

The Macroeconomics of Automation: Sorting Frictions and the Asymmetry Between Robots and AI*

Carlos Chávez[†]

February 2026

Abstract

Does automation rebuild work around the people it displaces, and does the answer differ for robots and AI? Reinstatement and sorting rates come from commuting-zone moments under design-based inference, the physical rate bracketed by Danish evidence. I model both frontiers advancing together, each reinstating tasks it destroys, workers sorting across the rest. Cognitive displacement scrambles sorting, a friction the physical frontier lacks. Robots deliver more productivity per unit of advance than AI in every sorting-parameter draw, in 98.5 percent adding the physical rate's sampling error, and under either estimate of AI's disputed reinstatement rate. The point-estimate gap is about a quarter. The steady-state ranking depends on how transitional sorting is, which I bound. AI raises the skill premium and does twice the labor-share damage. A planner subsidizes robot research at 1.4 to 2.1 times the AI rate across the central ninety percent, and a uniform automation tax locally lowers welfare.

JEL: E24, J23, J31, O33, O41

Keywords: Automation, reinstatement, productivity, labor share, robots, artificial intelligence, directed technical change, optimal policy

*I am grateful to James Heckman for guidance and support. All errors are my own. An earlier version of this paper was titled “The Macroeconomics of Asymmetric Automation: Frontier Technologies, Productivity, and the Labor Share.”

[†]Center for Economics of Human Development, University of Chicago. Email: cchavezpadilla@uchicago.edu

1 Introduction

When a technology automates human work, does it rebuild new work around the people it displaces? For robots and AI the answers differ, and the difference is large enough to reorder productivity, the labor share, and optimal innovation policy. Yet the two are usually analyzed as a single phenomenon called “automation,” one frontier advancing at one rate, with one parameter for how much complementary work the technology creates. Collapsing the two frontiers into one discards information that is quantitatively first-order.

Physical and cognitive automation differ in what displacement does to the match between workers and the tasks that get rebuilt. When a manufacturing firm installs a robot, the surrounding production process is reorganized: quality-control positions, maintenance roles, and programming tasks emerge alongside the displacement, and displaced production workers slot into them. When an AI system automates cognitive work (document review, data classification, routine analysis), the surrounding reorganization creates tasks at a similar rate, but the displacement scrambles the sorting of skilled workers across the remaining cognitive occupations, and the match that emerges is worse than the one it replaced. The question is how large this sorting difference is and what it implies for aggregate outcomes.

I build a two-frontier growth model in which physical and cognitive automation advance simultaneously, each governed by its own reinstatement rate and interacting with a shared labor market where workers sort across sectors according to comparative advantage. Productivity, labor shares, skill premia, and optimal policy all follow from a single growth decomposition that separates frontier-specific contributions from general equilibrium corrections.

The cognitive reinstatement rate, the number of complementary tasks AI creates per task it automates, is estimated by minimum distance on six wage and employment moments from 722 U.S. commuting zones (CZs), building on the two-frontier identification framework of [Chávez \(2026\)](#). The estimate is just under one: AI creates roughly ten complementary tasks for every ten it displaces. The same estimation identifies a sorting disruption parameter near two, capturing how AI degrades the matching of skilled workers to remaining cognitive tasks. Inference respects the shock structure of both exposure measures, which differ sharply: the robot side is effectively a single large natural experiment in European automotive robotization, while the AI side draws on roughly

eighty effective occupation-level shocks. The physical reinstatement rate is not estimated here. It is the commuting-zone estimate of [Chávez \(2026\)](#), recovered from the same U.S. data by a different moment system, and it lies within the range implied by [Humlum \(2022\)](#)’s firm-level analysis of Danish robot adopters. Humlum’s firms cut production-task value by a fifth yet expand total labor demand, so task creation roughly paces displacement, and the agreement of Danish firm-level, U.S. commuting-zone, German, and Spanish evidence suggests the reinstatement regime is a property of the technology, not of local institutions.

The headline finding is that robot automation is more productive per unit of frontier advance than AI automation, a ranking that holds under both this paper’s cognitive estimates and the near-zero estimates of [Chávez \(2026\)](#). At this paper’s estimates the gap is about a quarter, and it widens where AI rebuilds little. What carries the asymmetry depends on the calibration: at this paper’s estimates the two technologies reinstate at comparable rates and a cognitive-specific sorting friction accounts for the gap, while under Chávez’s estimates the reinstatement gap accounts for it. Either way the drag falls on the cognitive frontier. Robot dominance holds in all 200 bootstrap draws of the jointly identified sorting parameters, and in 98.5 percent of draws once the sampling error of the imported physical reinstatement rate is added, over the contemporaneous-to-medium-run horizon, the horizon over which research and development policy is set. Whether it persists into the steady state depends on how much of the measured sorting disruption is transitional adjustment and how much is permanent friction, a question the data cannot settle and the paper bounds instead.

Three macro results follow.

Supply-side dominance. The distributional consequences of asymmetric automation are determined almost entirely by the supply-side productivity asymmetry, not by aggregate demand channels. I solve the model with non-tradable services, endogenous labor supply and heterogeneous marginal propensities to consume, and compute each demand-side channel rather than bounding it. The total general equilibrium feedback is more than two orders of magnitude below the level at which it would matter, on both frontiers. On the robot frontier, two of the three additional demand-side channels dampen the feedback and the net demand-side correction is negative; on the cognitive frontier, the net correction is positive but negligible. The distributional consequences of asymmetric automation are therefore supply-side determined.

