

Melbourne Institute of
Applied Economic and Social Research

Working Paper Series

Managers and Gatekeepers in the Age of AI

Cassandra Merritt
Jacob Dominkski
Christipher Hoy
Yong Suk Lee

Working Paper 02/26

Managers as Gatekeepers in the Age of AI*

Cassandra Merritt[†] Jacob Dominski[‡] Christopher Hoy[§] Yong Suk Lee[¶]

March 26, 2026

Abstract

Artificial intelligence (AI) is frequently cast as a transformative technology that will raise productivity while displacing human work, yet organizational adoption remains uneven and aggregate effects are mixed. We examine whether middle managers contribute to this gap by acting as gatekeepers to AI adoption. In a pre-registered survey experiment of 2,000 managers in the United States and United Kingdom, respondents were randomly assigned to view videos summarizing recent evidence on AI's productivity benefits, its labor-displacing potential, or a placebo control. Exposure to information about labor displacement leads to a large reduction in intended AI adoption and advocacy (by 0.4–0.5 standard deviations) and a moderate reduction in staffing intentions (by 0.2 standard deviations). In contrast, information about productivity benefits has no significant average effect, although it increases advocacy among managers with low prior familiarity with AI. These findings indicate that middle managers' responses to the information environment shape both technology adoption and employment intentions. Rather than inducing substitution away from labor and toward AI, information about AI's labor-displacing potential leads managers to scale back both planned AI adoption and their staffing intentions.

JEL classification: J23, J24, O33, L2, M54, D83

Keywords: Artificial Intelligence, Technology adoption, Managers, Gatekeepers, Productivity, Labor displacement

*We thank Pádraig Slattery and Appo Nikiema for excellent research assistance. Lee thanks the Notre Dame Ethics Initiative for funding this project. This study is pre-registered under AEA RCT ID: AEARCTR-0017306 and was approved by the University of Notre Dame Institutional Review Board (Protocol ID: 25-08-9497).

[†]Keough School of Global Affairs, University of Notre Dame. Email: cjmerritt@nd.edu.

[‡]Institute for Ethics and the Common Good, University of Notre Dame. Email: jdomins2@nd.edu.

[§]University of Melbourne, and the World Bank. Email: choy@worldbank.org.

[¶]Keough School of Global Affairs, and Institute for Ethics and the Common Good, University of Notre Dame. Email: yong.s.lee@nd.edu.

1 Introduction

The rapid ascent of artificial intelligence (AI) in public discourse fuels growing tension between expectations of substantial productivity gains and concern about labor displacement. However, empirical evidence remains mixed. Experimental studies have documented considerable productivity gains at the individual level (Noy and Zhang, 2023; Brynjolfsson et al., 2023; Peng et al., 2023), but organizational-level impacts have been more modest (Challapally et al., 2025; Humlum and Vestergaard, 2025a), and adoption across organizations remains uneven (Bick et al., 2024; Kerr and Mertens, 2024; Korst et al., 2025; McClain and Sidoti, 2025; U.S. Chamber of Commerce and Teneo Research, 2025). Estimates of aggregate productivity gains diverge widely, from Acemoglu’s modest 0.07 percentage points per year in TFP (Acemoglu, 2024) to McKinsey’s projection of up to 3.4 percentage points annually when combining generative AI with other technologies (McKinsey Global Institute, 2023). The evidence on the impacts of AI on the labor-market are also mixed (Hampole et al., 2025; Dominski and Lee, 2025; Humlum and Vestergaard, 2025a; Brynjolfsson et al., 2025). This gap between AI’s promise to transform the economy and the more mixed organizational- and economy-wide evidence presents a key puzzle in understanding AI’s impact on the economy and labor market.

Senior executives typically advocate for rapid adoption of AI, hoping it will boost productivity and lower labor costs, but they can face significant implementation barriers (Yotzov et al., 2026; McKinsey, 2025; IBM, 2024). Managers, especially those with day-to-day operational authority, play a crucial role as organizational “gatekeepers,” shaping the pace of technology adoption and its implications for staff in their business units (Aghion and Tirole, 1997; Atkin et al., 2017). As a result, even when senior leadership pushes for greater AI use, managers can reshape, slow down, or stall deployment. This raises a key question: what role do managers play in mediating AI’s impact on the economy? We examine this by investigating how information about AI’s promised productivity gains and its labor-replacing potential affects managers’ plans for AI adoption and staffing.

We conducted a broadly representative survey experiment with 2000 managers in the United States and the United Kingdom. We collected extensive baseline information from respondents, including managers’ own AI perceptions, the prevalence of AI within their teams and firms, existing AI policies, managerial decision authority, and anticipated plans related to AI and hiring. We then randomly assigned managers to view one of three 2-minute videos: (i) a placebo control, providing generic facts on major AI firms; (ii) a treatment summarizing empirical evidence on AI’s potential productivity-enhancing effects (productivity treatment group); or (iii) a treatment summarizing empirical evidence on AI’s potential labor-displacing effects (labor treatment group). After viewing the video, we elicited their attitudes towards adopting/advocating for AI in their workplace and their intentions/advocacy for adjustments to staffing levels. We also measured a set of belief channels that may mediate treatment effects, including perceived benefits and risks, expected productivity gains, concerns about disruption and employee welfare, and ethical considerations.

Our baseline data reveal that AI adoption among surveyed managers and their firms is already

substantial and broadly positive. Over 70 percent of managers report positive sentiments toward AI, and we find similar shares across several other measures—team sentiments, executive support for AI use, expected AI adoption—all before viewing the survey’s informational video. Current use is well-established: around three-quarters of managers report their teams use AI at least weekly, and over 60 percent report that a majority of their team members are active users. These patterns are consistent with the rapid adoption trajectories documented across recent surveys of executives, firms, and workers. More than half of surveyed managers expect AI to be labor-augmenting, with only 13% expecting it to be labor-displacing.¹ In short, our respondents are neither unfamiliar with AI nor pessimistic about its effects—making the informational treatments a test of whether credible statements can shift the intentions and views of typical managers in the current labor market.

Managers respond to information about AI’s labor-displacing potential by substantially reducing their intent to advocate for AI adoption and lowering their staffing intentions. The labor treatment group reported 0.4–0.5 standard deviations lower intent to adopt or advocate for AI use compared to the placebo control group. To contextualize this magnitude, roughly two in three managers exposed to the labor treatment scored below the control-group average (compared to one in two in the control group). Mediation analysis of the labor treatment effects on managers’ AI adoption and advocacy suggests that the effect can be explained equally by managers decreasing their perceived benefits of AI and increasing their perceived risks. This treatment also resulted in 0.2 standard deviations lower staffing intentions, which is equivalent to about 58 percent of treated managers falling below the control-group average on staffing intentions (versus 50 percent in the control group).² The negative effects of the labor treatment are largely stable across manager and firm characteristics, suggesting a broadly shared pullback in both AI and staffing intentions. One exception is that there were significantly larger negative treatment effects across AI adoption and staffing intentions indices among respondents who had expected to hire more staff over the next year. These heterogeneous treatment effects are consistent with evidence on AI’s labor-displacement effects, leading managers to prioritize maintaining the status quo.

Information about AI’s productivity benefits does not have meaningful effects on AI adoption or staffing. We can rule out unconditional average effects exceeding 0.2 standard deviations for any outcome (we also find generally shrinking point estimates as we add controls). This result is likely due to managers being already informed and optimistic about AI’s productivity potential in their own workplace. In our heterogeneity analysis, we find strong direct relationships between baseline characteristics and our AI outcome indices. Managers with greater decision authority, more AI training and team use, and clearer firm-level AI policies show stronger AI adoption and advocacy. In terms of treatment effects, the productivity treatment has a significant positive effect on AI advocacy among managers with low team AI use, less technical authority, and less executive support for AI. These findings suggest that the productivity treatment shifts AI advocacy most where AI familiarity is weakest.

Our findings have several implications for understanding AI in today’s economy. First, managers

¹ In the remainder, 26% thought it would have mixed effects and 8% thought there would be no effects.

² This effect cannot be explained by the mediation analysis we conducted in this study.

matter. The narrative around AI adoption in organizations often centers on top-down executive initiatives or bottom-up worker experimentation. However, our findings show that middle managers can be decisive bottlenecks, or bridges, in AI adoption and staffing plans, independent of executive initiatives or firm-level policies.

Second, the information environment matters. Our findings indicate that singing the praises of AI's potential productivity benefits will do little to shift managers' perceptions, as many already hold optimistic views and have experimented with AI. On the other hand, information about AI's potential for labor displacement may trigger broad pullbacks in both AI advocacy and staffing intent. Initial media coverage and industry reports around AI often focused on its transformative potential. More recently, labor-displacement concerns have become the dominant narrative in AI-related reporting. This shift in messaging, and the society-wide concern over labor displacement, may influence AI adoption.

Third, where the manager stands in terms of AI familiarity matters. Managers with low AI familiarity, limited technical authority, and weak executive support are more sensitive to information about AI's potential benefits. These are precisely the population most likely to rely on external information to form expectations about a new technology. As AI diffuses unevenly across firms and to the Global South, this finding suggests differential reception of AI across contexts.

Finally, our findings speak directly to the evolving landscape in AI adoption. Figure 1 presents major survey estimates of AI adoption since 2023. It shows a clear upward trend but also substantial dispersion across the populations being surveyed. Executives report considerably higher levels of AI adoption compared to workers. Our estimate, drawn predominantly from middle managers, sits right in the middle, underscoring their role as gatekeepers whose decisions are shaped by the narratives surrounding AI and its implications for their organizations and workers, and who can meaningfully affect its trajectory.

Our paper contributes to several strands of the literature. To the best of our knowledge, this is the first cross-country randomized survey experiment of middle managers focusing on AI. As shown in Figure 1, to date, most empirical evidence about attitudes towards AI comes from broad surveys of individuals or firms. Large, reputable public agencies, think tanks, and research teams regularly document rising AI awareness and uneven uptake, with firm-level adoption climbing sharply since 2024, alongside mixed sentiments (Bick et al., 2024; Bonney et al., 2024; Crane et al., 2025). Those surveys are invaluable for levels and trends, but they are descriptive and typically focus on executives or workers rather than isolating the managerial layer that implements day-to-day change. Within management research, the closest studies are attitudinal models (Cao et al., 2021; Majrashi, 2024) and qualitative work on trust, risk, and acceptance (Daly et al., 2025; Wen et al., 2025). These illuminate mechanisms but rarely link information to managers' team-level adoption or hiring intentions in a representative experiment.

We contribute to this knowledge base in three ways. First, we generate causal evidence about how AI's potential productivity gains and labor-market effects influence the likelihood that managers adopt or advocate AI in their team, and whether they intend to hire, maintain, or reduce staff. Second, we pair those outcomes with belief mechanisms, allowing us to map which beliefs actually shift and whether belief shifts co-move with intentions. Third, we collect the organizational context that conditions

implementation: decision authority, baseline use, policies, function, and firm size. This allows us to document heterogeneity that matters for practice (who is responsive, under what constraints) and for policy (where complementary investments may be most needed). Together, this moves beyond questions of “who is using AI” to “what causes managers to implement or stall AI adoption, through which beliefs, and with what staffing implications,” while still providing valuable updated descriptive statistics and cross-country comparisons on AI use, organizational policies, managerial adoption, and staffing intentions.

Our paper also contributes to the literature on organizational barriers to technology adoption. Aghion and Tirole (1997) establish theoretically that real authority over decisions often resides not with senior leadership but with subordinates who control day-to-day operations and information flows. Atkin et al. (2017) demonstrate this empirically: in Pakistani soccer-ball factories, workers resisted a cost-saving technology by misinforming owners, resulting in puzzlingly low adoption despite clear net benefits. In the AI context, Tambe et al. (2019) document a substantial gap between AI’s promise and its reality in management, identifying ethical and legal complexity and adverse employee reactions as key barriers. Yet to the best of our knowledge no study has causally examined how managers themselves respond to information about AI. We provide the first experimental evidence that middle managers, as organizational gatekeepers, significantly adjust both AI adoption and staffing intentions in response to information about AI’s productivity and labor-displacement potential.

The rest of this paper is structured as follows: Section 2 summarizes the study methodology, including sampling, survey structure, treatment design and assignment, key outcomes, data cleaning, and empirical methods. Section 3 presents stylized facts based on managers’ responses. Section 4 presents the effects of our information treatments on our pre-specified main outcome indices and extends the analysis through heterogeneity and decomposition analyses. Section 5 concludes with discussions, implications, and future research avenues.

2 Study Design, Implementation, Methodology

2.1 Recruitment and Data Collection

Recruiting Respondents: The survey was administered through Qualtrics and distributed to managers through the broadly nationally representative samples found on Prolific, which handled participant recruitment, compensation, and country quotas. Respondents were recruited through Prolific notifications and were compensated \$4.40 for successful survey completion. Respondents were also entered into a lottery for five \$100 prizes. The survey was approximately 20-25 minutes long, with a median completion time of 21 minutes and 51 seconds and a median compensation \$12.08 per hour.

Our target population consisted of individuals currently in management roles who supervise at least one employee and make (or influence) decisions regarding their team’s work. In practice, this population spans first-line supervisors, middle managers, and senior leaders across functions, industries, and firm

sizes. The distribution of the realized sample by organizational role is reported in Appendix Table A.37.

Pre-randomization screeners determined eligibility, and only those who passed all checks proceeded to the randomized portion of the survey. We used two approaches to ensure that we reach managers and minimize misclassification. We targeted participants who answered “yes” to the following Prolific screener:

At work, do you have any supervisory responsibilities? In other words, do you have the authority to give instructions to subordinates?

At the start of the survey, we also confirmed that respondents currently supervise at least one employee. We also included the following additional inclusion criteria: currently employed (full-time or part-time), residing in the United States or the United Kingdom, and over 18 years of age. Anyone who failed to meet any of these inclusion criteria was automatically screened out and excluded from analysis.

Pre-analysis plan and registration: Prior to fielding the main survey, we wrote a pre-analysis plan specifying the primary outcomes, index construction, main specifications, heterogeneity analyses, and planned robustness checks. We pre-registered the study on the AEA RCT Registry (AEARCTR-0017306).

Representativeness: Because recruitment occurs on an online panel, the realized sample differs from the broader population of managers in two ways. Relative to the general population in the US and UK, Prolific participants tend to be more digitally engaged, comfortable with online research tasks, and more concentrated in urban/metro areas. In comparison to Prolific’s general participant pool, our manager sample tends to be older, more educated, and have more work experience than the platform’s average. These two sources of selection may also interact, shaping the demographic or occupational composition of Prolific’s manager population. We document the realized composition and report any imbalances in Section 3.

Data Collection: Data collection occurred in two stages. The pilot survey was fielded on Prolific to approximately 300 managers on September 25th, 2025, and reached full completion within 24 hours. Following the pilot, we refined the instrument before fielding the main wave from November 24th to November 25th, 2025. Data collection in both countries ran concurrently until sample targets were reached.

We aimed to recruit 2,000 managers, with 1,000 from the US and 1,000 from the UK. 2,100 people entered the survey, and 1995 people reached the randomization stage. Of these 995 were from the UK and 1000 were from the US. Of those that did not reach the randomization portion of the survey, fifteen did not give consent, two were under 18, and 88 did not have supervisory responsibilities.

2.2 Survey Structure

The study used a single online survey instrument consisting of a pre-treatment baseline questionnaire, a randomized informational video module, and a post-treatment questionnaire measuring outcomes and mechanisms.

The baseline questionnaire collected demographic and occupational information (e.g., age, gender,

education, industry, firm size, and managerial role) using items adapted from reputable instruments such as the U.S. Current Population Survey. It also measured team and unit characteristics, decision authority, and firm AI policies, AI support and attitudes, and baseline use. The full survey instrument can be found in Appendix Section C.1.

The randomized informational video module assigned each participant to one of three short videos: a neutral control, a productivity treatment, or a labor-market treatment. The control video described major AI developers and popular AI tools, without making claims about productivity or labor market effects of AI.³ The productivity treatment summarized recent research and reports on the productivity effects of AI tools (Peng et al., 2023; Noy and Zhang, 2023; Brynjolfsson et al., 2023; Goldman Sachs Global Investment Research, 2023; McKinsey & Company, 2023). Finally, the labor-market treatment video discussed evidence on task and occupational exposure to AI and its early labor-market effects (Eloundou et al., 2024; Acemoglu and Restrepo, 2020; Brynjolfsson et al., 2025; Hatzius et al., 2023; Cutter, 2025). Full scripts are provided in Appendix Section C.2. The videos⁴ were comparable in length (ranging from 1:55 to 2:18) and were designed to have a similar tone, format, production style, and structure.⁵ Each video was followed by a content-based attention check to assess basic engagement with the video. Time-on-page data were automatically recorded in Qualtrics to monitor exposure and engagement. Random assignment to the videos was implemented within the Qualtrics platform.

The post-treatment questionnaire elicited the study’s primary outcomes and the mechanisms underlying these beliefs. Primary outcomes include multiple 1–5 Likert items on AI adoption and expansion intentions, AI advocacy, hiring and layoff intentions, and hiring advocacy. Post-treatment belief items capture mechanisms such as perceived benefits and risks of AI, expected productivity gains, concerns about labor disruption, ethical salience, and views on employee welfare in AI decision-making. Questions on AI use, firm policies, outcomes, and belief channels were developed by the research team to measure theoretically relevant constructs specific to managerial decision-making. These items and the three informational videos have not been used in prior studies, but were piloted to assess item distributions, inter-item correlations, completion rates, and treatment responsiveness. Details about the outcome questions are outlined in Section 2.4.

2.3 Experimental Design

Randomization: Respondents were randomly assigned at the individual level to one of three video conditions (control, productivity, or labor-market). Assignment was implemented in Qualtrics imme-

³In the pilot, the control video described the historical development of AI. After reviewing pilot responses, we redesigned the control video for our main survey as a listing of generic facts about AI labs—like where they are headquartered and what models they are known for—to reduce the risk that the control condition shifted productivity beliefs.

⁴The videos are viewable at the following links:

Control: <https://tinyurl.com/MDHL2026-Control-Video>

Productivity: <https://tinyurl.com/MDHL2026-Productivity-Video>

Labor: <https://tinyurl.com/MDHL2026-Labor-Video>

⁵Additionally, all videos used the same narrator.

diately after completion of the pre-treatment baseline module and stratified by country, gender, and generational cohort. Treatment was allocated with probabilities of 40% to the control condition and 30% to each treatment condition. Recruitment quotas targeted approximately equal sample sizes in the United States and the United Kingdom. Random assignment was executed using Qualtrics' randomizer tool and recorded as an embedded variable in the survey data. Randomization occurred after respondents completed all pre-treatment baseline questions and immediately before the video page. This sequencing ensured that randomization was not affected by respondents who exited before completing the baseline section and that all randomized participants had provided comparable pre-treatment data.

2.4 Outcomes

Our primary outcomes are four standardized indices measuring managers' stated intentions regarding AI adoption and staffing. The *AI Adoption* index captures intentions to adopt or expand AI use in the respondent's team/unit. This index averages four items asking how likely the respondent is to begin adopting new AI, increase current AI use, reduce or slow AI use, and halt or avoid new AI adoption. Items use a 1–5 scale from “Not at all likely” to “Extremely likely.” The two “reduce/slow” and “halt/avoid” items are reverse-coded so that higher values indicate greater adoption intent.

The *AI Advocacy* index captures willingness to advocate for expanding AI use within the organization. This index averages two agreement items: advocating expanding AI use and advocating reducing AI use. Both are measured on 1–5 agreement scales. The “reduce” item is reverse-coded so that higher values represent stronger pro-AI advocacy.

The *Staffing Intentions* index measures expected headcount changes over the next year for the self-reported, three most common occupation groups in the respondent's team. For each group, respondents report expected headcount changes on a 1–5 scale (1 = major reduction, 3 = no change, 5 = major increase). Higher values indicate stronger hiring intentions.

The *Staffing Advocacy* index captures willingness to advocate for hiring vs. layoffs. This index averages two three-category advocacy items: one on whether the respondent would advocate for hiring additional employees and one on whether the respondent would advocate for layoffs or workforce reductions (reverse-coded). For each item, responses are coded so that higher values indicate more pro-hiring positions, while “Unsure” and “I would not feel comfortable voicing my opinion” are treated as missing.

We constructed each index in three steps. First, we standardized each component item using the mean and standard deviation of the pooled control group. Second, we averaged the standardized items of each index, requiring at least half of the component items to be non-missing. For those missing more than half the items, the index was set to missing. Third, we standardized the resulting index using the mean and standard deviation of the pooled control group. Responses of “Unsure” and “Not relevant” were treated as missing. Full item wording can be seen in the survey in Appendix Section C.1.

Our main analysis uses the sample of randomized respondents with non-missing values for all four primary outcome indices. Because index construction requires at least half of an index's component

items to be observed, some randomized respondents are excluded from this analytical sample when one or more primary indices cannot be constructed. We refer to this post-randomization exclusion as attrition. Appendix A.6 reports descriptive and regression-based evidence on attrition patterns, where we find little evidence of differential attrition and confirm our results are robust to our sample restriction.

In addition, we analyze post-treatment belief measures. Belief index construction, including component item wording, coding, and standardization, is described in Appendix Section E.2.

2.5 Empirical Methodology

2.5.1 Main Analysis

Our main experimental results are captured by estimating the following equation,

$$y_i = \alpha_c + \sum_{k \in \{1,2\}} \delta_k \text{arm}_{i,k} + x_i' \beta + \varepsilon_i, \quad (1)$$

where i , c , and k indexes individual, country, and treatment arm respectively. Our coefficient of interest is δ_k that capture the effect of treatment arm k (productivity or labor distress) relative to the neutral control video. α_c captures the country fixed effects. In our primary specification, we do not include any controls (dropping the $X\beta$ term), but in alternative specifications we include covariates to rule out confounding variation from imbalances in baseline characteristics.⁶ Across all specifications, we use heteroskedasticity robust standard errors.

Given the imbalance across the US and UK samples seen in Section 3, we also run our analysis in each country subsample to allow the possibility of distinct parameters as below.

$$y_i = \alpha_c + \sum_{k \in \{1,2\}} \delta_{c,k} \text{arm}_{i,c,k} + x_i' \beta_c + \varepsilon_i. \quad (2)$$

Note the subscripts c , indexing country, are retained in this notation to emphasize these specifications allow the estimated effects of the treatment arms and associations with the covariates to vary by country.

2.5.2 Heterogeneity Analysis

We estimate heterogeneous treatment effects by splitting each pre-treatment covariate of interest at the median. That is, for each covariate z_i , we define a dummy variable $Dz_i = \mathbf{1}[z_i > \text{median}(z)]$ and estimate,

⁶Seemingly random item non-response across increasing sets of controls results in little precision gains.

$$y_i = \alpha_c^H + \sum_{k \in \{1,2\}} \delta_k^H \text{arm}_{i,k} + \sum_{k \in \{1,2\}} \gamma_k^H \text{arm}_{i,k} \cdot Dz_i + \zeta Dz_i + x_i' \beta^H + \varepsilon_i^H, \quad (3)$$

where δ_k^H is the treatment effect for respondents below the median and $\delta_k^H + \gamma_k^H$ is the treatment effect for those above the median. We examine whether γ_k^H , the differential treatment effect between the two groups, is significant for different covariates of interest. ζ captures the mean difference in y between the two subgroups. All specifications pool both countries and include country fixed effects. As before, we use heteroskedasticity robust standard errors.

2.5.3 Decomposition Analysis

To examine which beliefs account for the estimated treatment effects, we follow Gelbach (2016). We perform this analysis separately for each treatment arm, restricting the sample to the relevant treatment group and the control group. For a given treatment arm k , the baseline specification is,

$$y_i = \alpha_c^B + \delta_k^B \text{arm}_{i,k} + x_i' \beta^B + \varepsilon_i^B, \quad (4)$$

and the full specification augments this with two post-treatment belief indices — perceived AI benefits ($m_{1,i}$) and perceived AI risks ($m_{2,i}$) — as candidate mediators,

$$y_i = \alpha_c^F + \delta_k^F \text{arm}_{i,k} + \theta_1 m_{1,i} + \theta_2 m_{2,i} + x_i' \beta^F + \varepsilon_i^F. \quad (5)$$

The difference $\hat{\delta}_k^B - \hat{\delta}_k^F$ decomposes exactly into mediator-specific contributions. For each mediator $j \in \{1, 2\}$, we estimate an auxiliary regression of the mediator on treatment assignment,

$$m_{j,i} = \alpha_c^j + \pi_{j,k} \text{arm}_{i,k} + x_i' \beta^j + \varepsilon_i^j. \quad (6)$$

The contribution of mediator j is $\hat{\theta}_j \times \hat{\pi}_{j,k}$, the product of the mediator's coefficient in the full outcome regression and the treatment's effect on that mediator, such that,

$$\hat{\delta}_k^B - \hat{\delta}_k^F = \hat{\theta}_1 \hat{\pi}_{1,k} + \hat{\theta}_2 \hat{\pi}_{2,k}. \quad (7)$$

The residual $\hat{\delta}_k^F$ is the unexplained direct effect. We conduct this analysis for each of our four primary outcome indices and both treatment arms. Across all specifications, we use heteroskedasticity robust standard errors.

3 Descriptive Findings

3.1 AI use over time based on leading surveys, including this study

There has been a rapid increase in AI use over the last three years, such that by the end of 2025, our survey along with evidence across the literature suggests a substantial share of firms and workers are using AI. Figure 1 plots AI adoption estimates from over 20 surveys fielded between 2023 and late 2025, primarily in the United States.⁷ Two features emerge immediately. First, there is a broadly positive trend across all respondent populations since 2023, consistent with the standard pattern that diffusion of new technologies is neither immediate nor uniform (Kalyani et al., 2025). Second, there is substantial cross-survey variance at any given point in time. This variance partly reflects the natural unevenness of technology diffusion, but it is further amplified by differences in the populations being surveyed—varying in organizational hierarchy and geography—and in how AI use is measured, whether as binary firm adoption, any individual use, weekly frequency, or task intensity.

Within this variation, a clear hierarchy is visible. The highest adoption rates are reported by senior decision-makers (Korst et al. 2023, 2024, 2025) and in a survey targeting workers in highly AI-exposed occupations (Humlum and Vestergaard 2025a,b). Firm-level and leadership-oriented measures consistently exceed individual worker self-reports of AI use at work. This gap raises a question central to our paper: whether middle managers serve as key gatekeepers in AI adoption, occupying a position between executive enthusiasm and the rate at which workers actually integrate AI into their tasks. Our survey’s estimates—plotted at the right edge of the figure—place managers’ reports of their teams’ AI use above the worker-level estimates but below the executive tier, consistent with this intermediary role.

3.2 Managers’ baseline perceptions and plans for AI Adoption

Managers in our sample report that AI use is already widespread and that adoption is likely to continue expanding. Figure 2 presents the distribution of baseline measures of AI adoption and use among the managers in our sample. The left column reports current AI use and attitudes; the right column reports the components and composite index of our primary AI adoption outcome.

The overall picture is one of established and expanding AI use. Approximately 70 percent of managers expect AI adoption in their team to expand over the coming year, with fewer than 3 percent anticipating any reduction. This optimism is supported by the organizational environment these managers describe: roughly 75 percent report moderate or high executive support for AI adoption, and over 76 percent view AI’s impact on their team’s work as at least moderately positive. Current use is similarly well-established. Nearly 75 percent of managers report that their team uses AI at least a few times per week, and over 60 percent report that more than half of their team members are active AI users. Taken together, these baseline distributions suggest that AI is neither nascent nor disappointing among the firms and teams

⁷This is a preliminary version of the figure and is not exhaustive of all related surveys on AI adoption. See the figure notes for source descriptions and frequency thresholds.

in our sample—managers are not reporting the experience of a “so-so technology” (Acemoglu, 2024; Acemoglu and Restrepo, 2019).

The right column shows the distribution of our standardized AI adoption index and its component measures. Managers express strong intent to adopt new AI tools and accelerate current use, while intent to reduce or halt AI use is rare—over 75 percent report being “not likely” to reduce, and over 83 percent report being “not likely” to halt. The resulting index is left-skewed, with the mass of the distribution concentrated at positive values. This skewness has implications for our experiment: if managers’ perspectives on AI productivity are already well-formed and positively oriented, information treatments emphasizing productivity gains may have limited room to shift adoption intentions further—a point we return to in Section 4.

3.3 Associations between manager characteristics, AI, and staffing intentions

Managers’ views and workplace context are strongly associated with AI adoption and advocacy, but far less predictive of staffing decisions. Figure 3 summarizes the cross-sectional associations, estimated within the control group, between managers’ pre-treatment characteristics and the four main outcome indices: AI adoption, AI advocacy, staffing intentions, and staffing advocacy. The left panel illustrates a clear pattern, which is that baseline AI exposure and pro-AI beliefs are strongly positively associated with AI-related outcomes. Managers who already expect greater future AI adoption, report stronger executive support for AI, perceive larger team-level AI impacts, hold more favorable team attitudes toward AI, and work in teams with more frequent AI use all exhibit substantially higher AI adoption and AI advocacy. These associations are substantially larger than other covariates and are consistently positive across both AI outcomes. By contrast, standard demographic characteristics explain much less of the variation: age, education, and political views have only modest correlations, while women report lower AI adoption and advocacy on average. Prior AI training is also strongly positively associated with AI outcomes, whereas prior experiences of unemployment—either personal or within the family—are negatively associated with them. Overall, the left panel suggests that managers’ AI intentions are most closely linked to their existing organizational AI environment and prior exposure to the technology, rather than to basic demographic characteristics alone.

The right panel shows that the correlates of staffing outcomes are weaker and less systematic. Most AI-related baseline variables that strongly predict AI adoption and advocacy show only small associations with staffing intentions and staffing advocacy, suggesting that managers do not mechanically map their AI views into employment plans. One notable exception is staffing plans for the next year, which is strongly positively associated with staffing intentions, as expected, while the respondent’s seniority and gender are related to staffing advocacy. AI training is positively associated with staffing intentions, but correlations with many other firm- and manager-level characteristics are close to zero. Taken together, Figure 3 suggests that AI-related beliefs and organizational context are highly predictive of managers’ willingness to adopt and promote AI, but much less predictive of their stated staffing decisions. This descriptive

evidence is consistent with the broader interpretation of managers as gatekeepers of AI diffusion: their views about AI are closely tied to prior exposure, team use, and leadership support, while staffing decisions appear to depend on a somewhat different and less tightly structured set of considerations.

4 Experimental Results

4.1 Main Results

Overall, we find stark differences in the impact of the two treatment arms, with managers exhibiting a strong pullback from AI and staffing intentions in the labor treatment—a marked negative reaction to AI’s labor displacing potential. Meanwhile, the productivity treatment has no apparent effect relative to the placebo control group. Figure 4 depicts estimates from our main specification⁸ where Subfigures 4a – 4d progress through our main outcomes—indices for AI adoption, AI advocacy, staffing intentions, and staffing advocacy—with the index on top and their component measures beneath. It is immediately apparent nearly all the point estimates for the AI productivity treatment (blue hollow markers) are statistically insignificant impacts of less than 10% of a standard deviation change.⁹

Meanwhile, the managers in the labor treatment (red solid markers) exhibit nearly a half standard deviation decrease in AI adoption and advocacy indices, seen in Figures 4a & 4b. To put this magnitude in context, roughly two in three managers exposed to the labor treatment scored below the control-group average (compared to one in two in the control group).¹⁰ This effect is substantial, especially in the context of similar cross-country, information-treatment survey experiments like Dechezleprêtre et al. (2025) and Alesina et al. (2018, 2022) that found information had a much smaller impact on preferences about other topics, related to climate change, immigration, and redistribution.

Moreover, the labor treatment also negatively affected managers’ staffing intentions, exhibited in Figure 4c. Interestingly, this strong pullback from staffing intentions is mainly driven by managers in the US. There is a much smaller and statistically insignificant reaction from UK managers. However, an increasing set of controls included in the specification leads to the estimated negative effect on UK managers’ staffing intentions increase in magnitude and the difference across countries becomes insignificant.¹¹

These main outcome effects are internally consistent within each index (across index components) and robust to controls. This is visible in each subfigure as the components below their respective index do not

⁸See equation 1. Our main specification restricts the sample to randomized respondents with non-missing values for all four primary outcome indices. Re-estimating each specification on the sample of randomized respondents with a non-missing value for the corresponding outcome yields very similar estimates, indicating that this sample restriction does not meaningfully affect the results. See Table A.62

⁹All outcomes are standardized relative to the placebo arm. See Appendix Section 2.4 for details on index construction.

¹⁰See Figure 5

¹¹See Appendix Table A.52. The convergence of estimates across samples particularly follows once controls include respondent and team AI postures and furthermore with firm level policies and pre-treatment plans.

exhibit divergent patterns.¹² Point estimates remain largely stable as we add increasingly expansive sets of controls, and are similar across our full analytic sample and the subsample with complete covariate/item survey response.¹³ Our checks corroborate two primary conclusions from our main results. First, the productivity treatment has little to no estimated effect on either managerial intentions or advocacy. And notably, the labor treatment prompts strong negative reactions on AI adoption and staffing decisions.

Regardless of what AI information was presented, managers did not have a detectable shift their willingness to advocate for workers. The point estimates are around 0.06 standard deviation declines in the staffing advocacy index for the productivity treated group, and essentially no change in the labor group. We speculate that staffing advocacy might implicitly be a more complicated function of managerial responsibilities that balance business unit objectives and ethical considerations. This raises the question of why staffing advocacy is unaffected on average, or relatively persistent for individuals regardless of information. This, and a broader unpacking of our main results, is explored in the next section.

4.2 Heterogeneous Treatment Effects

Our heterogeneity results suggest that the main treatment effects are broadly shared among respondents, as there is only substantial heterogeneity along three dimensions. This is shown in Figure 6, which illustrates subgroup treatment effects for five key pre-treatment characteristics: staffing plans for next year, team AI use frequency, respondent pre-treatment AI views, technology authority, and executive support for AI. Each Subfigure 6a – 6d correspond to one of the four primary outcome indices. Within each panel, the key characteristics are separated by “rows” that report the overall effect for the four distinct groups: the high (square markers) and below median (circle markers) subgroups of a given heterogeneity dimension for each labor (solid red markers) and productivity (hollow blue markers) treatment arms.¹⁴

The first notable dimension of heterogeneity relates to how the effect of the labor treatment varied based on the respondent’s expected staffing plans for the next year. This can be seen in the top row of each subfigure in Figure 6 (and Tables A.6–A.7), which shows how baseline plans for staffing corresponds to differential treatment effects on every outcome (apart from AI advocacy) from at least one information treatment arm. High baseline staffing plans are positively associated with AI adoption and advocacy, and naturally, very strongly associated with staffing intentions; furthermore, high baseline staffing plans exacerbate the effect of labor-displacing information on AI Adoption but have a smaller differential effect on staffing intentions relative its direct association. This suggests that managers with growing teams may react with greater caution in their decisions, but do not completely reverse their staffing plans.

¹²Appendix Tables A.46–A.49 respectively report estimates for treatment effects for each main index reported in column (1) and the index’s components (before standardization) in the following columns.

¹³Appendix Tables A.50–A.53 report estimates for each index across increasingly expansive control sets. These estimates are based on the subsample of respondents with no missing item responses across all variables in the largest control set. The raw estimated effects in this subsample are largely similar to those in the full analytic sample, mitigating the concern that compositional changes drive any conclusions.

¹⁴The corresponding interaction-model estimates and subgroup treatment effects are reported in AppendixSection A.3, Tables A.6–A.15. Definitions of the high and low subgroups for the primary heterogeneity dimensions are provided in Appendix Section E.1.

The second dimension where heterogeneous treatment effects exist is in regard to how the productivity treatment varied based on managers' prior exposure to AI. There is a strong differential response in AI advocacy in reaction to the productivity treatment—seemingly driven by unfamiliarity with AI. This can be seen in the second row of Figure 6 (and Table A.8) where the degree of heterogeneity is the manager's team's AI use frequency. With minimal AI use frequency, the estimated effects are around a 0.2 sd increase in AI advocacy—roughly half the effect of the labor arm on AI adoption/advocacy. However, this effect is completely erased if there is high AI use. There is a similar, but statistically insignificant pattern for the AI adoption outcome. This suggests that productivity treatment effects may be moderated by managers' prior familiarity with AI adoption.

The differential effect for the productivity treatment group on AI adoption and advocacy is most sharply explained by AI use, but managers' own sentiment towards AI, as well as their teams, or their firm leadership's support for AI, resonates with the interpretation that there is a higher impact for those with weaker priors on AI. This can be seen in the third row of each subfigure in Figure 6 (and in Tables A.10 and A.11), where the considered variable of heterogeneity is an indicator for positive pre-treatment sentiment towards AI. However, over 70% of our respondents report positive sentiments from both themselves and among their teams, and another 20% report neutral sentiments, which limits our ability to understand how negative sentiments interact with the informational treatments. However, the notable differential effect is visible in 6c, where non-positive sentiment corresponds to a near-zero point estimate in staffing intentions. This suggests the retrenchment behavior may be moderated by perceptions of AI's productive capabilities.

The final dimension of heterogeneity worth noting relates to executive support for AI, and its impact on staffing. This can be seen in the fifth row of each subfigure in Figure 6 (and Tables A.14 and A.15). In the two lower subfigures, it is evident the negative effect of the labor treatment on staffing intentions is driven more by the managers who report low executive support for AI. In fact, this characteristic corresponds to negative effects on staffing intentions and advocacy in both treatment arms, whereas managers with high executive support for AI respond with decreasing staffing intentions only in the labor treatment.¹⁵

These heterogeneity results and our descriptive statistics lead us to two suggestive takeaways. First, there may be little opportunity for such general information to have an impact when there is already prior knowledge and adoption of AI. While our results do not directly address this, information on AI productivity benefits may need to be increasingly specific to meaningfully shift managers' decisions, as their own specific experience, whether positive or negative, will likely override general information. We further speculate that AI advocacy is a more responsive margin than AI adoption. Information on general productivity benefits may leave AI managers uncertain about how to adopt/implement AI to realize benefits in their specific business functions. The other key takeaway is that employment and AI adoption are not intrinsically in tension within a firm or business unit as demonstrated by positive correlation between

¹⁵While we do not find the same statistically significant differential AI advocacy responses to the productivity treatment, there is some indication that lower support may correspond to greater opportunity for AI productivity benefits information to impact managers; for example, moving from the high executive-support-for-AI group to the low executive-support-for-AI group increases the estimated effect of AI productivity information on AI advocacy by 0.16 standard deviations.

more positive AI posture (from baseline adoption plans to executive support) and elevated staffing baseline plans and post-treatment intentions, alongside exaggerated retrenchment in staffing intentions in the labor arm among managers with positive AI sentiment.

4.3 Mechanism Analysis

Decomposing the labor-distress treatment effect estimates illustrates that changing perceptions in benefits and risks of AI account for manager's pull back from adoption and advocacy, but otherwise our estimated effects on other outcomes are not well explained by the channels measured in our survey. Figure 7 shows decompositions of all the treatment effects estimates on each primary outcome index. The decompositions are constructed and shown separately by treatment arm, based on the baseline, full, and auxiliary specifications in equations (4)–(7).¹⁶

The impacts of labor distress information on our AI outcome indices are well explained and sensible. Nearly 40% of the total effect on AI adoption (-0.403 sd) is accounted for by the AI benefits index, and another 30% is accounted for by changes in the AI risks index, leaving about 30% of this overall effect unexplained. Of the labor treatment's total effect on AI advocacy (-0.435 sd), again 40% is accounted for by the AI benefits index, and nearly 42% is attributed to changes in the AI risks index, leaving only about 18% of this overall effect unexplained. Table A.17 and Table A.18 further substantiate the intuitive interpretation one would presume from these channels and their contribution to the overall effects. In the full regressions for the AI labor treatment, the AI benefits index is positively associated with AI adoption and AI advocacy, while the AI risks index is negatively associated with both outcomes. In the auxiliary regressions, the AI labor treatment substantially lowers the AI benefits index and substantially raises the AI risks index. Given our sample's positive posture towards AI at baseline, it is perhaps unsurprising that information about AI's labor distress potential induces managers to recalibrate AI's perceived benefits and risks.

However, the labor treatment impact on staffing intentions is not well explained by any of the channels our survey measured. Of the two reported here (the strongest channels), less than 16% of this overall effect is accounted for by the AI benefits index, and less than 6% is accounted for by the AI risks index, leaving over a substantial portion unexplained. We speculate that, among many other potential beliefs not well captured by our measured channels, an overall sense of uncertainty could reconcile these negative effects on managers' decisions about AI adoption and staffing.¹⁷

¹⁶The corresponding decomposition estimates are reported in Table A.16. Full outcome regressions corresponding to equation (5) for all outcomes and both treatment arms are reported in Table A.17, and the auxiliary regressions corresponding to equation (6) for the AI benefits and AI risks indices are reported in Table A.18.

¹⁷We caution interpretation of the decompositions for the labor treatment's impact on staffing advocacy, and most productivity treatment impacts, given the overall estimated effects are insignificant. The effect on AI advocacy is marginally significant in the pooled regression, but again, not well explained by any measured channels. Table A.16 show AI benefits and risks are strongly associated with AI advocacy, but the auxiliary regressions in Table A.18 show little movement in either the AI benefits or AI risks indices under the AI Productivity treatment. This again resonates with our findings that managers have a positive posture towards AI at baseline, and AI productivity-enhancing information has little impact on their positive perception of AI.

5 Discussion and Conclusion

The rapid growth of artificial intelligence (AI) in public discourse has emphasized both productivity boosting and labor displacing potential. This information environment may be crucial in shaping organizational decisions and, by extension, broader societal diffusion. Our study shows that middle managers, as gatekeepers between executive strategy and team-level implementation, respond to this information in ways that meaningfully affect both AI adoption and staffing intentions.

