offer is higher than yours.

**B13/Q14.** You offer \$4 to a [non-quota] candidate with a \$5 reservation wage. If your competitors offer \$2 and \$3 to that same candidate, do you hire the candidate?

- Yes, because your offer is higher than the competitors' offers.
- No, because though your offer is higher than the competitors', it's lower than the reservation wage.

**B14/Q15.** You offer \$5 to a [non-quota] candidate with a \$4 reservation wage. If competitors offer \$2 and \$3 to that same candidate, do you hire the candidate?

- Yes, because your offer is higher than the competitors' offers and the reservation wage.
- No, because the competitors' offers are lower than the reservation wage.

## A.4 Policy Views Questions

All questions are selected from Gallup polls.

**1-2. Gender in Hiring Decisions.** *When companies consider gender as a factor in hiring decisions in order to increase the gender diversity of the workplace, do you think...*

1. *This makes the overall admissions process of these companies...* [More-Less Fair]
2. *That the candidates who are hired to these companies are...* [More-Less Qualified]

**3. DEI in college admissions.** *Which comes closer to your view about evaluating students for admission into a college or university...*

- a. Applicants should be admitted solely on the basis of merit, even if that results in few minority students being admitted.
- b. An applicant's racial and ethnic background should be considered to help promote diversity on college campuses, even if that means admitting some minority students who otherwise would not be admitted

**4. Gender quotas in boards of directors.** *Should companies impose gender quotas to increase diversity on their board of directors?* [Agree - Disagree]

**5. Equal opportunities for women.** *When it comes to giving women equal opportunities with men, do you think our country...* [Has gone too far - Has not gone far enough]