Labor share dynamics. Both automation frontiers lower the labor share, but with asymmetric

magnitudes. Robots reduce labor’s income share roughly half as much as AI per unit of frontier advance, and this ratio is the most tightly pinned down of the paper’s headline magnitudes. At this paper’s estimates, the gap traces to the same sorting friction, which degrades the productivity the cognitive frontier extracts from its tasks, so AI weighs more heavily on labor’s share. Both frontiers widen the skill premium, and the skill-premium response to AI is the paper’s sharpest empirical moment, significant under exposure-robust inference and the within-skill-bin permutation null, marginal under the harsher task-residualized null. Because both technologies lower the labor share, the margin that matters is which damages it least.

Optimal frontier-specific policy. Over the transition horizon, a social planner subsidizes robot R&D at between 1.4 and 2.1 times the rate of AI R&D across the central ninety percent of the estimated parameter set, reflecting an efficiency motive (robots generate more productivity per frontier advance) and an equity motive (robot-created tasks disproportionately benefit non-college workers, who have higher marginal utility in the welfare function). This tilt is a transition-horizon result. As AI matures its productivity content rises, and under every development path considered the ranking eventually reverses. What differs across paths is how much of the trajectory robots dominate, from about three quarters of it under the slowest path the current micro evidence favors, to about a quarter under the [Autor \(2024\)](#) scenario. A uniform automation tax is locally welfare-reducing: the optimal policy subsidizes both frontiers at every finite value of inequality aversion, so introducing a common tax moves both instruments away from their optima.

Three caveats constrain the interpretation. First, the cognitive frontier is identified from 2005–2019 American Community Survey (ACS) data using measures that largely capture routine cognitive automation, not GPT-era generative AI. The estimated reinstatement rate is a statement about historical cognitive displacement. I investigate what happens if AI learns to reinstate at higher rates. Second, the physical and cognitive reinstatement rates come from different estimations. The physical rate is the structural estimate of [Chávez \(2026\)](#) from commuting-zone wage moments, with firm-level Danish evidence serving only as an external bracket, while the cognitive rate is estimated here from the six-moment system. The difference in designs raises comparability concerns that the calibration section addresses with extensive sensitivity analysis. Third, the joint estimation of [Chávez \(2026\)](#) places the cognitive reinstatement rate near zero, well below the

estimate here, because it matches wage-level moments in place of the college employment share.¹ Carried through the macro model, that configuration *widens* robot dominance, because a lower cognitive reinstatement rate lowers AI’s productivity content. The estimates used here are therefore the conservative side of the disagreement.

This paper extends the task-based automation framework of [Acemoglu and Restrepo \(2018, 2022\)](#) in three directions: two frontiers with asymmetric reinstatement, a frontier-specific productivity decomposition that the single-frontier model cannot deliver, and frontier-specific optimal policy. Relative to [Hémous and Olsen \(2022\)](#), who model automation and augmentation as competing technologies, the contribution is the empirical mapping from observed reinstatement rates to macro objects. Closest to the quantitative exercise here is [Acemoglu \(2024\)](#), which prices the aggregate productivity gain from cognitive automation within a single-frontier task framework. The two-frontier structure is what allows the present paper to pose a question that framework cannot: not how large the gain from artificial intelligence is, but how it compares, per unit of advance, to the gain from the physical frontier. The three-sector extension with non-tradable services relates to the structural change literature ([Herrendorf et al., 2014](#); [Baumol, 1967](#)), and the dynamic model with heterogeneous savings to the macro-inequality literature ([Karabarbounis and Neiman, 2014](#); [Piketty, 2014](#)). The evidence on the physical frontier builds on [Graetz and Michaels \(2018\)](#) and [Acemoglu and Restrepo \(2020\)](#). The empirical motivation draws on [Webb \(2020\)](#); [Eloundou et al. \(2024\)](#); [Babina et al. \(2024\)](#), and the two-frontier identification framework extends [Chávez \(2026\)](#).

Section 2 presents the model and its three demand-side extensions as theory, with the quantitative channels deferred. Section 3 documents the reduced-form moments and their design-based inference, estimates the cognitive-frontier parameters by minimum distance, and imports the physical frontier from the commuting-zone estimate of [Chávez \(2026\)](#), which Danish firm-level evidence brackets but does not supply. Section 4 computes the general equilibrium feedback and derives the static results: supply-side dominance and uniqueness of equilibrium. Section 5 extends to contin-

¹The two estimates are design-specific rather than directly interchangeable. They come from different exposure measures, samples, controls, outcomes, covariance weights, and maintained mappings from cognitive exposure to wages and employment shares. Controlled bridges that swap the Language Model Occupational Exposure (LMOE) and Webb measures while holding the receiving paper’s structural mapping fixed do not close the gap: LMOE in this paper’s mapping raises rather than lowers δ_A , while Webb in the Chávez mapping produces a poor boundary fit. In this paper the AI wage and college-share moments jointly inform δ_A and φ . Robot productivity dominance holds under both calibrations. What differs is the channel that carries it, the sorting friction at the estimate here and the reinstatement gap at the near-zero rate.

uous time. Section 6 presents the positive macro results: productivity decomposition, labor share dynamics, and the threshold for AI reinstatement. Section 7 develops welfare and optimal policy. Section 8 concludes.

2 Model

The model extends the Acemoglu–Restrepo task framework to two automation frontiers with frontier-specific reinstatement. The production structure builds on Chávez (2026). This section presents it compactly and develops three extensions: non-tradable services, endogenous labor supply, and endogenous capital income distribution. The dynamic extension follows in Section 5.

2.1 Production

A final good Y combines physical and cognitive bundles via CES:

$$Y = \left[\mu Y_P^{\frac{\rho-1}{\rho}} + (1-\mu) Y_C^{\frac{\rho-1}{\rho}} \right]^{\frac{\rho}{\rho-1}}, \quad (1)$$

where $\rho > 0$ is the elasticity of substitution between physical and cognitive output.