Our central finding is that information about AI's labor-displacing potential prompts managers to substantially pull back from both AI adoption and staffing. We find a broad retreat of 0.4–0.5 standard deviations in AI adoption and advocacy and 0.2 standard deviations in staffing intentions. Mediation analysis suggests that this information simultaneously lowers managers' perceived benefits of AI and raises perceived risks, which together help explain the pullback from AI adoption and advocacy. Notably, the reduction in staffing intentions cannot be explained by any of our measured belief channels, suggesting that the labor treatment may trigger a more general retrenchment that extends beyond attitudes toward AI itself. Information about AI's productivity benefits, by contrast, does not shift intentions on average—consistent with managers already holding optimistic priors—though it does increase AI advocacy among managers with low AI familiarity, limited technical authority, and weak executive support.

Our findings also speak to the relationship between AI adoption and employment within firms. If managers viewed AI primarily as a substitute for labor, information about displacement should encourage AI adoption while reducing staffing. Instead, we find the opposite: displacement information leads managers to scale back both AI adoption and staffing simultaneously. Descriptive patterns in our baseline data reinforce this interpretation. For instance, managers who report higher AI adoption expectations also report higher staffing intentions—the only stronger association was baseline staffing plans. These associations are not causal, but together with the retrenchment behavior in the labor treatment, this survey documents instances that AI and human labor are not intrinsically at tension at the organizational level.

Our findings suggest that middle managers are part of the explanation for the gap between AI's individual-level promise and its more modest organizational-level impact, but several questions remain for future research. A first task is to establish whether the treatment effects we identify over managerial intentions map into realized organizational decisions, such as AI deployment, task reallocation, hiring, and workforce reductions. A second is to examine how managerial gatekeeping varies across occupations, business functions, and institutional environments, where the scope for AI adoption and the constraints on staffing decisions are likely to differ substantially. A third is to better isolate the mechanisms underlying managers' response to labor-displacement information, particularly the roles of uncertainty, perceived implementation risk, and concern for employee welfare. More broadly, future work should study which combinations of information, training, and firm-level support enable managers to adopt AI in ways that enhance productivity without undermining employment. Answering these questions would help clarify when managers act as barriers, when they act as bridges, and how their decisions shape the broader economic effects of AI.

References

- Acemoglu, Daron**, “The Simple Macroeconomics of AI,” *IMF Economic Review*, 2024.
- **and Pascual Restrepo**, “Automation and New Tasks: How Technology Displaces and Reinstates Labor,” *Journal of Economic Perspectives*, May 2019, 33 (2), 3–30.
- **and —**, “Robots and Jobs: Evidence from US Labor Markets,” *Journal of Political Economy*, 2020, 128 (6), 2188–2244.
- Aghion, Philippe and Jean Tirole**, “Formal and Real Authority in Organizations,” *Journal of Political Economy*, 1997, 105 (1), 1–29.
- Alesina, Alberto, Armando Miano, and Stefanie Stantcheva**, “Immigration and Redistribution,” *The Review of Economic Studies*, 03 2022, 90 (1), 1–39.
- **, Stefanie Stantcheva, and Edoardo Teso**, “Intergenerational Mobility and Preferences for Redistribution,” *American Economic Review*, February 2018, 108 (2), 521–54.
- Atkin, David, Azam Chaudhry, Shamyla Chaudry, Amit K. Khandelwal, and Eric Verhoogen**, “Organizational Barriers to Technology Adoption: Evidence from Soccer-Ball Producers in Pakistan,” *The Quarterly Journal of Economics*, 2017, 132 (3), 1101–1164.
- Barrero, Jose Maria, Nicholas Bloom, and Steven J. Davis**, “Why Working from Home Will Stick,” 2021. Describes the Survey of Working Arrangements and Attitudes (SWAA), running monthly since May 2020. AI module (December 2024) reported in Bick et al. (2024).
- Bick, Alexander, Adam Blandin, and David J. Deming**, “The Rapid Adoption of Generative AI,” September 2024.
- Bonney, Kathryn, Cory Breaux, Cathy Buffington, Emin Dinlersoz, Lucia S. Foster, Nathan Goldschlag, John C. Haltiwanger, Zachary Kroff, and Keith Savage**, “Tracking Firm Use of AI in Real Time: A Snapshot from the Business Trends and Outlook Survey,” April 2024.
- Brynjolfsson, Erik, Bharat Chandar, and Ruyu Chen**, “Canaries in the Coal Mine? Six Facts about the Recent Employment Effects of Artificial Intelligence,” November 2025. Version dated November 13, 2025.
- **, Danielle Li, and Lindsey R. Raymond**, “Generative AI at Work,” April 2023.
- Cao, Guangming, Yanqing Duan, John S. Edwards, and Yogesh K. Dwivedi**, “Understanding managers’ attitudes and behavioral intentions towards using artificial intelligence for organizational decision-making,” *Technovation*, August 2021, 106, 102312.
- Challapally, Aditya, Chris Pease, Ramesh Raskar, and Pradyumna Chari**, “The GenAI Divide: State of AI in Business 2025,” Technical Report, MIT NANDA (Networked Agents and Decentralized Architecture) July 2025.
- Crane, Leland, Michael Green, and Paul Soto**, “Measuring AI Uptake in the Workplace,” February 2025.

- Cutter, Chip**, “CEOs Start Saying the Quiet Part Out Loud: AI Will Wipe Out Jobs,” *The Wall Street Journal*, July 2025.
- Daly, Sarah J., Anna Wiewiora, and Greg Hearn**, “Shifting attitudes and trust in AI: Influences on organizational AI adoption,” *Technological Forecasting and Social Change*, June 2025, 215, 124108.
- Dechezleprêtre, Antoine, Adrien Fabre, Tobias Kruse, Bluebery Planterose, Ana Sanchez Chico, and Stefanie Stantcheva**, “Fighting Climate Change: International Attitudes toward Climate Policies,” *American Economic Review*, April 2025, 115 (4), 1258–1300.
- Dominski, Jacob and Yong Suk Lee**, “Advancing AI Capabilities and Evolving Labor Outcomes,” July 2025. arXiv:2507.08244.
- Eloundou, Tyna, Sam Manning, Pamela Mishkin, and Daniel Rock**, “GPTs are GPTs: Labor market impact potential of LLMs,” *Science*, June 2024, 384 (6702), 1306–1308.
- Federal Reserve Bank of Richmond**, “Automation and AI: What Does Adoption Look Like for Fifth District Businesses?,” June 2024. Survey of 211 CFOs in the Fifth Federal Reserve District, fielded May–June 2024.
- Gelbach, Jonah B.**, “When Do Covariates Matter? And Which Ones, and How Much?,” *Journal of Labor Economics*, 2016, 34 (2), 509–543.
- Goldman Sachs Global Investment Research**, “Generative AI could raise global GDP by 7%,” 2023. Goldman Sachs Global Investment Research.
- Hampole, Menaka, Dimitris Papanikolaou, Lawrence Schmidt, and Bryan Seegmiller**, “Artificial Intelligence and the Labor Market,” July 2025.
- Hatzius, Jan, Joseph Briggs, Devesh Kodnani, and Giovanni Pierdomenico**, “The Potentially Large Effects of Artificial Intelligence on Economic Growth,” March 2023. Global Economics Analyst.
- Humlum, Anders and Emilie Vestergaard**, “Large Language Models, Small Labor Market Effects,” May 2025. Survey of 25,241 Danish workers in 11 AI-exposed occupations, fielded November 2024.
- and —, “The Unequal Adoption of ChatGPT Exacerbates Existing Inequalities Among Workers,” *Proceedings of the National Academy of Sciences*, 2025, 122 (1), e2414972121. Survey of 18,109 Danish workers in 11 AI-exposed occupations, fielded November 2023–January 2024.
- IBM**, “AI in Action 2024 Report | IBM,” September 2024.
- Kalyani, Aakash, Nicholas Bloom, Marcela Carvalho, Tarek Hassan, Josh Lerner, and Ahmed Tahoun**, “The Diffusion of New Technologies*,” 2025, 140 (2), 1299–1365. eprint: <https://academic.oup.com/qje/article-pdf/140/2/1299/61489798/qjaf002.pdf>.
- Kerr, Emily and Karel Mertens**, “Texas Firms Using AI Report Little Impact on Employment,” June 2024. Texas Business Outlook Survey special questions, April 2024 (n=291). Additional waves: May 2025 (n=318), December 2025 (n=358).
- Korst, Jeremy, Stefano Puntoni, and Mary Purk**, “Growing Up: Navigating Generative AI’s Early Years – 2024 AI Adoption Report,” October 2024. Survey of 802 senior decision-makers at U.S. enterprises (1,000+ employees), fielded July 2024.

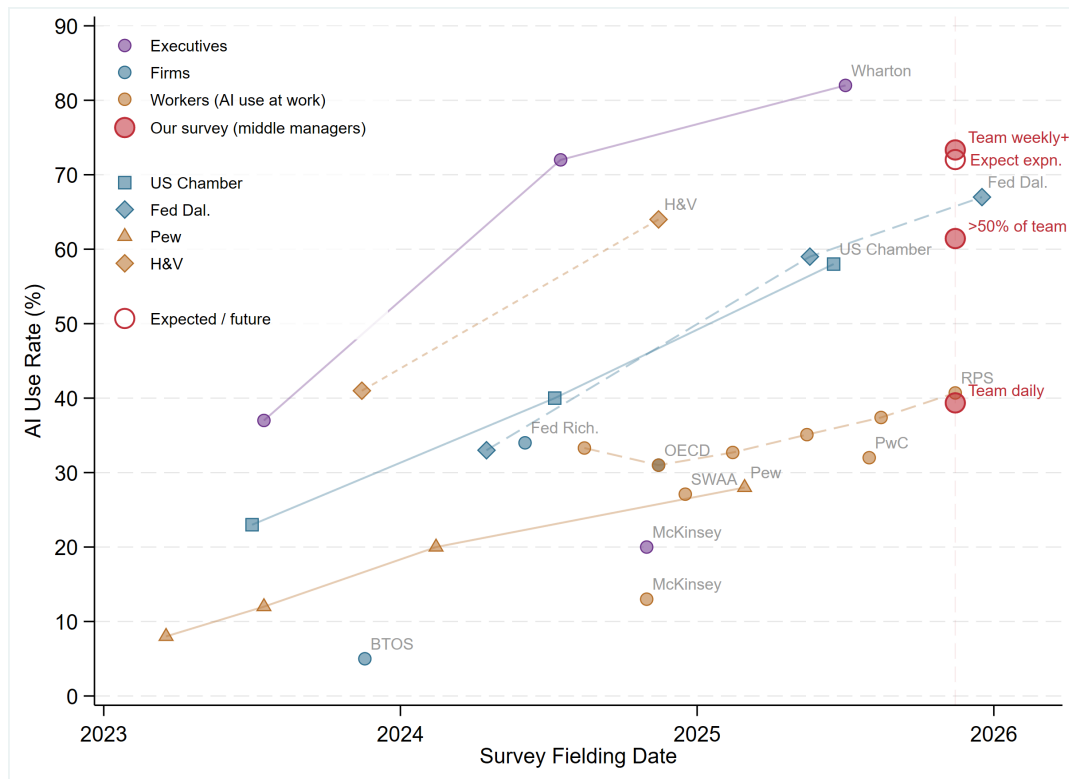
- , —, and **Prasanna Tambe**, “Accountable Acceleration: Gen AI Fast-Tracks into the Enterprise – 2025 AI Adoption Report,” October 2025. Survey of 801 senior decision-makers at U.S. enterprises (1,000+ employees), fielded June 26–July 11, 2025.
- , —, **Mary Purk, and Dan Ives**, “The Rise of Generative AI in the Enterprise,” 2023. Survey of 672 senior decision-makers at U.S. enterprises, fielded July 2023.
- Majrashi, Khalid**, “Determinants of Public Sector Managers’ Intentions to Adopt AI in the Workplace,” *International Journal of Public Administration in the Digital Age*, May 2024, 11.
- McClain, Colleen and Olivia Sidoti**, “34% of U.S. Adults Have Used ChatGPT, about Double the Share in 2023,” June 2025. American Trends Panel. Estimates from March 2023, July 2023, February 2024, and February/March 2025 waves.
- McKinsey**, “The State of AI: Global Survey 2025 | McKinsey,” 2025.
- McKinsey & Company**, “The economic potential of generative AI: The next productivity frontier,” 2023.
- , “Superagency in the Workplace: Empowering People to Unlock AI’s Full Potential at Work,” January 2025. Survey of 3,613 employees and 238 C-level executives across six countries, fielded October–November 2024.
- McKinsey Global Institute**, “The Economic Potential of Generative AI: The Next Productivity Frontier,” Technical Report, McKinsey & Company June 2023.
- Noy, Shakked and Whitney Zhang**, “Experimental evidence on the productivity effects of generative artificial intelligence,” *Science*, 2023.
- OECD**, “Generative AI and the SME Workforce: New Survey Evidence,” 2025. Telephone survey of 5,232 SMEs (≥250 employees) in Austria, Canada, Germany, Ireland, Japan, Korea, and the United Kingdom, fielded October 14–December 12, 2024.
- Peng, Sida, Pratyusha Kalluri, Yixin Zhang, Weizhu Chen et al.**, “Check Your Work: Investigating the Role of LLMs in Context of Scientific Writing,” February 2023. arXiv:2302.06590.
- PwC**, “Global Workforce Hopes and Fears Survey 2025,” 2025. Survey of 49,843 workers across 48 countries, fielded July–August 2025. No US-specific breakout available.
- Tambe, Prasanna, Peter Cappelli, and Valery Yakubovich**, “Artificial Intelligence in Human Resources Management: Challenges and a Path Forward,” *California Management Review*, 2019, 61 (4), 15–42.
- U.S. Chamber of Commerce and Teneo Research**, “Empowering Small Business: The Impact of Technology on U.S. Small Business,” 2023. Survey of 1,025 U.S. small businesses (≥250 employees), fielded June 15–July 9, 2023.
- , “The Impact of Technology on U.S. Small Business,” September 2024. Survey of 1,100 U.S. small businesses (≥250 employees), fielded June 21–July 16, 2024.
- , “Empowering Small Business: The Impact of Technology on U.S. Small Business,” 2025. Survey of 3,870 U.S. small businesses (≥250 employees), fielded June 6–26, 2025.

Wen, Yanjun, Jiale Wang, and Xiaoxi Chen, “Trust and AI weight: human-AI collaboration in organizational management decision-making,” *Frontiers in Organizational Psychology*, June 2025, 3. Publisher: Frontiers.

Yotzov, Ivan, Jose Maria Barrero, Nicholas Bloom, Philip Bunn, Steven J Davis, Kevin M Foster, Aaron Jalca, Brent H Meyer, Paul Mizen, Michael A Navarrete, Pawel Smietanka, Gregory Thwaites, and Ben Zhe Wang, “Firm Data on AI,” Working Paper 34836, National Bureau of Economic Research February 2026.

Figures

Figure 1: AI Use Rates Across Related Surveys



Notes: Each marker is a survey estimate of AI adoption. The x-axis is the approximate survey fielding date; the y-axis is the reported AI use rate. Color indicates the respondent population. Within each color, marker shapes and line styles distinguish multi-wave sources. Filled markers denote current use; hollow markers denote expected or future use. Frequency thresholds vary across surveys (ever used, any use, weekly, binary firm adoption); comparisons should be interpreted accordingly. **Wharton:** Korst et al. (2024, 2025). Senior decision-makers using generative AI at least weekly.

McKinsey: McKinsey & Company (2025). Share reporting AI use for more than 30% of daily tasks.

BTOS: Bonney et al. (2024). Firms reporting any AI use in the past two weeks, firm-weighted.

US Chamber: U.S. Chamber of Commerce and Teneo Research (2023, 2024, 2025). Small businesses reporting any AI use.

Fed Dal.: Kerr and Mertens (2024). Firms currently using any AI.

Fed Rich.: Federal Reserve Bank of Richmond (2024). Firms reporting AI use in task automation since January 2022.

OECD: OECD (2025). SMEs across seven countries that had adopted generative AI.

RPS: Bick et al. (2024). Employed adults reporting any generative AI use for work.

Pew: McClain and Sidoti (2025). Adults who have ever used ChatGPT for work tasks.

SWAA: Barrero et al. (2021). Employed adults reporting any generative AI use for work.

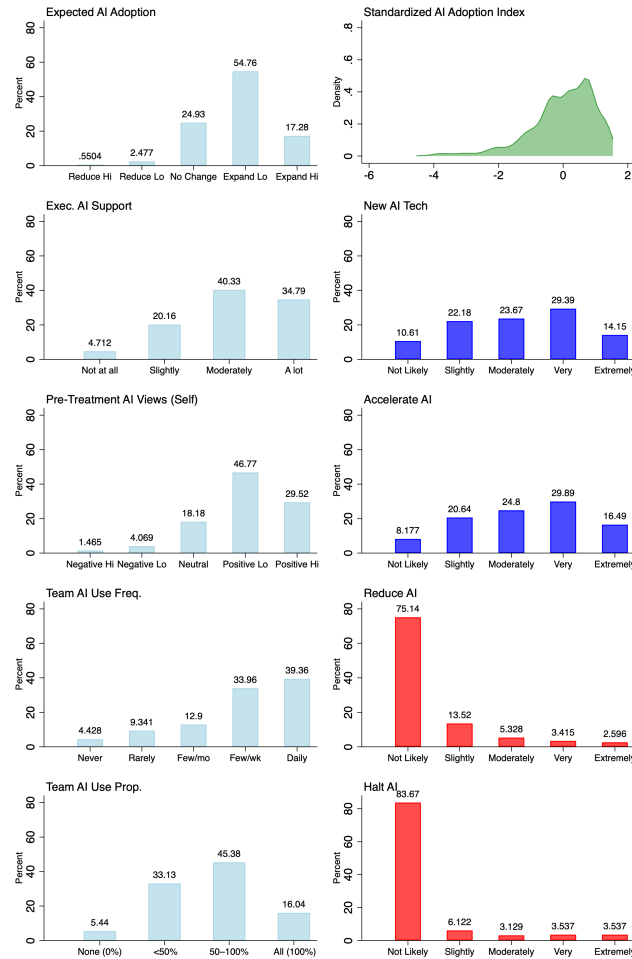
H&V: Humlum and Vestergaard (2025b,a). Workers in 11 AI-exposed occupations who have used AI for work.

PwC: PwC (2025). Workers using generative AI at least weekly, 48 countries.

Our survey: Pre-registered experiment of approximately 2,000 US and UK managers, fielded November 2025; collected multiple AI use measures.

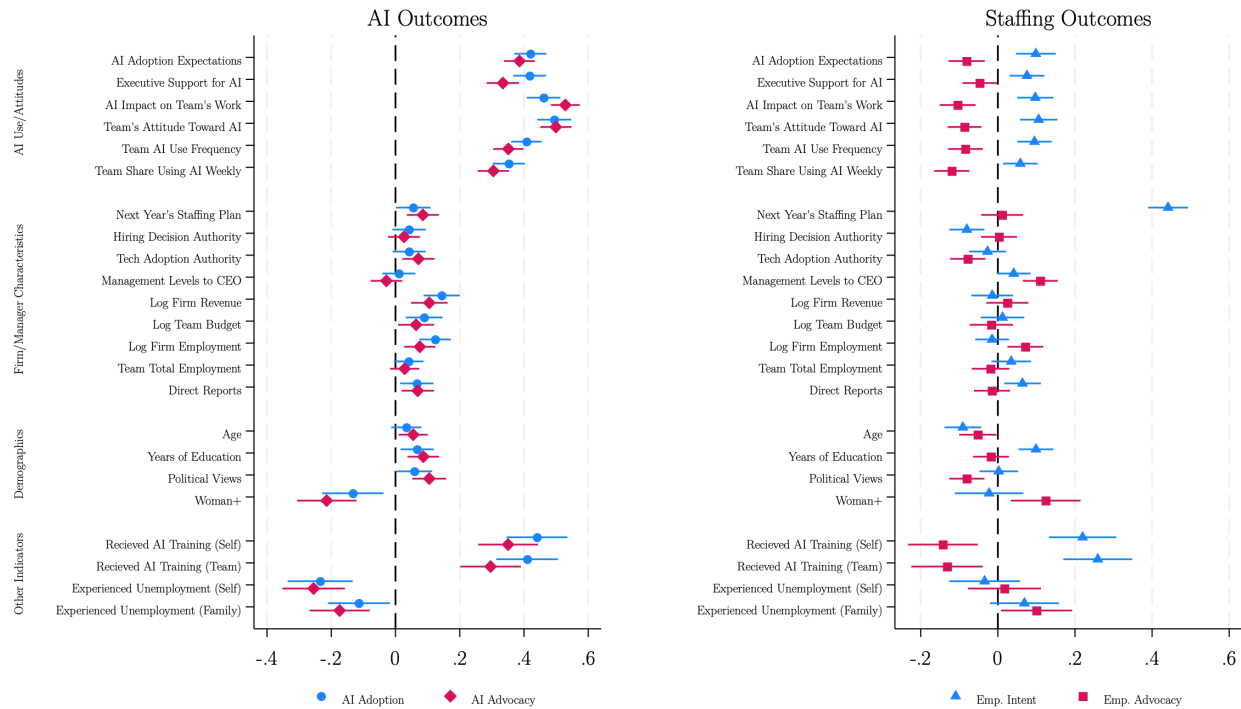
Note that the OECD marker overlaps with an RPS series marker. The OECD sample corresponds to firms, and the RPS sample to workers.

Figure 2: Distributions of Pre-Treatment AI-Related Measures and Control-Group AI Adoption Outcomes



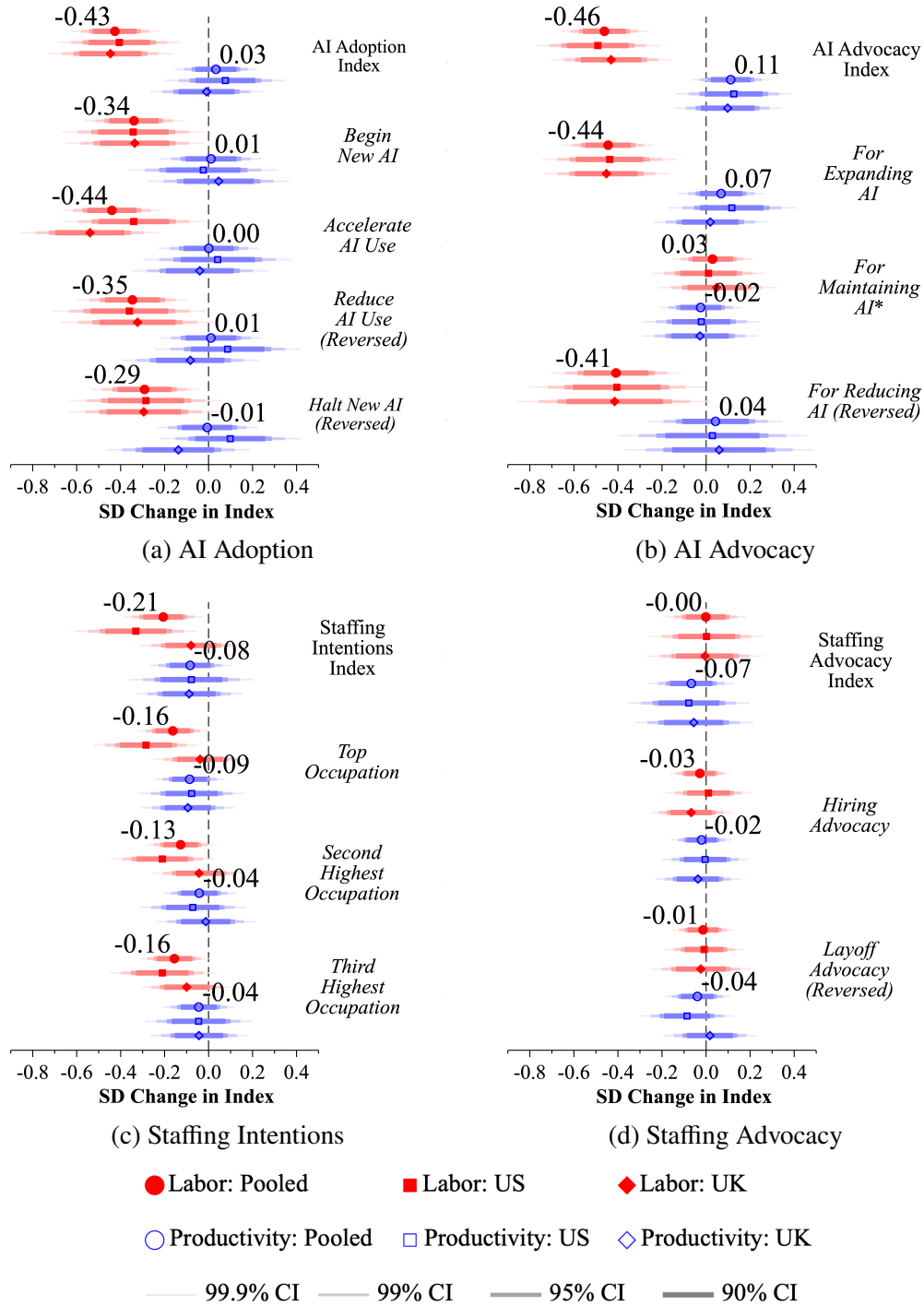
Notes: The left column of subfigures shows histograms of key measures collected at baseline related to AI use and adoption: Expected AI Adoption (5-point Likert scale); the firm's executive support for AI adoption (4-point Likert intensity scale); the respondent's perception of AI impact on their team's work (5-point Likert scale); the frequency of their team's AI use (5-point frequency scale) and proportion of their team using AI (4-point proportion scale). The right column of subfigures shows the distribution for our AI adoption index and the its contributing component measures. The index distribution is depicted as a (Gaussian) kernel density plot derived from 5-point Likert scale component measures about the respondent's intent to: begin new AI adoption, increase current AI use, reduce current AI use, and halt/avoid AI adoption. Outcome construction is described in Section 2.4.

Figure 3: Associations Between Manager Characteristics, Beliefs, and AI and Staffing Outcomes in the Control Group



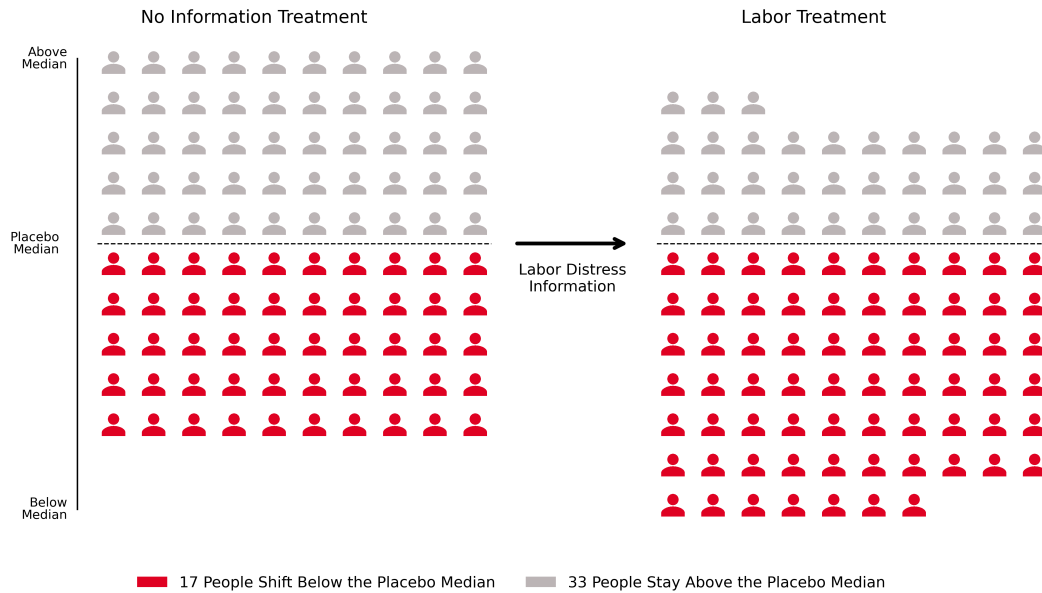
Notes: This figure reports descriptive bivariate associations, estimated using only respondents in the control group, between four standardized outcome indices and 23 pre-treatment or stable baseline-related predictors capturing AI beliefs and use, staffing expectations and authority, firm and team characteristics, demographics, training, and unemployment experience. In each regression, the dependent variable is a single standardized outcome index, and the independent variable is a single predictor, so the estimates show how each characteristic is associated with outcomes in the absence of treatment exposure. *AI Adoption* is a 4-item index measuring intentions to adopt or expand AI use in the respondent's team or unit. *AI Advocacy* is a 2-item index measuring willingness to advocate for expanding AI use within the organization. *Staffing Intentions* is a 3-item index measuring expected headcount changes over the next year for the respondent's three most common team occupation groups. *Staffing Advocacy* is a 2-item index measuring willingness to advocate for hiring versus layoffs. Indices are constructed in three steps: each component item is standardized using the pooled control-group mean and standard deviation, standardized items are averaged, requiring at least half of the component items to be non-missing, and the resulting index is standardized again using the pooled control-group mean and standard deviation. Predictors are grouped into four categories: AI use and attitudes, firm and manager characteristics, demographics, and other indicators. Most predictors are measured pre-treatment. Political views, unemployment experience, and selected firm-level characteristics are collected post-treatment to reduce attrition but are treated as stable, not plausibly affected by treatment. Continuous predictors are entered in standardized units, while woman+, AI training, and unemployment-experience measures are entered as binary indicators. Points denote coefficient estimates and horizontal lines denote 95% confidence intervals.

Figure 4: Effects of AI Labor and Productivity Treatments on Adoption and Hiring Outcomes

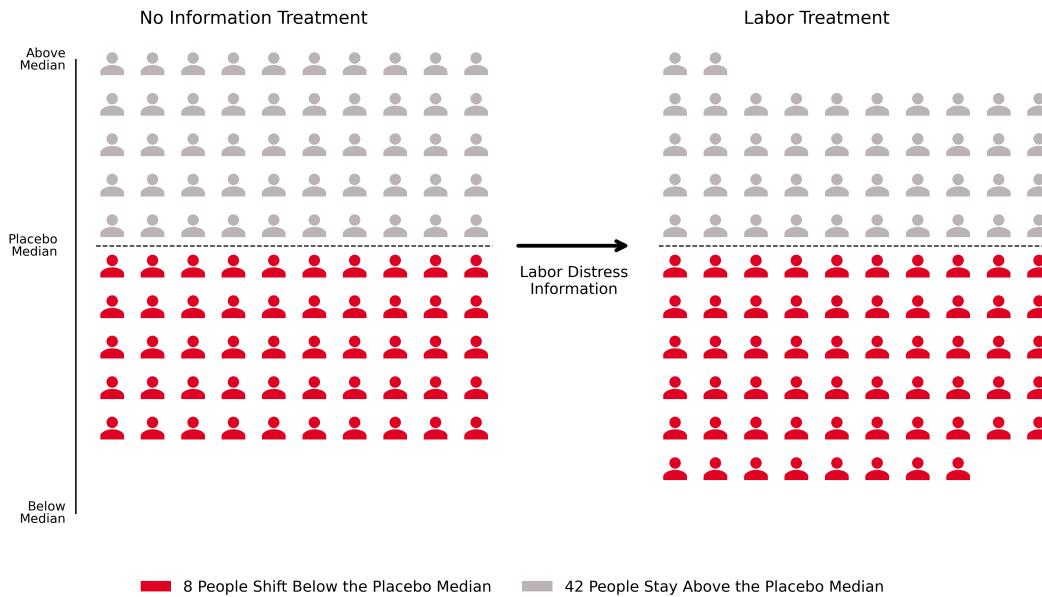


Notes: This figure reports treatment-effect estimates for each primary outcome index and its component items. Each panel corresponds to one of the four primary indices: AI Adoption, AI Advocacy, Staffing Intentions, and Staffing Advocacy. Within each panel, the top row reports the treatment effect on the full index, and the rows below report treatment effects for the component items used to construct that index. Item labels in italics denote component items. In Panel B, *For Maintaining AI Use* is shown for completeness but is not included in the construction of the *AI Advocacy* index. Outcome construction is described in Section 2.4. For each row, six estimates are shown: three for the AI Labor treatment arm and three for the AI Productivity treatment arm. Treatment arms are distinguished by color and fill, with red filled markers denoting AI Labor and blue hollow markers denoting AI Productivity. Within each treatment arm, the circle marker reports the pooled estimate from the full analytical sample, controlling for country fixed effects; the square marker reports the estimate from the U.S. subsample; and the diamond marker reports the estimate from the U.K. subsample. All estimates come from OLS regressions with an indicator for treatment assignment and robust standard errors. Pooled specifications include observations from both countries and country fixed effects, while country-specific specifications are estimated separately within each country. Horizontal intervals report 90%, 95%, 99%, and 99.9% confidence intervals; thicker and darker intervals correspond to lower confidence levels, while thinner and lighter intervals correspond to higher confidence levels.

Figure 5: Interpreting the Labor-Treatment Effects on AI Adoption and Staffing Intentions



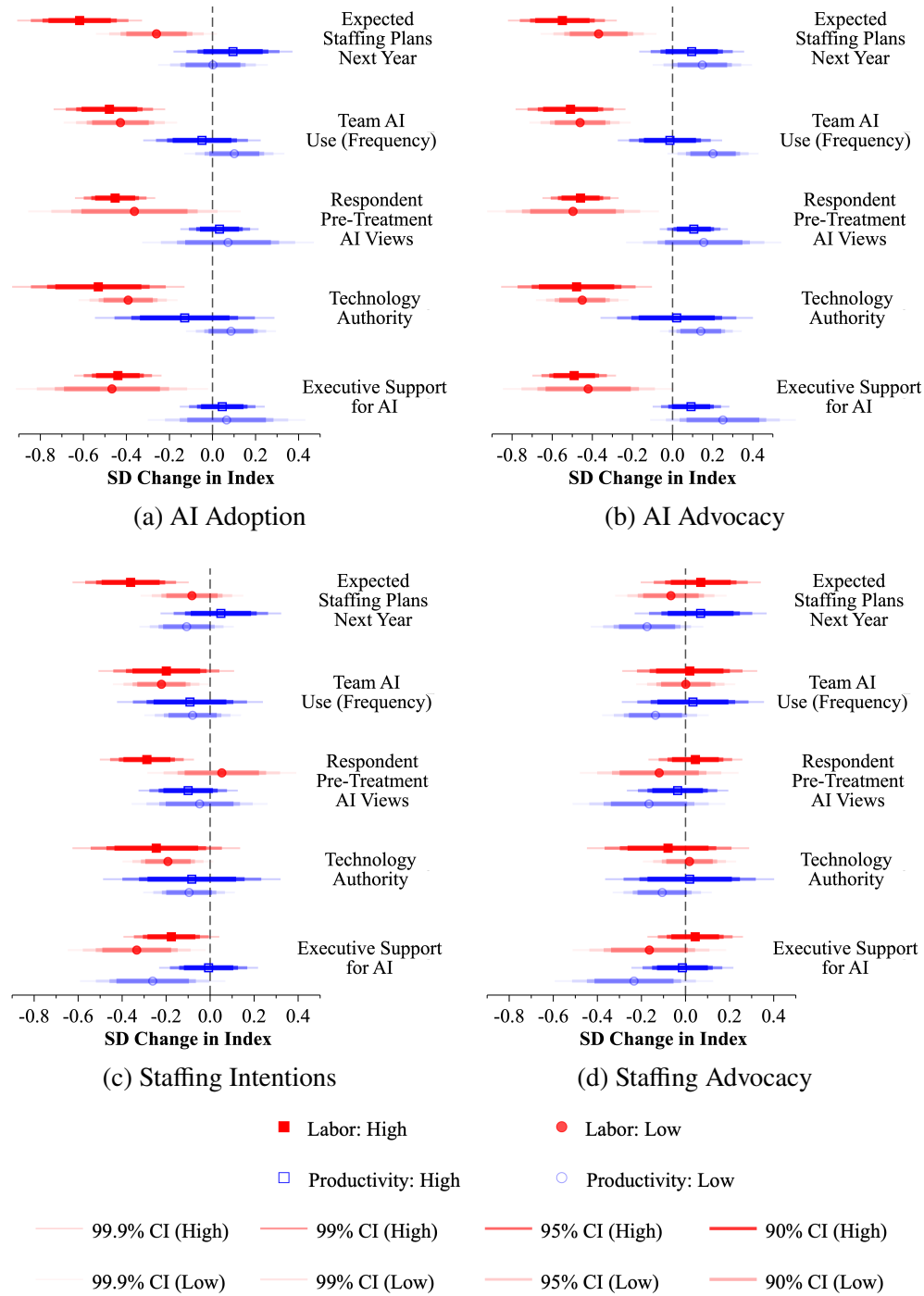
(a) AI Adoption



(b) Staffing Intentions

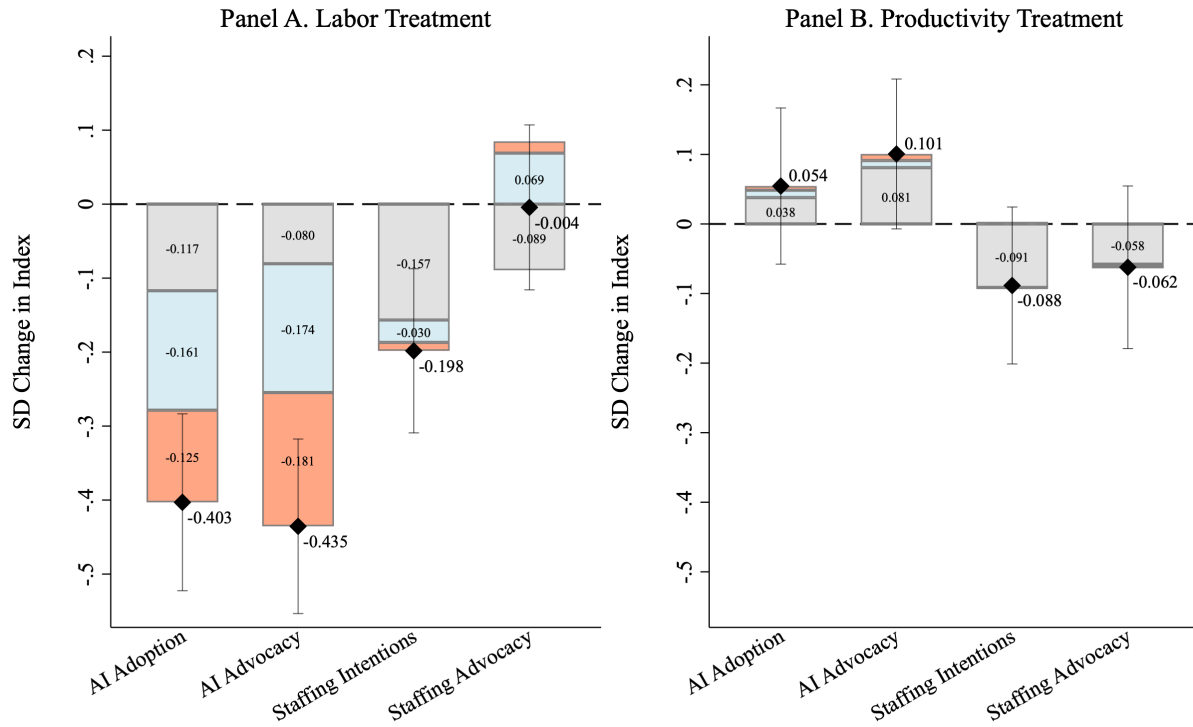
Notes: This figure translates the pooled labor-treatment estimates for *AI Adoption* and *Staffing Intentions* into an “out of 100” comparison using Cohen’s U_3 . In each panel, the left side depicts the placebo benchmark, and the right side depicts the implied distribution under the AI labor treatment. Red icons indicate the implied number of respondents below the placebo benchmark, while gray icons indicate those above it. Because the underlying indices are constructed from discrete survey items and are not exactly normally distributed, these figures should be interpreted as a visualization of effect size rather than as a literal reallocation of observed respondents. In the realized placebo-group data, the share below the placebo-group mean is not exactly 50 percent because the outcome indices are discrete and their distributions are skewed left. A complementary rank-based exercise that uses the placebo-group median after adding a small mean-zero jitter to break ties yields nearly identical results, suggesting that the visual magnitudes are not driven by the normality approximation.

Figure 6: Treatment Effect Heterogeneity Across Key Pre-Treatment Characteristics



Notes: This figure reports heterogeneous treatment-effect estimates for each of the four primary outcome indices: *AI Adoption*, *AI Advocacy*, *Staffing Intentions*, and *Staffing Advocacy*. Each panel corresponds to one primary index. Outcome construction is described in Section 2.4. Within each panel, the five rows report treatment effects by subgroup for five key heterogeneity dimensions identified in the CRF. In each row, four estimates are shown: two for the AI Labor treatment arm and two for the AI Productivity treatment arm. Treatment arms are distinguished by color and fill, with red filled markers denoting AI Labor and blue hollow markers denoting AI Productivity. Within each treatment arm, marker shape distinguishes subgroups: squares denote the high subgroup and circles denote the low subgroup. Thus, each row reports the estimated treatment effect for the high and low groups under each treatment arm. The construction of the high and low subgroup indicators is described in Appendix Section E.1. Estimates are obtained from the interacted heterogeneity specification in Equation 3, estimated separately for each outcome index and heterogeneity dimension. For each treatment arm, the low-group effect corresponds to the baseline treatment coefficient, while the high-group effect corresponds to the sum of the baseline treatment coefficient and the associated interaction coefficient. All specifications pool the sample across countries, include country fixed effects, and use robust standard errors. Horizontal intervals report 90%, 95%, 99%, and 99.9% confidence intervals. Within each subgroup, thicker and darker intervals correspond to lower confidence levels, while thinner and lighter intervals correspond to higher confidence levels.