Disaggregated task structure.

Each sector $j \in \{P, C\}$ produces output from a continuum of tasks. The task space comprises three segments:

- (i) *Automated tasks* $z \in [0, I_j]$: produced by capital with task-specific productivity $g_{Kj}(z)$.
- (ii) *Surviving labor tasks* $z \in (I_j, 1]$: produced by labor with task-specific productivity $g_{Lj}(z)$.
- (iii) *Reinstated tasks* $z \in (1, 1 + \delta_j I_j]$: new complementary tasks created at rate δ_j per unit of frontier advance, produced by labor.

Output in sector j aggregates all tasks via CES with elasticity $\sigma > 1$ (exponent $\varsigma \equiv (\sigma - 1)/\sigma$):

$$Y_j = \left[\int_0^{I_j} y_{Kj}(z)^\varsigma dz + \int_{I_j}^1 y_{Lj}(z)^\varsigma dz + \int_1^{1+\delta_j I_j} y_{Rj}(z)^\varsigma dz \right]^{1/\varsigma}, \quad (2)$$

where $y_{Kj}(z) = g_{Kj}(z) k_j(z)$ is output of automated task z using capital $k_j(z)$, $y_{Lj}(z) = g_{Lj}(z) \ell_j(z)$ is output of surviving task z using labor $\ell_j(z)$, and $y_{Rj}(z) = g_{Rj}(z) \ell_j^R(z)$ is output of reinstated task z using labor $\ell_j^R(z)$. Following [Acemoglu and Restrepo \(2018\)](#), the task productivities g_{Kj} , g_{Lj} , and g_{Rj} are defined in efficiency units, so that raising them to the task exponent yields the task content coefficient in the CES aggregator. The exact aggregator below carries the exponent $\sigma - 1$. Some expository passages use a distinct additive shorthand in ς ; Appendix [A](#) shows that it is generically unequal to the exact CES aggregate and it enters no quantitative computation.

Reinstated task productivities.

Reinstated tasks are created by entrepreneurs who draw a potential complementary task at frontier j with productivity γ from a sector-specific distribution G_j on $[\underline{\gamma}_j, \bar{\gamma}_j]$. An entrepreneur creates the task if and only if the value $v_j(\gamma)$ of operating it exceeds the creation cost $c(\gamma)$. In equilibrium, a measure $\delta_j I_j$ of tasks is created. The following lemma establishes that the heterogeneous productivities aggregate exactly under CES.

Lemma 1 (CES Aggregation of Reinstated Tasks). *Let reinstated tasks at frontier j have productivities $\{\gamma\}$ drawn from equilibrium distribution G_j^* (the distribution of γ conditional on entry). Under CES aggregation with exponent ς and optimal labor allocation across reinstated tasks, the reinstated-task contribution to the CES aggregator satisfies:*

$$\int_1^{1+\delta_j I_j} (g_{Rj}(z) \ell_j^R(z))^\varsigma dz = \delta_j I_j \cdot (\gamma_j B_j)^\varsigma \cdot \left(\frac{L_j^R}{\delta_j I_j} \right)^\varsigma, \quad (3)$$

where:

$$\gamma_j \equiv \left(\mathbb{E}_{G_j^*} [\gamma^{\sigma-1}] \right)^{1/(\sigma-1)} \quad (4)$$

is the $(\sigma - 1)$ -power mean of reinstated-task productivities, B_j is a sector-specific efficiency scalar, and L_j^R is total labor allocated to reinstated tasks.

The power mean γ_j in equation (4), with exponent $\sigma - 1$ per Lemma 2 of Appendix [A](#), is an exact aggregation object, with no approximation involved in this step. When reinstated tasks are homogeneous (G_j^* degenerate at a single γ_j), the power mean reduces to γ_j itself. The proof is in Appendix [A](#).

Compressed production function.

The firm in sector j allocates capital across automated tasks and labor across surviving and reinstated tasks to maximize Y_j . The following proposition derives the compressed form.

Proposition 1 (Task Aggregation). *Under optimal within-sector factor allocation:*

$$Y_j = \left[Z_j(I_j) + \tilde{\Gamma}_j(I_j, \delta_j) \cdot (L_j^{\text{eff}})^\varsigma \right]^{1/\varsigma}, \quad j \in \{P, C\}, \quad (5)$$

where:

$$Z_j(I_j) = \int_0^{I_j} g_{Kj}(z)^\varsigma \bar{k}_j(z)^\varsigma dz, \quad (6)$$

$$\tilde{\Gamma}_j(I_j, \delta_j) = \left[\int_{I_j}^1 g_{Lj}(z)^{\sigma-1} dz + \delta_j \cdot I_j \cdot (\gamma_j B_j)^{\sigma-1} \right]^{1/\sigma}, \quad (7)$$

with $\bar{k}_j(z)$ denoting optimal capital per automated task. Both Z_j and $\tilde{\Gamma}_j \cdot (L_j^{\text{eff}})^\varsigma$ carry units [output] $^\varsigma$: Z_j absorbs optimal capital quantities. $\tilde{\Gamma}_j$ collects the task-content terms that multiply aggregate effective labor raised to ς .

Proof. See Appendix A.