Figure 7: Gelbach Decomposition of Treatment Effects



Notes: This figure reports Gelbach decompositions of the treatment effects on the four primary outcome indices: *AI Adoption*, *AI Advocacy*, *Staffing Intentions*, and *Staffing Advocacy*. Panel A reports results for the AI Labor treatment, and Panel B reports results for the AI Productivity treatment. Within each panel, each bar corresponds to one outcome index and decomposes the total treatment effect into three components: an unexplained component, shown in gray; the component attributable to the AI Benefits index, shown in light blue; and the component attributable to the AI Risks index, shown in red-orange. The black diamond marks the total treatment effect, and the vertical line through the diamond reports its 95% confidence interval. For each outcome and treatment arm, the total effect is estimated from a regression of the outcome on the treatment indicator and country fixed effects, using robust standard errors. The decomposition then adds the AI Benefits and AI Risks indices to this specification. Following Gelbach (2016), the contribution of each mediator is computed as the product of: (i) the coefficient on that mediator in the regression that includes the treatment and both mediators; and (ii) the effect of the treatment on that mediator from an auxiliary regression of the mediator on the treatment and country fixed effects. The unexplained component is the residual difference between the total treatment effect and the sum of the two mediator-specific components. All specifications are estimated separately by treatment arm and pool the sample across countries with country fixed effects. Construction of the mechanism indices, including component item wording, coding, and standardization, is described in Appendix Section E.2. Outcome construction is described in Section 2.4.

Appendix

FOR ONLINE PUBLICATION

AI Adoption, Manager Beliefs, and Hiring Intentions

YONG SUK LEE

CASSANDRA MERRITT

JACOB DOMINSKI

CHRISTOPHER HOY

A Appendix: Tables and Figures

A.1 Balance Tables

Table A.1: Baseline Characteristic Balance Across Countries

	(1) US	(2) UK	(1)-(2) Difference
Panel A: Respondent Demographics			
Age	42.72 (11.90)	41.37 (10.89)	1.35** (0.53)
Years Schooling (Derived)	16.03 (2.11)	15.98 (2.23)	0.05 (0.10)
Woman+	0.46 (0.50)	0.48 (0.50)	-0.03 (0.02)
Full Time	0.86 (0.35)	0.83 (0.37)	0.03* (0.02)
Part Time	0.08 (0.27)	0.10 (0.30)	-0.02* (0.01)
Self Empl.	0.06 (0.24)	0.06 (0.24)	-0.00 (0.01)
Other Empl. Status	0.00 (0.07)	0.01 (0.09)	-0.00 (0.00)
<i>Politics (Left/Right)</i>	0.02 (1.22)	-0.19 (0.90)	0.21*** (0.05)
<i>Income (Numeric)</i>	89.81 (52.59)	61.28 (38.92)	28.53*** (2.14)
<i>Was Unempl. (Self)</i>	0.44 (0.50)	0.30 (0.46)	0.14*** (0.02)
<i>Was Unempl. (Fam.)</i>	0.57 (0.50)	0.40 (0.49)	0.17*** (0.02)
<i>Unempl. Dur. (Numeric)</i>	6.30 (4.14)	5.97 (4.29)	0.33 (0.33)
Panel B: Role & Team Characteristics			
Mng. Levels from Top (Num)	2.80 (1.58)	3.09 (1.72)	-0.29*** (0.08)
Direct Reports (Num)	8.38 (5.59)	7.19 (5.60)	1.19*** (0.26)
Team Total Empl. (Num)	24.48 (34.85)	26.03 (37.83)	-1.55 (1.69)
<i>Team Budget (Numeric,)</i>	3602.48 (11235.34)	4191.46 (13715.03)	-588.99 (702.05)
N	934	929	1,863

Notes: The sample is restricted to respondents with non-missing values for all four primary outcome indices. The table reports unweighted summary statistics by country. Columns (1) and (2) report country means with standard deviations in parentheses. Column (3) reports the difference in means (US–UK) with standard errors in parentheses from unequal-variance two-sided *t*-tests. Age is self-reported in whole years. Years of Schooling maps education categories to years. Woman+ is an indicator equal to one for respondents who don't identify as a Man. Full Time, Part Time, and Self Empl. and Other Empl. (e.g., unemployed, not in the labor force, or student) are indicators for employment status. Politics is a centered 5-point ideology measure from Very liberal to Very conservative. Income converts self-reported income brackets into midpoints (in thousands of dollars, with \$200k+ top-coded as \$250k). Was Unempl. (Self/Fam.) are indicators for whether the respondent (or a family member) has experienced unemployment. Unempl. Dur. converts categorical unemployment duration into midpoints in months. Mng. Levels from Top (Num), Direct Reports (Num), and Team Total Empl. (Num) convert categorical bins into midpoints. Team Budget converts budget brackets into midpoints expressed in thousands of dollars. Italicized variables were collected after the video module to reduce attrition prior to the randomized portion of the survey and because they are stable characteristics unlikely to be affected by treatment. “Unsure” and “Not relevant” responses are treated as missing. (***) $p < 0.01$, (**) $p < 0.05$, (*) $p < 0.10$.

Table A.2: Pre-Treatment AI Posture Balance Across Countries

	(1) US	(2) UK	(1)-(2) Difference
Panel A: Respondent			
Sentiment on AI	1.08 (0.90)	0.89 (0.84)	0.19*** (0.04)
AI Augments	0.55 (0.50)	0.49 (0.50)	0.06** (0.02)
AI Automates	0.11 (0.32)	0.14 (0.35)	-0.03* (0.02)
AI does Both	0.25 (0.43)	0.27 (0.44)	-0.01 (0.02)
AI does Neither	0.07 (0.25)	0.09 (0.28)	-0.02 (0.01)
Had AI Training	0.63 (0.48)	0.49 (0.50)	0.14*** (0.02)
Expect AI Adoption	0.89 (0.77)	0.82 (0.71)	0.07** (0.03)
Panel B: Team			
Sentiment on AI	1.02 (0.94)	0.79 (0.95)	0.23*** (0.04)
Had AI Training	0.57 (0.49)	0.44 (0.50)	0.13*** (0.02)
Team AI Use Prop. (Numeric)	0.61 (0.31)	0.56 (0.31)	0.04*** (0.01)
Team AI Use Freq. (Numeric)	0.60 (0.37)	0.57 (0.38)	0.03* (0.02)
Panel C: Firm			
Leadership AI Support	2.13 (0.83)	1.97 (0.88)	0.15*** (0.04)
AI Training	0.58 (0.49)	0.48 (0.50)	0.10*** (0.02)
AI Encouragement	0.60 (0.49)	0.55 (0.50)	0.05** (0.02)
Reqd AI Use	0.28 (0.45)	0.22 (0.42)	0.05** (0.02)
AI Hiring Freeze	0.15 (0.36)	0.10 (0.30)	0.05*** (0.02)
AI Guidance	0.74 (0.44)	0.70 (0.46)	0.03 (0.02)
No AI Initiatives	0.32 (0.47)	0.31 (0.46)	0.01 (0.02)
AI Restricted	0.23 (0.42)	0.27 (0.45)	-0.05** (0.02)
Other Firm AI Policy	0.08 (0.27)	0.07 (0.25)	0.01 (0.02)
N	934	929	1,863

Notes: The sample is restricted to respondents with non-missing values for all four primary outcome indices. The table reports unweighted summary statistics by country. Columns (1) and (2) report country means with standard deviations in parentheses. Column (3) reports the difference in means (US–UK) with standard errors in parentheses from unequal-variance two-sided *t*-tests. Sentiment on AI (Respondent) is a centered 5-point pre-treatment measure of the respondent’s sentiment toward AI, with higher values indicating more positive sentiment. AI Augments, AI Automates, AI does Both, and AI does Neither are indicator variables constructed from a pre-treatment categorical question on how the respondent views AI’s impact on work. Had AI Training (Respondent) is an indicator equal to one if the respondent reports having received AI training. Expected AI Adoption is a centered 5-point pre-treatment measure of expected AI adoption over the next year. Sentiment on AI (Team) is a centered 5-point pre-treatment measure of the respondent’s assessment of their team’s sentiment toward AI. Had AI Training (Team) is an indicator equal to one if the respondent reports that members of their team have received AI training. Team AI Use Prop. converts a categorical measure of the share of the team using AI into numeric values. Team AI Use Freq. converts a categorical measure of team AI use frequency into numeric values. Leadership AI Support is a pre-treatment measure of perceived executive support for AI. Firm AI policy variables are indicators for whether the respondent reports that their firm has each policy: AI Training, AI Encouragement, Required AI Use, AI Hiring Freeze, AI Guidance, No AI Initiatives, AI Restricted, and Other Firm AI Policy. “Unsure” and “Not relevant” responses are treated as missing. (***) $p < 0.01$, ** $p < 0.05$, * $p < 0.10$.

Table A.3: Pre-Treatment AI Posture Balance across US Survey Arms

	(1) Control	(2) AI Productivity	(3) AI Labor Distress	(2)-(1) Difference	(3)-(1) Difference
Panel A: Respondent					
Sentiment on AI	1.01 (0.98)	1.06 (0.93)	1.07 (0.85)	0.05 (0.07)	0.05 (0.07)
AI Augments	0.53 (0.50)	0.54 (0.50)	0.55 (0.50)	0.01 (0.04)	0.02 (0.04)
AI Automates	0.11 (0.32)	0.09 (0.28)	0.13 (0.34)	-0.03 (0.02)	0.01 (0.02)
AI does Both	0.25 (0.43)	0.26 (0.44)	0.24 (0.43)	0.01 (0.03)	-0.01 (0.03)
AI does Neither	0.08 (0.27)	0.11 (0.31)	0.06 (0.24)	0.03 (0.02)	-0.01 (0.02)
Had AI Training	0.65 (0.48)	0.57 (0.50)	0.63 (0.48)	-0.07** (0.04)	-0.02 (0.04)
Expect AI Adoption	0.90 (0.79)	0.84 (0.79)	0.86 (0.75)	-0.06 (0.06)	-0.04 (0.06)
Panel B: Team					
Sentiment on AI	0.95 (0.97)	1.01 (0.93)	0.99 (1.01)	0.06 (0.07)	0.03 (0.08)
Had AI Training	0.59 (0.49)	0.52 (0.50)	0.56 (0.50)	-0.08** (0.04)	-0.03 (0.04)
Team AI Use Prop. (Numeric)	0.59 (0.31)	0.57 (0.32)	0.62 (0.30)	-0.02 (0.02)	0.03 (0.02)
Team AI Use Freq. (Numeric)	0.57 (0.39)	0.56 (0.38)	0.62 (0.37)	-0.01 (0.03)	0.04 (0.03)
Panel C: Firm					
Leadership AI Support	2.09 (0.85)	2.05 (0.86)	2.18 (0.80)	-0.04 (0.07)	0.09 (0.06)
AI Training	0.60 (0.49)	0.50 (0.50)	0.59 (0.49)	-0.10** (0.04)	-0.01 (0.04)
AI Encouragement	0.61 (0.49)	0.54 (0.50)	0.61 (0.49)	-0.07* (0.04)	-0.00 (0.04)
Reqd AI Use	0.28 (0.45)	0.24 (0.43)	0.29 (0.46)	-0.04 (0.03)	0.01 (0.04)
AI Hiring Freeze	0.15 (0.36)	0.13 (0.33)	0.17 (0.38)	-0.02 (0.03)	0.03 (0.03)
AI Guidance	0.75 (0.43)	0.66 (0.47)	0.73 (0.45)	-0.09** (0.04)	-0.03 (0.03)
No AI Initiatives	0.32 (0.47)	0.35 (0.48)	0.33 (0.47)	0.03 (0.04)	0.00 (0.04)
AI Restricted	0.22 (0.42)	0.24 (0.43)	0.22 (0.42)	0.02 (0.03)	0.00 (0.03)
Other Firm AI Policy	0.11 (0.31)	0.07 (0.25)	0.06 (0.24)	-0.04 (0.03)	-0.04 (0.03)
N	402	295	303	697	705

Notes: The table reports unweighted summary statistics for the U.S. subsample. Columns (1)–(3) report arm means with standard deviations in parentheses for the Control, AI Productivity, and AI Labor Distress treatment conditions. Columns (4) and (5) report differences in means relative to the Control group (AI Productivity–Control and AI Labor Distress–Control), with standard errors in parentheses from unequal-variance two-sided t -tests. Sentiment on AI (Respondent) is a centered 5-point pre-treatment measure of the respondent’s sentiment toward AI, with higher values indicating more positive sentiment. AI Augments, AI Automates, AI does Both, and AI does Neither are indicator variables constructed from a pre-treatment categorical question on how the respondent views AI’s impact on work. Had AI Training (Respondent) is an indicator equal to one if the respondent reports having received AI training. Expected AI Adoption is a centered 5-point pre-treatment measure of expected AI adoption over the next year. Sentiment on AI (Team) is a centered 5-point pre-treatment measure of the respondent’s assessment of their team’s sentiment toward AI. Had AI Training (Team) is an indicator equal to one if the respondent reports that members of their team have received AI training. Team AI Use Prop. converts a categorical measure of the share of the team using AI into numeric values. Team AI Use Freq. converts a categorical measure of team AI use frequency into numeric values. Leadership AI Support is a pre-treatment measure of perceived executive support for AI. Firm AI policy variables are indicators for whether the respondent reports that their firm has each policy: AI Training, AI Encouragement, Required AI Use, AI Hiring Freeze, AI Guidance, No AI Initiatives, AI Restricted, and Other Firm AI Policy. “Unsure” and “Not relevant” responses are treated as missing. (***) $p < 0.01$, ** $p < 0.05$, * $p < 0.10$.

Table A.4: Pre-Treatment AI Posture Balance across UK Survey Arms

	(1) Control	(2) AI Productivity	(3) AI Labor Distress	(2)-(1) Difference	(3)-(1) Difference
Panel A: Respondent					
Sentiment on AI	0.93 (0.82)	0.87 (0.81)	0.82 (0.90)	-0.07 (0.06)	-0.11* (0.07)
AI Augments	0.50 (0.50)	0.47 (0.50)	0.49 (0.50)	-0.03 (0.04)	-0.02 (0.04)
AI Automates	0.13 (0.34)	0.15 (0.36)	0.13 (0.34)	0.02 (0.03)	0.00 (0.03)
AI does Both	0.27 (0.44)	0.26 (0.44)	0.26 (0.44)	-0.01 (0.03)	-0.01 (0.03)
AI does Neither	0.08 (0.27)	0.08 (0.28)	0.10 (0.30)	0.00 (0.02)	0.02 (0.02)
Had AI Training	0.47 (0.50)	0.48 (0.50)	0.47 (0.50)	0.01 (0.04)	0.00 (0.04)
Expect AI Adoption	0.82 (0.67)	0.83 (0.70)	0.81 (0.75)	0.01 (0.05)	-0.01 (0.06)
Panel B: Team					
Sentiment on AI	0.83 (0.91)	0.81 (0.88)	0.69 (1.03)	-0.02 (0.07)	-0.14* (0.08)
Had AI Training	0.42 (0.49)	0.42 (0.49)	0.44 (0.50)	0.00 (0.04)	0.02 (0.04)
Team AI Use Prop. (Numeric)	0.54 (0.31)	0.57 (0.32)	0.55 (0.31)	0.04 (0.02)	0.01 (0.02)
Team AI Use Freq. (Numeric)	0.56 (0.38)	0.55 (0.38)	0.56 (0.39)	-0.01 (0.03)	-0.00 (0.03)
Panel C: Firm					
Leadership AI Support	1.97 (0.87)	1.94 (0.89)	1.95 (0.89)	-0.04 (0.07)	-0.02 (0.07)
AI Training	0.49 (0.50)	0.42 (0.49)	0.46 (0.50)	-0.07* (0.04)	-0.03 (0.04)
AI Encouragement	0.59 (0.49)	0.51 (0.50)	0.50 (0.50)	-0.08** (0.04)	-0.09** (0.04)
Reqd AI Use	0.22 (0.42)	0.20 (0.40)	0.22 (0.42)	-0.03 (0.03)	0.00 (0.03)
AI Hiring Freeze	0.08 (0.28)	0.13 (0.33)	0.09 (0.29)	0.04 (0.03)	0.01 (0.02)
AI Guidance	0.69 (0.46)	0.71 (0.45)	0.68 (0.47)	0.02 (0.04)	-0.01 (0.04)
No AI Initiatives	0.32 (0.47)	0.34 (0.47)	0.29 (0.45)	0.02 (0.04)	-0.03 (0.04)
AI Restricted	0.27 (0.44)	0.26 (0.44)	0.29 (0.45)	-0.01 (0.04)	0.02 (0.04)
Other Firm AI Policy	0.03 (0.18)	0.07 (0.26)	0.09 (0.29)	0.04 (0.03)	0.06* (0.03)
N	400	295	300	695	700

Notes: The table reports unweighted summary statistics for the U.K. subsample. Columns (1)–(3) report arm means with standard deviations in parentheses for the Control, AI Productivity, and AI Labor Distress treatment conditions. Columns (4) and (5) report differences in means relative to the Control group (AI Productivity–Control and AI Labor Distress–Control), with standard errors in parentheses from unequal-variance two-sided t -tests. Sentiment on AI (Respondent) is a centered 5-point pre-treatment measure of the respondent’s sentiment toward AI, with higher values indicating more positive sentiment. AI Augments, AI Automates, AI does Both, and AI does Neither are indicator variables constructed from a pre-treatment categorical question on how the respondent views AI’s impact on work. Had AI Training (Respondent) is an indicator equal to one if the respondent reports having received AI training. Expected AI Adoption is a centered 5-point pre-treatment measure of expected AI adoption over the next year. Sentiment on AI (Team) is a centered 5-point pre-treatment measure of the respondent’s assessment of their team’s sentiment toward AI. Had AI Training (Team) is an indicator equal to one if the respondent reports that members of their team have received AI training. Team AI Use Prop. converts a categorical measure of the share of the team using AI into numeric values. Team AI Use Freq. converts a categorical measure of team AI use frequency into numeric values. Leadership AI Support is a pre-treatment measure of perceived executive support for AI. Firm AI policy variables are indicators for whether the respondent reports that their firm has each policy: AI Training, AI Encouragement, Required AI Use, AI Hiring Freeze, AI Guidance, No AI Initiatives, AI Restricted, and Other Firm AI Policy. “Unsure” and “Not relevant” responses are treated as missing. (***) $p < 0.01$, ** $p < 0.05$, * $p < 0.10$.

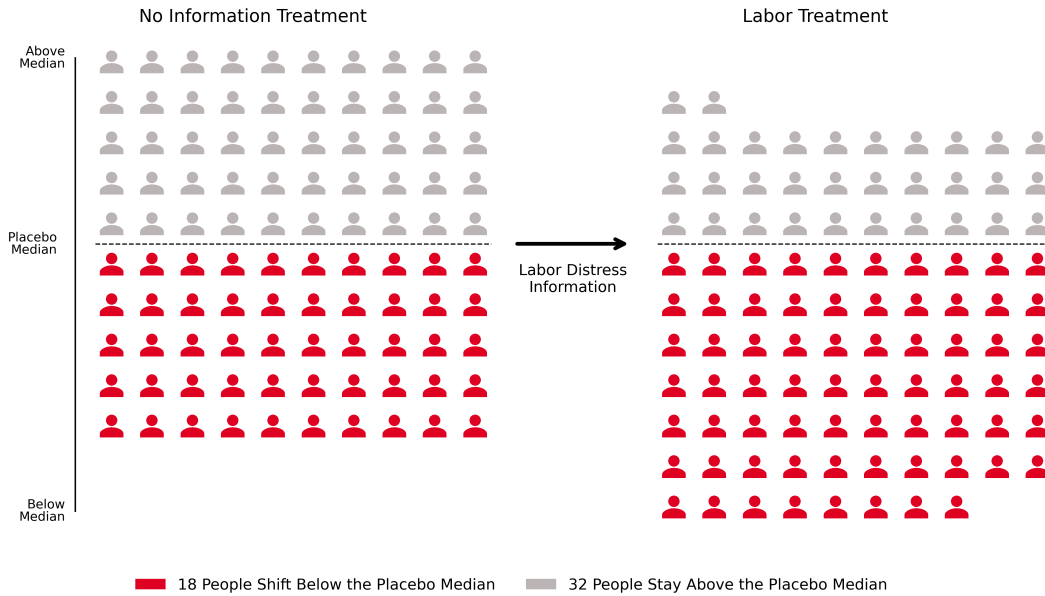
A.2 Main Treatment Effect Estimates

Table A.5: Treatment Effects on Main Indices

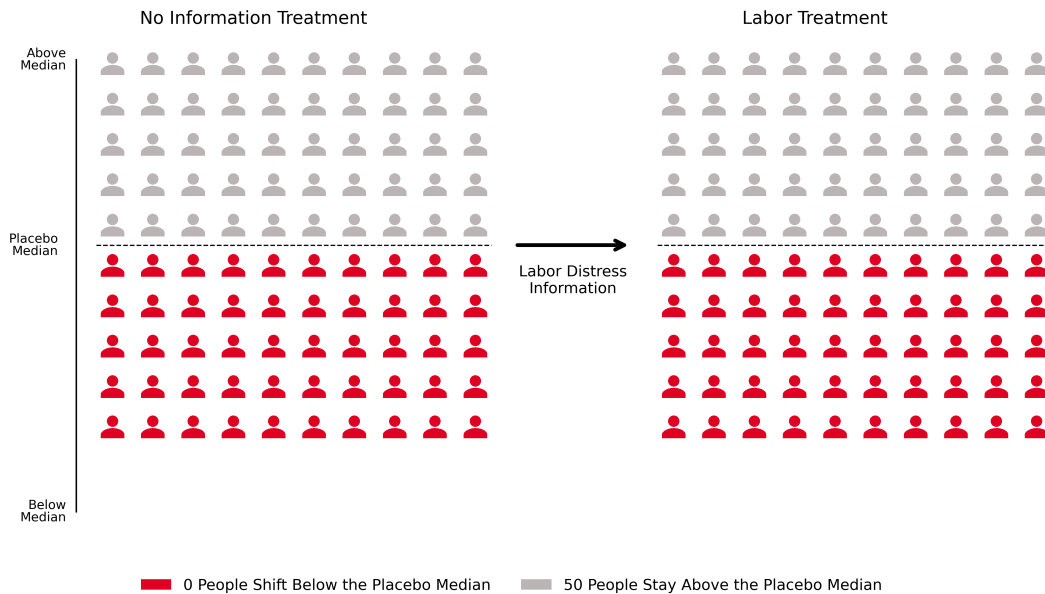
	(1)	(2)	(3)	(4)
	AI Adoption	AI Advocacy	Staffing Intentions	Staffing Advocacy
Panel A: Pooled (Country FE)				
AI Labor	-0.426*** (0.060)	-0.461*** (0.059)	-0.206*** (0.055)	-0.001 (0.056)
AI Productivity	0.033 (0.057)	0.112** (0.054)	-0.084 (0.056)	-0.067 (0.059)
R-squared	0.037	0.052	0.018	0.001
Observations	1863	1863	1863	1863
Panel B: US				
AI Labor	-0.406*** (0.085)	-0.491*** (0.085)	-0.331*** (0.084)	0.002 (0.078)
AI Productivity	0.076 (0.083)	0.126 (0.080)	-0.078 (0.085)	-0.078 (0.084)
R-squared	0.036	0.056	0.017	0.001
Observations	934	934	934	934
Panel C: UK				
AI Labor	-0.447*** (0.086)	-0.431*** (0.083)	-0.080 (0.070)	-0.005 (0.079)
AI Productivity	-0.009 (0.077)	0.098 (0.074)	-0.088 (0.073)	-0.056 (0.082)
R-squared	0.037	0.047	0.002	0.001
Observations	929	929	929	929
UK Diff, Labor	-0.040	0.060	0.251**	-0.007
UK Diff, Prod.	-0.086	-0.028	-0.010	0.023
R-squared	0.037	0.052	0.022	0.001
Observations	1863	1863	1863	1863

Notes: The table reports OLS estimates from equation (1) with no controls, where the unit of observation is an individual manager. The dependent variables are four standardized primary outcome indices. *AI Adoption* is a 4-item index measuring intentions to adopt or expand AI use in the respondent's team/unit (higher values indicate greater adoption intent). *AI Advocacy* is a 2-item index measuring willingness to advocate for expanding AI use within the organization (higher values indicate more pro-AI advocacy). *Staffing Intentions* is a 3-item index measuring expected headcount changes over the next year for the respondent's three most common team occupation groups (higher values indicate stronger hiring/expansion plans). *Staffing Advocacy* is a 2-item index measuring willingness to advocate for hiring versus layoffs (higher values indicate more pro-hiring advocacy). Indices are constructed in three steps: each component item is standardized using the pooled control-group mean and standard deviation, standardized items are averaged (requiring at least half of the component items to be non-missing), and the resulting index is standardized again using the pooled control-group mean and standard deviation. More information on index construction is described in Section 2.4. The omitted category is the neutral control video, so coefficients on *AI Productivity* and *AI Labor* compare each treatment condition to the control group. Panels report estimates separately for the U.S. and U.K. subsamples. The bottom rows report the difference in treatment effects between the U.K. and U.S. for each arm, along with R^2 and the number of observations. The sample is restricted to respondents with non-missing values for all four primary outcome indices. Robust standard errors are reported in parentheses. (***) $p < 0.01$, (**) $p < 0.05$, (*) $p < 0.10$.

Figure A.1: Interpreting the Labor-Treatment Effects on AI Advocacy and Staffing Advocacy



(a) AI Advocacy



(b) Staffing Advocacy

Notes: This figure translates the pooled labor-treatment estimates for *AI Advocacy* and *Staffing Advocacy* into an “out of 100” comparison using Cohen’s U_3 , which maps the estimated effect into the implied share of the treated distribution falling below the placebo benchmark. In each panel, the left side depicts the placebo benchmark, and the right side depicts the implied distribution under the AI labor treatment. Red icons indicate the implied number of respondents below the placebo benchmark, while gray icons indicate those above it. Because the underlying indices are constructed from discrete survey items and are not exactly normally distributed, these figures should be interpreted as a visualization of effect size rather than as a literal reallocation of observed respondents. In the realized placebo-group data, the share below the placebo-group mean is not exactly 50 percent because the outcome indices are discrete and somewhat skewed. A complementary rank-based exercise that uses the placebo-group median after adding a small mean-zero jitter to break ties yields nearly identical results, suggesting that the visual magnitudes are not driven by the normality approximation.

A.3 High-Low Heterogeneity Specifications

For each outcome, we report estimates based on the interacted heterogeneity specification in equation (3). In that specification, δ_k^H gives the treatment effect for respondents in the low subgroup, $\delta_k^H + \gamma_k^H$ gives the treatment effect for respondents in the high subgroup, γ_k^H captures the differential treatment effect between the high and low subgroups for treatment arm k , and ζ captures the mean difference in the outcome between the two subgroups. All specifications pool both countries, include country fixed effects, and use heteroskedasticity-robust standard errors. Construction of the high and low subgroup indicators is described in Appendix Section E.1.

For each heterogeneity dimension, the first table reports the coefficient estimates from equation (3) for each outcome. The second table reports subgroup-specific treatment effects implied by equation (3). In particular, for treatment arm k , the treatment effect in the low subgroup is δ_k^H , and the treatment effect in the high subgroup is $\delta_k^H + \gamma_k^H$. The bottom panel reports p-values for the null hypothesis that $\gamma_k^H = 0$, that is, for whether the treatment effect differs significantly between the high and low subgroups.

Table A.6: Pre-Treatment Expected Staffing Plans: Interaction Model Estimates

	(1)	(2)	(3)	(4)
	AI Adopt	AI Advoc.	Staff Int.	Staff Advoc.
AI Labor	-0.260*** (0.0852)	-0.369*** (0.0875)	-0.0826 (0.0708)	-0.0664 (0.0766)
AI Prod.	0.00208 (0.0773)	0.148** (0.0751)	-0.107 (0.0650)	-0.175** (0.0774)
High Staffing Expansion Plans	0.211*** (0.0741)	0.248*** (0.0738)	0.755*** (0.0682)	-0.0105 (0.0743)
AI Labor $\times$ High Staffing Expansion Plans	-0.357*** (0.123)	-0.181 (0.120)	-0.279*** (0.107)	0.136 (0.113)
AI Prod. $\times$ High Staffing Expansion Plans	0.0931 (0.114)	-0.0531 (0.109)	0.156 (0.106)	0.243** (0.119)
Observations	1826	1826	1826	1826
R-squared	0.048	0.059	0.149	0.006

Notes: This table reports heterogeneous treatment effect regressions for the four primary outcome indices using pre-treatment expected staffing plans next year as the heterogeneity dimension. Columns correspond to *AI Adoption*, *AI Advocacy*, *Staffing Intentions*, and *Staffing Advocacy*. The estimated specification includes indicators for assignment to the AI Productivity and AI Labor treatments, an indicator for *High Staffing Expansion Plans*, interactions between each treatment indicator and *High Staffing Expansion Plans*, and country fixed effects. The omitted category is the control group among respondents in the low staffing-expansion-plans subgroup. The row labeled *High Staffing Expansion Plans* reports the coefficient on the subgroup indicator. The interaction rows report how the treatment effects differ for respondents with high staffing expansion plans relative to respondents in the low subgroup. Expected staffing plans next year is constructed from three pre-treatment five-point items measuring expected staffing changes for the team's most common occupations. Respondents are classified into the high subgroup if their composite staffing-plan measure is strictly greater than 3, and into the low subgroup if it is less than or equal to 3, where 3 corresponds to "No Change" on the underlying scale. Subsequently, the high group represents respondents who, on average, expect to expand headcount in their three most common occupations. Construction of this subgroup measure is described in Appendix Section E.1. All specifications are estimated on the pooled sample across countries and include robust standard errors in parentheses. Outcome construction is described in Section 2.4. (***) $p < 0.01$, (**) $p < 0.05$, (*) $p < 0.10$.

Table A.7: Pre-Treatment Expected Staffing Plans: Subgroup Treatment Effects Relative to Control

	AI Adopt	AI Advoc.	Staff Int.	Staff Advoc.
AI Labor Effect (Low Staffing Expansion Plans)	-0.260*** (0.085)	-0.369*** (0.087)	-0.083 (0.071)	-0.066 (0.077)
AI Labor Effect (High Staffing Expansion Plans)	-0.617*** (0.088)	-0.549*** (0.082)	-0.361*** (0.080)	0.069 (0.083)
AI Productivity Effect (Low Staffing Expansion Plans)	0.002 (0.077)	0.148** (0.075)	-0.107 (0.065)	-0.175** (0.077)
AI Productivity Effect (High Staffing Expansion Plans)	0.095 (0.084)	0.095 (0.079)	0.049 (0.083)	0.069 (0.091)
p diff (Prod: High vs Low)	0.415	0.627	0.141	0.042
p diff (Labor: High vs Low)	0.004	0.133	0.009	0.230

Notes: This table reports subgroup-specific treatment effects implied by the interacted heterogeneity specification using pre-treatment expected staffing plans next year as the heterogeneity dimension. Columns correspond to *AI Adoption*, *AI Advocacy*, *Staffing Intentions*, and *Staffing Advocacy*. The row labeled *AI Labor Effect (Low Staffing Expansion Plans)* reports the estimated effect of the AI Labor treatment relative to the control group among respondents in the low staffing-expansion-plans subgroup. The row labeled *AI Labor Effect (High Staffing Expansion Plans)* reports the corresponding effect among respondents in the high subgroup. The next two rows report analogous treatment effects for the AI Productivity treatment within the low and high subgroups. These subgroup-specific effects are obtained from the interacted regression with treatment indicators, a *High Staffing Expansion Plans* indicator, interactions between each treatment indicator and *High Staffing Expansion Plans*, and country fixed effects. Expected staffing plans next year is constructed from three pre-treatment five-point items measuring expected staffing changes for the team's most common occupations. Respondents are classified into the high subgroup if their composite staffing-plan measure is strictly greater than 3, and into the low subgroup if it is less than or equal to 3, where 3 corresponds to "No Change" on the underlying scale. Subsequently, the high group represents respondents who, on average, expect to expand headcount in their three most common occupations. Construction of this subgroup measure is described in Appendix Section E.1. The bottom panel reports p-values for the difference in treatment effects between the high and low subgroups for each treatment arm. These p-values correspond to tests of whether the relevant interaction coefficient is equal to zero in the underlying interacted regression. All specifications are estimated on the pooled sample across countries with country fixed effects, and robust standard errors are reported in parentheses. Outcome construction is described in Section 2.4. (***) $p < 0.01$, (**) $p < 0.05$, (*) $p < 0.10$.

Table A.8: Team AI Use Frequency: Interaction Model Estimates

	(1) AI Adopt	(2) AI Advoc.	(3) Staff Int.	(4) Staff Advoc.
AI Labor	-0.428*** (0.0799)	-0.461*** (0.0766)	-0.221*** (0.0669)	0.00115 (0.0681)
AI Prod.	0.101 (0.0706)	0.202*** (0.0688)	-0.0796 (0.0666)	-0.136* (0.0732)
High AI Use Frequency	0.736*** (0.0656)	0.695*** (0.0677)	0.198*** (0.0755)	-0.169** (0.0761)
AI Labor × High AI Use Frequency	-0.0513 (0.112)	-0.0477 (0.113)	0.0221 (0.115)	0.0180 (0.115)
AI Prod. × High AI Use Frequency	-0.150 (0.109)	-0.215** (0.105)	-0.0120 (0.121)	0.170 (0.122)
Observations	1852	1852	1852	1852
R-squared	0.134	0.137	0.028	0.005

Notes: This table reports heterogeneous treatment effect regressions for the four primary outcome indices using team AI use frequency as the heterogeneity dimension. Columns correspond to *AI Adoption*, *AI Advocacy*, *Staffing Intentions*, and *Staffing Advocacy*. The estimated specification includes indicators for assignment to the AI Productivity and AI Labor treatments, an indicator for *High AI Use Frequency*, interactions between each treatment indicator and *High AI Use Frequency*, and country fixed effects. The omitted category is the control group among respondents in the low team-AI-use-frequency subgroup. The row labeled *High AI Use Frequency* reports the coefficient on the subgroup indicator. The interaction rows report how the treatment effects differ for respondents in the high team-AI-use-frequency subgroup relative to respondents in the low subgroup. Team AI use frequency is measured using five ordered categories: Never, Rarely, Few/month, Few/week, and Daily. Respondents are classified into the high subgroup if they report Daily AI use on their team, and into the low subgroup otherwise. Construction of this subgroup measure is described in Appendix Section E.1. All specifications are estimated on the pooled sample across countries and include robust standard errors in parentheses. Outcome construction is described in Section 2.4. (***) $p < 0.01$, (**) $p < 0.05$, (*) $p < 0.10$.

Table A.9: Team AI Use Frequency: Subgroup Treatment Effects Relative to Control

	AI Adopt	AI Advoc.	Staff Int.	Staff Advoc.
AI Labor Effect (Low AI Use Frequency)	-0.428*** (0.080)	-0.461*** (0.077)	-0.221*** (0.067)	0.001 (0.068)
AI Labor Effect (High AI Use Frequency)	-0.479*** (0.079)	-0.509*** (0.083)	-0.199** (0.094)	0.019 (0.093)
AI Productivity Effect (Low AI Use Frequency)	0.101 (0.071)	0.202*** (0.069)	-0.080 (0.067)	-0.136* (0.073)
AI Productivity Effect (High AI Use Frequency)	-0.049 (0.083)	-0.013 (0.079)	-0.092 (0.101)	0.033 (0.098)
p diff (Prod: High vs Low)	0.167	0.040	0.921	0.165
p diff (Labor: High vs Low)	0.647	0.673	0.848	0.876

Notes: This table reports subgroup-specific treatment effects implied by the interacted heterogeneity specification using team AI use frequency as the heterogeneity dimension. Columns correspond to *AI Adoption*, *AI Advocacy*, *Staffing Intentions*, and *Staffing Advocacy*. The row labeled *AI Labor Effect (Low AI Use Frequency)* reports the estimated effect of the AI Labor treatment relative to the control group among respondents in the low team-AI-use-frequency subgroup. The row labeled *AI Labor Effect (High AI Use Frequency)* reports the corresponding effect among respondents in the high subgroup. The next two rows report analogous treatment effects for the AI Productivity treatment within the low and high subgroups. These subgroup-specific effects are obtained from the interacted regression with treatment indicators, a *High Team AI Use Frequency* indicator, interactions between each treatment indicator and *High Team AI Use Frequency*, and country fixed effects. Team AI use frequency is measured using five ordered categories: Never, Rarely, Few/month, Few/week, and Daily. Respondents are classified into the high subgroup if they report Daily AI use on their team, and into the low subgroup otherwise. Construction of this subgroup measure is described in Appendix Section E.1. The bottom panel reports p-values for the difference in treatment effects between the high and low subgroups for each treatment arm. These p-values correspond to tests of whether the relevant interaction coefficient is equal to zero in the underlying interacted regression. All specifications are estimated on the pooled sample across countries with country fixed effects, and robust standard errors are reported in parentheses. Outcome construction is described in Section 2.4. (***) $p < 0.01$, (**) $p < 0.05$, (*) $p < 0.10$.

Table A.10: Respondent's Pre-Treatment AI Views: Interaction Model Estimates

	(1)	(2)	(3)	(4)
	AI Adopt	AI Advoc.	Staff Int.	Staff Advoc.
AI Labor	-0.363** (0.150)	-0.496*** (0.130)	0.0532 (0.103)	-0.119 (0.109)
AI Prod.	0.0722 (0.121)	0.155 (0.117)	-0.0478 (0.0936)	-0.166 (0.105)
Positive Self AI Attitudes	1.005*** (0.0913)	1.080*** (0.0877)	0.256*** (0.0745)	-0.319*** (0.0787)
AI Labor × Positive Self AI Attitudes	-0.0897 (0.160)	0.0376 (0.142)	-0.341*** (0.122)	0.164 (0.127)
AI Prod. × Positive Self AI Attitudes	-0.0398 (0.133)	-0.0495 (0.128)	-0.0521 (0.116)	0.130 (0.126)
Observations	1843	1843	1843	1843
R-squared	0.188	0.244	0.025	0.011

Notes: This table reports heterogeneous treatment effect regressions for the four primary outcome indices using respondents' pre-treatment AI views as the heterogeneity dimension. Columns correspond to *AI Adoption*, *AI Advocacy*, *Staffing Intentions*, and *Staffing Advocacy*. The estimated specification includes indicators for assignment to the AI Productivity and AI Labor treatments, an indicator for *Positive Self AI Attitudes*, interactions between each treatment indicator and *Positive Self AI Attitudes*, and country fixed effects. The omitted category is the control group among respondents in the low pre-treatment-AI-views subgroup. The row labeled *Positive Self AI Attitudes* reports the coefficient on the subgroup indicator. The interaction rows report how the treatment effects differ for respondents in the high pre-treatment-AI-views subgroup relative to respondents in the low subgroup. Respondents' pre-treatment AI views are measured on a five-point scale ranging from *Negative High* to *Positive High*. Respondents are classified into the high subgroup if they report *Positive Low* or *Positive High* views towards AI, and into the low subgroup if they report *Negative High*, *Negative Low*, or *Neutral* views towards AI. Construction of this subgroup measure is described in Appendix Section E.1. All specifications are estimated on the pooled sample across countries and include robust standard errors in parentheses. Outcome construction is described in Section 2.4. (***) $p < 0.01$, (**) $p < 0.05$, (*) $p < 0.10$.

Table A.11: Respondent's Pre-Treatment AI Views: Subgroup Treatment Effects Relative to Control

	AI Adopt	AI Advoc.	Staff Int.	Staff Advoc.
AI Labor Effect (Not Positive Self Attitudes)	-0.363** (0.150)	-0.496*** (0.130)	0.053 (0.103)	-0.119 (0.109)
AI Labor Effect (Positive Self Attitudes)	-0.452*** (0.057)	-0.459*** (0.058)	-0.287*** (0.065)	0.045 (0.065)
AI Productivity Effect (Not Positive Self Attitudes)	0.072 (0.121)	0.155 (0.117)	-0.048 (0.094)	-0.166 (0.105)
AI Productivity Effect (Positive Self Attitudes)	0.032 (0.055)	0.106** (0.052)	-0.100 (0.068)	-0.036 (0.070)
p diff (Prod: High vs Low)	0.765	0.699	0.654	0.304
p diff (Labor: High vs Low)	0.575	0.791	0.005	0.196

Notes: This table reports subgroup-specific treatment effects implied by the interacted heterogeneity specification using respondent's pre-treatment AI views as the heterogeneity dimension. Columns correspond to *AI Adoption*, *AI Advocacy*, *Staffing Intentions*, and *Staffing Advocacy*. The row labeled *AI Labor Effect (Not Positive Self Attitudes)* reports the estimated effect of the AI Labor treatment relative to the control group among respondents in the low pre-treatment-AI-views subgroup. The row labeled *AI Labor Effect (Positive Self Attitudes)* reports the corresponding effect among respondents in the high subgroup. The next two rows report analogous treatment effects for the AI Productivity treatment within the low and high subgroups. These subgroup-specific effects are obtained from the interacted regression with treatment indicators, a *High Pre-Treatment AI Views* indicator, interactions between each treatment indicator and *High Pre-Treatment AI Views*, and country fixed effects. Respondents' pre-treatment AI views are measured on a five-point scale ranging from *Negative High* to *Positive High*. Respondents are classified into the high subgroup if they report *Positive Low* or *Positive High* views towards AI, and into the low subgroup if they report *Negative High*, *Negative Low*, or *Neutral* views towards AI. Construction of this subgroup measure is described in Appendix Section E.1. The bottom panel reports p-values for the difference in treatment effects between the high and low subgroups for each treatment arm. These p-values correspond to tests of whether the relevant interaction coefficient is equal to zero in the underlying interacted regression. All specifications are estimated on the pooled sample across countries with country fixed effects, and robust standard errors are reported in parentheses. Outcome construction is described in Section 2.4. (*** p<0.01, ** p<0.05, * p<0.10).

Table A.12: Technology Authority Heterogeneity: Interaction Model Estimates

	(1) AI Adopt	(2) AI Advoc.	(3) Staff Int.	(4) Staff Advoc.
AI Labor	-0.392*** (0.0697)	-0.450*** (0.0700)	-0.192*** (0.0626)	0.0180 (0.0640)
AI Prod.	0.0859 (0.0633)	0.141** (0.0619)	-0.0957 (0.0631)	-0.106 (0.0682)
High Tech Authority	0.112 (0.0832)	0.116 (0.0830)	-0.0586 (0.0849)	-0.175** (0.0811)
AI Labor × High Tech Authority	-0.139 (0.140)	-0.0284 (0.134)	-0.0529 (0.132)	-0.0973 (0.129)
AI Prod. × High Tech Authority	-0.215 (0.141)	-0.120 (0.131)	0.0124 (0.138)	0.125 (0.135)
Observations	1851	1851	1851	1851
R-squared	0.038	0.052	0.019	0.008

Notes: This table reports heterogeneous treatment effect regressions for the four primary outcome indices using technology authority as the heterogeneity dimension. Columns correspond to *AI Adoption*, *AI Advocacy*, *Staffing Intentions*, and *Staffing Advocacy*. The estimated specification includes indicators for assignment to the AI Productivity and AI Labor treatments, an indicator for *High Tech Authority*, interactions between each treatment indicator and *High Tech Authority*, and country fixed effects. The omitted category is the control group among respondents in the low technology-authority subgroup. The row labeled *High Tech Authority* reports the coefficient on the subgroup indicator. The interaction rows report how the treatment effects differ for respondents in the high technology-authority subgroup relative to respondents in the low subgroup. Technology authority captures where decisions about team technology adoption are made. Respondents are classified into the high subgroup if technology decisions are made mostly or entirely at the team level, and into the low subgroup if such decisions are made only or mostly at the corporate level or are shared across levels. Construction of this subgroup measure is described in Appendix Section E.1. All specifications are estimated on the pooled sample across countries and include robust standard errors in parentheses. Outcome construction is described in Section 2.4. (*** p<0.01, ** p<0.05, * p<0.10).

Table A.13: Technology Authority: Subgroup Treatment Effects Relative to Control

	AI Adopt	AI Advoc.	Staff Int.	Staff Advoc.
AI Labor Effect (Low Tech Authority)	-0.392*** (0.070)	-0.450*** (0.070)	-0.192*** (0.063)	0.018 (0.064)
AI Labor Effect (High Tech Authority)	-0.531*** (0.121)	-0.479*** (0.114)	-0.245** (0.116)	-0.079 (0.112)
AI Productivity Effect (Low Tech Authority)	0.086 (0.063)	0.141** (0.062)	-0.096 (0.063)	-0.106 (0.068)
AI Productivity Effect (High Tech Authority)	-0.129 (0.127)	0.021 (0.115)	-0.083 (0.122)	0.019 (0.116)
p diff (Prod: High vs Low)	0.129	0.360	0.928	0.356
p diff (Labor: High vs Low)	0.323	0.832	0.687	0.450

Notes: This table reports subgroup-specific treatment effects implied by the interacted heterogeneity specification using technology authority as the heterogeneity dimension. Columns correspond to *AI Adoption*, *AI Advocacy*, *Staffing Intentions*, and *Staffing Advocacy*. The row labeled *AI Labor Effect (Low Tech Authority)* reports the estimated effect of the AI Labor treatment relative to the control group among respondents in the low technology-authority subgroup. The row labeled *AI Labor Effect (High Tech Authority)* reports the corresponding effect among respondents in the high subgroup. The next two rows report analogous treatment effects for the AI Productivity treatment within the low and high subgroups. These subgroup-specific effects are obtained from the interacted regression with treatment indicators, a *High Technology Authority* indicator, interactions between each treatment indicator and *High Technology Authority*, and country fixed effects. Technology authority captures where decisions about team technology adoption are made. Respondents are classified into the high subgroup if technology decisions are made mostly or entirely at the team level, and into the low subgroup if such decisions are made only or mostly at the corporate level or are shared across levels. Construction of this subgroup measure is described in Appendix Section E.1. The bottom panel reports p-values for the difference in treatment effects between the high and low subgroups for each treatment arm. These p-values correspond to tests of whether the relevant interaction coefficient is equal to zero in the underlying interacted regression. All specifications are estimated on the pooled sample across countries with country fixed effects, and robust standard errors are reported in parentheses. Outcome construction is described in Section 2.4. (***) $p < 0.01$, ** $p < 0.05$, * $p < 0.10$.

Table A.14: Executive Support: Interaction Model Estimates

	(1)	(2)	(3)	(4)
	AI Adopt	AI Advoc.	Staff Int.	Staff Advoc.
AI Labor	-0.467*** (0.135)	-0.421*** (0.129)	-0.334*** (0.0950)	-0.164 (0.106)
AI Prod.	0.0655 (0.111)	0.251** (0.110)	-0.261*** (0.100)	-0.234** (0.109)
High Executive Support	0.837*** (0.0846)	0.788*** (0.0836)	-0.00142 (0.0770)	-0.211*** (0.0786)
AI Labor × High Executive Support	0.0274 (0.149)	-0.0698 (0.144)	0.157 (0.116)	0.208* (0.125)
AI Prod. × High Executive Support	-0.0202 (0.126)	-0.159 (0.124)	0.254** (0.121)	0.220* (0.130)
Observations	1825	1825	1825	1825
R-squared	0.156	0.143	0.025	0.004

Notes: This table reports heterogeneous treatment effect regressions for the four primary outcome indices using executive support for AI as the heterogeneity dimension. Columns correspond to *AI Adoption*, *AI Advocacy*, *Staffing Intentions*, and *Staffing Advocacy*. The estimated specification includes indicators for assignment to the AI Productivity and AI Labor treatments, an indicator for *High Executive Support*, interactions between each treatment indicator and *High Executive Support*, and country fixed effects. The omitted category is the control group among respondents in the low executive-support-for-AI subgroup. The row labeled *High Executive Support* reports the coefficient on the subgroup indicator. The interaction rows report how the treatment effects differ for respondents in the high executive-support-for-AI subgroup relative to respondents in the low subgroup. Executive support for AI captures respondents' perceptions of how supportive organizational leadership is of AI adoption. Respondents are classified into the high subgroup if they perceive leadership as moderately or strongly supportive of AI adoption, and into the low subgroup if they perceive leadership support as limited. Construction of this subgroup measure is described in Appendix Section E.1. All specifications are estimated on the pooled sample across countries and include robust standard errors in parentheses. Outcome construction is described in Section 2.4. (***) $p < 0.01$, ** $p < 0.05$, * $p < 0.10$.

Table A.15: Executive Support: Subgroup Treatment Effects Relative to Control

	AI Adopt	AI Advoc.	Staff Int.	Staff Advoc.
AI Labor Effect (Low Executive Support)	-0.467*** (0.135)	-0.421*** (0.129)	-0.334*** (0.095)	-0.164 (0.106)
AI Labor Effect (High Executive Support)	-0.440*** (0.062)	-0.491*** (0.063)	-0.176*** (0.066)	0.044 (0.066)
AI Productivity Effect (Low Executive Support)	0.065 (0.111)	0.251** (0.110)	-0.261*** (0.100)	-0.234** (0.109)
AI Productivity Effect (High Executive Support)	0.045 (0.060)	0.092 (0.058)	-0.007 (0.068)	-0.014 (0.070)
p diff (Prod: High vs Low)	0.873	0.202	0.036	0.090
p diff (Labor: High vs Low)	0.854	0.627	0.174	0.095

Notes: This table reports subgroup-specific treatment effects implied by the interacted heterogeneity specification using executive support for AI as the heterogeneity dimension. Columns correspond to *AI Adoption*, *AI Advocacy*, *Staffing Intentions*, and *Staffing Advocacy*. The row labeled *AI Labor Effect (Low Executive Support)* reports the estimated effect of the AI Labor treatment relative to the control group among respondents in the low executive-support-for-AI subgroup. The row labeled *AI Labor Effect (High Executive Support)* reports the corresponding effect among respondents in the high subgroup. The next two rows report analogous treatment effects for the AI Productivity treatment within the low and high subgroups. These subgroup-specific effects are obtained from the interacted regression with treatment indicators, a *High Executive Support for AI* indicator, interactions between each treatment indicator and *High Executive Support for AI*, and country fixed effects. Executive support for AI captures respondents' perceptions of how supportive organizational leadership is of AI adoption. Respondents are classified into the high subgroup if they perceive leadership as moderately or strongly supportive of AI adoption, and into the low subgroup if they perceive leadership support as limited. Construction of this subgroup measure is described in Appendix Section E.1. The bottom panel reports p-values for the difference in treatment effects between the high and low subgroups for each treatment arm. These p-values correspond to tests of whether the relevant interaction coefficient is equal to zero in the underlying interacted regression. All specifications are estimated on the pooled sample across countries with country fixed effects, and robust standard errors are reported in parentheses. Outcome construction is described in Section 2.4. (***) $p < 0.01$, ** $p < 0.05$, * $p < 0.10$.

A.4 Gelbach Decompositions

Table A.16: Gelbach Decomposition with Benefits and Risks

	(1)	(2)	(3)	(4)
	AI Adoption	AI Advocacy	Staffing Intentions	Staffing Advocacy
Panel A: AI Labor Treatment				
Total treatment effect	-0.403 (0.061)	-0.435 (0.060)	-0.198 (0.057)	-0.004 (0.057)
Unexplained component	-0.117 (0.055)	-0.080 (0.049)	-0.157 (0.058)	-0.089 (0.059)
Perceived Benefit of AI	-0.161	-0.174	-0.030	0.069
Perceived Risk of AI	-0.125	-0.181	-0.011	0.016
Panel B: AI Productivity Treatment				
Total treatment effect	0.054 (0.057)	0.101 (0.055)	-0.088 (0.058)	-0.062 (0.060)
Unexplained component	0.038 (0.050)	0.081 (0.045)	-0.091 (0.057)	-0.058 (0.059)
Perceived Benefit of AI	0.011	0.010	0.003	-0.004
Perceived Risk of AI	0.006	0.009	-0.000	-0.001

Notes: This table reports Gelbach decompositions of the treatment effects on the four primary outcome indices: *AI Adoption*, *AI Advocacy*, *Staffing Intentions*, and *Staffing Advocacy*. Each panel compares one treatment arm to the control group. Columns correspond to outcome indices. The first row reports the total treatment effect, defined as the coefficient on the treatment indicator from the baseline specification that includes only the treatment indicator and country fixed effects. The second row reports the unexplained component, defined as the coefficient on the treatment indicator from the full specification that adds the AI Benefits and AI Risks indices to the baseline specification. The remaining rows report each mechanism's contribution to the explained component. Following Gelbach (2016), each mechanism contribution is computed as the product of: (i) the coefficient on that mechanism index in the full outcome regression; and (ii) the coefficient on the treatment indicator from an auxiliary regression of that mechanism index on the treatment indicator and country fixed effects. All specifications are estimated separately by treatment arm and outcome, restricting the sample in each panel to respondents in the relevant treatment arm and the control group. Robust standard errors are reported in parentheses for the total treatment effect and unexplained component only. Outcome construction is described in Section 2.4. Construction of the AI Benefits and AI Risks indices is described in the Appendix Section E.2. (***) $p < 0.01$, (**) $p < 0.05$, (*) $p < 0.10$.

Table A.17: Full Regressions with Benefits and Risks

	(1)	(2)	(3)	(4)
	AI Adoption	AI Advocacy	Staffing Intentions	Staffing Advocacy
Panel A: AI Labor Treatment				
AI Labor	-0.117** (0.055)	-0.080 (0.049)	-0.157*** (0.058)	-0.089 (0.059)
Perceived Benefits of AI	0.389*** (0.031)	0.419*** (0.027)	0.073** (0.032)	-0.166*** (0.028)
Perceived Risks of AI	-0.245*** (0.030)	-0.356*** (0.027)	-0.022 (0.031)	0.031 (0.029)
R-squared	0.283	0.410	0.026	0.034
Observations	1259	1259	1259	1259
Panel B: AI Productivity Treatment				
AI Productivity	0.038 (0.050)	0.081* (0.045)	-0.091 (0.057)	-0.058 (0.059)
Perceived Benefits of AI	0.388*** (0.030)	0.376*** (0.028)	0.114*** (0.031)	-0.140*** (0.030)
Perceived Risks of AI	-0.199*** (0.029)	-0.309*** (0.026)	0.005 (0.032)	0.021 (0.031)
R-squared	0.240	0.333	0.034	0.022
Observations	1248	1248	1248	1248

Notes: This table reports the full outcome regressions used in the Gelbach decomposition exercise. Each panel compares one treatment arm to the control group, and columns correspond to the four primary outcome indices: *AI Adoption*, *AI Advocacy*, *Staffing Intentions*, and *Staffing Advocacy*. In each column, the reported specification regresses the indicated outcome index on a treatment indicator, the AI Benefits index, the AI Risks index, and country fixed effects. The first row reports the coefficient on the treatment indicator from this full specification. The next two rows report the coefficients on the AI Benefits and AI Risks indices. All specifications are estimated separately by treatment arm and outcome, restricting the sample in each panel to respondents in the relevant treatment arm and the control group. Robust standard errors are reported in parentheses. Outcome construction is described in Section 2.4. Construction of the AI Benefits and AI Risks indices is described in Appendix Section E.2. (***) $p < 0.01$, (**) $p < 0.05$, (*) $p < 0.10$.

Table A.18: Auxiliary Regressions for Benefits and Risks

	(1)	(2)
	Perceived Benefits of AI	Perceived Risks of AI
Panel A: AI Labor Treatment		
AI Labor	-0.415*** (0.060)	0.508*** (0.058)
R-squared	0.041	0.058
Observations	1259	1259
Panel B: AI Productivity Treatment		
AI Productivity	0.027 (0.058)	-0.030 (0.058)
R-squared	0.006	0.001
Observations	1248	1248

Notes: This table reports the auxiliary regressions used in the Gelbach decomposition exercise, in which the AI Benefits and AI Risks indices are the dependent variables. Each panel compares one treatment arm to the control group. The two columns report results for the AI Benefits and AI Risks indices, respectively. In each column, the reported specification regresses the indicated mechanism index on a treatment indicator and country fixed effects. The reported coefficient is therefore the estimated effect of the treatment on the corresponding mechanism index. All specifications are estimated separately by treatment arm, restricting the sample in each panel to respondents in the relevant treatment arm and the control group. Robust standard errors are reported in parentheses. Construction of the AI Benefits and AI Risks indices is described in Appendix Section E.2. (***) $p < 0.01$, (**) $p < 0.05$, (*) $p < 0.10$.

A.5 Bivariate Associations Within the Control Group

Table A.19: Bivariate Associations Between Pre-Treatment Characteristics and Main Outcomes Within the Control Group

	(1) AI Adoption	(2) AI Advocacy	(3) Staffing Intentions	(4) Staffing Advocacy
AI Adoption Expectations	0.421*** (0.037)	0.388*** (0.037)	0.105*** (0.040)	-0.080** (0.038)
Executive Support for AI	0.451*** (0.038)	0.386*** (0.039)	0.038 (0.036)	-0.104*** (0.034)
AI Impact on Team's Work	0.487*** (0.037)	0.552*** (0.033)	0.118*** (0.035)	-0.126*** (0.034)
Team's Attitude Toward AI	0.515*** (0.040)	0.538*** (0.039)	0.082** (0.039)	-0.099*** (0.034)
Team AI Use Frequency	0.438*** (0.034)	0.392*** (0.035)	0.086** (0.035)	-0.113*** (0.035)
Team Share Using AI Weekly	0.372*** (0.038)	0.373*** (0.037)	0.045 (0.036)	-0.153*** (0.037)
Next Year's Staffing Plan	0.143*** (0.038)	0.164*** (0.039)	0.485*** (0.041)	-0.083** (0.042)
Hiring Decision Authority	0.021 (0.037)	0.021 (0.037)	-0.048 (0.036)	-0.006 (0.037)
Tech Adoption Authority	0.082** (0.039)	0.091** (0.039)	-0.001 (0.039)	-0.074** (0.036)
Management Levels to CEO	-0.027 (0.040)	-0.062 (0.038)	0.019 (0.036)	0.060* (0.037)
Log Firm Revenue	0.116*** (0.044)	0.106** (0.042)	-0.014 (0.042)	0.022 (0.043)
Log Team Budget	0.048 (0.042)	0.034 (0.042)	-0.001 (0.045)	-0.007 (0.044)
Log Firm Employment	0.065* (0.038)	0.041 (0.037)	-0.021 (0.035)	0.070* (0.037)
Team Total Employment	0.060** (0.030)	0.053* (0.030)	0.068* (0.037)	-0.002 (0.035)
Direct Reports	0.042 (0.041)	0.063 (0.040)	0.021 (0.039)	-0.007 (0.039)
Age	0.036 (0.032)	0.027 (0.034)	-0.043 (0.037)	-0.038 (0.037)
Years of Education	0.066* (0.039)	0.083** (0.037)	0.139*** (0.035)	-0.025 (0.037)
Political Views	0.098** (0.039)	0.114*** (0.039)	-0.043 (0.039)	-0.051 (0.034)
Woman+	-0.125* (0.073)	-0.100 (0.073)	0.001 (0.073)	0.151** (0.072)
Received AI Training (Self)	0.525*** (0.073)	0.465*** (0.074)	0.207*** (0.071)	-0.192*** (0.073)
Received AI Training (Team)	0.468*** (0.074)	0.412*** (0.075)	0.242*** (0.073)	-0.189** (0.074)
Experienced Unemployment (Self)	-0.186** (0.080)	-0.327*** (0.080)	-0.105 (0.077)	0.002 (0.077)
Experienced Unemployment (Family)	-0.141* (0.073)	-0.252*** (0.073)	0.002 (0.073)	0.058 (0.074)

Notes: This table reports descriptive bivariate associations between the four primary outcome indices and 23 pre-treatment or stable baseline-related predictors, estimated using only respondents in the control group. Each column corresponds to one of the four primary indices: *AI Adoption*, *AI Advocacy*, *Staffing Intentions*, and *Staffing Advocacy*. In each row, the dependent variable is the indicated outcome index, and the independent variable is a single predictor. Outcome construction is described in Section 2.4. Most predictors are measured pre-treatment. Political views, unemployment experience, and selected firm-level characteristics are collected post-treatment to reduce attrition but are treated as stable and not plausibly affected by treatment. Continuous predictors are entered in standardized units, while woman+, AI training, and unemployment-experience measures are entered as binary indicators. Robust standard errors are reported in parentheses. Sample sizes vary slightly across predictors due to small differences in item nonresponse. (***) $p < 0.01$, (**) $p < 0.05$, (*) $p < 0.10$.

A.6 Attrition

Table A.20: Attrition Regressions

	Pooled	US	UK
AI Productivity	0.023* (0.014)	0.035* (0.020)	0.011 (0.019)
AI Labor	0.015 (0.013)	0.020 (0.018)	0.010 (0.019)
Constant	0.055*** (0.010)	0.050*** (0.011)	0.060*** (0.012)
Joint test p-value	0.206	0.177	0.799
R-squared	0.002	0.003	0.000
Observations	1995	1000	995

Notes: The table reports OLS estimates from regressions where the dependent variable is an indicator equal to one if a randomized respondent is excluded from the final analytical sample. The omitted category is the control arm, so coefficients on *AI Productivity* and *AI Labor* measure differences in attrition rates relative to control. The pooled column includes a country fixed effect. The row labeled *Joint test p-value* reports the p-value for the null hypothesis that the coefficients on *AI Productivity* and *AI Labor* are jointly equal to zero. The sample is restricted to respondents who reach randomization. Robust standard errors are reported in parentheses. (***) $p < 0.01$, ** $p < 0.05$, * $p < 0.10$.

Table A.21: Completed Primary Indices Among Randomized Respondents

Completed Primary Indices	Pooled				United States				United Kingdom			
	Freq.	Control	Labor	Prod	Freq.	Control	Labor	Prod	Freq.	Control	Labor	Prod
0	0	0.0	0.0	0.0	0	0.0	0.0	0.0	0	0.0	0.0	0.0
1	2	0.0	0.2	0.2	0	0.0	0.0	0.0	2	0.0	0.3	0.3
2	13	0.6	0.5	0.8	7	1.0	0.3	0.7	6	0.2	0.7	1.0
3	117	4.9	6.3	6.8	59	4.0	6.6	7.8	58	5.8	6.0	5.8
4	1863	94.5	93.0	92.2	934	95.0	93.1	91.5	929	94.0	93.0	92.9
Frequency	1995	802	603	590	1000	402	303	295	995	400	300	295

Notes: The table reports the distribution of completed primary indices among all respondents who reached randomization. Each row corresponds to the number of completed primary indices. Within each sample and treatment arm, entries report the percentage of randomized respondents with that number of completed indices. The frequency columns report the total number of randomized respondents in each sample with the corresponding number of completed indices. The bottom row reports the total number of randomized respondents in each sample and treatment arm.

Table A.22: Missing Main Outcome Indices by Treatment Arm

Index	Control	Labor	Productivity
AI Adoption Plans	1.2	2.0	1.7
AI Advocacy	0.7	1.3	1.5
Employment Intentions	0.5	1.2	1.5
Employment Advocacy	3.6	3.3	4.2
Frequency	802	603	590

Notes: The table reports the share of randomized respondents in each treatment arm who are missing each main outcome index. Columns are ordered as Control, AI Labor, and AI Productivity. The bottom row reports the number of randomized respondents in each arm.

Table A.23: Treatment Effects on Unsure Responses for Index Components

Component	Pooled		United States		United Kingdom	
	Labor	Prod	Labor	Prod	Labor	Prod
New AI Tech	0.016 (0.012)	0.009 (0.011)	0.014 (0.017)	0.005 (0.016)	0.019 (0.016)	0.013 (0.016)
Accelerate AI	0.009 (0.010)	0.005 (0.009)	0.006 (0.015)	0.007 (0.015)	0.013 (0.013)	0.004 (0.011)
Reduce AI	0.008 (0.012)	0.000 (0.011)	0.006 (0.015)	-0.003 (0.014)	0.010 (0.018)	0.004 (0.017)
Halt AI	0.006 (0.011)	0.014 (0.012)	0.013 (0.014)	0.014 (0.015)	-0.002 (0.017)	0.013 (0.019)
Advocacy 1	-0.014* (0.008)	-0.007 (0.008)	-0.014 (0.010)	-0.007 (0.012)	-0.013 (0.011)	-0.006 (0.012)
Advocacy 2	0.007 (0.007)	0.015* (0.008)	0.007 (0.010)	0.021* (0.012)	0.007 (0.010)	0.008 (0.010)
Advocacy 3	-0.002 (0.007)	0.008 (0.008)	-0.009 (0.010)	0.018 (0.014)	0.005 (0.010)	-0.001 (0.009)
Employment Occupation 1	0.015** (0.006)	0.010* (0.005)	0.008 (0.007)	0.009 (0.008)	0.021** (0.009)	0.011 (0.007)
Employment Occupation 2	0.007 (0.005)	0.009 (0.005)	0.002 (0.007)	0.003 (0.007)	0.011 (0.007)	0.014* (0.008)
Employment Occupation 3	0.014* (0.008)	0.005 (0.008)	0.016 (0.012)	-0.000 (0.010)	0.013 (0.012)	0.010 (0.012)
Hire Advocacy	-0.017* (0.010)	-0.008 (0.011)	0.002 (0.015)	0.003 (0.015)	-0.036*** (0.013)	-0.019 (0.015)
Layoff Advocacy	0.007 (0.008)	0.023** (0.010)	0.007 (0.012)	0.032** (0.015)	0.007 (0.012)	0.014 (0.013)

Notes: Each row reports treatment effects for an indicator equal to one if a randomized respondent selected the unsure response for the corresponding index component item. Columns 1 and 2 report coefficients from pooled regressions that include country fixed effects. Columns 3 and 4 report coefficients from regressions estimated on the United States sample. Columns 5 and 6 report coefficients from regressions estimated on the United Kingdom sample. Robust standard errors are reported in parentheses. *** p<0.01, ** p<0.05, * p<0.10).

A.7 Binary Component Outcomes

Table A.24: Treatment Effects on AI Adoption Item Responses

	(1)	(2)	(3)	(4)
	Begin New AI	Accelerate AI	Reduce AI	Halt New AI
Panel A: Top Category				
AI Labor	-0.065*** (0.017)	-0.074*** (0.018)	0.022** (0.011)	0.003 (0.011)
AI Productivity	0.023 (0.021)	0.023 (0.022)	-0.003 (0.009)	-0.010 (0.010)
Control Mean	0.141	0.165	0.026	0.035
R-squared	0.019	0.017	0.005	0.002
Observations	1799	1831	1803	1803
Panel B: Top Two Categories				
AI Labor	-0.122*** (0.027)	-0.141*** (0.027)	0.032** (0.015)	0.038** (0.016)
AI Productivity	-0.004 (0.028)	-0.001 (0.028)	-0.005 (0.013)	-0.015 (0.014)
Control Mean	0.435	0.464	0.060	0.071
R-squared	0.016	0.018	0.004	0.007
Observations	1799	1831	1803	1803

Notes: The table reports OLS estimates from pooled linear probability models with country fixed effects, where the unit of observation is an individual manager. Each column corresponds to a binary outcome based on one AI adoption survey item. In Panel A, the dependent variable equals one if the respondent selected the top response category only (*Extremely Likely*). In Panel B, the dependent variable equals one if the respondent selected one of the top two response categories (*Very Likely* or *Extremely Likely*). The omitted category is the neutral control video. For the *Reduce AI* and *Halt New AI* items, higher values reflect greater willingness to scale back AI use. These items are reverse-coded when constructing the primary AI adoption index, but they are kept in their original direction here so coefficients can be interpreted directly as percentage-point changes in the probability of selecting those responses. The bottom rows report the control-group mean, R^2 , and the number of observations. Sample sizes vary slightly across columns because of item-level nonresponse. Robust standard errors are reported in parentheses. (*** $p < 0.01$, ** $p < 0.05$, * $p < 0.10$).

Table A.25: Treatment Effects on AI Advocacy Item Responses

	(1)	(2)	(3)
	Advocate to Expand AI	Advocate to Maintain AI	Advocate to Reduce AI
Panel A: Top Category			
AI Labor	-0.094*** (0.021)	-0.001 (0.018)	0.044*** (0.015)
AI Productivity	0.012 (0.024)	-0.003 (0.018)	-0.002 (0.013)
Control Mean	0.232	0.124	0.055
R-squared	0.017	0.005	0.009
Observations	1837	1840	1843
Panel B: Top Two Categories			
AI Labor	-0.176*** (0.027)	0.041 (0.028)	0.080*** (0.023)
AI Productivity	0.032 (0.026)	-0.001 (0.028)	-0.035* (0.020)
Control Mean	0.687	0.561	0.166
R-squared	0.034	0.008	0.016
Observations	1837	1840	1843

Notes: The table reports OLS estimates from pooled linear probability models with country fixed effects, where the unit of observation is an individual manager. Each column corresponds to a binary outcome based on one AI advocacy survey item. In Panel A, the dependent variable equals one if the respondent selected the top response category only (*Strongly Agree*). In Panel B, the dependent variable equals one if the respondent selected one of the top two response categories (*Agree* or *Strongly Agree*). The omitted category is the neutral control video. The *Reduce AI* advocacy item is reverse-coded when constructing the primary AI advocacy index, but it is kept in its original direction here so coefficients can be interpreted directly as percentage-point changes in the probability of advocating to reduce AI use. The *Maintain AI* item is shown here for completeness, but it is not used in the construction of the primary AI advocacy index. The bottom rows report the control-group mean, R^2 , and the number of observations. Sample sizes vary slightly across columns because of item-level nonresponse. Robust standard errors are reported in parentheses. (***) $p < 0.01$, ** $p < 0.05$, * $p < 0.10$.

Table A.26: Treatment Effects on Staffing Intentions: Expansion Responses

	(1)	(2)	(3)
	Top Occupation	Second Highest Occupation	Third Highest Occupation
Panel A: Top Category			
AI Labor	-0.044** (0.018)	-0.007 (0.016)	-0.043*** (0.015)
AI Productivity	-0.019 (0.019)	0.044** (0.018)	0.009 (0.018)
Control Mean	0.142	0.096	0.107
R-squared	0.015	0.017	0.024
Observations	1859	1861	1829
Panel B: Top Two Categories			
AI Labor	-0.070** (0.028)	-0.058** (0.027)	-0.068*** (0.026)
AI Productivity	-0.041 (0.028)	-0.028 (0.027)	-0.036 (0.027)
Control Mean	0.497	0.410	0.353
R-squared	0.006	0.012	0.014
Observations	1859	1861	1829

Notes: The table reports OLS estimates from pooled linear probability models with country fixed effects, where the unit of observation is an individual manager. Columns correspond to the respondent's three most common occupation groups. Each outcome is based on responses to the question asking how the respondent would adjust the current headcount in their team or business unit for each of these occupation groups within the next year, if they had the authority to do so. In Panel A, the dependent variable equals one if the respondent selected the strongest expansion response category (*Increase Headcount a Lot*). In Panel B, the dependent variable equals one if the respondent selected one of the top two expansion response categories (*Increase Headcount Slightly* or *Increase Headcount a Lot*). The omitted category is the neutral control video. The bottom rows report the control-group mean, R^2 , and the number of observations. Sample sizes vary slightly across columns because of item-level nonresponse. Robust standard errors are reported in parentheses. (***) $p < 0.01$, ** $p < 0.05$, * $p < 0.10$.

Table A.27: Treatment Effects on Staffing Intentions: Reduction Responses

	(1)	(2)	(3)
	Top Occupation	Second Highest Occupation	Third Highest Occupation
Panel A: Top Category			
AI Labor	0.008 (0.007)	0.015* (0.009)	0.008 (0.009)
AI Productivity	0.003 (0.007)	0.010 (0.008)	0.001 (0.009)
Control Mean	0.013	0.016	0.023
R-squared	0.003	0.002	0.001
Observations	1859	1861	1829
Panel B: Top Two Categories			
AI Labor	0.041** (0.016)	0.048*** (0.018)	0.036* (0.020)
AI Productivity	0.023 (0.016)	0.048*** (0.018)	0.017 (0.019)
Control Mean	0.073	0.090	0.121
R-squared	0.004	0.006	0.004
Observations	1859	1861	1829

Notes: The table reports OLS estimates from pooled linear probability models with country fixed effects, where the unit of observation is an individual manager. Columns correspond to the respondent's three most common occupation groups. Each outcome is based on responses to the question asking how the respondent would adjust the current headcount in their team or business unit for each of these occupation groups within the next year, if they had the authority to do so. In Panel A, the dependent variable equals one if the respondent selected the strongest reduction response category (*Reduce Headcount a Lot*). In Panel B, the dependent variable equals one if the respondent selected one of the top two reduction response categories (*Reduce Headcount Slightly* or *Reduce Headcount a Lot*). The omitted category is the neutral control video. The bottom rows report the control-group mean, R^2 , and the number of observations. Sample sizes vary slightly across columns because of item-level nonresponse. Robust standard errors are reported in parentheses. (***) $p < 0.01$, (**) $p < 0.05$, (*) $p < 0.10$.

Table A.28: Treatment Effects on Staffing Advocacy Item Responses

	(1)	(2)
	Advocate for Hiring	Advocate for Layoffs
Panel A: Top Category		
AI Labor	-0.005 (0.029)	-0.003 (0.011)
AI Productivity	-0.000 (0.029)	0.016 (0.013)
Control Mean	0.504	0.042
R-squared	0.000	0.002
Observations	1771	1742

Notes: The table reports OLS estimates from pooled linear probability models with country fixed effects, where the unit of observation is an individual manager. Each column corresponds to a binary staffing advocacy outcome. The dependent variable equals one if the respondent selected the top response category for that item (*Advocate For*). The omitted category is the neutral control video. The *Advocate Layoffs* item is reverse-coded when constructing the primary staffing advocacy index, but it is kept in its original direction here, so coefficients are directly interpretable as percentage-point changes in the probability of advocating for layoffs. The bottom rows report the control-group mean, R^2 , and the number of observations. Sample sizes vary slightly across columns because of item-level nonresponse. Robust standard errors are reported in parentheses. (***) $p < 0.01$, ** $p < 0.05$, * $p < 0.10$).

Table A.29: Treatment Effects on Binary Index Outcomes

	(1)	(2)	(3)	(4)
	AI Adoption Index	AI Advocacy Index	Staffing Intentions Index	Staffing Advocacy Index
Panel A: Above Control-Group Median				
AI Labor	-0.169*** (0.026)	-0.152*** (0.024)	-0.094*** (0.027)	-0.012 (0.026)
AI Productivity	0.009 (0.028)	0.038 (0.027)	-0.048* (0.028)	-0.030 (0.027)
Control Mean	0.443	0.354	0.488	0.668
R-squared	0.028	0.029	0.011	0.001
Observations	1863	1863	1863	1863

Notes: The table reports OLS estimates from pooled linear probability models with country fixed effects, where the unit of observation is an individual manager. Each column corresponds to a binary version of one of the four primary standardized indices. For each index, the dependent variable equals one if the respondent's index value is above the control-group 50th-percentile cutoff, where the cutoff is computed within the neutral control video group. Because the indices are discrete and exhibit substantial bunching, this rule does not necessarily split the control group evenly. The omitted category is the neutral control video. The bottom rows report the control-group mean, R^2 , and the number of observations. Sample sizes vary slightly across columns because of item-level nonresponse in the components used to construct the indices. Robust standard errors are reported in parentheses. (***) $p < 0.01$, ** $p < 0.05$, * $p < 0.10$).

A.8 Additional Balance Tables

Table A.30: Baseline Characteristic Balance Across US Survey Arms

	(1)	(2)	(3)	(2)-(1)	(3)-(1)
	Control	AI Productivity	AI Labor Distress	Difference	Difference
Panel A: Respondent Demographics					
Age	42.65 (12.01)	42.35 (12.12)	42.48 (11.23)	-0.31 (0.93)	-0.18 (0.88)
Years Schooling (Derived)	15.90 (2.03)	16.10 (2.16)	16.00 (2.22)	0.20 (0.16)	0.11 (0.16)
Woman+	0.46 (0.50)	0.46 (0.50)	0.46 (0.50)	0.00 (0.04)	-0.00 (0.04)
Full Time	0.85 (0.36)	0.86 (0.35)	0.86 (0.34)	0.01 (0.03)	0.01 (0.03)
Part Time	0.08 (0.27)	0.07 (0.26)	0.07 (0.26)	-0.01 (0.02)	-0.01 (0.02)
Self Empl.	0.06 (0.24)	0.06 (0.25)	0.06 (0.24)	0.00 (0.02)	-0.00 (0.02)
Other Empl. Status	0.01 (0.09)	0.00 (0.06)	0.00 (0.06)	-0.00 (0.01)	-0.00 (0.01)
<i>Politics (Left/Right)</i>	-0.02 (1.23)	0.00 (1.23)	0.02 (1.21)	0.03 (0.09)	0.05 (0.09)
<i>Income (Numeric)</i>	87.08 (51.24)	90.78 (54.35)	88.99 (52.21)	3.70 (4.07)	1.90 (3.94)
<i>Was Unempl. (Self)</i>	0.43 (0.50)	0.44 (0.50)	0.48 (0.50)	0.01 (0.04)	0.05 (0.04)
<i>Was Unempl. (Fam.)</i>	0.57 (0.50)	0.57 (0.50)	0.57 (0.50)	-0.01 (0.04)	-0.00 (0.04)
<i>Unempl. Dur. (Numeric)</i>	6.12 (4.14)	6.17 (4.09)	6.69 (4.28)	0.06 (0.48)	0.57 (0.48)
Panel B: Role & Team Characteristics					
Mng. Levels from Top (Num)	2.84 (1.56)	2.77 (1.67)	2.77 (1.60)	-0.07 (0.12)	-0.07 (0.12)
Direct Reports (Num)	8.36 (5.46)	8.42 (5.78)	8.13 (5.68)	0.06 (0.43)	-0.23 (0.42)
Team Total Empl. (Num)	26.99 (38.93)	23.16 (31.88)	22.30 (33.47)	-3.84 (2.69)	-4.69* (2.73)
<i>Team Budget (Numeric,)</i>	3136.82 (9914.09)	4392.66 (12179.84)	3213.40 (11284.43)	1255.85 (1008.77)	76.58 (933.71)
N	402	295	303	697	705

Notes: The table reports unweighted summary statistics for the U.S. subsample. Columns (1)–(3) report arm means with standard deviations in parentheses for the Control, AI Productivity, and AI Labor Distress treatment conditions. Columns (4) and (5) report differences in means relative to the Control group (AI Productivity–Control and AI Labor Distress–Control), with standard errors in parentheses from unequal-variance two-sided *t*-tests. Age is self-reported in whole years. Years of Schooling maps education categories to years. Woman+ is an indicator equal to one for respondents who don't identify as a Man. Full Time, Part Time, and Self Empl. and Other Empl. (e.g., unemployed, not in the labor force, or student) are indicators for employment status. Politics is a centered 5-point ideology measure from Very liberal to Very conservative. Income converts self-reported income brackets into midpoints (in thousands of dollars, with \$200k+ top-coded as \$250k). Was Unempl. (Self/Fam.) are indicators for whether the respondent (or a family member) has experienced unemployment. Unempl. Dur. converts categorical unemployment duration into midpoints in months. Mng. Levels from Top (Num), Direct Reports (Num), and Team Total Empl. (Num) convert categorical bins into midpoints. Team Budget converts budget brackets into midpoints expressed in thousands of dollars. Italicized variables were collected after the video module to reduce attrition prior to the randomized portion of the survey and because they are stable characteristics unlikely to be affected by treatment. “Unsure” and “Not relevant” responses are treated as missing. (***) $p < 0.01$, ** $p < 0.05$, * $p < 0.10$.

Table A.31: Baseline Characteristic Balance Across UK Survey Arms

	(1)	(2)	(3)	(2)-(1)	(3)-(1)
	Control	AI Productivity	AI Labor Distress	Difference	Difference
Panel A: Respondent Demographics					
Age	41.16 (10.87)	41.61 (10.96)	41.47 (10.98)	0.46 (0.84)	0.31 (0.84)
Years Schooling (Derived)	15.90 (2.19)	15.94 (2.26)	16.07 (2.24)	0.04 (0.17)	0.17 (0.17)
Woman+	0.49 (0.50)	0.49 (0.50)	0.50 (0.50)	0.00 (0.04)	0.01 (0.04)
Full Time	0.84 (0.36)	0.81 (0.39)	0.80 (0.40)	-0.03 (0.03)	-0.05 (0.03)
Part Time	0.09 (0.28)	0.12 (0.32)	0.13 (0.33)	0.03 (0.02)	0.04* (0.02)
Self Empl.	0.07 (0.25)	0.05 (0.22)	0.07 (0.25)	-0.01 (0.02)	0.00 (0.02)
Other Empl. Status	0.01 (0.09)	0.02 (0.13)	0.01 (0.10)	0.01 (0.01)	0.00 (0.01)
<i>Politics (Left/Right)</i>	-0.18 (0.92)	-0.13 (0.87)	-0.28 (0.91)	0.06 (0.07)	-0.09 (0.07)
<i>Income (Numeric)</i>	60.50 (35.48)	61.06 (41.31)	61.09 (40.72)	0.56 (2.99)	0.59 (2.95)
<i>Was Unempl. (Self)</i>	0.28 (0.45)	0.31 (0.46)	0.35 (0.48)	0.04 (0.04)	0.08** (0.04)
<i>Was Unempl. (Fam.)</i>	0.39 (0.49)	0.41 (0.49)	0.41 (0.49)	0.02 (0.04)	0.02 (0.04)
<i>Unempl. Dur. (Numeric)</i>	6.16 (4.45)	5.81 (4.40)	6.02 (3.99)	-0.34 (0.63)	-0.14 (0.58)
Panel B: Role & Team Characteristics					
Mng. Levels from Top (Num)	3.06 (1.75)	3.25 (1.76)	2.96 (1.65)	0.19 (0.13)	-0.11 (0.13)
Direct Reports (Num)	7.25 (5.51)	7.06 (5.68)	6.82 (5.62)	-0.20 (0.43)	-0.44 (0.43)
Team Total Empl. (Num)	25.44 (36.68)	24.33 (36.26)	28.52 (41.40)	-1.11 (2.80)	3.08 (3.01)
<i>Team Budget (Numeric,)</i>	4237.41 (14301.60)	5867.20 (16179.84)	2231.63 (8491.41)	1629.79 (1514.67)	-2005.79* (1085.07)
N	400	295	300	695	700

Notes: The table reports unweighted summary statistics for the U.K. subsample. Columns (1)–(3) report arm means with standard deviations in parentheses for the Control, AI Productivity, and AI Labor Distress treatment conditions. Columns (4) and (5) report differences in means relative to the Control group (AI Productivity–Control and AI Labor Distress–Control), with standard errors in parentheses from unequal-variance two-sided *t*-tests. Age is self-reported in whole years. Years of Schooling maps education categories to years. Woman+ is an indicator equal to one for respondents who don't identify as a Man. Full Time, Part Time, and Self Empl. and Other Empl. (e.g., unemployed, not in the labor force, or student) are indicators for employment status. Politics is a centered 5-point ideology measure from Very liberal to Very conservative. Income converts self-reported income brackets into midpoints (in thousands of dollars, with \$200k+ top-coded as \$250k). Was Unempl. (Self/Fam.) are indicators for whether the respondent (or a family member) has experienced unemployment. Unempl. Dur. converts categorical unemployment duration into midpoints in months. Mng. Levels from Top (Num), Direct Reports (Num), and Team Total Empl. (Num) convert categorical bins into midpoints. Team Budget converts budget brackets into midpoints expressed in thousands of dollars. Italicized variables were collected after the video module to reduce attrition prior to the randomized portion of the survey and because they are stable characteristics unlikely to be affected by treatment. “Unsure” and “Not relevant” responses are treated as missing. (***) $p < 0.01$, ** $p < 0.05$, * $p < 0.10$.