Three features merit comment. First, $Z_j(I_j)$ is not a pure technology parameter: it contains optimal capital deployment $\bar{k}_j(z)$ at each automated task. In general equilibrium, $\bar{k}_j(z)$ depends on the rental rate, so Z_j is endogenous to factor prices. The notation $Z_j(I_j)$ suppresses this dependence. All comparative statics in Sections 6–7 account for the full equilibrium response.²

Second, the surviving-task integral and the reinstatement term enter the bracket of $\tilde{\Gamma}_j$ additively because both represent task-content terms that multiply $(L_j^{\text{eff}})^\varsigma$ after optimal labor allocation. The equal-effective-labor-share result under CES (Appendix A, Step 2) is what permits this additive separation inside the aggregate.

Third, the reinstatement term is proportional to $\delta_j \cdot I_j$, the total measure of reinstated tasks, not to δ_j alone. The entry condition (10) makes this product invariant to I_j except through the

²Equation (7) is the exact CES aggregate derived in Appendix A, equation (A.9), and is what enters every quantitative computation. Where the text refers to the two terms separately, it uses the additive shorthand $\int_{I_j}^1 g_{Lj}(z)^\varsigma dz + \delta_j I_j (\gamma_j B_j)^\varsigma$, display notation that Appendix A shows is generically unequal to (7). Appendix G.8 reports the productivity decomposition under the joint estimates of Ch  vez (2026), with $\pi_R/\pi_A \in [1.08, 2.7]$ across configurations.

equilibrium objects $(v_j, W_{k(j)}^d)$ (Appendix A, Step 5), so the stock of reinstated tasks responds to frontier advance only through equilibrium task values, and the interaction between δ_j and I_j generates the nonlinearities exploited in the productivity decomposition (Section 6).

Effective labor aggregates non-college (S) and college (H) workers via CES with elasticity $\varepsilon_w > 1$:

$$L_j^{\text{eff}} = \left[(a_S^j \cdot \ell_S^j)^{\frac{\varepsilon_w - 1}{\varepsilon_w}} + (a_H^j \cdot \ell_H^j)^{\frac{\varepsilon_w - 1}{\varepsilon_w}} \right]^{\frac{\varepsilon_w}{\varepsilon_w - 1}}, \quad (8)$$

with skill–sector productivities satisfying comparative advantage: non-college workers have comparative advantage in physical production, college workers in cognitive production ($\kappa \equiv (a_S^P/a_H^P)/(a_S^C/a_H^C) > 1$). Sector prices satisfy $p_P/p_C = [\mu/(1 - \mu)](Y_C/Y_P)^{1/\rho}$.

Frontier and sector indices correspond one-to-one: $R \equiv P$ for the robot frontier operating on physical tasks, and $A \equiv C$ for the AI frontier operating on cognitive tasks. The endogenous automation mix $\omega \equiv I_R/(I_R + I_A)$, the robot share of total automation investment, and the reinstatement rates (δ_R, δ_A) are determined by the two-frontier framework of [Chávez \(2026\)](#):

$$\omega = \omega\left(\frac{W_S}{W_H}, \frac{r_R}{r_A}\right), \quad \frac{\partial \omega}{\partial (W_S/W_H)} > 0, \quad (9)$$

$$\delta_j = \frac{1}{I_j} \left[\frac{\theta_j}{c_0(1 + \eta)} (v_j - W_{k(j)}^d) \right]^{1/\eta}, \quad (10)$$

where θ_j is the task complementarity parameter at frontier j , v_j is the value of a reinstated task, and $W_{k(j)}^d$ is the post-displacement wage of the skill group $k(j)$ that frontier j 's automation displaces, following the comparative-advantage assignment of equation (8): $k(R) = S$ and $k(A) = H$. Here r_R and r_A are the rental rates on robot and AI capital, and $c_0 > 0$ and $\eta > 0$ are the scale and the curvature of the task-creation cost function $c(\cdot)$ introduced above. The two frontiers can differ both in the number of feasible complementary tasks per unit of advance and in the value margin $v_j - W_{k(j)}^d$ that governs whether those tasks are created. What enters the macro model is the *realized* reinstatement rate that this margin produces in equilibrium, and the rates estimated below are comparable across frontiers. The physical rate is taken from the estimation in [Chávez \(2026\)](#) (Section 3.4) and the cognitive rate is estimated here.

2.2 Three-Sector Extension: Non-Tradable Services

The quantitative results of Section 6 come from the two-sector core. This extension is solved separately, in Appendix B.1, and it is what produces the computed feedback channels of Section 4.1. I extend the production structure to three sectors: physical (P), cognitive (C), and non-tradable services (N). The final good becomes:

$$Y = \left[\mu_P Y_P^{\frac{\rho-1}{\rho}} + \mu_C Y_C^{\frac{\rho-1}{\rho}} + \mu_N Y_N^{\frac{\rho-1}{\rho}} \right]^{\frac{\rho}{\rho-1}}, \quad \mu_P + \mu_C + \mu_N = 1. \quad (11)$$

Physical and cognitive sectors retain the two-frontier task structure. Non-tradable services are produced with labor only ($Y_N = A_N \cdot L_N^{\text{eff}}$) and are not automated: $I_N = 0$, $\delta_N = 0$. This is the Baumol sector. Accordingly the three-sector weights (μ_P, μ_C, μ_N) are not separately calibrated. The two-sector core uses the physical share $\mu = 0.45$, and of the three-sector weights only $\mu_N = 0.35$ enters, as the scale of the demand-side channels computed in Section 4.1.

With non-homothetic preferences, non-college workers spend a larger share of income on local services ($\alpha_S^N > \alpha_H^N$). AI displacement that reduces $W_S L_S$ directly reduces service demand, releasing non-college workers to the tradable sectors and depressing W_S further. The channel operates through a local demand multiplier and enters the feedback through labor reallocation, bypassing the CES price dampening that limits the composition channel. Section 4 computes it.