Table A.32: Manager Role and Function Balance Across Countries

	(1) US	(2) UK	(1)-(2) Difference
Panel A: Role			
Supervisor	0.19 (0.39)	0.30 (0.46)	-0.11*** (0.02)
Low-level	0.28 (0.45)	0.30 (0.46)	-0.02 (0.02)
Mid-level	0.32 (0.47)	0.28 (0.45)	0.04* (0.02)
High-level	0.14 (0.35)	0.08 (0.28)	0.06*** (0.01)
Executive	0.06 (0.24)	0.03 (0.18)	0.03*** (0.01)
Other	0.01 (0.10)	0.01 (0.09)	0.00 (0.00)
Panel B: Function			
Sales/BD	0.10 (0.30)	0.08 (0.26)	0.03* (0.01)
Marketing	0.05 (0.21)	0.03 (0.17)	0.02* (0.01)
Customer	0.16 (0.36)	0.16 (0.37)	-0.01 (0.02)
Product	0.04 (0.21)	0.03 (0.16)	0.02** (0.01)
Engineering	0.05 (0.21)	0.05 (0.22)	-0.00 (0.01)
Data/Analytics	0.03 (0.18)	0.05 (0.23)	-0.02** (0.01)
IT/Cyber	0.14 (0.35)	0.08 (0.27)	0.07*** (0.01)
Operations	0.10 (0.30)	0.12 (0.32)	-0.02 (0.01)
Manufacturing	0.05 (0.22)	0.04 (0.19)	0.01 (0.01)
Finance	0.09 (0.29)	0.08 (0.27)	0.01 (0.01)
HR	0.06 (0.24)	0.05 (0.22)	0.01 (0.01)
Legal	0.02 (0.14)	0.03 (0.18)	-0.01* (0.01)
Other	0.11 (0.31)	0.20 (0.40)	-0.09*** (0.02)
N	934	929	1,863

Notes: The sample is restricted to respondents with non-missing values for all four primary outcome indices. The table reports unweighted summary statistics by country. Columns (1) and (2) report country means with standard deviations in parentheses. Column (3) reports the difference in means (US–UK) with standard errors in parentheses from unequal-variance two-sided *t*-tests. Role reports indicator variables for the respondent’s management level. Function reports indicator variables for the respondent’s business function categories. “Unsure” and “Not relevant” responses are treated as missing. (***) $p < 0.01$, ** $p < 0.05$, * $p < 0.10$.

Table A.33: Manager Role and Function Balance Across US Survey Arms

	(1)	(2)	(3)	(2)-(1)	(3)-(1)
	Control	AI Productivity	AI Labor Distress	Difference	Difference
Panel A: Role					
Supervisor	0.20 (0.40)	0.18 (0.38)	0.19 (0.40)	-0.03 (0.03)	-0.01 (0.03)
Low-level	0.31 (0.46)	0.26 (0.44)	0.24 (0.43)	-0.05 (0.03)	-0.07** (0.03)
Mid-level	0.29 (0.46)	0.33 (0.47)	0.34 (0.47)	0.04 (0.04)	0.04 (0.04)
High-level	0.12 (0.33)	0.15 (0.35)	0.15 (0.36)	0.02 (0.03)	0.03 (0.03)
Executive	0.06 (0.25)	0.06 (0.25)	0.07 (0.25)	-0.00 (0.02)	0.00 (0.02)
Other	0.00 (0.07)	0.02 (0.13)	0.01 (0.11)	0.01 (0.01)	0.01 (0.01)
Panel B: Function					
Sales/BD	0.09 (0.29)	0.10 (0.30)	0.10 (0.30)	0.01 (0.02)	0.00 (0.02)
Marketing	0.05 (0.22)	0.05 (0.21)	0.04 (0.20)	-0.00 (0.02)	-0.01 (0.02)
Customer	0.15 (0.36)	0.19 (0.39)	0.13 (0.34)	0.04 (0.03)	-0.02 (0.03)
Product	0.05 (0.22)	0.03 (0.18)	0.04 (0.20)	-0.02 (0.02)	-0.01 (0.02)
Engineering	0.05 (0.21)	0.04 (0.20)	0.05 (0.22)	-0.01 (0.02)	0.00 (0.02)
Data/Analytics	0.03 (0.18)	0.04 (0.20)	0.02 (0.15)	0.01 (0.01)	-0.01 (0.01)
IT/Cyber	0.14 (0.34)	0.13 (0.33)	0.16 (0.37)	-0.01 (0.03)	0.03 (0.03)
Operations	0.10 (0.30)	0.10 (0.30)	0.11 (0.31)	-0.00 (0.02)	0.00 (0.02)
Manufacturing	0.04 (0.21)	0.04 (0.21)	0.07 (0.25)	-0.00 (0.02)	0.02 (0.02)
Finance	0.09 (0.29)	0.09 (0.28)	0.09 (0.29)	-0.00 (0.02)	-0.00 (0.02)
HR	0.07 (0.25)	0.05 (0.23)	0.06 (0.23)	-0.01 (0.02)	-0.01 (0.02)
Legal	0.02 (0.16)	0.02 (0.13)	0.02 (0.13)	-0.01 (0.01)	-0.01 (0.01)
Other	0.11 (0.31)	0.12 (0.32)	0.12 (0.32)	0.01 (0.02)	0.01 (0.02)
N	402	295	303	697	705

Notes: The table reports unweighted summary statistics for the U.S. subsample. Columns (1)–(3) report arm means with standard deviations in parentheses for the Control, AI Productivity, and AI Labor Distress treatment conditions. Columns (4) and (5) report differences in means relative to the Control group (AI Productivity–Control and AI Labor Distress–Control), with standard errors in parentheses from unequal-variance two-sided *t*-tests. Role reports indicator variables for the respondent’s management level. Function reports indicator variables for the respondent’s business function categories. “Unsure” and “Not relevant” responses are treated as missing. (***) $p < 0.01$, ** $p < 0.05$, * $p < 0.10$.

Table A.34: Manager Role and Function Balance Across UK Survey Arms

	(1)	(2)	(3)	(2)-(1)	(3)-(1)
	Control	AI Productivity	AI Labor Distress	Difference	Difference
Panel A: Role					
Supervisor	0.31 (0.46)	0.33 (0.47)	0.28 (0.45)	0.02 (0.04)	-0.02 (0.03)
Low-level	0.28 (0.45)	0.29 (0.46)	0.31 (0.46)	0.01 (0.03)	0.03 (0.04)
Mid-level	0.30 (0.46)	0.26 (0.44)	0.26 (0.44)	-0.03 (0.03)	-0.03 (0.03)
High-level	0.07 (0.26)	0.07 (0.26)	0.10 (0.30)	-0.00 (0.02)	0.02 (0.02)
Executive	0.03 (0.17)	0.04 (0.19)	0.03 (0.18)	0.01 (0.01)	0.00 (0.01)
Other	0.01 (0.10)	0.01 (0.08)	0.01 (0.10)	-0.00 (0.01)	0.00 (0.01)
Panel B: Function					
Sales/BD	0.07 (0.25)	0.07 (0.25)	0.09 (0.29)	0.00 (0.02)	0.03 (0.02)
Marketing	0.04 (0.18)	0.03 (0.16)	0.03 (0.18)	-0.01 (0.01)	-0.00 (0.01)
Customer	0.17 (0.37)	0.19 (0.39)	0.15 (0.36)	0.02 (0.03)	-0.01 (0.03)
Product	0.03 (0.16)	0.03 (0.17)	0.02 (0.13)	0.00 (0.01)	-0.01 (0.01)
Engineering	0.06 (0.24)	0.05 (0.21)	0.03 (0.17)	-0.01 (0.02)	-0.03** (0.02)
Data/Analytics	0.05 (0.22)	0.06 (0.23)	0.05 (0.22)	0.01 (0.02)	0.00 (0.02)
IT/Cyber	0.07 (0.25)	0.07 (0.26)	0.09 (0.29)	0.00 (0.02)	0.03 (0.02)
Operations	0.13 (0.33)	0.10 (0.30)	0.12 (0.32)	-0.03 (0.02)	-0.01 (0.02)
Manufacturing	0.03 (0.16)	0.05 (0.22)	0.03 (0.18)	0.03* (0.02)	0.01 (0.01)
Finance	0.09 (0.29)	0.07 (0.26)	0.08 (0.28)	-0.02 (0.02)	-0.01 (0.02)
HR	0.06 (0.23)	0.06 (0.24)	0.04 (0.20)	0.01 (0.02)	-0.01 (0.02)
Legal	0.03 (0.18)	0.03 (0.18)	0.04 (0.19)	0.00 (0.01)	0.00 (0.01)
Other	0.20 (0.40)	0.18 (0.39)	0.22 (0.42)	-0.01 (0.03)	0.03 (0.03)
N	400	295	300	695	700

Notes: The table reports unweighted summary statistics for the U.K. subsample. Columns (1)–(3) report arm means with standard deviations in parentheses for the Control, AI Productivity, and AI Labor Distress treatment conditions. Columns (4) and (5) report differences in means relative to the Control group (AI Productivity–Control and AI Labor Distress–Control), with standard errors in parentheses from unequal-variance two-sided *t*-tests. Role reports indicator variables for the respondent’s management level. Function reports indicator variables for the respondent’s business function categories. “Unsure” and “Not relevant” responses are treated as missing. (***) $p < 0.01$, ** $p < 0.05$, * $p < 0.10$.

Table A.35: Decision Authority Balance Across Countries

	(1) US	(2) UK	(1)-(2) Difference
Panel A: Tech Authority			
Corp. only	0.12 (0.32)	0.17 (0.37)	-0.05*** (0.02)
Corp. mostly	0.32 (0.47)	0.35 (0.48)	-0.03 (0.02)
Shared	0.27 (0.44)	0.26 (0.44)	0.00 (0.02)
Team mostly	0.20 (0.40)	0.14 (0.35)	0.06*** (0.02)
Team only	0.10 (0.29)	0.07 (0.26)	0.02 (0.01)
Panel B: Hiring Authority			
Corp. only	0.05 (0.22)	0.08 (0.27)	-0.03** (0.01)
Corp. mostly	0.20 (0.40)	0.22 (0.41)	-0.01 (0.02)
Shared	0.34 (0.47)	0.29 (0.46)	0.05** (0.02)
Team mostly	0.26 (0.44)	0.25 (0.43)	0.01 (0.02)
Team only	0.14 (0.35)	0.16 (0.37)	-0.02 (0.02)
N	934	929	1,863

Notes: The sample is restricted to respondents with non-missing values for all four primary outcome indices. The table reports unweighted summary statistics by country. Columns (1) and (2) report country means with standard deviations in parentheses. Column (3) reports the difference in means (US–UK) with standard errors in parentheses from unequal-variance two-sided *t*-tests. Tech Authority reports indicator variables for the respondent’s level of authority over technology adoption decisions. Hiring Authority reports indicator variables for the respondent’s level of authority over hiring decisions. “Unsure” and “Not relevant” responses are treated as missing. (***) $p < 0.01$, (**) $p < 0.05$, (*) $p < 0.10$.

Table A.36: Decision Authority Balance Across US Survey Arms

	(1)	(2)	(3)	(2)-(1)	(3)-(1)
	Control	AI Productivity	AI Labor Distress	Difference	Difference
Panel A: Tech Authority					
Corp. only	0.14 (0.35)	0.14 (0.35)	0.08 (0.27)	-0.00 (0.03)	-0.07*** (0.02)
Corp. mostly	0.30 (0.46)	0.29 (0.46)	0.37 (0.48)	-0.00 (0.04)	0.07* (0.04)
Shared	0.25 (0.43)	0.30 (0.46)	0.27 (0.45)	0.05 (0.03)	0.02 (0.03)
Team mostly	0.21 (0.41)	0.16 (0.37)	0.19 (0.40)	-0.05* (0.03)	-0.02 (0.03)
Team only	0.10 (0.30)	0.10 (0.30)	0.09 (0.29)	0.00 (0.02)	-0.01 (0.02)
Panel B: Hiring Authority					
Corp. only	0.07 (0.25)	0.04 (0.21)	0.06 (0.23)	-0.02 (0.02)	-0.01 (0.02)
Corp. mostly	0.21 (0.41)	0.21 (0.41)	0.18 (0.38)	0.00 (0.03)	-0.04 (0.03)
Shared	0.33 (0.47)	0.31 (0.46)	0.38 (0.49)	-0.02 (0.04)	0.05 (0.04)
Team mostly	0.23 (0.42)	0.30 (0.46)	0.25 (0.43)	0.07** (0.03)	0.02 (0.03)
Team only	0.16 (0.36)	0.14 (0.34)	0.14 (0.34)	-0.02 (0.03)	-0.02 (0.03)
N	402	295	303	697	705

Notes: The table reports unweighted summary statistics for the U.S. subsample. Columns (1)–(3) report arm means with standard deviations in parentheses for the Control, AI Productivity, and AI Labor Distress treatment conditions. Columns (4) and (5) report differences in means relative to the Control group (AI Productivity–Control and AI Labor Distress–Control), with standard errors in parentheses from unequal-variance two-sided t -tests. Tech Authority reports indicator variables for the respondent’s level of authority over technology adoption decisions. Hiring Authority reports indicator variables for the respondent’s level of authority over hiring decisions. “Unsure” and “Not relevant” responses are treated as missing. (***) $p < 0.01$, (**) $p < 0.05$, (*) $p < 0.10$.

Table A.37: Decision Authority Balance Across UK Survey Arms

	(1)	(2)	(3)	(2)-(1)	(3)-(1)
	Control	AI Productivity	AI Labor Distress	Difference	Difference
Panel A: Tech Authority					
Corp. only	0.18 (0.38)	0.15 (0.36)	0.18 (0.38)	-0.02 (0.03)	-0.00 (0.03)
Corp. mostly	0.34 (0.47)	0.36 (0.48)	0.37 (0.48)	0.02 (0.04)	0.03 (0.04)
Shared	0.26 (0.44)	0.30 (0.46)	0.25 (0.43)	0.04 (0.03)	-0.01 (0.03)
Team mostly	0.15 (0.36)	0.11 (0.32)	0.13 (0.34)	-0.04 (0.03)	-0.02 (0.03)
Team only	0.07 (0.26)	0.08 (0.26)	0.07 (0.26)	0.00 (0.02)	0.00 (0.02)
Panel B: Hiring Authority					
Corp. only	0.10 (0.29)	0.07 (0.26)	0.09 (0.29)	-0.02 (0.02)	-0.00 (0.02)
Corp. mostly	0.21 (0.41)	0.26 (0.44)	0.19 (0.39)	0.05 (0.03)	-0.02 (0.03)
Shared	0.30 (0.46)	0.28 (0.45)	0.30 (0.46)	-0.01 (0.03)	-0.00 (0.03)
Team mostly	0.24 (0.43)	0.23 (0.42)	0.26 (0.44)	-0.01 (0.03)	0.02 (0.03)
Team only	0.15 (0.36)	0.16 (0.36)	0.17 (0.37)	0.00 (0.03)	0.01 (0.03)
N	400	295	300	695	700

Notes: The table reports unweighted summary statistics for the U.K. subsample. Columns (1)–(3) report arm means with standard deviations in parentheses for the Control, AI Productivity, and AI Labor Distress treatment conditions. Columns (4) and (5) report differences in means relative to the Control group (AI Productivity–Control and AI Labor Distress–Control), with standard errors in parentheses from unequal-variance two-sided *t*-tests. Tech Authority reports indicator variables for the respondent’s level of authority over technology adoption decisions. Hiring Authority reports indicator variables for the respondent’s level of authority over hiring decisions. “Unsure” and “Not relevant” responses are treated as missing. (***) $p < 0.01$, ** $p < 0.05$, * $p < 0.10$.

Table A.38: Firm Characteristics and Ownership Balance Across Countries

	(1) US	(2) UK	(1)-(2) Difference
Panel A: Firm Characteristics			
Firm Size (Empl., Numeric)	2149.98 (3969.79)	3321.71 (4591.45)	-1171.74*** (200.68)
Firm Age (Numeric)	20.55 (9.64)	23.14 (9.21)	-2.59*** (0.44)
Annual Firm Revenue (Numeric,)	1498.12 (3848.10)	1894.82 (4448.72)	-396.70* (227.13)
Panel B: Ownership			
Private	0.65 (0.48)	0.55 (0.50)	0.09*** (0.02)
Government	0.10 (0.30)	0.22 (0.41)	-0.12*** (0.02)
Public	0.13 (0.34)	0.09 (0.29)	0.04*** (0.01)
Non-profit	0.08 (0.28)	0.11 (0.31)	-0.02* (0.01)
Worker-Owned	0.03 (0.17)	0.01 (0.11)	0.02*** (0.01)
Other	0.01 (0.08)	0.01 (0.10)	-0.00 (0.00)
N	934	929	1,863

Notes: The sample is restricted to respondents with non-missing values for all four primary outcome indices. The table reports unweighted summary statistics by country. Columns (1) and (2) report country means with standard deviations in parentheses. Column (3) reports the difference in means (US–UK) with standard errors in parentheses from unequal-variance two-sided t -tests. Firm Size (Numeric), Firm Age (Numeric), and Firm Revenue (Numeric) convert categorical bins into midpoints. Ownership reports indicator variables for the firm ownership categories. “Unsure” and “Not relevant” responses are treated as missing. (***) $p < 0.01$, ** $p < 0.05$, * $p < 0.10$).

Table A.39: Firm Characteristics and Ownership Balance Across US Survey Arms

	(1)	(2)	(3)	(2)-(1)	(3)-(1)
	Control	AI Productivity	AI Labor Distress	Difference	Difference
Panel A: Firm Characteristics					
Firm Size (Empl., Numeric)	2102.38 (3870.10)	2030.57 (3830.61)	2261.95 (4179.43)	-71.81 (296.74)	159.56 (309.92)
Firm Age (Numeric)	20.04 (9.70)	20.30 (10.04)	20.99 (9.48)	0.26 (0.77)	0.95 (0.73)
Annual Firm Revenue (Numeric,)	1485.26 (3844.36)	1511.35 (3814.17)	1335.90 (3694.62)	26.09 (331.00)	-149.36 (315.10)
Panel B: Ownership					
Private	0.66 (0.47)	0.65 (0.48)	0.63 (0.48)	-0.01 (0.04)	-0.03 (0.04)
Government	0.08 (0.27)	0.09 (0.29)	0.12 (0.32)	0.01 (0.02)	0.04 (0.02)
Public	0.14 (0.35)	0.12 (0.32)	0.13 (0.34)	-0.03 (0.03)	-0.01 (0.03)
Non-profit	0.07 (0.26)	0.11 (0.31)	0.08 (0.28)	0.03 (0.02)	0.01 (0.02)
Worker-Owned	0.03 (0.16)	0.03 (0.17)	0.03 (0.17)	0.00 (0.01)	0.00 (0.01)
Other	0.01 (0.10)	0.00 (0.06)	0.01 (0.08)	-0.01 (0.01)	-0.00 (0.01)
N	402	295	303	697	705

Notes: The table reports unweighted summary statistics for the U.S. subsample. Columns (1)–(3) report arm means with standard deviations in parentheses for the Control, AI Productivity, and AI Labor Distress treatment conditions. Columns (4) and (5) report differences in means relative to the Control group (AI Productivity–Control and AI Labor Distress–Control), with standard errors in parentheses from unequal-variance two-sided *t*-tests. Firm Size (Numeric), Firm Age (Numeric), and Firm Revenue (Numeric) convert categorical bins into midpoints. Ownership reports indicator variables for the firm ownership categories. “Unsure” and “Not relevant” responses are treated as missing. (***) $p < 0.01$, (**) $p < 0.05$, (*) $p < 0.10$.

Table A.40: Firm Characteristics and Ownership Balance Across UK Survey Arms

	(1)	(2)	(3)	(2)-(1)	(3)-(1)
	Control	AI Productivity	AI Labor Distress	Difference	Difference
Panel A: Firm Characteristics					
Firm Size (Empl., Numeric)	3133.02 (4522.76)	3436.84 (4588.16)	3322.79 (4637.70)	303.81 (354.87)	189.77 (354.25)
Firm Age (Numeric)	22.46 (9.47)	23.56 (9.11)	22.84 (9.31)	1.10 (0.72)	0.38 (0.73)
Annual Firm Revenue (Numeric,)	1694.26 (4162.88)	2109.97 (4653.01)	1984.99 (4590.28)	415.71 (426.82)	290.73 (415.51)
Panel B: Ownership					
Private	0.55 (0.50)	0.56 (0.50)	0.54 (0.50)	0.01 (0.04)	-0.00 (0.04)
Government	0.23 (0.42)	0.23 (0.42)	0.22 (0.41)	-0.00 (0.03)	-0.02 (0.03)
Public	0.10 (0.30)	0.08 (0.27)	0.10 (0.30)	-0.02 (0.02)	-0.00 (0.02)
Non-profit	0.10 (0.30)	0.09 (0.29)	0.13 (0.33)	-0.00 (0.02)	0.03 (0.02)
Worker-Owned	0.01 (0.12)	0.02 (0.13)	0.01 (0.10)	0.00 (0.01)	-0.01 (0.01)
Other	0.01 (0.09)	0.02 (0.13)	0.01 (0.08)	0.01 (0.01)	-0.00 (0.01)
N	400	295	300	695	700

Notes: The table reports unweighted summary statistics for the U.K. subsample. Columns (1)–(3) report arm means with standard deviations in parentheses for the Control, AI Productivity, and AI Labor Distress treatment conditions. Columns (4) and (5) report differences in means relative to the Control group (AI Productivity–Control and AI Labor Distress–Control), with standard errors in parentheses from unequal-variance two-sided *t*-tests. Firm Size (Numeric), Firm Age (Numeric), and Firm Revenue (Numeric) convert categorical bins into midpoints. Ownership reports indicator variables for the firm ownership categories. “Unsure” and “Not relevant” responses are treated as missing. (***) $p < 0.01$, (**) $p < 0.05$, (*) $p < 0.10$.

Table A.41: Firm Industry Balance Across Countries

	(1)	(2)	(1)-(2)
	US	UK	Difference
Panel A: Industry			
Health/Social	0.12 (0.32)	0.16 (0.37)	-0.04*** (0.02)
Prof/Sci/Tech	0.12 (0.33)	0.13 (0.34)	-0.00 (0.02)
Education	0.09 (0.29)	0.13 (0.33)	-0.03** (0.01)
Finance/Insurance	0.13 (0.33)	0.09 (0.29)	0.03** (0.01)
Retail Trade	0.09 (0.29)	0.07 (0.26)	0.02 (0.01)
Manufacturing	0.08 (0.27)	0.08 (0.27)	-0.00 (0.01)
Information	0.10 (0.29)	0.05 (0.21)	0.05*** (0.01)
Other Services	0.05 (0.21)	0.05 (0.21)	-0.00 (0.01)
Arts/Ent./Rec.	0.04 (0.20)	0.03 (0.17)	0.01 (0.01)
Construction	0.04 (0.20)	0.03 (0.18)	0.01 (0.01)
Public Admin.	0.01 (0.09)	0.06 (0.23)	-0.05*** (0.01)
Accom/Food Serv.	0.04 (0.18)	0.02 (0.15)	0.01 (0.01)
Transport/Wareh.	0.02 (0.14)	0.03 (0.18)	-0.01* (0.01)
Real Estate	0.02 (0.14)	0.02 (0.13)	0.01 (0.01)
Wholesale Trade	0.02 (0.13)	0.01 (0.09)	0.01 (0.01)
Utilities	0.01 (0.08)	0.02 (0.13)	-0.01** (0.01)
Ag./Forest/Fish	0.01 (0.10)	0.01 (0.10)	-0.00 (0.00)
Company Mgmt.	0.01 (0.12)	0.01 (0.08)	0.01 (0.00)
Admin./Waste	0.01 (0.07)	0.01 (0.09)	-0.00 (0.00)
Mining/Oil/Gas	0.00 (0.05)	0.00 (0.05)	-0.00 (0.00)
N	933	929	1,862

Notes: The sample is restricted to respondents with non-missing values for all four primary outcome indices. The table reports unweighted summary statistics by country. Columns (1) and (2) report country means with standard deviations in parentheses. Column (3) reports the difference in means (US–UK) with standard errors in parentheses from unequal-variance two-sided *t*-tests. Industry reports indicator variables for the firm industry categories. “Unsure” and “Not relevant” responses are treated as missing. (***) $p < 0.01$, (**) $p < 0.05$, (*) $p < 0.10$.

Table A.42: Firm Industry Balance Across US Survey Arms

	(1)	(2)	(3)	(2)-(1)	(3)-(1)
	Control	AI Productivity	AI Labor Distress	Difference	Difference
Panel A: Industry					
Health/Social	0.09 (0.29)	0.15 (0.36)	0.11 (0.32)	0.06** (0.03)	0.02 (0.02)
Prof/Sci/Tech	0.12 (0.33)	0.13 (0.33)	0.13 (0.33)	0.00 (0.03)	0.00 (0.03)
Education	0.10 (0.30)	0.09 (0.28)	0.09 (0.29)	-0.01 (0.02)	-0.00 (0.02)
Finance/Insurance	0.14 (0.35)	0.13 (0.33)	0.10 (0.29)	-0.02 (0.03)	-0.05* (0.02)
Retail Trade	0.08 (0.27)	0.10 (0.30)	0.12 (0.32)	0.02 (0.02)	0.04* (0.02)
Manufacturing	0.07 (0.26)	0.06 (0.24)	0.09 (0.28)	-0.01 (0.02)	0.01 (0.02)
Information	0.10 (0.31)	0.09 (0.28)	0.10 (0.29)	-0.02 (0.02)	-0.01 (0.02)
Other Services	0.04 (0.20)	0.05 (0.22)	0.05 (0.21)	0.01 (0.02)	0.01 (0.02)
Arts/Ent./Rec.	0.05 (0.23)	0.05 (0.21)	0.03 (0.17)	-0.01 (0.02)	-0.03* (0.01)
Construction	0.04 (0.20)	0.03 (0.17)	0.06 (0.24)	-0.01 (0.01)	0.02 (0.02)
Public Admin.	0.01 (0.10)	0.01 (0.08)	0.01 (0.08)	-0.00 (0.01)	-0.00 (0.01)
Accom/Food Serv.	0.04 (0.20)	0.05 (0.22)	0.03 (0.17)	0.01 (0.02)	-0.01 (0.01)
Transport/Wareh.	0.02 (0.16)	0.02 (0.15)	0.02 (0.14)	-0.00 (0.01)	-0.01 (0.01)
Real Estate	0.02 (0.15)	0.01 (0.12)	0.03 (0.16)	-0.01 (0.01)	0.00 (0.01)
Wholesale Trade	0.02 (0.14)	0.01 (0.10)	0.02 (0.13)	-0.01 (0.01)	-0.00 (0.01)
Utilities	0.01 (0.10)	0.00 (0.06)	0.00 (0.06)	-0.01 (0.01)	-0.01 (0.01)
Ag./Forest/Fish	0.01 (0.10)	0.01 (0.10)	0.01 (0.10)	0.00 (0.01)	-0.00 (0.01)
Company Mgmt.	0.01 (0.11)	0.01 (0.08)	0.02 (0.14)	-0.01 (0.01)	0.01 (0.01)
Admin./Waste	0.01 (0.10)	0.00 (0.06)	0.00 (0.06)	-0.01 (0.01)	-0.01 (0.01)
Mining/Oil/Gas	0.00 (0.00)	0.00 (0.06)	0.00 (0.06)	0.00 (0.00)	0.00 (0.00)
N	402	294	303	696	705

Notes: The table reports unweighted summary statistics for the U.S. subsample. Columns (1)–(3) report arm means with standard deviations in parentheses for the Control, AI Productivity, and AI Labor Distress treatment conditions. Columns (4) and (5) report differences in means relative to the Control group (AI Productivity–Control and AI Labor Distress–Control), with standard errors in parentheses from unequal-variance two-sided *t*-tests. Industry reports indicator variables for the firm industry categories. “Unsure” and “Not relevant” responses are treated as missing. (***) $p < 0.01$, ** $p < 0.05$, * $p < 0.10$.

Table A.43: Firm Industry Balance Across UK Survey Arms

	(1)	(2)	(3)	(2)-(1)	(3)-(1)
	Control	AI Productivity	AI Labor Distress	Difference	Difference
Panel A: Industry					
Health/Social	0.14 (0.34)	0.20 (0.40)	0.17 (0.37)	0.07** (0.03)	0.03 (0.03)
Prof/Sci/Tech	0.10 (0.30)	0.14 (0.35)	0.13 (0.34)	0.04 (0.03)	0.03 (0.02)
Education	0.12 (0.33)	0.12 (0.32)	0.14 (0.35)	-0.01 (0.03)	0.01 (0.03)
Finance/Insurance	0.10 (0.30)	0.08 (0.27)	0.09 (0.29)	-0.02 (0.02)	-0.01 (0.02)
Retail Trade	0.07 (0.25)	0.07 (0.25)	0.09 (0.28)	0.00 (0.02)	0.02 (0.02)
Manufacturing	0.05 (0.22)	0.10 (0.30)	0.07 (0.26)	0.05** (0.02)	0.02 (0.02)
Information	0.06 (0.23)	0.03 (0.16)	0.05 (0.22)	-0.03** (0.02)	-0.01 (0.02)
Other Services	0.07 (0.25)	0.05 (0.21)	0.03 (0.16)	-0.02 (0.02)	-0.04*** (0.02)
Arts/Ent./Rec.	0.04 (0.20)	0.02 (0.14)	0.03 (0.18)	-0.02* (0.01)	-0.01 (0.01)
Construction	0.03 (0.16)	0.02 (0.15)	0.04 (0.20)	-0.00 (0.01)	0.01 (0.01)
Public Admin.	0.07 (0.26)	0.05 (0.21)	0.06 (0.23)	-0.02 (0.02)	-0.01 (0.02)
Accom/Food Serv.	0.03 (0.16)	0.03 (0.16)	0.02 (0.15)	0.00 (0.01)	-0.00 (0.01)
Transport/Wareh.	0.04 (0.20)	0.03 (0.17)	0.02 (0.15)	-0.01 (0.01)	-0.02 (0.01)
Real Estate	0.03 (0.16)	0.01 (0.12)	0.01 (0.08)	-0.01 (0.01)	-0.02** (0.01)
Wholesale Trade	0.02 (0.13)	0.00 (0.00)	0.01 (0.08)	-0.02*** (0.01)	-0.01 (0.01)
Utilities	0.02 (0.15)	0.01 (0.10)	0.02 (0.13)	-0.01 (0.01)	-0.01 (0.01)
Ag./Forest/Fish	0.01 (0.10)	0.01 (0.10)	0.01 (0.11)	0.00 (0.01)	0.00 (0.01)
Company Mgmt.	0.00 (0.05)	0.01 (0.12)	0.01 (0.08)	0.01 (0.01)	0.00 (0.01)
Admin./Waste	0.01 (0.09)	0.01 (0.10)	0.01 (0.08)	0.00 (0.01)	-0.00 (0.01)
Mining/Oil/Gas	0.00 (0.00)	0.01 (0.08)	0.00 (0.00)	0.01 (0.00)	0.00 (0.00)
N	400	295	300	695	700

Notes: The table reports unweighted summary statistics for the U.K. subsample. Columns (1)–(3) report arm means with standard deviations in parentheses for the Control, AI Productivity, and AI Labor Distress treatment conditions. Columns (4) and (5) report differences in means relative to the Control group (AI Productivity–Control and AI Labor Distress–Control), with standard errors in parentheses from unequal-variance two-sided *t*-tests. Industry reports indicator variables for the firm industry categories. “Unsure” and “Not relevant” responses are treated as missing. (***) $p < 0.01$, (**) $p < 0.05$, (*) $p < 0.10$.

A.9 Additional Results on Index Components, Control Sets, and Heterogeneity Dimensions

Table A.44: Control Group Raw Means: AI Outcome Index Components

	(1)	(2)	(3)
	Pooled	US	UK
Panel A: AI Adoption Components			
New AI Tech	3.14 (1.22)	3.24 (1.24)	3.04 (1.20)
Accelerate AI	3.26 (1.19)	3.26 (1.22)	3.25 (1.17)
Reduce AI	1.45 (0.93)	1.49 (0.98)	1.40 (0.88)
Halt AI	1.37 (0.97)	1.41 (1.01)	1.33 (0.92)
Panel B: AI Advocacy Components			
For AI Expansion	3.74 (1.06)	3.72 (1.11)	3.75 (1.00)
For AI Maintain	3.48 (0.99)	3.58 (0.96)	3.38 (1.01)
For AI Reduction	2.16 (1.18)	2.14 (1.20)	2.18 (1.15)
N	758	382	376

Notes: The sample is restricted to respondents with non-missing values for all four primary outcome indices. The table reports unweighted summary statistics for the Control condition only, pooling respondents across both countries in Column (1) and separately by country in Columns (2) and (3). Each row reports the component mean with the standard deviation in parentheses. Panel A reports the four AI adoption component items. Panel B reports the AI advocacy component items. The item measuring advocacy for maintaining existing AI use is reported for completeness but is not included in the construction of the main AI Advocacy Index. “Unsure” and “Not relevant” responses are treated as missing.

Table A.45: Control Group Raw Means: Staffing Outcome Index Components

	(1)	(2)	(3)
	Pooled	US	UK
Panel A: Staffing Intention Components			
Post Staffing for Occ 1	3.55 (0.85)	3.64 (0.91)	3.46 (0.78)
Post Staffing for Occ 2	3.40 (0.82)	3.52 (0.88)	3.28 (0.74)
Post Staffing for Occ 3	3.32 (0.87)	3.44 (0.93)	3.20 (0.79)
Panel B: Staffing Advocacy Components			
Hiring Advocacy	2.37 (0.70)	2.35 (0.75)	2.40 (0.66)
Layoff Advocacy	1.28 (0.53)	1.26 (0.54)	1.29 (0.53)
N	758	382	376

Notes: The sample is restricted to respondents with non-missing values for all four primary outcome indices. The table reports unweighted summary statistics for the Control condition only, pooling respondents across both countries in Column (1) and separately by country in Columns (2) and (3). Each row reports the component mean with the standard deviation in parentheses. Panel A reports the staffing intention component items. Panel B reports the staffing advocacy component items. “Unsure” and “Not relevant” responses are treated as missing.

Table A.46: Treatment Effects on AI Adoption Index and Components

	(1)	(2)	(3)	(4)	(5)
	AI Adoption Index	New AI Tech	Accelerate AI	Reduce AI	Halt AI
Panel A: Pooled (Country FE)					
AI Labor	-0.426*** (0.060)	-0.339*** (0.068)	-0.440*** (0.068)	0.312*** (0.059)	0.239*** (0.059)
AI Productivity	0.033 (0.057)	0.011 (0.070)	0.001 (0.069)	-0.053 (0.052)	-0.060 (0.052)
Constant	0.025 (0.044)	3.232*** (0.054)	3.306*** (0.053)	1.480*** (0.042)	1.393*** (0.043)
Panel B: US					
AI Labor	-0.406*** (0.085)	-0.343*** (0.097)	-0.340*** (0.098)	0.332*** (0.086)	0.251*** (0.087)
AI Productivity	0.076 (0.083)	-0.024 (0.102)	0.041 (0.103)	-0.117 (0.076)	-0.127* (0.074)
Constant	0.006 (0.052)	3.243*** (0.064)	3.264*** (0.063)	1.493*** (0.051)	1.408*** (0.052)
Panel C: UK					
AI Labor	-0.447*** (0.086)	-0.336*** (0.095)	-0.540*** (0.095)	0.292*** (0.082)	0.227*** (0.081)
AI Productivity	-0.009 (0.077)	0.046 (0.097)	-0.040 (0.093)	0.013 (0.072)	0.009 (0.073)
Constant	-0.007 (0.050)	3.041*** (0.063)	3.253*** (0.061)	1.401*** (0.047)	1.332*** (0.048)
UK Diff, Labor	-0.040	0.007	-0.199	-0.040	-0.024
UK Diff, Prod.	-0.086	0.070	-0.081	0.131	0.136
R-squared	0.037	0.022	0.029	0.026	0.017
Observations	1863	1799	1831	1803	1803

Notes: The table reports OLS estimates from equation (1) with no controls, where the unit of observation is an individual manager. Column (1) reports treatment effects on the standardized *AI Adoption* index described in Section 2.4. Columns (2)–(5) report treatment effects on the raw component items that make up the index, measured on their original 1–5 scales and coded in their original direction (i.e., not reverse-coded). The omitted category is the neutral control video. Coefficients on *AI Productivity* and *AI Labor* therefore compare each treatment condition to the control group. Panel A reports pooled estimates with country fixed effects. Panels B and C report estimates separately for the U.S. and U.K. subsamples. The bottom rows report the difference in treatment effects between the U.K. and U.S. for each arm, along with R^2 and the number of observations. The sample is restricted to respondents with non-missing values for all four primary outcome indices. Component-item regressions may have fewer observations than the index regressions because indices are defined when at least half of their component items are non-missing. Robust standard errors are reported in parentheses. (***) $p < 0.01$, ** $p < 0.05$, * $p < 0.10$.

Table A.47: Treatment Effects on AI Advocacy Index and Components

	(1)	(2)	(3)	(4)
	AI Advocacy Index	For AI Expansion	For AI Maintain	For AI Reduction
Panel A: Pooled (Country FE)				
AI Labor	-0.461*** (0.059)	-0.445*** (0.063)	0.029 (0.056)	0.414*** (0.068)
AI Productivity	0.112** (0.054)	0.068 (0.059)	-0.025 (0.056)	-0.141** (0.065)
Constant	-0.005 (0.044)	3.738*** (0.048)	3.574*** (0.042)	2.166*** (0.052)
Panel B: US				
AI Labor	-0.491*** (0.085)	-0.437*** (0.093)	0.011 (0.078)	0.487*** (0.098)
AI Productivity	0.126 (0.080)	0.117 (0.088)	-0.022 (0.080)	-0.113 (0.098)
Constant	-0.000 (0.052)	3.722*** (0.057)	3.579*** (0.049)	2.136*** (0.062)
Panel C: UK				
AI Labor	-0.431*** (0.083)	-0.452*** (0.086)	0.048 (0.080)	0.340*** (0.093)
AI Productivity	0.098 (0.074)	0.019 (0.078)	-0.027 (0.080)	-0.169* (0.088)
Constant	0.000 (0.051)	3.752*** (0.052)	3.382*** (0.052)	2.180*** (0.060)
UK Diff, Labor	0.060	-0.015	0.037	-0.147
UK Diff, Prod.	-0.028	-0.098	-0.006	-0.056
R-squared	0.052	0.040	0.009	0.036
Observations	1863	1837	1840	1843

Notes: The table reports OLS estimates from equation (1) with no controls, where the unit of observation is an individual manager. Column (1) reports treatment effects on the standardized *AI Advocacy* index described in Section 2.4. Columns (2)–(4) report treatment effects on the underlying raw advocacy items, measured on their original 1–5 scales and coded in their original direction (i.e., not reverse-coded). The *AI Advocacy* index is constructed from advocacy for expanding AI use and advocacy for reducing AI use. Advocacy for maintaining current AI use is reported here, but is not used to construct the *AI Advocacy Index*. The omitted category is the neutral control video. Coefficients on *AI Productivity* and *AI Labor* therefore compare each treatment condition to the control group. Panel A reports pooled estimates with country fixed effects. Panels B and C report estimates separately for the U.S. and U.K. subsamples. The bottom rows report the difference in treatment effects between the U.K. and U.S. for each arm, along with R^2 and the number of observations. The sample is restricted to respondents with non-missing values for all four primary outcome indices. Component-item regressions may have fewer observations than the index regressions because indices are defined when at least half of their component items are non-missing. Robust standard errors are reported in parentheses. (***) $p < 0.01$, ** $p < 0.05$, * $p < 0.10$.

Table A.48: Treatment Effects on Staffing Intentions Index and Components

	(1)	(2)	(3)	(4)
	Staffing Intentions Index	Post Staffing for Occ 1	Post Staffing for Occ 2	Post Staffing for Occ 3
Panel A: Pooled (Country FE)				
AI Labor	-0.206*** (0.055)	-0.163*** (0.048)	-0.127*** (0.048)	-0.156*** (0.048)
AI Productivity	-0.084 (0.056)	-0.086* (0.048)	-0.043 (0.050)	-0.045 (0.050)
Constant	0.105** (0.045)	3.610*** (0.038)	3.490*** (0.037)	3.418*** (0.039)
Panel B: US				
AI Labor	-0.331*** (0.084)	-0.285*** (0.072)	-0.210*** (0.071)	-0.210*** (0.073)
AI Productivity	-0.078 (0.085)	-0.077 (0.072)	-0.072 (0.073)	-0.046 (0.074)
Constant	0.141** (0.056)	3.644*** (0.047)	3.524*** (0.045)	3.435*** (0.048)
Panel C: UK				
AI Labor	-0.080 (0.070)	-0.039 (0.064)	-0.044 (0.063)	-0.100 (0.062)
AI Productivity	-0.088 (0.073)	-0.094 (0.064)	-0.013 (0.068)	-0.044 (0.067)
Constant	-0.143*** (0.045)	3.463*** (0.040)	3.277*** (0.038)	3.196*** (0.041)
UK Diff, Labor	0.251**	0.246**	0.167*	0.111
UK Diff, Prod.	-0.010	-0.017	0.059	0.002
R-squared	0.022	0.015	0.016	0.020
Observations	1863	1859	1861	1829

Notes: The table reports OLS estimates from equation (1) with no controls, where the unit of observation is an individual manager. Column (1) reports treatment effects on the standardized *Staffing Intentions* index described in Section 2.4. Columns (2)–(4) report treatment effects on the raw, underlying component items measuring expected headcount changes over the next year for the respondent’s first, second, and third most common occupation groups on their team (each measured on a 1–5 scale from major reduction to major increase, where higher values indicate stronger hiring intentions). The omitted category is the neutral control video. Coefficients on *AI Productivity* and *AI Labor* therefore compare each treatment condition to the control group. Panel A reports pooled estimates with country fixed effects. Panels B and C report estimates separately for the U.S. and U.K. subsamples. The bottom rows report the difference in treatment effects between the U.K. and U.S. for each arm, along with R^2 and the number of observations. The sample is restricted to respondents with non-missing values for all four primary outcome indices. Component-item regressions may have fewer observations than the index regressions because indices are defined when at least half of their component items are non-missing. Robust standard errors are reported in parentheses. (***) $p < 0.01$, (**) $p < 0.05$, (*) $p < 0.10$.

Table A.49: Treatment Effects on Staffing Advocacy Index and Components

	(1)	(2)	(3)
	Staffing Advocacy Index	Hiring Advocacy	Layoff Advocacy
Panel A: Pooled (Country FE)			
AI Labor	-0.001 (0.056)	-0.028 (0.041)	-0.018 (0.030)
AI Productivity	-0.067 (0.059)	-0.020 (0.041)	0.039 (0.033)
Constant	-0.007 (0.044)	2.367*** (0.032)	1.280*** (0.024)
Panel B: US			
AI Labor	0.002 (0.078)	0.011 (0.059)	0.016 (0.043)
AI Productivity	-0.078 (0.084)	-0.005 (0.059)	0.066 (0.048)
Constant	-0.004 (0.053)	2.351*** (0.039)	1.262*** (0.029)
Panel C: UK			
AI Labor	-0.005 (0.079)	-0.068 (0.057)	-0.052 (0.043)
AI Productivity	-0.056 (0.082)	-0.037 (0.058)	0.012 (0.044)
Constant	0.004 (0.050)	2.397*** (0.035)	1.293*** (0.028)
UK Diff, Labor	-0.007	-0.078	-0.068
UK Diff, Prod.	0.023	-0.032	-0.054
R-squared	0.001	0.001	0.003
Observations	1863	1771	1742

Notes: The table reports OLS estimates from equation (1) with no controls, where the unit of observation is an individual manager. Column (1) reports treatment effects on the standardized *Staffing Advocacy* index described in Section 2.4. Columns (2)–(3) report treatment effects on the raw component items that make up the index: advocacy for hiring additional employees and advocacy for layoffs or workforce reductions. Both components are measured in their original direction (not reverse-coded). The omitted category is the neutral control video. Coefficients on *AI Productivity* and *AI Labor* therefore compare each treatment condition to the control group. Panel A reports pooled estimates with country fixed effects. Panels B and C report estimates separately for the U.S. and U.K. subsamples. The bottom rows report the difference in treatment effects between the U.K. and U.S. for each arm, along with R^2 and the number of observations. The sample is restricted to respondents with non-missing values for all four primary outcome indices. Component-item regressions may have fewer observations than the index regressions because indices are defined when at least half of their component items are non-missing. Robust standard errors are reported in parentheses. (***) $p < 0.01$, (**) $p < 0.05$, (*) $p < 0.10$.

Table A.50: Treatment Effects on AI Adoption Index (Robustness to Controls)

	(1)	(2)	(3)	(4)	(5)
	Baseline	+ Indiv	+ Firm	+ Beliefs	+ Policies
Panel A: Pooled (Country FE)					
AI Labor	-0.443*** (0.076)	-0.433*** (0.077)	-0.439*** (0.077)	-0.438*** (0.066)	-0.422*** (0.064)
AI Productivity	0.017 (0.074)	0.016 (0.073)	0.010 (0.073)	-0.001 (0.063)	0.006 (0.061)
Constant	0.063 (0.054)	0.132* (0.076)	0.154** (0.076)	0.157** (0.078)	0.195** (0.078)
Panel B: US					
AI Labor	-0.438*** (0.102)	-0.432*** (0.100)	-0.435*** (0.102)	-0.489*** (0.092)	-0.470*** (0.089)
AI Productivity	0.045 (0.104)	0.029 (0.103)	0.030 (0.103)	-0.009 (0.093)	0.002 (0.089)
Constant	0.053 (0.063)	0.177* (0.094)	0.202** (0.094)	0.206* (0.112)	0.263** (0.111)
Panel C: UK					
AI Labor	-0.449*** (0.117)	-0.432*** (0.118)	-0.438*** (0.119)	-0.368*** (0.092)	-0.361*** (0.092)
AI Productivity	-0.017 (0.104)	-0.016 (0.103)	-0.023 (0.103)	-0.008 (0.083)	-0.012 (0.080)
Constant	0.141** (0.064)	0.175** (0.084)	0.170** (0.086)	0.246*** (0.086)	0.259*** (0.085)
Individual Characteristics		Y	Y	Y	Y
Firm Characteristics			Y	Y	Y
Beliefs / Training				Y	Y
Firm Policies / Authority					Y
UK Diff, Labor	-0.011	0.003	0.007	0.125	0.138
UK Diff, Prod.	-0.061	-0.043	-0.048	-0.008	-0.025
R-squared	0.040	0.064	0.072	0.321	0.368
Observations	1100	1100	1100	1100	1100

Notes: The table reports OLS estimates from equation (1), where the unit of observation is an individual manager. Panel A reports pooled estimates with country fixed effects, while Panels B and C report estimates separately for the U.S. and U.K., respectively. The dependent variable is the standardized *AI adoption* index described in Section 2.4. The omitted category is the neutral control video, with coefficients on *AI Productivity* and *AI Labor* comparing each treatment condition to the control group. Columns add controls sequentially. Column (1) includes no controls. Column (2) adds individual characteristics (age, years of schooling, log income; all demeaned), a gender indicator, political ideology (5-point scale; standardized), and indicators for respondent and family unemployment experience. Column (3) additionally adds demeaned firm and team characteristics: log firm revenue, log firm employment, log team budget, team size, and number of direct reports. Column (4) additionally adds pre-treatment beliefs and training: respondent and team AI sentiment (5-point Likert scale, demeaned), a categorical pre-treatment belief about whether AI augments/automates work, and indicators for respondent and team AI training. Column (5) additionally adds baseline organizational context: a standardized pre-treatment staffing-plan index (constructed from expected headcount changes for up to three occupation groups), technology and hiring-decision authority (5-point Likert scale, standardized), perceived executive support for AI (4-point Likert scale, standardized), and expected AI adoption next year (5-point Likert scale, standardized). The sample is restricted to respondents with non-missing values for all four primary outcome indices and non-missing values for the full control set. The bottom panel indicates which control groups are included in each column and reports the difference in treatment effects between the U.K. and U.S. (U.K.–U.S.), estimated from a pooled interacted specification that allows treatment effects to vary by country. The reported R^2 and number of observations in the bottom panel also come from this pooled interacted specification and therefore do not correspond to the separate country-specific regressions reported in Panels B and C. Robust standard errors are reported in parentheses. (***) $p < 0.01$, (**) $p < 0.05$, (*) $p < 0.10$.

Table A.51: Treatment Effects on AI Advocacy Index (Robustness to Controls)

	(1)	(2)	(3)	(4)	(5)
	Baseline	+ Indiv	+ Firm	+ Beliefs	+ Policies
Panel A: Pooled (Country FE)					
AI Labor	-0.443*** (0.075)	-0.434*** (0.075)	-0.439*** (0.076)	-0.447*** (0.062)	-0.433*** (0.062)
AI Productivity	0.077 (0.071)	0.076 (0.070)	0.069 (0.070)	0.048 (0.059)	0.057 (0.058)
Constant	0.087 (0.054)	0.167** (0.073)	0.176** (0.074)	0.245*** (0.075)	0.256*** (0.074)
Panel B: US					
AI Labor	-0.454*** (0.100)	-0.453*** (0.100)	-0.453*** (0.101)	-0.520*** (0.087)	-0.502*** (0.086)
AI Productivity	0.063 (0.099)	0.049 (0.099)	0.046 (0.099)	-0.008 (0.084)	0.010 (0.080)
Constant	0.095 (0.064)	0.168* (0.091)	0.173* (0.093)	0.235** (0.105)	0.284*** (0.103)
Panel C: UK					
AI Labor	-0.428*** (0.115)	-0.413*** (0.115)	-0.423*** (0.115)	-0.345*** (0.088)	-0.333*** (0.089)
AI Productivity	0.095 (0.100)	0.092 (0.099)	0.079 (0.100)	0.109 (0.085)	0.119 (0.084)
Constant	0.130* (0.067)	0.259*** (0.083)	0.252*** (0.083)	0.371*** (0.084)	0.351*** (0.084)
Individual Characteristics		Y	Y	Y	Y
Firm Characteristics			Y	Y	Y
Beliefs / Training				Y	Y
Firm Policies / Authority					Y
UK Diff, Labor	0.026	0.036	0.034	0.173	0.183
UK Diff, Prod.	0.032	0.041	0.038	0.112	0.104
R-squared	0.047	0.080	0.087	0.366	0.392
Observations	1100	1100	1100	1100	1100

Notes: The table reports OLS estimates from equation (1), where the unit of observation is an individual manager. Panel A reports pooled estimates with country fixed effects, while Panels B and C report estimates separately for the U.S. and U.K., respectively. The dependent variable is the standardized *AI Advocacy* index described in Section 2.4. The omitted category is the neutral control video, with coefficients on *AI Productivity* and *AI Labor* comparing each treatment condition to the control group. Columns add controls sequentially. Column (1) includes no controls. Column (2) adds individual characteristics (age, years of schooling, log income; all demeaned), a gender indicator, political ideology (5-point scale; standardized), and indicators for respondent and family unemployment experience. Column (3) additionally adds demeaned firm and team characteristics: log firm revenue, log firm employment, log team budget, team size, and number of direct reports. Column (4) additionally adds pre-treatment beliefs and training: respondent and team AI sentiment (5-point Likert scale, demeaned), a categorical pre-treatment belief about whether AI augments/automates work, and indicators for respondent and team AI training. Column (5) additionally adds baseline organizational context: a standardized pre-treatment staffing-plan index (constructed from expected headcount changes for up to three occupation groups), technology and hiring-decision authority (5-point Likert scale, standardized), perceived executive support for AI (4-point Likert scale, standardized), and expected AI adoption next year (5-point Likert scale, standardized). The sample is restricted to respondents with non-missing values for all four primary outcome indices and non-missing values for the full control set. The bottom panel indicates which control groups are included in each column and reports the difference in treatment effects between the U.K. and U.S. (U.K.–U.S.), estimated from a pooled interacted specification that allows treatment effects to vary by country. The reported R^2 and number of observations in the bottom panel also come from this pooled interacted specification and therefore do not correspond to the separate country-specific regressions reported in Panels B and C. Robust standard errors are reported in parentheses. (***) $p < 0.01$, (**) $p < 0.05$, (*) $p < 0.10$.

Table A.52: Treatment Effects on Staffing Intentions Index (Robustness to Controls)

	(1) Baseline	(2) + Indiv	(3) + Firm	(4) + Beliefs	(5) + Policies
Panel A: Pooled (Country FE)					
AI Labor	-0.251*** (0.075)	-0.269*** (0.075)	-0.269*** (0.075)	-0.270*** (0.073)	-0.309*** (0.067)
AI Productivity	-0.038 (0.077)	-0.046 (0.076)	-0.050 (0.076)	-0.047 (0.076)	-0.003 (0.066)
Constant	0.143** (0.060)	0.165** (0.076)	0.154** (0.076)	0.017 (0.091)	0.056 (0.080)
Panel B: US					
AI Labor	-0.359*** (0.106)	-0.372*** (0.107)	-0.366*** (0.107)	-0.375*** (0.104)	-0.342*** (0.092)
AI Productivity	-0.086 (0.110)	-0.108 (0.107)	-0.118 (0.107)	-0.110 (0.107)	-0.045 (0.095)
Constant	0.191*** (0.073)	0.233** (0.097)	0.208** (0.097)	0.020 (0.124)	0.062 (0.108)
Panel C: UK					
AI Labor	-0.104 (0.102)	-0.173* (0.102)	-0.173* (0.102)	-0.185* (0.102)	-0.287*** (0.098)
AI Productivity	0.021 (0.105)	0.020 (0.102)	0.019 (0.103)	-0.008 (0.105)	-0.005 (0.094)
Constant	-0.166*** (0.060)	-0.189*** (0.069)	-0.191*** (0.070)	-0.239*** (0.091)	-0.094 (0.084)
Individual Characteristics		Y	Y	Y	Y
Firm Characteristics			Y	Y	Y
Beliefs / Training				Y	Y
Firm Policies / Authority					Y
UK Diff, Labor	0.256*	0.250*	0.253*	0.254*	0.074
UK Diff, Prod.	0.107	0.137	0.145	0.128	0.074
R-squared	0.026	0.050	0.054	0.089	0.292
Observations	1100	1100	1100	1100	1100

Notes: The table reports OLS estimates from equation (1), where the unit of observation is an individual manager. Panel A reports pooled estimates with country fixed effects, while Panels B and C report estimates separately for the U.S. and U.K., respectively. The dependent variable is the standardized, 3-item *Staffing Intentions* index described in Section 2.4. The omitted category is the neutral control video, with coefficients on *AI Productivity* and *AI Labor* comparing each treatment condition to the control group. Columns add controls sequentially. Column (1) includes no controls. Column (2) adds individual characteristics (age, years of schooling, log income; all demeaned), a gender indicator, political ideology (5-point scale; standardized), and indicators for respondent and family unemployment experience. Column (3) additionally adds demeaned firm and team characteristics: log firm revenue, log firm employment, log team budget, team size, and number of direct reports. Column (4) additionally adds pre-treatment beliefs and training: respondent and team AI sentiment (5-point Likert scale, demeaned), a categorical pre-treatment belief about whether AI augments/automates work, and indicators for respondent and team AI training. Column (5) additionally adds baseline organizational context: a standardized pre-treatment staffing-plan index (constructed from expected headcount changes for up to three occupation groups), technology and hiring-decision authority (5-point Likert scale, standardized), perceived executive support for AI (4-point Likert scale, standardized), and expected AI adoption next year (5-point Likert scale, standardized). The sample is restricted to respondents with non-missing values for all four primary outcome indices and non-missing values for the full control set. The bottom panel indicates which control groups are included in each column and reports the difference in treatment effects between the U.K. and U.S. (U.K.–U.S.), estimated from a pooled interacted specification that allows treatment effects to vary by country. The reported R^2 and number of observations in the bottom panel also come from this pooled interacted specification and therefore do not correspond to the separate country-specific regressions reported in Panels B and C. Robust standard errors are reported in parentheses. (***) $p < 0.01$, (**) $p < 0.05$, (*) $p < 0.10$.

Table A.53: Treatment Effects on Staffing Advocacy Index (Robustness to Controls)

	(1)	(2)	(3)	(4)	(5)
	Baseline	+ Indiv	+ Firm	+ Beliefs	+ Policies
Panel A: Pooled (Country FE)					
AI Labor	-0.083 (0.076)	-0.091 (0.076)	-0.097 (0.077)	-0.094 (0.076)	-0.105 (0.077)
AI Productivity	-0.042 (0.079)	-0.041 (0.078)	-0.043 (0.078)	-0.052 (0.078)	-0.051 (0.077)
Constant	-0.042 (0.057)	0.004 (0.077)	0.035 (0.076)	0.209** (0.094)	0.224** (0.095)
Panel B: US					
AI Labor	-0.019 (0.098)	-0.020 (0.099)	-0.024 (0.099)	-0.019 (0.099)	-0.011 (0.100)
AI Productivity	-0.108 (0.109)	-0.120 (0.109)	-0.115 (0.109)	-0.135 (0.109)	-0.130 (0.108)
Constant	-0.045 (0.067)	-0.013 (0.093)	0.015 (0.093)	0.274** (0.126)	0.286** (0.128)
Panel C: UK					
AI Labor	-0.177 (0.120)	-0.219* (0.122)	-0.223* (0.121)	-0.223* (0.121)	-0.257** (0.124)
AI Productivity	0.039 (0.113)	0.040 (0.111)	0.023 (0.112)	0.025 (0.114)	0.034 (0.114)
Constant	-0.052 (0.070)	-0.049 (0.092)	-0.066 (0.092)	0.025 (0.117)	0.050 (0.122)
Individual Characteristics		Y	Y	Y	Y
Firm Characteristics			Y	Y	Y
Beliefs / Training				Y	Y
Firm Policies / Authority					Y
UK Diff, Labor	-0.158	-0.184	-0.177	-0.182	-0.225
UK Diff, Prod.	0.146	0.161	0.152	0.162	0.164
R-squared	0.004	0.024	0.034	0.053	0.062
Observations	1100	1100	1100	1100	1100

Notes: The table reports OLS estimates from equation (1), where the unit of observation is an individual manager. Panel A reports pooled estimates with country fixed effects, while Panels B and C report estimates separately for the U.S. and U.K., respectively. The dependent variable is the standardized *Staffing Advocacy* index described in Section 2.4. The omitted category is the neutral control video, with coefficients on *AI Productivity* and *AI Labor* comparing each treatment condition to the control group. Columns add controls sequentially. Column (1) includes no controls. Column (2) adds individual characteristics (age, years of schooling, log income; all demeaned), a gender indicator, political ideology (5-point scale; standardized), and indicators for respondent and family unemployment experience. Column (3) additionally adds demeaned firm and team characteristics: log firm revenue, log firm employment, log team budget, team size, and number of direct reports. Column (4) additionally adds pre-treatment beliefs and training: respondent and team AI sentiment (5-point Likert scale, demeaned), a categorical pre-treatment belief about whether AI augments/automates work, and indicators for respondent and team AI training. Column (5) additionally adds baseline organizational context: a standardized pre-treatment staffing-plan index (constructed from expected headcount changes for up to three occupation groups), technology and hiring-decision authority (5-point Likert scale, standardized), perceived executive support for AI (4-point Likert scale, standardized), and expected AI adoption next year (5-point Likert scale, standardized). The sample is restricted to respondents with non-missing values for all four primary outcome indices and non-missing values for the full control set. The bottom panel indicates which control groups are included in each column and reports the difference in treatment effects between the U.K. and U.S. (U.K.–U.S.), estimated from a pooled interacted specification that allows treatment effects to vary by country. The reported R^2 and number of observations in the bottom panel also come from this pooled interacted specification and therefore do not correspond to the separate country-specific regressions reported in Panels B and C. Robust standard errors are reported in parentheses. (***) $p < 0.01$, (**) $p < 0.05$, (*) $p < 0.10$.

Table A.54: Hiring Authority: Interaction Model Estimates

	(1)	(2)	(3)	(4)
	AI Adopt	AI Advoc.	Staff Int.	Staff Advoc.
AI Labor	-0.470*** (0.0770)	-0.433*** (0.0762)	-0.156** (0.0698)	0.0436 (0.0726)
AI Prod.	0.0424 (0.0727)	0.0784 (0.0694)	-0.0476 (0.0735)	-0.107 (0.0780)
High Hiring Authority	-0.00358 (0.0757)	-0.0148 (0.0762)	-0.145* (0.0742)	0.0356 (0.0740)
AI Labor × High Hiring Authority	0.101 (0.124)	-0.0676 (0.122)	-0.109 (0.113)	-0.106 (0.113)
AI Prod. × High Hiring Authority	-0.0242 (0.116)	0.0807 (0.112)	-0.0824 (0.114)	0.0904 (0.118)
Observations	1860	1860	1860	1860
R-squared	0.037	0.052	0.028	0.003

Notes: This table reports heterogeneous treatment effect regressions for the four primary outcome indices using hiring authority as the heterogeneity dimension. Columns correspond to *AI Adoption*, *AI Advocacy*, *Staffing Intentions*, and *Staffing Advocacy*. The estimated specification includes indicators for assignment to the AI Productivity and AI Labor treatments, an indicator for *High Hiring Authority*, interactions between each treatment indicator and *High Hiring Authority*, and country fixed effects. The omitted category is the control group among respondents in the low hiring-authority subgroup. The row labeled *High Hiring Authority* reports the coefficient on the subgroup indicator. The interaction rows report how the treatment effects differ for respondents in the high hiring-authority subgroup relative to respondents in the low subgroup. Hiring authority captures where decisions about team hiring are made. Respondents are classified into the high subgroup if hiring decisions are made primarily at the team level, and into the low subgroup if such decisions are made primarily at the corporate level or are shared across levels. Construction of this subgroup measure is described in Appendix Section E.1. All specifications are estimated on the pooled sample across countries and include robust standard errors in parentheses. Outcome construction is described in Section 2.4. (***) $p < 0.01$, ** $p < 0.05$, * $p < 0.10$.

Table A.55: Hiring Authority: Subgroup Treatment Effects Relative to Control

	AI Adopt	AI Advoc.	Staff Int.	Staff Advoc.
AI Labor Effect (Low Hiring Authority)	-0.470*** (0.077)	-0.433*** (0.076)	-0.156** (0.070)	0.044 (0.073)
AI Labor Effect (High Hiring Authority)	-0.369*** (0.097)	-0.501*** (0.096)	-0.265*** (0.088)	-0.062 (0.087)
AI Productivity Effect (Low Hiring Authority)	0.042 (0.073)	0.078 (0.069)	-0.048 (0.073)	-0.107 (0.078)
AI Productivity Effect (High Hiring Authority)	0.018 (0.091)	0.159* (0.087)	-0.130 (0.087)	-0.017 (0.088)
p diff (Prod: High vs Low)	0.835	0.469	0.469	0.443
p diff (Labor: High vs Low)	0.417	0.580	0.336	0.349

Notes: This table reports subgroup-specific treatment effects implied by the interacted heterogeneity specification using hiring authority as the heterogeneity dimension. Columns correspond to *AI Adoption*, *AI Advocacy*, *Staffing Intentions*, and *Staffing Advocacy*. The row labeled *AI Productivity effect (Low Hiring Authority)* reports the estimated effect of the AI Productivity treatment relative to the control group among respondents in the low-hiring-authority subgroup. The row labeled *AI Productivity effect (High Hiring Authority)* reports the corresponding effect among respondents in the high subgroup. The next two rows report analogous treatment effects for the AI Labor treatment within the low and high subgroups. These subgroup-specific effects are obtained from the interacted regression with treatment indicators, a *High Hiring Authority* indicator, interactions between each treatment indicator and *High Hiring Authority*, and country fixed effects. Hiring authority captures where decisions about team hiring are made. Respondents are classified into the high subgroup if hiring decisions are made primarily at the team level, and into the low subgroup if such decisions are made primarily at the corporate level or are shared across levels. Construction of this subgroup measure is described in Appendix Section E.1. The bottom panel reports p-values for the difference in treatment effects between the high and low subgroups for each treatment arm. These p-values correspond to tests of whether the relevant interaction coefficient is equal to zero in the underlying interacted regression. All specifications are estimated on the pooled sample across countries with country fixed effects, and robust standard errors are reported in parentheses. Outcome construction is described in Section 2.4. (***) $p < 0.01$, ** $p < 0.05$, * $p < 0.10$.

Table A.56: Team AI Use (Proportion): Interaction Model Estimates

	(1)	(2)	(3)	(4)
	AI Adopt	AI Advoc.	Staff Int.	Staff Advoc.
AI Labor	-0.450*** (0.103)	-0.372*** (0.101)	-0.208** (0.0811)	-0.0782 (0.0860)
AI Prod.	0.0702 (0.0917)	0.243*** (0.0941)	-0.164* (0.0836)	-0.0965 (0.0869)
High Team AI Use	0.660*** (0.0744)	0.656*** (0.0746)	0.0498 (0.0716)	-0.299*** (0.0725)
AI Labor × High Team AI Use	0.0308 (0.123)	-0.123 (0.122)	0.000525 (0.111)	0.135 (0.113)
AI Prod. × High Team AI Use	-0.0744 (0.115)	-0.215* (0.113)	0.127 (0.113)	0.0676 (0.118)
Observations	1820	1820	1820	1820
R-squared	0.122	0.116	0.021	0.014

Notes: This table reports heterogeneous treatment effect regressions for the four primary outcome indices using team AI use proportion as the heterogeneity dimension. Columns correspond to *AI Adoption*, *AI Advocacy*, *Staffing Intentions*, and *Staffing Advocacy*. The estimated specification includes indicators for assignment to the AI Productivity and AI Labor treatments, an indicator for *High Team AI Use*, interactions between each treatment indicator and *High Team AI Use*, and country fixed effects. The omitted category is the control group among respondents in the low team-AI-use subgroup. The row labeled *High Team AI Use* reports the coefficient on the subgroup indicator. The interaction rows report how the treatment effects differ for respondents in the high team-AI-use subgroup relative to respondents in the low subgroup. Team AI use captures the share of team members who use AI at least once per week. Respondents are classified into the high subgroup if at least half of their team members use AI weekly, and into the low subgroup if fewer than half do. Construction of this subgroup measure is described in Appendix Section E.1. All specifications are estimated on the pooled sample across countries and include robust standard errors in parentheses. Outcome construction is described in Section 2.4. (***) $p < 0.01$, (**) $p < 0.05$, (*) $p < 0.10$.

Table A.57: Team AI Use (Proportion): Subgroup Treatment Effects Relative to Control

	AI Adopt	AI Advoc.	Staff Int.	Staff Advoc.
AI Labor Effect (Low Team AI Use)	-0.450*** (0.103)	-0.372*** (0.101)	-0.208** (0.081)	-0.078 (0.086)
AI Labor Effect (High Team AI Use)	-0.419*** (0.068)	-0.496*** (0.068)	-0.207*** (0.075)	0.057 (0.074)
AI Productivity Effect (Low Team AI Use)	0.070 (0.092)	0.243*** (0.094)	-0.164* (0.084)	-0.096 (0.087)
AI Productivity Effect (High Team AI Use)	-0.004 (0.068)	0.029 (0.063)	-0.036 (0.076)	-0.029 (0.079)
p diff (Prod: High vs Low)	0.516	0.058	0.261	0.566
p diff (Labor: High vs Low)	0.802	0.312	0.996	0.234

Notes: This table reports subgroup-specific treatment effects implied by the interacted heterogeneity specification using team AI use proportion as the heterogeneity dimension. Columns correspond to *AI Adoption*, *AI Advocacy*, *Staffing Intentions*, and *Staffing Advocacy*. The row labeled *AI Productivity effect (Low Team AI Use)* reports the estimated effect of the AI Productivity treatment relative to the control group among respondents in the low team-AI-use subgroup. The row labeled *AI Productivity effect (High Team AI Use)* reports the corresponding effect among respondents in the high subgroup. The next two rows report analogous treatment effects for the AI Labor treatment within the low and high subgroups. These subgroup-specific effects are obtained from the interacted regression with treatment indicators, a *High Team AI Use* indicator, interactions between each treatment indicator and *High Team AI Use*, and country fixed effects. Team AI use captures the share of team members who use AI at least once per week. Respondents are classified into the high subgroup if at least half of their team members use AI weekly, and into the low subgroup if fewer than half do. Construction of this subgroup measure is described in Appendix Section E.1. The bottom panel reports p-values for the difference in treatment effects between the high and low subgroups for each treatment arm. These p-values correspond to tests of whether the relevant interaction coefficient is equal to zero in the underlying interacted regression. All specifications are estimated on the pooled sample across countries with country fixed effects, and robust standard errors are reported in parentheses. Outcome construction is described in Section 2.4. (***) $p < 0.01$, (**) $p < 0.05$, (*) $p < 0.10$.

Table A.58: Team AI Attitudes: Interaction Model Estimates

	(1)	(2)	(3)	(4)
	AI Adopt	AI Advoc.	Staff Int.	Staff Advoc.
AI Labor	-0.339*** (0.127)	-0.447*** (0.116)	-0.306*** (0.0992)	-0.0368 (0.0983)
AI Prod.	0.0843 (0.115)	0.248** (0.113)	-0.0973 (0.102)	-0.0981 (0.105)
Positive Team AI Attitudes	1.020*** (0.0847)	1.038*** (0.0825)	0.127* (0.0771)	-0.230*** (0.0805)
AI Labor × Positive Team AI Attitudes	-0.0965 (0.140)	0.000807 (0.130)	0.138 (0.119)	0.0410 (0.119)
AI Prod. × Positive Team AI Attitudes	-0.0711 (0.128)	-0.193 (0.125)	0.0324 (0.123)	0.0396 (0.127)
Observations	1838	1838	1838	1838
R-squared	0.201	0.230	0.025	0.009

Notes: This table reports heterogeneous treatment effect regressions for the four primary outcome indices using team AI attitudes as the heterogeneity dimension. Columns correspond to *AI Adoption*, *AI Advocacy*, *Staffing Intentions*, and *Staffing Advocacy*. The estimated specification includes indicators for assignment to the AI Productivity and AI Labor treatments, an indicator for *Positive Team AI Attitudes*, interactions between each treatment indicator and *Positive Team AI Attitudes*, and country fixed effects. The omitted category is the control group among respondents in the low team-AI-attitudes subgroup. The row labeled *Positive Team AI Attitudes* reports the coefficient on the subgroup indicator. The interaction rows report how the treatment effects differ for respondents in the positive-team-AI-attitudes subgroup relative to respondents in the low subgroup. Team AI attitudes capture respondents' pre-treatment perceptions of their teams' baseline orientation toward AI. Respondents are classified into the high subgroup if they perceive their team as having a positive attitude toward AI, and into the low subgroup if they perceive their team as neutral or negative toward AI. Construction of this subgroup measure is described in Appendix Section E.1. All specifications are estimated on the pooled sample across countries and include robust standard errors in parentheses. Outcome construction is described in Section 2.4. (***) $p < 0.01$, (**) $p < 0.05$, (*) $p < 0.10$.

Table A.59: Team AI Attitudes: Subgroup Treatment Effects Relative to Control

	AI Adopt	AI Advoc.	Staff Int.	Staff Advoc.
AI Labor Effect (Not Positive Team Attitudes)	-0.339*** (0.127)	-0.447*** (0.116)	-0.306*** (0.099)	-0.037 (0.098)
AI Labor Effect (Positive Team Attitudes)	-0.435*** (0.058)	-0.446*** (0.059)	-0.167** (0.066)	0.004 (0.067)
AI Productivity Effect (Not Positive Team Attitudes)	0.084 (0.115)	0.248** (0.113)	-0.097 (0.102)	-0.098 (0.105)
AI Productivity Effect (Positive Team Attitudes)	0.013 (0.057)	0.056 (0.053)	-0.065 (0.068)	-0.059 (0.071)
p diff (Prod: High vs Low)	0.579	0.123	0.792	0.755
p diff (Labor: High vs Low)	0.492	0.995	0.245	0.730

Notes: This table reports subgroup-specific treatment effects implied by the interacted heterogeneity specification using team AI attitudes as the heterogeneity dimension. Columns correspond to *AI Adoption*, *AI Advocacy*, *Staffing Intentions*, and *Staffing Advocacy*. The row labeled *AI Productivity effect (Low Team AI Attitudes)* reports the estimated effect of the AI Productivity treatment relative to the control group among respondents in the low team-AI-attitudes subgroup. The row labeled *AI Productivity effect (High Team AI Attitudes)* reports the corresponding effect among respondents in the high subgroup. The next two rows report analogous treatment effects for the AI Labor treatment within the low and high subgroups. These subgroup-specific effects are obtained from the interacted regression with treatment indicators, a *Positive Team AI Attitudes* indicator, interactions between each treatment indicator and *Positive Team AI Attitudes*, and country fixed effects. Team AI attitudes capture respondents' pre-treatment perceptions of their teams' baseline orientation toward AI. Respondents are classified into the high subgroup if they perceive their team as having a positive attitude toward AI, and into the low subgroup if they perceive their team as neutral or negative toward AI. Construction of this subgroup measure is described in Appendix Section E.1. The bottom panel reports p-values for the difference in treatment effects between the high and low subgroups for each treatment arm. These p-values correspond to tests of whether the relevant interaction coefficient is equal to zero in the underlying interacted regression. All specifications are estimated on the pooled sample across countries with country fixed effects, and robust standard errors are reported in parentheses. Outcome construction is described in Section 2.4. (***) $p < 0.01$, (**) $p < 0.05$, (*) $p < 0.10$.

Table A.60: Pre-Treatment AI Adoption Plans: Interaction Model Estimates

	(1)	(2)	(3)	(4)
	AI Adopt	AI Advoc.	Staff Int.	Staff Advoc.
AI Labor	-0.517*** (0.123)	-0.603*** (0.114)	-0.169* (0.101)	0.0478 (0.0978)
AI Prod.	0.0312 (0.109)	0.207* (0.109)	-0.0435 (0.0928)	-0.116 (0.0995)
High AI Adoption Plans	0.820*** (0.0810)	0.764*** (0.0813)	0.203*** (0.0762)	-0.142* (0.0771)
AI Labor × High AI Adoption Plans	0.109 (0.137)	0.180 (0.131)	-0.0341 (0.121)	-0.0676 (0.119)
AI Prod. × High AI Adoption Plans	-0.00704 (0.124)	-0.138 (0.123)	-0.0481 (0.116)	0.0616 (0.123)
Observations	1817	1817	1817	1817
R-squared	0.170	0.169	0.024	0.006

Notes: This table reports heterogeneous treatment effect regressions for the four primary outcome indices using pre-treatment AI adoption plans as the heterogeneity dimension. Columns correspond to *AI Adoption*, *AI Advocacy*, *Staffing Intentions*, and *Staffing Advocacy*. The estimated specification includes indicators for assignment to the AI Productivity and AI Labor treatments, an indicator for *High AI Adoption Plans*, interactions between each treatment indicator and *High AI Adoption Plans*, and country fixed effects. The omitted category is the control group among respondents in the low AI-adoption-plans subgroup. The row labeled *High AI Adoption Plans* reports the coefficient on the subgroup indicator. The interaction rows report how the treatment effects differ for respondents in the high AI-adoption-plans subgroup relative to respondents in the low subgroup. Pre-treatment AI adoption plans capture respondents' expectations about how AI adoption on their team will change over the next year. Respondents are classified into the high subgroup if they expect AI adoption to expand, and into the low subgroup if they expect adoption to remain stable or decline. Construction of this subgroup measure is described in Appendix Section E.1. All specifications are estimated on the pooled sample across countries and include robust standard errors in parentheses. Outcome construction is described in Section 2.4. (***) $p < 0.01$, ** $p < 0.05$, * $p < 0.10$.

Table A.61: Pre-Treatment AI Adoption Plans: Subgroup Treatment Effects Relative to Control

	AI Adopt	AI Advoc.	Staff Int.	Staff Advoc.
AI Labor Effect (Low AI Adoption Plans)	-0.517*** (0.123)	-0.603*** (0.114)	-0.169* (0.101)	0.048 (0.098)
AI Labor Effect (High AI Adoption Plans)	-0.407*** (0.062)	-0.423*** (0.063)	-0.203*** (0.067)	-0.020 (0.068)
AI Productivity Effect (Low AI Adoption Plans)	0.031 (0.109)	0.207* (0.109)	-0.043 (0.093)	-0.116 (0.100)
AI Productivity Effect (High AI Adoption Plans)	0.024 (0.059)	0.069 (0.058)	-0.092 (0.070)	-0.055 (0.073)
p diff (Prod: High vs Low)	0.955	0.265	0.679	0.618
p diff (Labor: High vs Low)	0.426	0.168	0.779	0.570

Notes: This table reports subgroup-specific treatment effects implied by the interacted heterogeneity specification using pre-treatment AI adoption plans as the heterogeneity dimension. Columns correspond to *AI Adoption*, *AI Advocacy*, *Staffing Intentions*, and *Staffing Advocacy*. The row labeled *AI Productivity effect (Low AI Adoption Plans)* reports the estimated effect of the AI Productivity treatment relative to the control group among respondents in the low AI-adoption-plans subgroup. The row labeled *AI Productivity effect (High AI Adoption Plans)* reports the corresponding effect among respondents in the high subgroup. The next two rows report analogous treatment effects for the AI Labor treatment within the low and high subgroups. These subgroup-specific effects are obtained from the interacted regression with treatment indicators, a *High AI Adoption Plans* indicator, interactions between each treatment indicator and *High AI Adoption Plans*, and country fixed effects. Pre-treatment AI adoption plans capture respondents' expectations about how AI adoption on their team will change over the next year. Respondents are classified into the high subgroup if they expect AI adoption to expand, and into the low subgroup if they expect adoption to remain stable or decline. Construction of this subgroup measure is described in Appendix Section E.1. The bottom panel reports p-values for the difference in treatment effects between the high and low subgroups for each treatment arm. These p-values correspond to tests of whether the relevant interaction coefficient is equal to zero in the underlying interacted regression. All specifications are estimated on the pooled sample across countries with country fixed effects, and robust standard errors are reported in parentheses. Outcome construction is described in Section 2.4. (***) $p < 0.01$, ** $p < 0.05$, * $p < 0.10$.

Table A.62: Treatment Effects on Main Indices Using Outcome-Specific Samples

	(1)	(2)	(3)	(4)
	AI Adoption	AI Advocacy	Staffing Intentions	Staffing Advocacy
Panel A: Pooled (Country FE)				
AI Labor	-0.438*** (0.060)	-0.463*** (0.058)	-0.184*** (0.054)	-0.006 (0.055)
AI Productivity	0.025 (0.055)	0.109** (0.053)	-0.072 (0.055)	-0.069 (0.057)
R-squared	0.037	0.052	0.016	0.001
Observations	1963	1972	1975	1921
Panel B: US				
AI Labor	-0.415*** (0.085)	-0.492*** (0.083)	-0.294*** (0.082)	0.009 (0.077)
AI Productivity	0.053 (0.081)	0.111 (0.078)	-0.045 (0.083)	-0.071 (0.082)
R-squared	0.034	0.054	0.014	0.001
Observations	981	990	991	965
Panel C: UK				
AI Labor	-0.461*** (0.083)	-0.435*** (0.080)	-0.072 (0.069)	-0.022 (0.078)
AI Productivity	-0.002 (0.075)	0.107 (0.071)	-0.098 (0.071)	-0.067 (0.081)
R-squared	0.040	0.049	0.002	0.001
Observations	982	982	984	956
UK Diff, Labor	-0.046	0.057	0.221**	-0.032
UK Diff, Prod.	-0.055	-0.004	-0.053	0.005
R-squared	0.037	0.052	0.019	0.001
Observations	1963	1972	1975	1921

Notes: The table reports OLS estimates from equation (1) with no controls, where the unit of observation is an individual manager. The dependent variables are four standardized primary outcome indices. *AI Adoption* is a 4-item index measuring intentions to adopt or expand AI use in the respondent's team/unit (higher values indicate greater adoption intent). *AI Advocacy* is a 2-item index measuring willingness to advocate for expanding AI use within the organization (higher values indicate more pro-AI advocacy). *Staffing Intentions* is a 3-item index measuring expected headcount changes over the next year for the respondent's three most common team occupation groups (higher values indicate stronger hiring/expansion plans). *Staffing Advocacy* is a 2-item index measuring willingness to advocate for hiring versus layoffs (higher values indicate more pro-hiring advocacy). Indices are constructed in three steps: each component item is standardized using the pooled control-group mean and standard deviation, standardized items are averaged (requiring at least half of the component items to be non-missing), and the resulting index is standardized again using the pooled control-group mean and standard deviation. More information on index construction is described in Section 2.4. The omitted category is the neutral control video, so coefficients on *AI Productivity* and *AI Labor* compare each treatment condition to the control group. Panels report estimates separately for the U.S. and U.K. subsamples. The bottom rows report the difference in treatment effects between the U.K. and U.S. for each arm, along with R^2 and the number of observations. Unlike Table A.5, each regression is estimated on the sample of randomized respondents with a non-missing value for the specific outcome index used in that column, so the number of observations can vary across columns. Robust standard errors are reported in parentheses. (***) $p < 0.01$, (**) $p < 0.05$, (*) $p < 0.10$.