A natural concern is that AI could automate services ($I_N > 0$), which would change the sign of this channel. If AI automates customer service, food ordering, or personal care scheduling, the non-tradable sector shrinks through automation itself and the labor reallocation channel weakens. Appendix E evaluates a parametric extension in which the service sector is partially automated. The supply-side ranking survives partial service automation.

2.3 Endogenous Labor Supply

Worker i of skill k participates if $W_k \geq \underline{w}_k(i)$, where reservation wages are drawn from $F_k(\cdot)$. The participation elasticity is $\varepsilon_k^{LS} \equiv W_k f_k(W_k) / F_k(W_k)$, calibrated for each skill group from Chetty et al. (2013) (Table 3).

The participation margin *dampens* the wage response (some workers exit when wages fall) but *amplifies* the income response. Which force wins is a quantitative question, and Section 4 answers

it by solving the margin: the dampening prevails, and participation reduces the total feedback.

2.4 Endogenous Capital Income Distribution

Automation generates capital income from returns to automated tasks (Π^{auto}) and reinstatement rents (Π^{reinst}). With heterogeneous marginal propensities to consume ($m_S > m_H > m_K$), a shift toward AI raises the capital income share and lowers the aggregate MPC, feeding back through the non-tradable channel.

Together with the demand composition channel, which operates through the relative price of the two final goods, Section 4 computes all four at the estimated parameters, shows that their net contribution is an order of magnitude smaller than the direct supply-side effects, and states the formal supply-side dominance result.

3 Estimation and Calibration

This section documents the reduced-form moments and their design-based inference, estimates the cognitive-frontier parameters by minimum distance, calibrates the physical frontier from the commuting-zone estimate of [Chávez \(2026\)](#), which external evidence brackets rather than supplies, and closes with the mapping from commuting-zone estimates to national parameters.

3.1 Reduced-Form Moments

The structural parameters ($\Lambda_R, \delta_A, \varphi, \varepsilon_H$) are identified from a minimum distance estimator (MDE) that matches the model’s predicted cross-sectional moments to their empirical counterparts. Here ε_H is the college local labor-supply elasticity, which governs the incidence of local demand shocks on college wages through the share $\eta_H = \varepsilon_D / (\varepsilon_D + \varepsilon_H)$; it is distinct from the college participation elasticity ε_H^{LS} of Section 2.3. This subsection documents the six reduced-form coefficients that serve as inputs.

I regress three commuting-zone-level outcomes, log change in non-college wages ($\Delta \ln W_{NC}$), log change in college wages ($\Delta \ln W_C$), and change in college employment share (Δs_C), on robot and AI exposure, using 722 commuting zones, long differences from pooled 2005–2007 baseline to pooled 2015–2019 endpoint ACS samples. Robot exposure is the EU5 Bartik instrument, built from

1993–2014 robot adoption quantities in Denmark, Finland, France, Italy, and Sweden (Acemoglu and Restrepo, 2020). AI exposure is the Webb (2020) AI exposure index in Bartik percentiles. All regressions include 38 baseline demographic and labor-market composition shares, the baseline farm-occupation employment share, and 8 census region dummies, and are weighted by baseline commuting-zone employment.³

Table 1 reports the results. Commuting zones with greater robot exposure experienced lower non-college wage growth ($\hat{\beta}_{R,NC} = -0.0079$, $t = -4.25$, precisely estimated under conventional inference, though Section 3.2 shows the design-based robot-side precision is far lower) and moderately lower college wage growth ($\hat{\beta}_{R,C} = -0.0055$, $t = -1.99$). AI exposure is associated with higher college wage growth ($\hat{\beta}_{AI,C} = +0.0052$, conventional $t = +2.07$, but exposure-robust $t = 1.46$, so it is not significant under the design-based inference this paper prefers), lower non-college wage growth ($\hat{\beta}_{AI,NC} = -0.0069$, $t = -2.88$, which survives exposure-robust inference at $t = -2.24$ but clears neither permutation null, $p = 0.10$ within skill bins and $p = 0.16$ on task residuals), and a rising college employment share ($+0.0030$, conventional $t = +2.29$, though it does not survive the within-bin permutation null). The sharpest AI moment is the skill-premium response in column (4): $+0.0121$, with $t = +4.51$ conventionally and $+3.60$ under exposure-robust inference. It is the only moment that clears a permutation null at conventional levels ($p = 0.014$ within skill bins), though it is marginal against the task-residual null ($p = 0.068$). The premium is the model’s central distributional prediction, and it is estimated more precisely than either wage moment because common commuting-zone-level shocks difference out of the contrast. This pattern, robots depressing wages for both groups while AI redistributes toward college workers, identifies Λ_R , δ_A , and φ in the structural model.

The college-share moments carry the expected signs (robots reduce, AI increases). The AI college-share coefficient survives exposure-robust inference ($t = 2.11$) but not the within-skill-bin permutation null ($p = 0.38$), so it is the weakest of the three AI moments on design-based inference. It nonetheless remains part of the joint minimum-distance system: the wage and share moments jointly inform δ_A and φ through the full 6×6 cross-equation variance-covariance matrix. Leave-one-moment-out estimates keep δ_A near one when either AI wage coefficient is removed, but removing

³The correlation between the robot and AI exposure instruments is $r = -0.067$, so the two exposure measures vary nearly independently across commuting zones, which is what allows the robot and AI gradients to be separated. That variation speaks to Λ_R , δ_A and φ , not to δ_R .

the college-share block sharply increases uncertainty and boundary incidence. The share moments therefore matter for precision without uniquely carrying the point estimate.