B Appendix: Alternative Heterogeneity and Mechanism Analysis

B.1 Preliminary Causal Random Forest Results

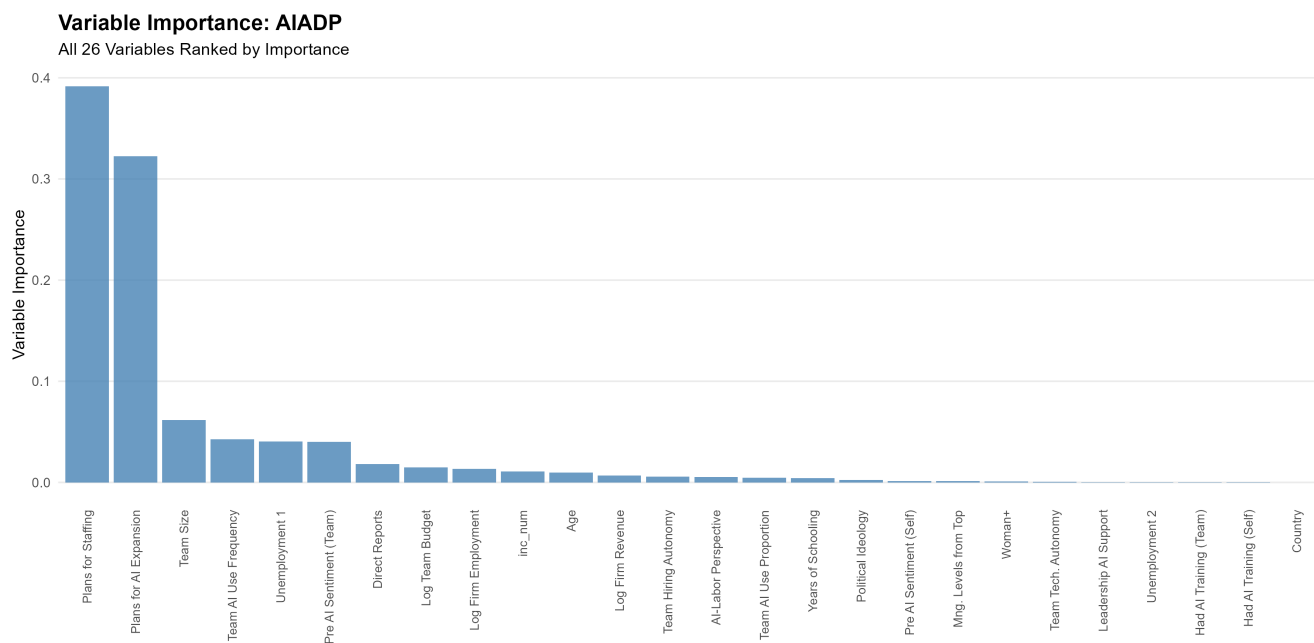


Figure B.2: Variable Importance for AI Adoption Index

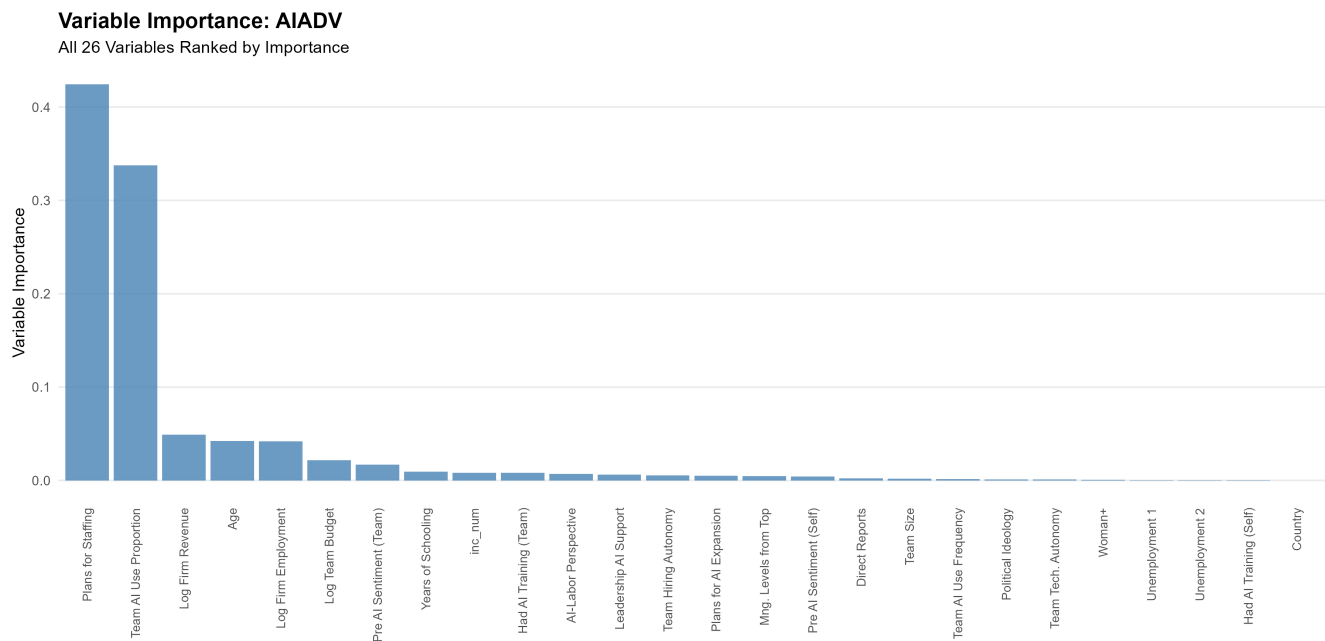


Figure B.3: Variable Importance for AI Advocacy Index

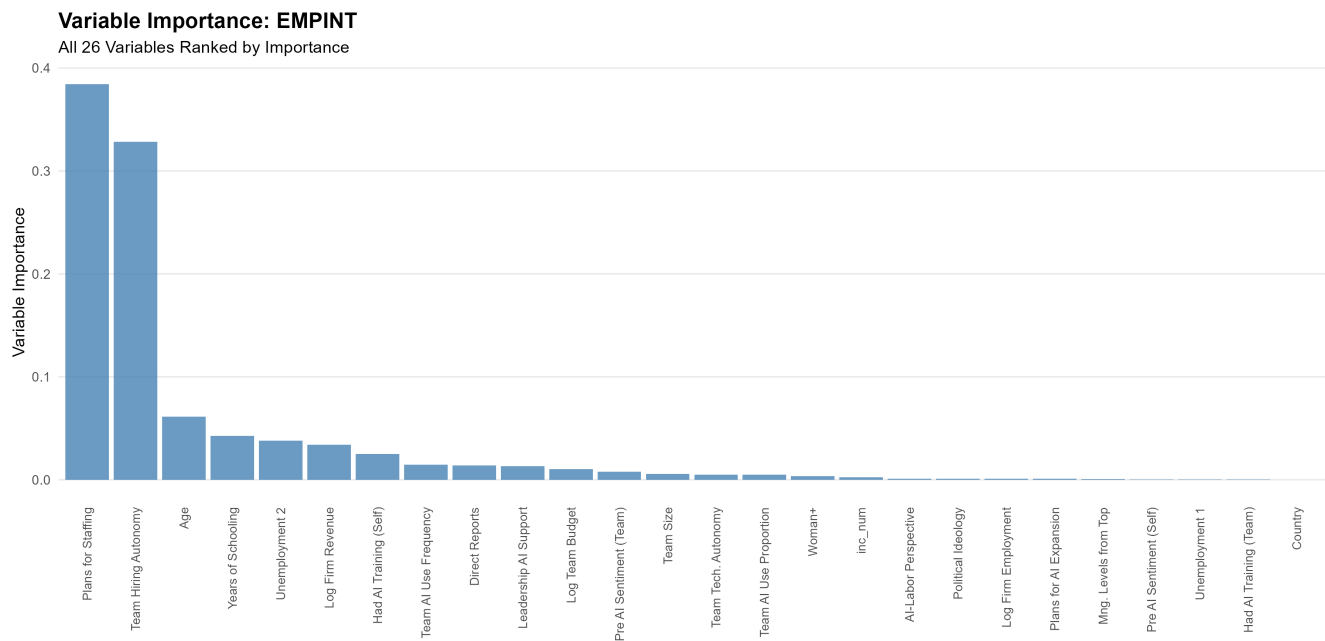


Figure B.4: Variable Importance for Staffing Intentions Index

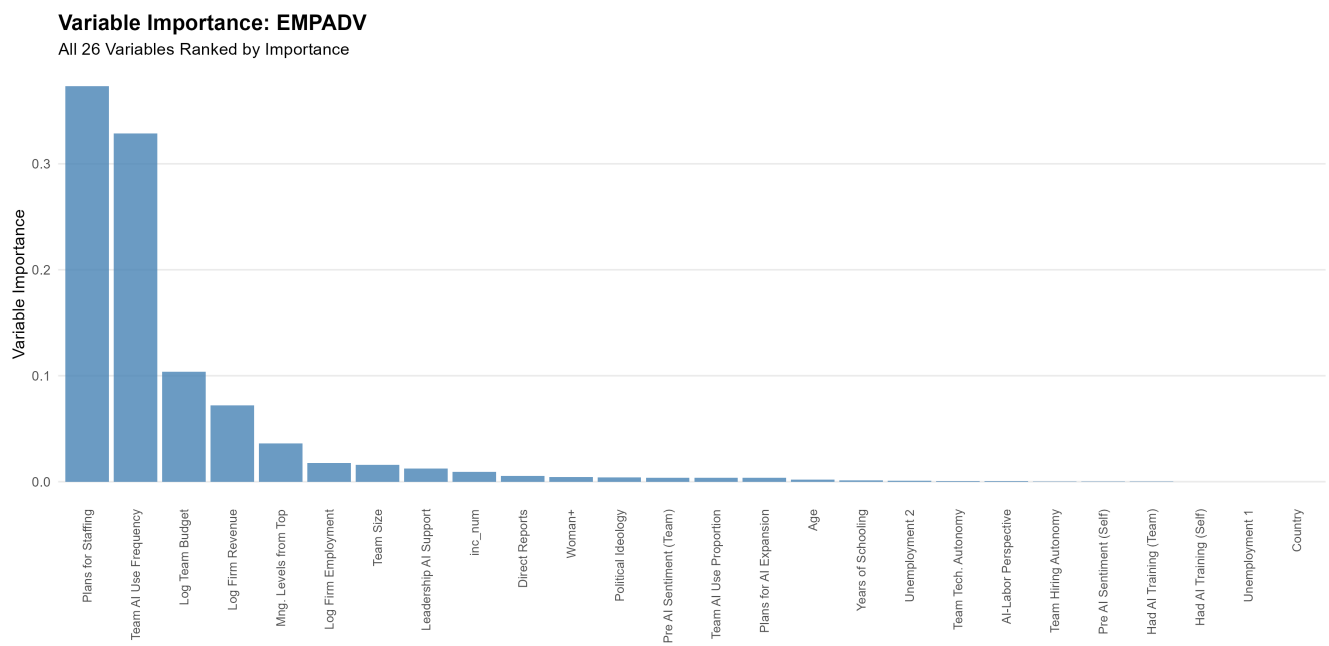
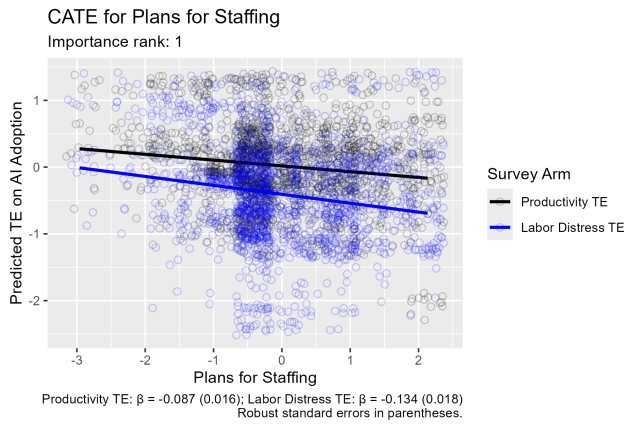
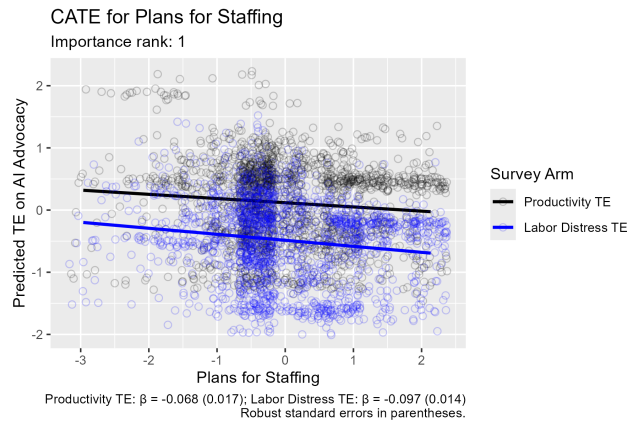


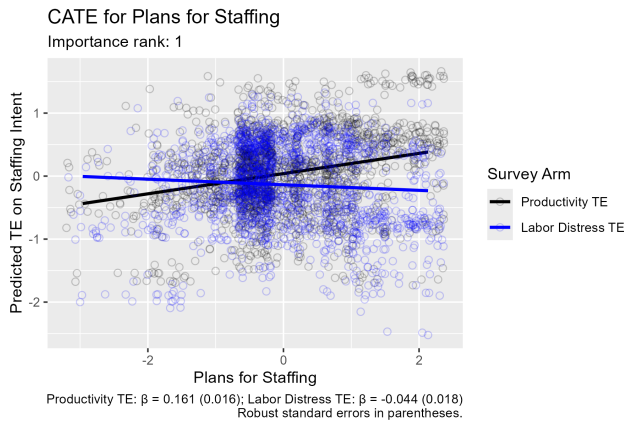
Figure B.5: Variable Importance for Staffing Advocacy Index



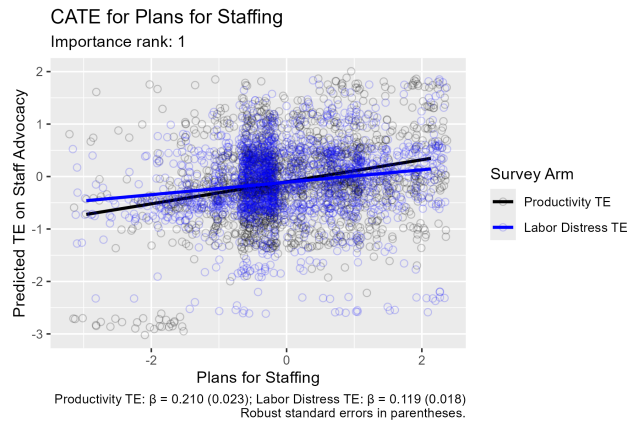
(a) AI Adoption



(b) AI Advocacy



(c) Staffing Plans



(d) Staffing Advocacy

Figure B.6: Predicted CATE on Main Indices based on Pre-Treatment Staffing Plans

B.2 Continuous Heterogeneity Effect Estimates

Table B.63: All Main Treatment Effect Heterogeneity from Decision Authority (Continuous)

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
	AI Adopt, Tech	AI Adopt, Hiring	AI Advoc., Tech	AI Advoc., Hiring	Staff Int., Tech	Staff Int., Hiring	Staff Advoc., Tech	Staff Advoc., Hiring
AI labor	-0.428*** (0.0604)	-0.431*** (0.0604)	-0.458*** (0.0595)	-0.463*** (0.0597)	-0.204*** (0.0552)	-0.201*** (0.0550)	-0.00499 (0.0556)	0.00162 (0.0557)
AI Prod.	0.0332 (0.0568)	0.0315 (0.0566)	0.110** (0.0544)	0.111** (0.0543)	-0.0921 (0.0562)	-0.0838 (0.0561)	-0.0722 (0.0587)	-0.0688 (0.0585)
Z1 level	0.0800** (0.0388)	0.0211 (0.0370)	0.0922** (0.0387)	0.0210 (0.0371)	-0.00978 (0.0387)	-0.0495 (0.0357)	-0.0742** (0.0360)	-0.00598 (0.0367)
Z × AI Labor	-0.0513 (0.0647)	0.0585 (0.0650)	0.0134 (0.0627)	-0.0141 (0.0624)	-0.0366 (0.0583)	-0.0491 (0.0540)	-0.0206 (0.0562)	-0.0275 (0.0562)
Z × AI Prod.	-0.0862 (0.0598)	0.0265 (0.0606)	-0.0793 (0.0580)	0.0524 (0.0558)	-0.0624 (0.0586)	-0.0594 (0.0577)	0.00639 (0.0578)	0.0623 (0.0587)
Z mean	2.715	3.217	2.715	3.217	2.715	3.217	2.715	3.217
Z sd	1.159	1.136	1.159	1.136	1.159	1.136	1.159	1.136
Observations	1851	1860	1851	1860	1851	1860	1851	1860
R-squared	0.040	0.039	0.057	0.053	0.020	0.025	0.007	0.002

Notes: The table reports OLS estimates of heterogeneous treatment effects from the interacted specification in equation (3), where the unit of observation is an individual manager. The dependent variables are the four standardized primary outcome indices described in Section 2.4. The omitted category is the neutral control video, so coefficients on *AI Prod.* and *AI Labor* reflect treatment effects relative to the control group. Each outcome is estimated twice using two alternative continuous, pre-treatment authority measures (labeled *Z1 level* in the table). In the *Tech* specifications, *Z1* is the respondent's technology decision authority (5-point scale), and in the *Hiring* specifications, *Z1* is the respondent's hiring decision authority (5-point scale); higher values indicate greater decision authority. The authority measures are standardized for estimation, so the interaction terms (*Z1 × AI Prod.* and *Z1 × AI Labor*) reflect how treatment effects change with a one-standard-deviation increase in the corresponding authority measure. Rows labeled *Z mean* and *Z sd* report the unweighted mean and standard deviation of the corresponding authority measure on its original 1–5 scale for the estimation sample in each column. All specifications include country fixed effects. Robust standard errors are reported in parentheses. The sample is restricted to respondents with non-missing values for all four primary outcome indices and non-missing values for the corresponding authority measure. (***) $p < 0.01$, ** $p < 0.05$, * $p < 0.10$.

Table B.64: All Main Treatment Effect Heterogeneity from AI Use

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
	AI Adopt, Prop.	AI Adopt, Freq.	AI Advoc., Prop.	AI Advoc., Freq.	Staff Int., Prop.	Staff Int., Freq.	Staff Advoc., Prop.	Staff Advoc., Freq.
AI Labor	-0.486*** (0.140)	-0.445*** (0.119)	-0.310** (0.135)	-0.460*** (0.112)	-0.186* (0.109)	-0.240*** (0.0921)	-0.103 (0.115)	-0.0347 (0.0954)
AI Prod.	0.137 (0.124)	0.133 (0.101)	0.341*** (0.125)	0.246** (0.101)	-0.217* (0.111)	-0.110 (0.0922)	-0.151 (0.115)	-0.174* (0.0997)
Z1 Level	1.205*** (0.123)	1.159*** (0.0903)	1.213*** (0.121)	1.039*** (0.0921)	0.119 (0.115)	0.223** (0.0930)	-0.496*** (0.119)	-0.299*** (0.0934)
Z1 x AI Labor	0.0738 (0.203)	-0.0277 (0.159)	-0.254 (0.196)	-0.0509 (0.154)	-0.0388 (0.177)	0.0461 (0.143)	0.190 (0.181)	0.0785 (0.145)
Z1 x AI Prod.	-0.206 (0.187)	-0.174 (0.144)	-0.407** (0.180)	-0.238* (0.142)	0.222 (0.179)	0.0402 (0.148)	0.169 (0.184)	0.182 (0.151)
Z mean	0.584	0.581	0.584	0.581	0.584	0.581	0.584	0.581
Z sd	0.310	0.378	0.310	0.378	0.310	0.378	0.310	0.378
Observations	1820	1852	1820	1852	1820	1852	1820	1852
R-squared	0.150	0.189	0.140	0.172	0.023	0.027	0.015	0.009

Notes: The table reports OLS estimates of heterogeneous treatment effects from the interacted specification in equation (3), where the unit of observation is an individual manager. The dependent variables are the four standardized primary outcome indices described in Table A.5 and Section 2.4. The omitted category is the neutral control video, so coefficients on *AI Prod.* and *AI Labor* reflect treatment effects relative to the control group. Each outcome is estimated using two alternative continuous, team-level heterogeneity measures of pre-treatment AI use (labeled *Z1 level* in the table). In the *Prop.* specifications, *Z1* is *Team AI Use Prop.*, the share of the respondent's team/unit using AI at least weekly, coded numerically as 0 (none), 0.25 (<50%), 0.75 (50–100%), and 1 (100%). In the *Freq.* specifications, *Z1* is *Team AI Use Freq.*, how often the team/unit uses AI at work, coded numerically as 0 (never), 0.05 (rarely), 0.1 (a few times per month), 0.5 (a few times per week), and 1 (daily). The interaction terms (*Z1 × AI Prod.* and *Z1 × AI Labor*) capture how treatment effects vary with baseline team AI use under each measure. Rows labeled *Z mean* and *Z sd* report the unweighted mean and standard deviation of the corresponding team-level AI use measure (Prop. or Freq.) for the estimation sample in each column. All specifications include country fixed effects. Robust standard errors are reported in parentheses. The sample is restricted to respondents with non-missing values for all four primary outcome indices and non-missing values for the corresponding team AI use measure. (***) $p < 0.01$, ** $p < 0.05$, * $p < 0.10$.

Table B.65: All Main Treatment Effect Heterogeneity from AI Attitudes (Continuous)

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
	AI Adopt, Self	AI Adopt, Team	AI Advoc., Self	AI Advoc., Team	Staff Int., Self	Staff Int., Team	Staff Advoc., Self	Staff Advoc., Team
AI labor	-0.422*** (0.0543)	-0.414*** (0.0527)	-0.457*** (0.0509)	-0.451*** (0.0509)	-0.207*** (0.0551)	-0.205*** (0.0550)	0.00400 (0.0557)	-0.00519 (0.0558)
AI Prod.	0.0345 (0.0497)	0.00865 (0.0500)	0.110** (0.0463)	0.0852* (0.0478)	-0.0891 (0.0565)	-0.0775 (0.0567)	-0.0653 (0.0586)	-0.0651 (0.0590)
Z1 level	0.489*** (0.0369)	0.519*** (0.0402)	0.555*** (0.0331)	0.544*** (0.0387)	0.113*** (0.0351)	0.0739* (0.0385)	-0.126*** (0.0337)	-0.0995*** (0.0337)
Z × AI Labor	-0.0513 (0.0666)	-0.0564 (0.0625)	0.000667 (0.0538)	-0.0425 (0.0560)	-0.0920 (0.0594)	0.0389 (0.0575)	0.0657 (0.0588)	0.0215 (0.0495)
Z × AI Prod.	-0.0325 (0.0603)	-0.0160 (0.0640)	-0.0703 (0.0518)	-0.104* (0.0590)	0.00504 (0.0571)	0.0223 (0.0620)	0.00722 (0.0557)	0.0215 (0.0605)
Z mean	3.988	3.905	3.988	3.905	3.988	3.905	3.988	3.905
Z sd	0.879	0.953	0.879	0.953	0.879	0.953	0.879	0.953
Observations	1843	1838	1843	1838	1843	1838	1843	1838
R-squared	0.229	0.253	0.312	0.285	0.027	0.027	0.012	0.008

Notes: The table reports OLS estimates of heterogeneous treatment effects from the interacted specification in equation (3), where the unit of observation is an individual manager. The dependent variables are the four standardized primary outcome indices described in Section 2.4. The omitted category is the neutral control video, so coefficients on *AI Prod.* and *AI Labor* reflect treatment effects relative to the control group. All specifications include country fixed effects, and robust standard errors are reported in parentheses. The heterogeneity measure (*Z1 level*) is pre-treatment AI attitude, measured on a 1–5 scale. Columns labeled *Self* use the respondent's own AI attitude (*pre_ai_views*), while columns labeled *Team* use the respondent's perception of their team/unit's AI attitude (*teamatt*). For estimation, the corresponding attitude measure is standardized, so the interaction terms (*Z1 × AI Prod.* and *Z1 × AI Labor*) indicate how treatment effects change with a one-standard-deviation increase in AI attitude. Rows labeled *Z mean* and *Z sd* report the unweighted mean and standard deviation of the corresponding attitude measure on its original 1–5 scale for the estimation sample in each column. The sample is restricted to respondents with non-missing values for all four primary outcome indices and a non-missing value of the corresponding AI attitude measure. (***) $p < 0.01$, ** $p < 0.05$, * $p < 0.10$.

Table B.66: All Main Treatment Effect Heterogeneity from Firm Executive Support for AI (Continuous)

	(1)	(2)	(3)	(4)
	AI Adopt	AI Advoc.	Staff Int.	Staff Advoc.
AI labor	-0.454*** (0.0556)	-0.479*** (0.0568)	-0.215*** (0.0550)	-0.00635 (0.0560)
AI Prod.	0.0383 (0.0515)	0.122** (0.0508)	-0.0719 (0.0569)	-0.0679 (0.0593)
Z1 level	0.452*** (0.0377)	0.388*** (0.0385)	0.0295 (0.0358)	-0.104*** (0.0344)
Z × AI Labor	-0.0273 (0.0650)	-0.0666 (0.0634)	0.0201 (0.0537)	0.0982* (0.0538)
Z × AI Prod.	-0.0669 (0.0575)	-0.0856 (0.0567)	0.107* (0.0562)	0.0898 (0.0559)
Z mean	3.052	3.052	3.052	3.052
Z sd	0.858	0.858	0.858	0.858
Observations	1825	1825	1825	1825
R-squared	0.199	0.162	0.027	0.005

Notes: The table reports OLS estimates of heterogeneous treatment effects from the interacted specification in equation (3), where the unit of observation is an individual manager. The dependent variables are the four standardized primary outcome indices described in Section 2.4. The omitted category is the neutral control video, so coefficients on *AI Prod.* and *AI Labor* reflect treatment effects relative to the control group. All specifications include country fixed effects, and robust standard errors are reported in parentheses. The heterogeneity measure (*Z1 level*) is pre-treatment perceived executive support for AI, measured on a 1–4 scale. For estimation, the executive-support measure is standardized, so the interaction terms (*Z1 × AI Prod.* and *Z1 × AI Labor*) indicate how treatment effects change with a one-standard-deviation increase in perceived executive support. Rows labeled *Z mean* and *Z sd* report the unweighted mean and standard deviation of the executive-support measure on its original 1–4 scale for the estimation sample. The sample is restricted to respondents with non-missing values for all four primary outcome indices and a non-missing executive-support measure. (***) $p < 0.01$, ** $p < 0.05$, * $p < 0.10$.

Table B.67: All Main Treatment Effect Heterogeneity from Baseline AI Adoption Plans

	(1)	(2)	(3)	(4)
	AI Adopt	AI Advoc.	Staff Int.	Staff Advoc.
AI labor	-0.433*** (0.0553)	-0.468*** (0.0551)	-0.192*** (0.0557)	-0.00200 (0.0561)
AI Prod.	0.0332 (0.0517)	0.114** (0.0508)	-0.0763 (0.0568)	-0.0732 (0.0594)
ZI level	0.420*** (0.0373)	0.390*** (0.0368)	0.0985** (0.0398)	-0.0797** (0.0375)
Z × AI Labor	0.0188 (0.0634)	0.0440 (0.0598)	-0.0227 (0.0639)	-0.0118 (0.0576)
Z × AI Prod.	-0.0218 (0.0570)	-0.0541 (0.0544)	0.00530 (0.0629)	0.00634 (0.0590)
Z mean	3.857	3.857	3.857	3.857
Z sd	0.742	0.742	0.742	0.742
Observations	1817	1817	1817	1817
R-squared	0.197	0.193	0.026	0.007

Notes: The table reports OLS estimates of heterogeneous treatment effects from the interacted specification in equation (3), where the unit of observation is an individual manager. The dependent variables are the four standardized primary outcome indices described in Section 2.4. The omitted category is the neutral control video, so coefficients on *AI Prod.* and *AI Labor* reflect treatment effects relative to the control group. All specifications include country fixed effects and robust standard errors are reported in parentheses. The heterogeneity measure (*ZI level*) is a pre-treatment measure of expected AI adoption over the next year, recorded on a 1–5 Likert scale (higher values indicate stronger expected adoption). For estimation, this measure is standardized, so the interaction terms (*ZI × AI Prod.* and *ZI × AI Labor*) indicate how treatment effects change with a one-standard-deviation increase in baseline AI adoption plans. Rows labeled *Z mean* and *Z sd* report the unweighted mean and standard deviation of the adoption-plan measure on its original 1–5 scale for the estimation sample. The sample is restricted to respondents with non-missing values for all four primary outcome indices and a non-missing value of the AI adoption-plan measure. (***) $p < 0.01$, ** $p < 0.05$, * $p < 0.10$.

Table B.68: All Main Treatment Effect Heterogeneity from Baseline Staffing Plans

	(1)	(2)	(3)	(4)
	AI Adoption Index	AI Advocacy Index	Staffing Index	Staff Advocacy Index
AI Labor	-0.428*** (0.0611)	-0.458*** (0.0605)	-0.232*** (0.0516)	-0.00613 (0.0562)
AI Prod.	0.0402 (0.0571)	0.120** (0.0546)	-0.0333 (0.0495)	-0.0641 (0.0592)
(Pre) Staffing Plan	0.142*** (0.0389)	0.169*** (0.0396)	0.481*** (0.0414)	-0.0824* (0.0428)
AI Labor × (Pre) Staffing Plan	-0.212*** (0.0652)	-0.119* (0.0616)	-0.166*** (0.0637)	0.159** (0.0652)
AI Prod. × (Pre) Staffing Plan	-0.0447 (0.0663)	-0.101* (0.0595)	0.0410 (0.0624)	0.152** (0.0686)
Z mean	3.330	3.330	3.330	3.330
Z sd	0.787	0.787	0.787	0.787
Observations	1826	1826	1826	1826
R-squared	0.048	0.062	0.212	0.007

Notes: The table reports OLS estimates of heterogeneous treatment effects from the interacted specification in equation (3), where the unit of observation is an individual manager. The dependent variables are the four standardized primary outcome indices described in Table A.5 and Section 2.4. The omitted category is the neutral control video, so coefficients on *AI Prod.* and *AI Labor* reflect treatment effects relative to the control group. The heterogeneity measure (*Pre*) Staffing Plan captures managers' baseline staffing outlook prior to treatment. It is constructed from respondents' expected headcount changes over the next year for the three most common occupation groups on their team. Each occupation-group item is standardized, the standardized items are averaged (requiring at least two non-missing occupation groups), and the resulting average is standardized again. Higher values indicate stronger baseline plans to expand headcount. The interaction terms *AI Prod. × (Pre) Staffing Plan* and *AI Labor × (Pre) Staffing Plan* indicate how treatment effects change with baseline staffing plans. Rows labeled *Z mean* and *Z sd* report the unweighted mean and standard deviation of the non-standardized staffing-plan composite (average of the raw occupation-group items, requiring at least two non-missing) for the estimation sample in each column. All specifications include country fixed effects. Robust standard errors are reported in parentheses. The sample is restricted to respondents with non-missing values for all four primary outcome indices and a non-missing Staffing Plan index. (***) $p < 0.01$, ** $p < 0.05$, * $p < 0.10$.

B.3 Imai Decomposition Analysis

Table B.69: AI Benefits Mediation Analysis on Main Indices

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)
	AI Benefits	AI Adopt.	AI Adopt.	AI Advoc.	AI Advoc.	Staff Intent	Staff Intent	Staff Advoc.	Staff Advoc.
AI Labor	-0.413*** (0.059)	-0.424*** (0.060)	-0.238*** (0.055)	-0.460*** (0.059)	-0.253*** (0.051)	-0.205*** (0.055)	-0.168*** (0.055)	-0.000 (0.056)	-0.063 (0.055)
AI Productivity	0.035 (0.058)	0.033 (0.057)	0.018 (0.051)	0.112** (0.054)	0.094** (0.048)	-0.087 (0.056)	-0.090 (0.056)	-0.062 (0.058)	-0.057 (0.058)
AI Benefits			0.451*** (0.025)		0.501*** (0.023)		0.091*** (0.024)		-0.152*** (0.023)
Constant	0.057 (0.044)	0.023 (0.044)	-0.002 (0.039)	-0.007 (0.044)	-0.035 (0.038)	0.103** (0.045)	0.098** (0.045)	-0.006 (0.044)	0.002 (0.043)
Observations	1861.000	1861.000	1861.000	1861.000	1861.000	1861.000	1861.000	1861.000	1861.000

Notes: The table reports OLS estimates from equation (??), where the unit of observation is an individual manager. The dependent variables are the four standardized primary outcome indices described in Section 2.4. The omitted category is the neutral control video. Coefficients on *AI Productivity* and *AI Labor* compare each treatment condition to the control group. For each outcome, odd-numbered columns correspond to the baseline specification (first line of equation (??)), and the adjacent even-numbered columns add the mediator (second line of equation (??)). For each outcome, comparing a baseline column to the immediately following column (which adds the mediator) shows how treatment-effect estimates change when conditioning on the mediator. All specifications include country fixed effects. The *AI Benefits* index is constructed from two post-treatment belief items capturing whether the respondent views AI as more beneficial and whether expected productivity benefits have increased, with higher values representing higher perceived AI benefits. Each item is standardized using the pooled control-group mean and standard deviation, averaged (requiring at least one non-missing component), and the resulting average is standardized again. Robust standard errors are reported in parentheses. The sample is restricted to respondents with non-missing values for all four primary outcome indices and a non-missing *AI Benefits* index. (***) $p < 0.01$, (**) $p < 0.05$, (*) $p < 0.10$.

Table B.70: AI Risks Causal Mediation Analysis on Main Indices

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)
	AI Risks	AI Adopt.	AI Adopt.	AI Advoc.	AI Advoc.	Staff Intent	Staff Intent	Staff Advoc.	Staff Advoc.
AI Labor	0.515*** (0.057)	-0.427*** (0.060)	-0.250*** (0.058)	-0.460*** (0.060)	-0.229*** (0.054)	-0.204*** (0.055)	-0.189*** (0.056)	-0.004 (0.056)	-0.042 (0.057)
AI Productivity	-0.023 (0.057)	0.033 (0.057)	0.026 (0.054)	0.112** (0.054)	0.101** (0.049)	-0.084 (0.056)	-0.085 (0.056)	-0.067 (0.059)	-0.065 (0.058)
AI Risks			-0.343*** (0.026)		-0.449*** (0.023)		-0.029 (0.025)		0.074*** (0.025)
Constant	0.014 (0.045)	0.025 (0.044)	0.030 (0.041)	-0.004 (0.044)	0.002 (0.039)	0.105** (0.045)	0.105** (0.045)	-0.006 (0.044)	-0.007 (0.044)
Observations	1860.000	1860.000	1860.000	1860.000	1860.000	1860.000	1860.000	1860.000	1860.000

Notes: The table reports OLS estimates from equation (??), where the unit of observation is an individual manager. The dependent variables are the four standardized primary outcome indices described in Section 2.4. The omitted category is the neutral control video. Coefficients on *AI Productivity* and *AI Labor* compare each treatment condition to the control group. For each outcome, odd-numbered columns correspond to the baseline specification (first line of equation (??)), and the adjacent even-numbered columns add the mediator (second line of equation (??)). For each outcome, comparing a baseline column to the immediately following column (which adds the mediator) shows how treatment-effect estimates change when conditioning on the mediator. All specifications include country fixed effects. The *AI Risk* index is constructed from two post-treatment belief items capturing whether the respondent views AI as more risky and whether concerns about AI-driven labor disruption have increased, with higher values representing higher perceived AI risks. Each item is standardized using the pooled control-group mean and standard deviation, averaged (requiring at least one non-missing component), and the resulting average is standardized again. Robust standard errors are reported in parentheses. The sample is restricted to respondents with non-missing values for all four primary outcome indices and a non-missing *AI Risk* index. (***) $p < 0.01$, (**) $p < 0.05$, (*) $p < 0.10$.

Table B.71: Ethical Concerns Causal Mediation Analysis on Main Indices

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)
	Ethical Concerns	AI Adopt.	AI Adopt.	AI Advoc.	AI Advoc.	Staff Intent	Staff Intent	Staff Advoc.	Staff Advoc.
AI Labor	0.157*** (0.053)	-0.424*** (0.060)	-0.415*** (0.060)	-0.463*** (0.059)	-0.442*** (0.059)	-0.209*** (0.055)	-0.232*** (0.055)	-0.002 (0.056)	-0.027 (0.055)
AI Productivity	-0.004 (0.055)	0.036 (0.057)	0.036 (0.057)	0.112** (0.054)	0.111** (0.054)	-0.086 (0.056)	-0.085 (0.056)	-0.060 (0.058)	-0.059 (0.057)
Ethical Concerns			-0.060** (0.029)		-0.137*** (0.027)		0.145*** (0.026)		0.164*** (0.027)
Constant	0.015 (0.043)	0.025 (0.044)	0.026 (0.044)	-0.003 (0.044)	-0.001 (0.044)	0.104** (0.045)	0.102** (0.045)	-0.005 (0.044)	-0.008 (0.044)
Observations	1858.000	1858.000	1858.000	1858.000	1858.000	1858.000	1858.000	1858.000	1858.000

Notes: The table reports OLS estimates from equation (??), where the unit of observation is an individual manager. The dependent variables are the four standardized primary outcome indices described in Section 2.4. The omitted category is the neutral control video. Coefficients on *AI Productivity* and *AI Labor* compare each treatment condition to the control group. For each outcome, odd-numbered columns correspond to the baseline specification (first line of equation (??)), and the adjacent even-numbered columns add the mediator (second line of equation (??)). For each outcome, comparing a baseline column to the immediately following column (which adds the mediator) shows how treatment-effect estimates change when conditioning on the mediator. All specifications include country fixed effects. The *AI Ethics* index is constructed from two post-treatment belief items capturing whether ethical considerations about workforce impacts have become more important in the respondent's decision-making and whether firms should increase consideration for employee welfare when making AI adoption decisions, with higher values representing greater ethical salience and concern for worker welfare. Each item is standardized using the pooled control-group mean and standard deviation, averaged (requiring at least one non-missing component), and the resulting average is standardized again. Robust standard errors are reported in parentheses. The sample is restricted to respondents with non-missing values for all four primary outcome indices and a non-missing *AI Ethics* index. (***) $p < 0.01$, (**) $p < 0.05$, (*) $p < 0.10$.

Table B.72: AI Costs Causal Mediation Analysis on Main Indices

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)
	AI Costs	AI Adopt.	AI Adopt.	AI Advoc.	AI Advoc.	Staff Intent	Staff Intent	Staff Advoc.	Staff Advoc.
AI Labor	-0.094 (0.058)	-0.406*** (0.061)	-0.398*** (0.061)	-0.439*** (0.060)	-0.436*** (0.060)	-0.198*** (0.056)	-0.183*** (0.056)	-0.001 (0.057)	-0.001 (0.057)
AI Productivity	-0.012 (0.058)	0.056 (0.057)	0.057 (0.057)	0.103* (0.055)	0.103* (0.055)	-0.087 (0.058)	-0.085 (0.057)	-0.065 (0.060)	-0.065 (0.060)
AI Costs			0.084*** (0.030)		0.040 (0.027)		0.156*** (0.024)		0.004 (0.023)
Constant	0.003 (0.045)	0.030 (0.045)	0.029 (0.045)	0.010 (0.044)	0.010 (0.044)	0.099** (0.046)	0.099** (0.046)	-0.001 (0.045)	-0.001 (0.045)
Observations	1786.000	1786.000	1786.000	1786.000	1786.000	1786.000	1786.000	1786.000	1786.000

Notes: The table reports OLS estimates from equation (??), where the unit of observation is an individual manager. The dependent variables are the four standardized primary outcome indices described in Section 2.4. The omitted category is the neutral control video. Coefficients on *AI Productivity* and *AI Labor* compare each treatment condition to the control group. For each outcome, odd-numbered columns correspond to the baseline specification (first line of equation (??)), and the adjacent even-numbered columns add the mediator (second line of equation (??)). For each outcome, comparing a baseline column to the immediately following column (which adds the mediator) shows how treatment-effect estimates change when conditioning on the mediator. All specifications include country fixed effects. The *AI Costs* index is constructed from four post-treatment items measuring perceived implementation costs of adopting AI in the respondent's organization: initial investment requirements, time required to realize returns, expected workforce reorganization, and expected process/production reorganization, with higher values representing higher perceived implementation costs. Each item is standardized using the pooled control-group mean and standard deviation, averaged (requiring at least one non-missing component), and the resulting average is standardized again. Robust standard errors are reported in parentheses. The sample is restricted to respondents with non-missing values for all four primary outcome indices and a non-missing *AI Costs* index. (***) $p < 0.01$, ** $p < 0.05$, * $p < 0.10$.

Table B.73: Information Impact Causal Mediation Analysis on Main Indices

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)
	Info. Impact	AI Adopt.	AI Adopt.	AI Advoc.	AI Advoc.	Staff Intent	Staff Intent	Staff Advoc.	Staff Advoc.
AI Labor	-0.245*** (0.060)	-0.426*** (0.060)	-0.385*** (0.062)	-0.461*** (0.059)	-0.411*** (0.061)	-0.206*** (0.055)	-0.188*** (0.055)	-0.001 (0.056)	-0.015 (0.056)
AI Productivity	-0.208*** (0.064)	0.032 (0.057)	0.066 (0.054)	0.110** (0.054)	0.152*** (0.051)	-0.084 (0.056)	-0.069 (0.057)	-0.069 (0.059)	-0.081 (0.059)
Info. Impact			0.167*** (0.026)		0.204*** (0.024)		0.074*** (0.022)		-0.057*** (0.021)
Constant	0.151*** (0.045)	0.024 (0.044)	-0.001 (0.043)	-0.006 (0.044)	-0.037 (0.042)	0.105** (0.045)	0.093** (0.045)	-0.007 (0.044)	0.002 (0.044)
Observations	1862.000	1862.000	1862.000	1862.000	1862.000	1862.000	1862.000	1862.000	1862.000

Notes: The table reports OLS estimates from equation (??), where the unit of observation is an individual manager. The dependent variables are the four standardized primary outcome indices described in Section 2.4. The omitted category is the neutral control video. Coefficients on *AI Productivity* and *AI Labor* compare each treatment condition to the control group. For each outcome, odd-numbered columns correspond to the baseline specification (first line of equation (??)), and the adjacent even-numbered columns add the mediator (second line of equation (??)). For each outcome, comparing a baseline column to the immediately following column (which adds the mediator) shows how treatment-effect estimates change when conditioning on the mediator. All specifications include country fixed effects. The *Information Impact* index is constructed from four post-treatment items capturing respondents' evaluations of the video's informational content: perceived novelty, perceived credibility, perceived bias (reverse-coded so higher values indicate less perceived bias), and agreement with the video's message, with higher values representing greater informational impact and perceived credibility. Each item is standardized using the pooled control-group mean and standard deviation, averaged (requiring at least one non-missing component), and the resulting average is standardized again. Robust standard errors are reported in parentheses. The sample is restricted to respondents with non-missing values for all four primary outcome indices and a non-missing *Information Impact* index. (***) $p < 0.01$, ** $p < 0.05$, * $p < 0.10$.

C Appendix: Survey Instrument, Treatment Scripts, & Final Sample

C.1 Survey Instrument

Consent form:

Thank you for considering participating in this research project related to artificial intelligence in the workplace. We are academic researchers studying how managers use AI and how information about AI might shape adoption and workforce planning.

Your participation in this survey is entirely voluntary. You can choose to not take part, as well as stop the survey, or withdraw from the survey at any time without penalty. We expect that the entire survey will take approximately 15-25 minutes to complete. You will be compensated for your participation in the survey. To ensure high-quality data, our panel provider may review responses for inattentiveness per its standard policy.

You will receive the surveys listed incentive and additionally be entered into a random drawing for five (5) separate \$100 prizes (USD or local equivalent). If you are selected as a winner, you will have the option at the end of the survey to donate part, all, or none of your \$100 prize to one or more nonprofit organizations listed at the end of the survey. Any donation you choose to make will be paid to the organization(s) by the research team, and you will receive the remainder of your prize (i.e., \$100 minus any donation). Participation in the donation decision is voluntary and does not affect your usual survey compensation.

The risks of participating in this study are minimal. However, as with any online survey, there is a small possibility of loss of confidentiality in the event of a data breach, even though the research team does not receive identifying information and security measures are in place. Some survey questions may feel sensitive, and you may feel mild discomfort when reflecting on these topics. You are free to skip any question or stop the survey at any time.

Individual survey results will not be made public and will be used by the research team only to further develop the research project. The data received by the research team will be anonymized. Employers won't see your responses. We do not ask for employer names or any direct identifiers. Any summary survey content that may be included in an academic publication or other academic outlets will also be based on the anonymized data, so that respondents can not be identified.

If you have any questions related to the survey or study, please contact the Principal Investigator, Yong Lee, at yong.s.lee@nd.edu or 574-631-7964.

For questions about your rights as a research participant and to discuss problems, complaints, or concerns about this research study, please contact Notre Dame Research Compliance at

574-631-1461 or at compliance@nd.edu.