Figure 1 visualizes the two headline relationships. The robot–non-college wage gradient is steep and approximately linear across the support. The AI–college wage relationship is positive but noisier, consistent with the larger standard errors in column (2).

Appendix H reports three robustness checks. The first two are estimated on the completed control set of Table 1, and the third drops the farm share. Adding the Autor et al. (2013) China import shock to the baseline specification leaves all six moments virtually unchanged (Panel A), indicating that robot and AI exposure capture channels distinct from trade competition. With augmented controls that include log population and baseline tech employment share (Panel B), the AI coefficient on college wages attenuates from 0.0052 to 0.0022 and loses even its conventional significance. It was in any case not significant under the design-based inference this paper prefers. This reflects the endogeneity of these controls: commuting zones with high AI-task concentration are precisely the large, tech-intensive metropolitan areas whose size and sectoral composition are themselves shaped by AI exposure. These controls absorb the cross-sectional variation carried by AI-task concentration and are therefore excluded from the preferred specification. The baseline specification in Table 1 is therefore preferred for the MDE. Outlier and design-based diagnostics for these moments are in Section 3.2.

The baseline control set comprises 38 pre-determined composition shares plus the baseline farm-occupation employment share. Farm share belongs to the same class as the other 38, baseline, pre-determined, and compositional, and unlike the log-population and tech-share controls rejected above it does not absorb the exposure variation. A Rotemberg decomposition of the AI exposure measure under this set places the top identifying weight on computer systems analysts, the occupation the model’s mechanism says should carry an AI moment. Dropping the farm share leaves the AI college-wage coefficient nearly unchanged (+0.0052 against +0.0056) while weakening the non-college and college-share moments. The farm-dropped specification is reported as Panel C of Table H.1, and all moments feeding the estimation use the completed set.

Table 1: Reduced-Form Effects of Robot and AI Exposure

	Dependent Variable			
	$\Delta \ln W_{NC}$	$\Delta \ln W_C$	Δs_C	$\Delta \ln(W_C/W_{NC})$
	(1)	(2)	(3)	(4)
<i>Panel A: Conventional inference</i>				
Robot exposure (EU5)	−0.0079*** (0.0019)	−0.0055** (0.0028)	−0.0016 (0.0011)	—
AI exposure (Webb pctl)	−0.0069*** (0.0024)	0.0052** (0.0025)	0.0030** (0.0013)	0.0121*** (0.0027)
<i>Panel B: Design-based inference</i>				
Robot: permutation p (19 shocks)	0.11	0.27	0.46	—
AI: exposure-robust t (AKM)	−2.24	+1.46	+2.11	+3.60
AI: within-bin permutation p	0.10	0.17	0.38	0.014
AI: task-residual permutation p	0.16	0.31	0.09	0.068
Completed controls + region FEs	Yes	Yes	Yes	Yes
Observations	722	722	722	722

Notes: Each column reports a weighted OLS regression of the indicated outcome on both exposures under the completed control specification of Section 3.1: 38 baseline composition shares, the baseline farm-occupation employment share, and 8 census region dummies, weighted by baseline employment. Robot exposure is the EU5 Bartik over 1993–2014 adoption quantities. AI exposure is the Webb (2020) index in Bartik percentiles. Column (4) is the skill-premium response, a linear combination of columns (1)–(2) whose information the joint estimation already uses. Panel B holds each estimate to its design: permutation over the 19 industry shocks on the robot side, where the coarse permutation distribution limits resolution to roughly the ten-percent level, and exposure-robust standard errors (Adão et al., 2019), abbreviated AKM in the row labels, with skill-bin-constrained permutation nulls over 338 occupation scores on the AI side. Robust standard errors in parentheses. Coefficients and standard errors are displayed rounded to four decimals. Reported t -statistics are computed from the unrounded estimates and need not equal the ratio of the displayed figures.

* $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.

The binned scatter plots underlying the two headline gradients show that neither relationship

is driven by a handful of commuting zones.

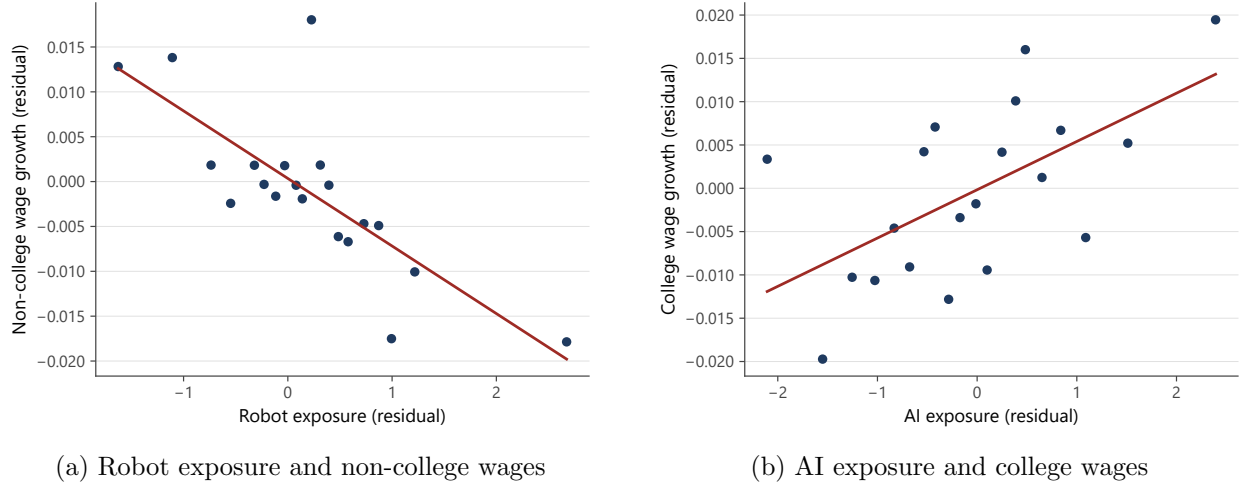


Figure 1: Partial Correlation Plots: Technology Exposure and Wages

Notes: Each panel displays a binscatter with 20 quantile bins, weighted by baseline commuting-zone employment. Panel (a) plots $\Delta \ln W_{NC}$ against EU5 robot exposure, residualized on AI exposure and baseline controls. Panel (b) plots $\Delta \ln W_C$ against Webb AI exposure, residualized on robot exposure and baseline controls.