Please select your response to the consent:

- I agree and give my consent to participate in this research project. I understand that participation is voluntary and that I may withdraw my consent at any time without penalty.

- I do not agree to participate and will be excluded from the remainder of the questions.

Block 1: About Respondent¹

1. What is your age? Please enter your age as a whole number. [Required]
[blank field] → (if younger than 18, end survey)
2. Which generation do you consider yourself to be part of?
 - Generation Z / Zoomers
 - Millennials / Generation Y
 - Generation X
 - Baby Boomers
 - Silent Generation
 - Unsure
3. What is your gender? [Required]
 - Male
 - Female
 - Non-binary
 - Prefer not to answer
4. Which category best describes your highest level of education?
 - Eighth grade or less
 - Some high school
 - High school diploma or equivalent (completion of secondary school, e.g., GED, A-levels, or similar)
 - Trade school or vocational training
 - Some college
 - Associate's degree or equivalent (typically 2 years of post-secondary education, such as a technical college or junior college program)
 - Bachelor's degree or equivalent (typically 3–4 years of university study)
 - Master's Degree (MA, MS, MBA)
 - Professional Degree (JD, MD)
 - Doctoral Degree
 - Prefer not to say

¹ * Indicates questions for which answer options are randomized. ** Indicates questions for which the answer scale (e.g., Strongly disagree → Strongly agree) is randomly flipped. The orientation of the scale is consistent for each respondent throughout the survey, but may differ across respondents. Regardless of randomization, all "Unsure," "Prefer not to say," and "Other" answer options always appear at the end of the list.

5. Which of the following best describes your current employment status?
- I am employed as a full-time employee
 - I am employed as a part-time employee
 - I am self-employed or a small business owner
 - I am unemployed and looking for work
 - I am unemployed and not looking for work
 - I am a student
6. In what country do you currently reside?*
- United States
 - United Kingdom
 - Other → (end survey)
7. At work, do you have any supervisory responsibilities? In other words, do you have the authority to give instructions to subordinates? [Required]
- Yes
 - No → (end survey)
8. What best describes your managerial role?
- Team Lead or Shift Supervisor – Oversees a small team, typically without full hiring/firing authority
 - Front-Line Manager – Directly manages individual workers; responsible for daily operations and performance
 - Mid-Level Manager – Manages other managers or multiple teams; involved in strategic planning and resource allocation
 - Director or Senior Manager – Oversees multiple departments or large organizational units
 - Executive – Sets organizational strategy and direction at the highest level
 - Other (please specify): _____
9. What is your current job title?
- [Blank Box]
10. Briefly describe your typical workday and main responsibilities (in 50 words or less)
- [Blank Box]

Block 2: About Your Team/Business Unit

We will next ask about your team/business unit, i.e., the immediate group you work with day-to-day.

11. Which category best describes your team / business unit?*

- ☐ Sales / Business Development
- ☐ Marketing / Communications
- ☐ Customer Success / Service
- ☐ Product Management / Category Management
- ☐ Engineering / R&D / Innovation
- ☐ Data & Analytics / Business Insights
- ☐ IT / Infrastructure / Cybersecurity
- ☐ Operations / Supply Chain / Logistics
- ☐ Manufacturing / Production / Quality
- ☐ Finance / Accounting
- ☐ HR / People / Talent
- ☐ Legal / Compliance / Risk
- ☐ Other – please specify

12. What is the total number of employees in your team / business unit?

- ☐ 1-5
- ☐ 6-10
- ☐ 11-20
- ☐ 21-50
- ☐ 51-100
- ☐ Over 100

13. What is the most common occupation in your team / business unit?*

- ☐ Managers
- ☐ Professional
- ☐ Technical
- ☐ Administrative/Clerical
- ☐ Sales
- ☐ Production/Operations
- ☐ Manual Support (janitorial, security)
- ☐ Other (please specify)

14. What is the second most common occupation in your team/unit?*

- Managers
- Professional
- Technical
- Administrative/Clerical
- Sales
- Production/Operations
- Manual Support (janitorial, security)
- Other (please specify)

15. What is the third most common occupation in your team/unit?*

- Managers
- Professional
- Technical
- Administrative/Clerical
- Sales
- Production/Operations
- Manual Support (janitorial, security)
- Other (please specify)

Block 3: Decision Authority, Span of Control, and Internal Dynamics

16. Who makes key decisions about hiring permanent employees in your team / business unit (i.e., the immediate group you work with day-to-day)?**

- The decision is made entirely by corporate leadership, with no input from my team / business unit.
- The decision is made mostly by corporate leadership, with limited input from my team / business unit.
- The decision is equally shared between corporate leadership and my team / business unit.
- The decision is made mostly by my team / business unit, with limited input from corporate leadership.
- The decision is made entirely by my team / business unit, with no input from corporate leadership.

² The occupation selected as the most common occupation will not appear in this selection.

³ The occupations selected as the most and second most common occupations will not appear in this selection.

17. Who has the final authority to approve adopting new technologies (such as AI) in your team / business unit (i.e., the immediate group you work with day-to-day)?**
- The decision is made entirely by corporate leadership, with no input from my team / business unit.
 - The decision is made mostly by corporate leadership, with limited input from my team / business unit.
 - The decision is equally shared between corporate leadership and my team / business unit.
 - The decision is made mostly by my team / business unit, with limited input from corporate leadership.
 - The decision is made entirely by my team / business unit, with no input from corporate leadership.
18. Approximately how many employees do you manage (i.e., report directly to you)?
- 1
 - 2-3
 - 4-6
 - 7-10
 - 11-15
 - More than 15
19. How many management levels are there between you and your organization's CEO or top executive?
- 0
 - 1
 - 2-3
 - 4-5
 - More than 5

Block 4: Firm Policy

Definition of AI: For the purpose of this survey, AI (Artificial Intelligence) refers to technologies that enable machines or software applications to perform tasks typically requiring human intelligence. This includes chatbots (e.g., ChatGPT, Claude, Gemini, etc.), image and video generating AI (e.g., Dall-E, Midjourney, etc.) AI-embedded software applications (e.g., speech-to-text applications, AI-embedded office applications), machine learning algorithms, data analytics, and similar technologies that are commonly used in businesses today.

20. To what extent do top executives in your firm support and encourage the adoption of AI?^{**}

- ☐ A lot
- ☐ Moderately
- ☐ Slightly
- ☐ Not at all
- ☐ Unsure

21. Do any of the following AI-related policies or initiatives currently exist in your company?⁴

[Yes; No; *Unsure*]

- ☐ Formal training programs for AI technologies
- ☐ Active encouragement or incentives for adopting AI
- ☐ Mandatory use of AI tools or technologies
- ☐ Hiring freeze or workforce reduction mandates linked explicitly to AI adoption
- ☐ Guidance or communication about AI initiatives
- ☐ No formal AI initiatives
- ☐ Restrictions or discouragement around AI tool use
- ☐ Other (please specify)

⁴ Here, the order of statements is randomized. The order of scale points is consistently randomly flipped.

22. To what degree is your team / business unit (i.e., the immediate group you work with day-to-day) currently using AI tools in the following areas?⁵

[Widespread Use – Fully integrated into daily work across the function; Routine Use – Used regularly by some staff for relevant tasks; Piloting – Limited adoption of AI or proof-of-concepts in progress; Exploring – Learning about possibilities; no formal adoption yet; None – Not in use; Not Applicable]

- Customer service
- Marketing or communications
- Software or product development
- HR or talent management
- Legal or compliance
- Finance
- Operations
- Strategy or analytics
- IT or technical support
- Training or learning & development
- Research and development
- Other:_____

23. Approximately how often does your team / business unit use AI at work?**

- Every day
- A few times per week
- A few times per month
- Rarely (less than once a month)
- Never
- Unsure

24. Approximately what percent of your team / business unit uses AI at work at least once a week?**

- 100%, that is, everyone
- 50 to 100%
- Less than 50%
- None
- Unsure

⁵ Here, the order of statements is randomized. The order of scale points is consistently randomly flipped.

25. Overall, what best describes your team / business unit's attitude toward AI use at work today? **

- Very positive
- Somewhat positive
- Neither negative nor positive
- Somewhat negative
- Very negative
- Unsure

26. Have you or your teammates received any training in AI? ⁶

[Yes; No; Unsure]

- Myself
- My team

27. Based on your current understanding of your organization's plans, how do you expect AI use to change in your team / business unit's core area (i.e., its primary function) in the next year? **

- Expand substantially – Significant increase in scope, adoption, or investment
- Expand moderately – Gradual or partial increase
- Stay about the same – No major changes expected
- Reduce moderately – Gradual or partial decrease
- Reduce substantially – Significant reduction in scope or use
- Unsure

28. Based on your current understanding of your organization's plans, how do you expect headcount to change in your team / business unit for the three most common roles in the next year? ⁷

[Increase a lot; Increase slightly; No change; Reduce slightly; Reduce a lot; Unsure]

- Occupation group 1 (top occupation group identified in question 13)
- Occupation group 2 (second most prevalent occupation group identified in question 14)
- Occupation group 3 (third most prevalent occupation group identified in question 15)

⁶ Here, the order of the statements and scale points is consistently randomly flipped.

⁷ Here, the order of scale points is consistently randomly flipped.

29. Overall, how do you personally view the impact of AI on your team's work today?**

- Very Positive
- Somewhat Positive
- Neither negative nor positive
- Somewhat Negative
- Very Negative
- Unsure

30. Thinking about your team / business unit, AI is more likely to...**

- Primarily complement/assist workers' tasks
- Both complement and replace about equally
- Primarily replace/substitute some tasks or roles
- Have no meaningful effect
- Unsure

Block 5: Randomized Informational Video

Control: History of AI

Treatment 1: Productivity

Treatment 2: Labor Displacement

Block 6: Post-treatment Key Outcomes (AI Adoption and Hiring Intentions and Advocacy)

31. Please indicate how likely you are to take the following actions within your team / business unit in the next year:⁸

[Extremely likely; Very likely; Moderately likely; Slightly likely; Not at all likely; Unsure]

- Begin adopting new AI technologies not currently in use.
- Increase or accelerate your current use of AI technologies.
- Reduce or slow down your current use of AI technologies.
- Completely halt or avoid new AI adoption.

⁸ Here, both the order of statements and the order of scale points is consistently randomly flipped.

32. To what extent do you agree or disagree with the following:⁷

[*Strongly agree; Agree; Neither agree nor disagree; Disagree; Strongly disagree; Unsure*]

- I would advocate for expanding AI use within my team / business unit, or initiating adoption if it has not yet been implemented.
- I would advocate for maintaining the current level of AI use within my team / business unit
- I would advocate for reducing AI use within my team / business unit

**** If the respondent strongly agrees or agrees they will advocate for expanding AI use within their team or unit ****

33. For what reasons would you advocate expanding AI use within your team / business unit? (Select all that apply)*

- Cost savings from using AI
- Better work quality from using AI
- Time savings or increased production efficiency
- Replacement of certain job functions or tasks
- Other (please specify)

**** If the respondent strongly agrees or agrees they will advocate for adopting fewer AI technologies or maintaining current use ****

34. For what reasons would you advocate for adopting fewer AI technologies or delaying adoption? (Select all that apply)*

- High implementation or maintenance costs
- Concerns about accuracy, reliability, or quality of AI outputs
- Potential negative impact on jobs or employee roles
- Ethical or legal concerns (e.g., bias, privacy, compliance)
- Lack of necessary skills or training within the team
- Other (please specify)

35. If you had the authority or ability to do so, how would you adjust current head count in your team / business unit for the three most common roles within the next year?⁹

[Increase a lot; Increase slightly; No change; Reduce slightly; Reduce a lot; Not relevant]

- Occupation group 1 (top occupation group identified in question 13)
- Occupation group 2 (top occupation group identified in question 14)
- Occupation group 3 (top occupation group identified in question 15)

36. Would your hiring or layoff decisions differ specifically for any of the following groups?¹⁰

[Yes; No; Unsure]

- Entry-level positions
- Workers over 50

**** If the respondents hiring or layoff decisions differ for entry-level positions ****

37. Please briefly explain how your hiring or layoff decisions would differ for entry-level positions and why (in 50 words or less):

[Blank Box]

**** If the respondents hiring or layoff decisions differ for workers over 50 ****

38. Please briefly explain how your hiring or layoff decisions would differ for workers over 50 and why (in 50 words or less):

[Blank Box]

⁹ Here, the order of scale points is consistently randomly flipped.

¹⁰ Here, both the order of statements and the order of scale points is consistently randomly flipped.

39. Suppose you are attending an important annual meeting to discuss your team or department's AI and workforce strategy for the upcoming year. In that meeting...

Would you advocate for or against hiring additional employees in your team/unit in the next year?**

- Advocate **for** hiring additional employees
- Neither advocate for nor against hiring additional employees
- Advocate **against** hiring additional employees
- I would not feel comfortable voicing my opinion on hiring
- Unsure

40. Why? In 50 words or less, what is the main reason for your answer?

[Blank Box]

41. Suppose you are attending an important annual meeting to discuss your team or department's AI and workforce strategy for the upcoming year. In that meeting...

Would you advocate for or against layoffs or workforce reductions in your team/unit in the next year?**

- Advocate **for** layoffs or workforce reductions
- Neither advocate for nor against layoffs
- Advocate **against** layoffs or workforce reductions
- I would not feel comfortable voicing my opinion on layoffs or workforce reductions
- Unsure

42. Why? In 50 words or less, what is the main reason for your answer?

[Blank Box]

Block 7: Post-treatment Mechanisms (Channels)

43. Please indicate your level of agreement with the following statements about your reasoning after seeing the informational video:¹¹

[*Strongly agree; Agree; Neither agree nor disagree; Disagree; Strongly disagree; Unsure*]

- AI seems more beneficial than I previously thought.
- AI seems more harmful or risky than I previously thought.
- My expectations about future productivity benefits from AI have increased.
- My concerns about AI-driven labor disruption have increased.
- Ethical considerations about workforce impacts have become more important in my decision-making.
- Firms should increase their consideration for employee welfare when making decisions about AI adoption.

44. Please briefly explain in your own words how the informational video **specifically** influenced your thinking regarding AI adoption and hiring:

[*Blank Box*]

45. Please indicate your level of agreement with the following statements about the informational video you watched:¹⁰

[*Strongly agree; Agree; Neither agree nor disagree; Disagree; Strongly disagree; Unsure*]

- The video provided new information that I was not previously aware of.
- The statistics cited in the video are credible and believable.
- The video's information is biased.
- I agreed with the video's message.

46. Please rate the expected magnitude of the following requirements for effective AI adoption in your organization:¹⁰

[*Very high; High; Moderate; Low; Very low; Unsure*]

- Initial investment required to begin AI adoption
- Time required to realize returns on effective AI adoption
- Changes required in tasks, roles, and team / business unit organization for effective AI adoption
- Changes required in production or operational processes for effective AI adoption

¹¹ Here, both the order of statements and the order of scale points is consistently randomly flipped.

Block 8: Post-treatment Individual Characteristics

47. What was your total personal income, before taxes, last year?

- \$0-\$9,999
- \$10,000-\$14,999
- \$15,000-\$19,999
- \$20,000-\$29,999
- \$30,000-\$39,999
- \$40,000-\$49,999
- \$50,000-\$69,999
- \$70,000-\$89,999
- \$90,000-\$109,999
- \$110,000-\$149,999
- \$150,000-\$199,999
- \$200,000 +

48. In political views, where do you see yourself on the liberal / conservative spectrum?**

- Very liberal
- Liberal
- Moderate
- Conservative
- Very Conservative

49. Have you ever been unemployed during your work life? Note: Excluding temporarily dormant employment (e.g., longer periods of sickness).**

- Yes
- No
- Unsure

*** If the respondent has experienced a period of unemployment in their work life ***

50. Which best describes the length of your longest period of unemployment?

- Less than a month
- 1 month
- 2 months-5 months
- 6 months-1 year
- Over 1 year

51. Has anyone in your immediate family been unemployed during your life? Note:
Excluding temporarily dormant employment (e.g., longer periods of sickness).**

- ☐ Yes
- ☐ No
- ☐ Unsure

Block 9: Post-treatment Firm Characteristics

52. Approximately how many people are employed at your organization (including all locations and business units)?

- Fewer than 10 employees
- 10-24 employees
- 25-49 employees
- 50-99 employees
- 100-249 employees
- 250-499 employees
- 500-999 employees
- 1,000-4,999 employees
- 5,000-9,999 employees
- 10,000 or more employees
- Unsure

53. How long has your organization been in operation?

- Fewer than 1 year
- 1-4 years
- 5-9 years
- 10-24 years
- 25 years or more
- Unsure

54. Which of the following best describes the structure of your organization?**

- Privately Owned
- Publicly Traded
- Worker Owned (e.g., co-op)
- Government Owned
- Non-profit organization
- Other (please specify)

55. Which of the following best describes the primary industry in which your organization operates? (If your organization operates in multiple industries, please select the one most relevant to your role.)

[Drop-down menu with 20 Naics Sectors]

56. What is the approximate annual revenue of your entire firm (company-wide):

- ☐ Less than \$99,999
- ☐ \$100,000 - \$499,999
- ☐ \$500,000 - \$999,999
- ☐ \$1 million – \$4.9 million
- ☐ \$5 million – \$9.9 million
- ☐ \$10 million – \$24.9 million
- ☐ \$25 million – \$49.9 million
- ☐ \$50 million – \$99.9 million
- ☐ \$100 million – \$249.9 million
- ☐ \$250 million – \$499.9 million
- ☐ \$500 million – \$999.9 million
- ☐ \$1 billion – \$4.9 billion
- ☐ \$5 billion – \$9.9 billion
- ☐ \$10 billion or more
- ☐ Unsure / Prefer not to say

57. Approximate total annual budget allocated to your specific unit or department:

- ☐ Less than \$10,000
- ☐ \$10,000 - \$24,999
- ☐ \$25,000 - \$49,999
- ☐ \$50,000 - \$99,999
- ☐ \$100,000 – \$249,999
- ☐ \$250,000 – \$499,999
- ☐ \$500,000 – \$999,999
- ☐ \$1 million – \$2.49 million
- ☐ \$2.5 million – \$4.9 million
- ☐ \$5 million – \$9.9 million
- ☐ \$10 million – \$24.9 million
- ☐ \$25 million – \$49.9 million
- ☐ \$50 million or more
- ☐ Unsure / Prefer not to say

Block 10: Effort check, Debrief, and Thank You

58. It is vital to our study that we only include responses from people who devoted their full attention to this study. Otherwise, years of effort (the researchers' and the time of other participants) could be wasted. You will receive credit for this study, no matter what; however, please tell us how much effort you put forth towards this study.**

- I put forth almost no effort
- I put forth very little effort
- I put forth some effort
- I put forth quite a bit of effort
- I put forth a lot of effort

59. Also, often there are several distractions present during studies (other people, TV, music, etc.). Please indicate how much attention you paid to this survey. Again, you will receive credit no matter what. We appreciate your honesty!**

- I gave this study almost no attention
- I gave this study very little attention
- I gave this study some of my attention
- I gave this study most of my attention
- I gave this study my full attention

60. *This section does not affect your standard survey compensation.*

By taking this survey, you are automatically enrolled in a lottery to win \$100. In a few weeks, you will know whether you won the \$100. The payment will be made to you in the same way as your regular survey pay, so no further action is required on your part. In case you won, would you be willing to donate part or all of your \$100 gain towards non-profit organizations that work directly on issues related to AI development or supporting workers? If you are one of the lottery winners, the additional compensation will be distributed according to your elected allocations below. We will pay your desired donation amount to the non-profit(s) of your choosing. Please enter how to allocate your prospective \$100 gain in whole dollar amounts.

Non-Profit Organizations

Allen Institute for Artificial Intelligence: The Allen Institute develops foundational AI research and innovation to deliver real-world impact through large-scale open models, data, robotics, conservation, and beyond.

FHI 360: FHI 360's economic opportunity programs partner with local institutions, organizations and businesses to equip job seekers with in-demand skills through training and certification and guide them to open positions.

Please note, if the sum of your allocations is more than \$100, we will distribute the \$100 according to the implied proportions. If no allocation is entered, we will equally distribute the proceeds.

- Donate to the Allen Institute for Artificial Intelligence [Blank Box]
- Donate to FHI 360 [Blank Box]
- Paid to you personally [Blank Box]

Thank you for participating. This study aims to understand how managers think about AI, workforce planning, and hiring decisions. All of your responses will remain strictly confidential and anonymous. Your input is essential to improving our understanding of how organizations and managers respond to new technologies. Thank you again for your time and insights.

If you have any questions related to the survey, or would like to withdraw from the study, please contact the Principal Investigator, Yong Lee, at yong.s.lee@nd.edu or 574-631-7964.

For questions about your rights as a research participant and to discuss problems, complaints, or concerns about this research study, please contact Notre Dame Research Compliance at 574-631-1461 or at compliance@nd.edu.

C.2 Treatment Scripts

Control Group: AI Firms

Artificial Intelligence, or AI, refers to technologies that enable machines or software to perform tasks typically requiring human intelligence. Today, a number of organizations around the world develop and provide AI tools for various applications.

OpenAI, based in San Francisco, California, created ChatGPT. It's credited with catalyzing widespread interest in AI and was first released in late 2022 ([The New York Times 2022](#)). OpenAI also collaborates closely with Microsoft, whose Copilot tools integrate AI into applications like Word, Excel, and Outlook.

Anthropic, also based in San Francisco, develops the Claude family of AI assistants. Claude's coding capabilities are especially popular among software developers.

Google DeepMind, headquartered in London, England, is Google's AI research and development laboratory. They are responsible for developing Google's AI products, such as Gemini and Flow, and supporting the integration of AI features into Google's search engine ([Google DeepMind 2025](#)).

Meta, also based in California, has developed the Llama family of AI models. Compared to other major labs, Meta openly shares critical components of its Llama models, enabling developers to customize them for their own use cases.

DeepSeek, an AI developer based in China, has drawn international attention for creating skillful AI systems at a fraction of the cost of other major AI models ([CNN 2025](#)).

Outside of large technology companies, universities and startups also contribute to AI research and development. Tech companies in other countries are developing AI models as well.

Together, these firms and institutions represent a diverse and evolving AI landscape, shaping the development of this technology.

Manipulation check:

Based on the information you saw, what was the main message communicated?*

- AI technology and capabilities have been advancing over time. (**Correct**)
- AI is a form of machine learning used primarily for marketing.
- AI was first invented by a group of companies in the 1990s.

Treatment 1: Productivity Benefits of AI

Artificial Intelligence, or AI, refers to technologies that enable machines or software to perform tasks typically requiring human intelligence, and it has been developing rapidly in recent years.

With this development, AI is transforming productivity across industries. Managers and firms adopting AI experience substantial productivity improvements. Recent rigorous studies demonstrate clear and measurable benefits:

- In software development, programmers using AI-assistants complete coding tasks 56% faster compared to their peers without AI. (Peng et al. 2023)
- In professional writing, AI-driven writing assistance tools generate 40% time savings. (Noy & Zhang 2023)
- In the customer support field, implementing AI chatbots and support tools has increased productivity by 14%, allowing employees to focus on higher-value tasks. (Brynjolfsson, Li, and Raymond 2023)

Leading economic analyses highlight the significant broader impact of AI:

- Goldman Sachs predicts AI adoption could boost global GDP by \$7 trillion, or approximately 7%, over the next decade. (Goldman Sachs 2023)
- McKinsey & Company forecasts annual economic benefits ranging from \$2.6 to \$4.4 trillion, driven by efficiency improvements across various business processes. (McKinsey & Company 2023).

Experts widely agree that current productivity gains from AI are just the beginning. As AI technologies become more advanced, accessible, and integrated, firms are expected to achieve even greater efficiency, task optimization, and cost savings. Adopting AI today could position your team and organization at the forefront of productivity and competitiveness.

Manipulation check:

Based on the information you saw, what was the main message communicated?*

- AI adoption is being mandated by corporate policy.
- AI can significantly improve productivity and efficiency across industries. (**Correct**)
- AI is not yet ready for widespread use in most business functions.

Treatment 2: Negative Labor Market Impacts of AI

Artificial Intelligence, or AI, refers to technologies that enable machines or software to perform tasks typically requiring human intelligence, and it has been developing rapidly in recent years.

With this development, AI is leading to significant workforce restructuring. The spread of AI is already impacting employment, with numerous companies announcing layoffs and hiring freezes. Recent rigorous studies demonstrate clear and measurable changes to the labor market:

- Leading research suggests that AI is capable of automating over half of the tasks performed in 46% of all current occupations. (Elondou et al. 2024)
- Studies of earlier automation technology estimate that a single industrial robot reduces the local employment-to-population ratio and wages by about one percent. (Acemoglu & Restrepo 2020)
- After the release of ChatGPT, early career workers in the most AI exposed occupations experienced a 13% decline in employment compared to their least exposed peers. (Brynjolfsson et al. 2025)

Leading economic analyses highlight the significant broader impact of AI:

- Goldman Sachs estimates that approximately 300 million jobs globally could be exposed to automation by AI, with white-collar and knowledge-based roles facing the greatest risk. (Goldman Sachs 2023)
- The CEO of Ford anticipates that AI will replace over half of all white collar jobs in America (Wall Street Journal 2025).

Experts emphasize that AI's rapid development and integration into the workplace will continue to reshape employment landscapes, leading to job displacement, increased uncertainty for workers, and potentially extensive unemployment in many sectors. Adopting AI today could displace certain roles within your team and risk unemployment for affected employees.

Manipulation check:

Based on the information you saw, what was the main message communicated?*

- AI is mostly used for marketing and customer service tasks.
- AI has the potential to cause widespread job losses and increase the risk of labor market disruption. (**Correct**)
- AI is not expected to affect white-collar employment.

C.3 Analytic Sample and Exclusions

We begin with all respondents who entered the Qualtrics survey. Eligibility was enforced prior to randomization (age ≥ 18 , residence in the United States or the United Kingdom, and current supervisory responsibility). Respondents who failed any eligibility check were exited from the survey before the video module and therefore were never assigned to a treatment condition. All analyses are conducted on respondents who were assigned to one of the three video conditions.

Our main analyses use a consistent analytic sample across the four primary outcomes. Specifically, we restrict attention to respondents with non-missing values for all four primary indices (AI Adoption, AI Advocacy, Staffing Intentions, and Staffing Advocacy). This restriction holds the estimation sample fixed across outcome regressions and simplifies comparisons across outcomes.

Of the 1995 who reached the randomization portion of the survey, 1863 had non-missing values for all four primary indices. Primary indices are treated as missing if over half of an index's items are missing for that respondent.

For those missing primary indices, 21 were only missing the AI adoption index, 15 were only missing the AI advocacy index, 10 were only missing the staffing intentions index, and 70 were only missing the staffing advocacy index. The remaining 16 individuals missing a primary index were missing more than one primary index. Exactly 66 people in each country are missing a primary index and are excluded from the data set.

Of the 1863 individuals who reached randomization and have values for the four primary indices, 934 were from the US, while the remaining 929 were from the UK. Unless otherwise noted, all main results use the analytic sample of 1,863 respondents.

D Appendix: Balance Checks

As noted in Section 2, our analytic sample of managers in the United States and United Kingdom was collected through a survey platform. While the platform's survey panel is meant to be nationally representative, conditioning potential respondents and filtering sampled respondents by managers does not preserve national representativeness of the target population of interest (i.e., all managers in the respective countries), but our sample statistics demonstrate broad coverage across demographics, business functions, and industries. This section presents descriptive statistics for three purposes: 1. to convey the qualities and coverage of our analytic sample to inform the extent of this study's external validity; 2. to establish new stylized facts on AI perceptions and adoptions within and across the US and UK; and 3. to report on the balance of baseline respondent characteristics within countries to cement the internal validity of our randomized experiment.

Table A.1 presents descriptive statistics and US–UK comparisons on baseline respondent characteristics.¹⁸ Overall, demographic and role/team characteristics are more similar than different across the two countries. Managers in our sample are mostly middle-aged adults (30-50 years old), well educated (some college to masters), and around half of respondents identify as a man. There are some minor cross-country differences, as sampled US managers are slightly older, more likely full-time, and notably more likely to

¹⁸Appendix tables in Section ?? show survey statistics for respondent's role, function, and firm characteristics, respectively.

have experienced unemployment personally or through a family member.¹⁹ In both countries, on average, responding managers are about three levels away from the top of their organization, with US managers slightly closer. The characteristics of respondent's role / business unit varies greatly throughout our analytic sample reflecting expansive coverage, and between countries, US managers only had an average of 1 more direct report while overall team / business unit size in employment and budget (nominally)²⁰ are similar. Overall, these observable demographic and role characteristics reflect some differences between countries in our analytic sample, but more so establish wide coverage of national manager populations.

Following the same format, Table A.2 presents statistics on survey respondents' own, team's, and firm's AI posture. In both countries, we find that respondents baseline sentiments on AI lean positive and more than 75% have perspectives that includes AI capacity to augment human labor. However, UK respondents are slightly less bullish, and these patterns (sign and difference) are further reflected in respondents' team and firm AI-related postures, from team sentiment and use, and firm-level AI policies, trainings, and executive support.

Without establishing the national representativeness of survey respondents, it is speculation to use these subsample differences as stylized facts on cross country comparisons. However, these descriptive statistics support the likely notion that there are real differences, and that information on AI may have heterogeneous effects on managers decisions between the US and UK.

While there are significant differences between countries, reassuringly, we find little imbalance between randomized survey arms within each country. This naturally follows from our survey instrument's stratified randomization by country (and age & gender). Descriptive statistics on AI posture within country between arms are presented in Tables A.3 and A.4 for the US and UK, respectively.²¹ Within the US, we see our respondents treated with the AI productivity information report less AI training (self and team) and less AI positive policies at their firms (e.g. related to AI training, encouragement, and guidance); all of which are correlated outcomes with similar extents of imbalance (about 20% of the standard deviation for the control arm's corresponding response). Within the UK, we see respondents in either treatment arm report lower firm AI encouragement relative to control respondents and those treated with the AI induced labor distress information report lower self and team AI sentiments. Overall, these tables suggest that the stratified survey arm randomization did not perfectly balance respondents across all characteristics, and any unconditional differences between arms within country should be checked for robustness to controls. Otherwise, the descriptive statistics exhibit managers with a wide variety of baseline characteristics, cross-country differences within our sample, and predominantly balance characteristics within national subsamples across treatment arms.

¹⁹Income and self-reported political ideological alignment (liberal-conservative) are also statistically different, but structurally different across national contexts.

²⁰Given the team budget size was asked in USD, this may imply UK teams have a larger budget per person. This characteristic may also be structurally different across national contexts.

²¹Appendix tables A.30 and A.31 report baseline respondent/team characteristics for the US and UK, respectively.

E Appendix: Additional Grouping and Index Definitions

E.1 Heterogeneity Subgroup Definitions

Expected Staffing Plans Next Year: We measure respondents' pre-treatment staffing plans based on three five-point items that ask how staffing is expected to change for the most common occupations on the respondent's team, ranging from "Reduce Hi" (1) to "Expand Hi" (5), with "No Change" (3) as the midpoint. We construct a composite staffing plan measure by taking the mean of the three occupation-specific responses, requiring at least two of the three items to be non-missing. In our estimation sample ($N = 1,826$ for the composite), the median value of the composite is 3.00, and the distribution is concentrated around 3 (34.28% of observations take the value 3.00), with additional mass at values above 3.

We define the high pre-treatment staffing plan group as respondents with a composite value strictly greater than 3, which corresponds to 47.04% of the sample. The low group includes respondents with a composite value less than or equal to 3, which corresponds to 52.96% of the sample. Conceptually, membership in the high group indicates that the respondent anticipates net staffing expansion (on average across the team's most common occupations), while the low group captures respondents expecting no change or net reductions. We use the > 3 cutoff because 3 represents "No Change" on the underlying scale, so values above 3 reflect expected expansion rather than stable staffing.

Team AI Use Frequency: We measure the frequency of AI use on respondents' teams using five ordered categories: Never, Rarely, Few/month, Few/week, and Daily. In our estimation sample ($N = 1,852$), the median response is "Few/week." The distribution of the underlying numeric coding is: 0 (Never; 4.43%), 0.05 (Rarely; 9.34%), 0.1 (Few/month; 12.90%), 0.5 (Few/week; 33.96%), and 1 (Daily; 39.36%).

We define the high AI use frequency group as respondents reporting Daily use (coded as 1), which corresponds to 39.36% of the sample. The low AI use frequency group includes all responses below Daily (Never through Few/week; coded 0 through 0.5), which corresponds to 60.64% of the sample. Conceptually, membership in the high group captures teams that use AI on a daily basis.

Respondent's Pre-Treatment AI Views: We measure the respondents' pre-treatment attitudes toward AI using a five-point sentiment scale: "Negative Hi" (1), "Negative Lo" (2), "Neutral" (3), "Positive Lo" (4), and "Positive Hi" (5). In our estimation sample ($N = 1,843$), the median response is 4 ("Positive Lo"). The distribution is: Negative Hi (1.47%), Negative Lo (4.07%), Neutral (18.18%), Positive Lo (46.77%), and Positive Hi (29.52%).

We define the high self AI attitude group as respondents reporting values 4–5 ("Positive Lo" or "Positive Hi"), which corresponds to 76.29% of the sample. The low self AI attitude group includes values 1–3 (Negative Hi/Lo or Neutral), which corresponds to 23.71% of the sample. Conceptually, membership in the high group indicates a positive baseline orientation toward AI prior to treatment, while

the low group captures respondents who are neutral or negative toward AI.

Technology Authority: We measure team technology decision authority using a five-point scale: “Corporate only” (1), “Corporate mostly” (2), “Shared” (3), and “Team mostly” (4) “Team only” (5). In our estimation sample ($N = 1,851$), the median response is 3 (“Shared”). The distribution is: Corporate only (14.37%), Corporate mostly (33.71%), Shared (26.47%), Team mostly (16.91%), and Team only (8.54%).

We define the high tech authority group as respondents reporting values 4–5 (“Team mostly” or “Team only”), which corresponds to 25.45% of the sample. The low tech authority group includes values 1–3 (“Corporate only,” “Corporate mostly,” or “Shared”), which corresponds to 74.55% of the sample. Conceptually, membership in the high group indicates that technology decisions are primarily made at the team level rather than centralized at corporate leadership or shared evenly across levels.

Executive Support for AI: We measure perceived executive support for AI adoption using a four-point scale: “Not at all” (1), “Slightly” (2), “Moderately” (3), and “A lot” (4). In our estimation sample ($N = 1,825$), the median response is 3 (“Moderately”). The distribution is: Not at all (4.71%), Slightly (20.16%), Moderately (40.33%), and A lot (34.79%).

We define the high executive support group as respondents reporting values 3–4 (“Moderately” or “A lot”), which corresponds to 75.12% of the sample. The low executive support group includes values 1–2 (“Not at all” or “Slightly”), which corresponds to 24.88% of the sample. Conceptually, membership in the high group indicates that respondents perceive leadership as broadly supportive of AI adoption, while the low group captures environments where leadership support is limited.

Hiring Authority: We measure team hiring decision authority using a five-point scale: “Corporate only” (1), “Corporate mostly” (2), “Shared” (3), “Team mostly” (4), and “Team only” (5). In our estimation sample ($N = 1,860$), the median response is 3 (“Shared”). The distribution is: Corporate only (6.61%), Corporate mostly (20.97%), Shared (31.77%), Team mostly (25.43%), and Team only (15.22%).

We define the high hiring authority group as respondents reporting values 4–5 (“Team mostly” or “Team only”), which corresponds to 40.65% of the sample. The low hiring authority group includes values 1–3 (“Corporate only,” “Corporate mostly,” or “Shared”), which corresponds to 59.35% of the sample. Conceptually, membership in the high group indicates that hiring decisions are primarily made at the team level rather than centralized at corporate leadership or shared evenly across levels.

Team AI Use Proportions: We measure team AI use as the proportion of team members who use AI at least once per week. In our estimation sample ($N = 1,820$), the median response is 0.75. The distribution of the underlying numeric coding is: 0 (5.44%), 0.25 (33.13%), 0.75 (45.38%), and 1 (16.04%). In the categorical version of this measure, responses correspond to None (0%), < 50%, 50–100%, and All (100%).

We define the high team AI use group as respondents reporting 0.75 or 1 (at least 50% of team members use AI weekly), which corresponds to 61.42% of the sample. The low team AI use group includes responses below 0.75 (less than 50% using AI weekly), which corresponds to 38.58% of the sample. Conceptually, membership in the high group captures teams where 50% or more of the team members use AI weekly.

Team AI Attitudes: We measure respondents' pre-treatment perceptions of their teams' attitudes toward AI using a five-point sentiment scale: "Negative Hi" (1), "Negative Lo" (2), "Neutral" (3), "Positive Lo" (4), and "Positive Hi" (5). In our estimation sample ($N = 1,838$), the median response is 4 ("Positive Lo"). The distribution is: Negative Hi (2.23%), Negative Lo (6.37%), Neutral (18.06%), Positive Lo (45.32%), and Positive Hi (28.02%).

We define the high team AI attitude group as respondents reporting values 4–5 ("Positive Lo" or "Positive Hi"), which corresponds to 73.34% of the sample. The low team AI attitude group includes values 1–3 (Negative Hi/Lo or Neutral), which corresponds to 26.66% of the sample. Conceptually, membership in the high group indicates that the respondent perceives their team as having a positive baseline orientation toward AI prior to treatment, while the low group captures teams perceived as neutral or negative.

Pre-Treatment AI Adoption Plans: We measure respondents' pre-treatment AI adoption plans for the next year using a five-point scale: "Reduce Hi" (1), "Reduce Lo" (2), "No Change" (3), "Expand Lo" (4), and "Expand Hi" (5). In our estimation sample ($N = 1,817$), the median response is 4 ("Expand Lo"). The distribution is: Reduce Hi (0.55%), Reduce Lo (2.48%), No Change (24.93%), Expand Lo (54.76%), and Expand Hi (17.28%).

We define the high pre-treatment adoption plan group as respondents reporting values 4–5 ("Expand Lo" or "Expand Hi"), which corresponds to 72.04% of the sample. The low group includes values 1–3 (Reduce Hi/Lo or No Change), which corresponds to 27.96% of the sample. Conceptually, membership in the high group indicates that the respondent expects AI adoption on their team to expand over the next year, while the low group captures respondents expecting stable or reduced adoption.

E.2 Construction of Belief Indices

This appendix describes the construction of the post-treatment mechanism indices used in the analysis. We construct five indices: *AI Benefits*, *AI Risks*, *Ethics*, *Costs*, and *Information Impact*. All indices are designed so that higher values correspond to more of the underlying concept. In the decomposition analysis reported in the main text, we focus on the *AI Benefits* and *AI Risks* indices. We also construct the *Ethics*, *Costs*, and *Information Impact* indices for completeness and exploratory analysis, but these additional indices contribute little in the decomposition exercises and are omitted from the reported figures and tables for parsimony.

For the *AI Benefits*, *AI Risks*, *Ethics*, and *Information Impact* indices, component items are measured on five-point Likert scales, with response options ranging from strongly agree to strongly disagree. For the *Costs* index, component items are measured on a five-point scale ranging from very high to very low. In all cases, an “unsure” response is treated as missing.

Each index is constructed in three steps. First, each component item is coded so that higher values reflect more of the intended construct. All items are straight-coded except for the *Information Impact* item on perceived bias, which is reverse-coded so that higher values correspond to lower perceived bias and therefore greater informational impact. Second, each component item is standardized using the pooled control-group mean and standard deviation. Third, the standardized component items are averaged to form a raw index, requiring at least one non-missing component item. The resulting raw index is then standardized again using the pooled control-group mean and standard deviation. As a result, each final index is centered and scaled relative to the control-group distribution, and higher values consistently indicate a greater level of the underlying mechanism.

The AI Benefits Index: The *AI Benefits* index is intended to capture the extent to which the treatment increases perceived upside or productivity-related benefits from AI. It is constructed from the following two post-treatment items:

1. *AI seems more beneficial than I previously thought.*
2. *My expectations about future productivity benefits from AI have increased.*

Both items are straight-coded. Higher values of the index indicate a greater perceived increase in the benefits of AI.

AI Risks Index: The *AI Risks* index is intended to capture the extent to which the treatment increases perceived harms, risks, or labor-displacement concerns associated with AI. It is constructed from the following two post-treatment items:

1. *AI seems more harmful or risky than I previously thought.*
2. *My concerns about AI-driven labor disruption have increased.*

Both items are straight-coded. Higher values of the index indicate greater perceived risk from AI.

Ethics Index: The *Ethics* index is intended to capture the extent to which respondents place greater weight on ethical or employee-welfare considerations when thinking about AI adoption. It is constructed from the following two post-treatment items:

1. *Ethical considerations about workforce impacts have become more important in my decision-making.*
2. *Firms should increase their consideration for employee welfare when making decisions about AI adoption.*

Both items are straight-coded. Higher values of the index indicate greater emphasis on ethical or employee-welfare considerations in AI-related decision-making.

Costs Index: The *Costs* index is intended to capture perceived implementation costs or barriers associated with adopting AI. It is constructed from the following four post-treatment items:

1. *Initial investment required to begin AI adoption*
2. *Time required to realize returns on effective AI adoption*
3. *Changes required in tasks, roles, and team / business unit organization for effective AI adoption*
4. *Changes required in production or operational processes for effective AI adoption*

All four items are straight-coded. Higher values indicate higher perceived costs, longer implementation horizons, or greater required organizational or operational change associated with effective AI adoption.

Information Impact Index: The *Information Impact* index is intended to capture the extent to which respondents view the informational video treatment as novel, credible, persuasive, and not biased. It is constructed from the following four post-treatment items:

1. *The video provided new information that I was not previously aware of.*
2. *The statistics cited in the video are credible and believable.*
3. *The video's information is biased.*
4. *I agreed with the video's message.*

The first, second, and fourth items are straight-coded. The third item, *The video's information is biased*, is reverse-coded so that higher values correspond to lower perceived bias. Higher values of the index therefore indicate greater perceived informational impact, including higher novelty, credibility, agreement, and lower perceived bias.