3.2 What Identifies the Robot Moments

The robot moments are estimated on the employment-weighted cross-section, so the estimand is the wage response of the average worker, and the estimates should be read as statements about where employment is, not about the average commuting zone. The exposure measure is a shift-share object: nineteen industry-level European adoption shocks spread over baseline industry employment shares. The identifying claim is that the shocks, not the shares, are as good as random (Borusyak et al., 2022), which makes the industry count the relevant sample size for inference. The effective number of shocks, the inverse Herfindahl of the exposure weights, is 3.05. The AI exposure measure aggregates 338 occupation scores with far more dispersed weights, and the same inverse-Herfindahl computation gives 77.9 effective occupation shocks, the roughly eighty referenced in the introduction and the sample size the exposure-robust AI inference exploits.

That number is small because the identification is concentrated, and the concentration has a name. A Rotemberg decomposition (Goldsmith-Pinkham et al., 2020) attributes 93.9 percent of the identifying covariance in the non-college wage moment to a single industry, automotive, and removing the automotive shock alone moves the coefficient from -0.0079 to $+0.0021$ ($t = +0.40$),

reversing it to an insignificant positive. This is the instrument working as built, since industrial robots concentrate in automotive production by construction of the International Federation of Robotics (IFR) data, and the same fact surfaces geographically as the leave-Michigan-out margin in the state-deletion grid. The EU5 instrument is, in effect, the European automotive robotization experiment. Published shift-share designs share this structure, including the China-shock literature this design descends from (Goldsmith-Pinkham et al., 2020), and the consequence is that conventional standard errors, which treat 722 commuting zones as independent observations, overstate the precision of what is closer to one large natural experiment. Excluding the identifying shock removes the design itself, so that result belongs here, to the identification discussion, and the robustness ladder below tests perturbations within the design.

Table 1 therefore reports two tiers of inference side by side. All estimates use the completed control specification, the baseline composition shares augmented with the baseline farm-occupation employment share, whose derivation is in the methods discussion of Section 3.1. Conventionally, the non-college wage moment is estimated at $t = -4.25$, and within the design it is sign-stable under trimming of three-standard-deviation residual outliers ($t = -4.23$), winsorized exposure ($t = -2.71$), log exposure, deletion of the agriculture-and-energy states outright, and every leave-one-state-out subsample, where the weakest margin is Michigan at $t = -1.85$, negative and marginal. The Michigan margin restates the automotive concentration geographically. Two within-design perturbations move it materially, and both appear here, not in a note. Unweighted, the coefficient is attenuated and imprecise ($t = -1.10$), which under the Solon et al. (2015) reading indicates effect heterogeneity correlated with size: the displacement response is concentrated among large exposed commuting zones. And under permutation of the nineteen industry shocks, the design-based test the shift-share structure calls for, the two-sided p -value is on the order of ten percent, and the coarse resolution of a nineteen-shock permutation distribution bounds how sharply this can be stated in either direction. The college wage and college employment share moments do not survive the shock-level test at all (permutation p of 0.27 and 0.46), so the design-based content of the robot block lives in the non-college wage response.

What these moments identify, and what they do not, follows the same two-tier logic. The robot block pins down Λ_R , the scale of the displacement shock. Table 2 reports the shock-level standard error (0.0193) as the estimate’s headline uncertainty. A conventional CZ-level standard error would

be 0.0030, overstating precision by a factor of 6.4 by treating 722 commuting zones as independent draws from a design with three effective shocks. The reinstatement rate δ_R never lived in these moments: it is the estimate of [Chávez \(2026\)](#), bracketed by the firm-level evidence in [Humlum \(2022\)](#), and no conclusion about reinstatement asymmetry rests on the automotive experiment. The restriction $\psi \equiv 1$ is maintained as an assumption on construction grounds: an industry-level demand shock is skill-neutral by construction (Section 3.1), and the formal difference test that could probe the restriction rests on the robot college-wage moment, which has essentially no power under design-based inference.

3.3 Minimum-Distance Estimation

The cognitive reinstatement rate $\hat{\delta}_A = 0.99$ (bootstrap SE = 0.19) and sorting disruption $\hat{\varphi} = 1.99$ (bootstrap SE = 0.31) are estimated by minimum distance, matching the six reduced-form moments from Table 1 (three from robot exposure, three from AI exposure) to four structural parameters (Λ_R , δ_A , φ , ε_H). The two-frontier identification framework builds on [Chávez \(2026\)](#), adapted to the macro production structure of this paper. The restricted specification imposes symmetric robot demand shocks ($\psi \equiv 1$), yielding two overidentifying restrictions. The Hansen J -test ($J = 0.012$, $p = 0.99$, two degrees of freedom) does not reject the model, though under the design-based moment variances the two overidentifying restrictions carry almost no information. The restriction is maintained as an assumption, on the grounds that an industry-level demand shock is skill-neutral by construction, and not because the data discriminate against the alternative. The unrestricted specification is not separately estimated. Any difference test would rest on the robot college-wage moment, which has almost no power under design-based inference, so it would carry no evidentiary weight.

Table 2: Estimated Parameters (Minimum Distance)

Parameter	Value	Source	Identification
<i>Estimated here (CZ inputs: Section 3.1; MDE: Section 3.3)</i>			
$\hat{\delta}_A$	0.99 (0.19)	MDE, 6 moments	AI $\rightarrow (\Delta W_{NC}, \Delta W_C, \Delta s_C)$
$\hat{\varphi}$	1.99 (0.31)	MDE, 6 moments	AI skill-premium (wage–wage) contrast
$\hat{\varepsilon}_H$	1.38 (2.91)	MDE, 6 moments	Robot/AI $\rightarrow \Delta s_C$
$\hat{\Lambda}_R$	0.0085 (0.019)	MDE, 6 moments	Robot $\rightarrow (\Delta W_{NC}, \Delta W_C)$

Notes: Bootstrap SE in parentheses (200 replications under the two-way variance-covariance matrix (VCV): industry-level resampling for the robot moments, CZ-level for the AI moments). $(\Lambda_R, \delta_A, \varphi, \varepsilon_H)$ estimated by minimum distance from six reduced-form moments (Table 1), 722 U.S. CZs, pooled 2005–2007 to 2015–2019 ACS long differences, restricted specification $\psi \equiv 1$. Hansen $J = 0.01$ ($p = 0.99$, 2 overid), uninformative under the design-based moment variances. A one-step estimator with a fixed diagonal weight matrix agrees with the two-step estimates on δ_A and φ to the second decimal. Identification framework follows Chávez (2026), adapted to the macro model. Values are from the baseline model with Cobb–Douglas final-good aggregation and the consistent Z_j convention. Table G.3 reports sensitivity to CES aggregation at the calibrated $\rho = 0.80$. ε_H weakly identified: 90% bootstrap percentile interval $[-0.60, 6.26]$ (asymmetric because it is read from the draws, not a normal approximation).

3.4 The Physical Reinstatement Rate

The physical reinstatement rate $\delta_R = 1.14$ (point-identified with bootstrap SE 0.18, externally validated by the Humlum bracket $[0.60, 2.00]$) is the minimum-distance estimate of Chávez (2026) from U.S. commuting-zone wage moments. External evidence brackets it. Humlum (2022)’s firm-level analysis of robot adoption in Danish manufacturing implies a creation-to-displacement value ratio between 0.6, counting only robot-complementary tech tasks, and 2.0, attributing the entire net labor-demand expansion of adopters to task creation, and the estimate sits inside that range (Appendix D.3 derives the bracket). Danish institutions, stronger employment protection, union co-determination, and active labor market policies, create incentives for firms to redeploy their workers, so the Danish bracket may sit above the U.S. rate. The bracket validates δ_R rather than supplying it, so this bears on the external check and not on the estimate the model uses. The sensitivity grid in Table 4 makes the exposure explicit: parity occurs at $\delta_R \approx 0.61$ at the baseline estimates, so the headline survives across most of the Humlum range and reaches approximate parity at the narrow bound (ratio 0.99 at 0.60, marginally below one). That bound is not a live

possibility for the estimate. It sits three standard errors below $\delta_R = 1.14$ and outside the bootstrap 95 percent interval $[0.84, 1.53]$, so near-parity at the narrow Humlum bound describes how wide the external bracket is, not the sampling uncertainty in the parameter the model uses. U.S. firms adopting industrial robots exhibit comparable patterns of internal task reallocation ([Acemoglu and Restrepo, 2020](#)), and European evidence corroborates reinstatement across institutional settings ([Koch et al., 2021](#); [Dauth et al., 2021](#)).

Table 3: Calibrated Parameters

Parameter	Value	Source	Identification
<i>Physical frontier (imported estimate)</i>			
δ_R	1.14 (0.18)	Chávez (2026)	CZ wage moments, point-identified
<i>Production structure</i>			
σ	2.00	Acemoglu and Restrepo (2020)	Task elasticity
ε_w	1.50	CZ wage differentials	Skill elasticity
κ	1.80	Sectoral employment	Comparative advantage
μ, ρ	0.45, 0.80	Output shares, structural change lit.	Physical share (two-good aggregator), goods subst.
ξ_R, ξ_A	0.45, 0.50	Assigned (no external estimate)	Results insensitive, App. G.4
$\varepsilon_D, \varepsilon_S^{\text{mig}}$	3.0, 0.78	Notowidigdo (2020)	Incidence share $\eta_S = 0.794$
<i>Three-sector extension and demand-side</i>			
μ_N	0.35 [†]	Service employment	Non-tradable share
s_N^S	0.45 [†]	CPS service employment	Non-college employment in services
α_S^N, α_H^N	0.40, 0.25	CEX data	Service expenditure shares
$\varepsilon_S^{LS}, \varepsilon_H^{LS}$	0.25, 0.10	Chetty et al. (2013)	Participation elasticities
m_S, m_H, m_K	0.75, 0.45, 0.20	Dynan et al. (2004)	Marginal propensities
<i>Dynamic model</i>			
σ_c	0.50	Standard macro	Intertemporal elasticity
ρ_S, ρ_H	0.06, 0.03	Wealth data	Time preference
d_A, d_R	0.10, 0.10	Standard	Capital depreciation, symmetric baseline

[†] These two service-sector anchors are recorded from the data, but the extended model cannot reproduce both at once: at its equilibrium the service share of value added is 0.15 and the share of non-college workers employed in services is 0.27, and no service-sector scale satisfies both targets simultaneously. Appendix B.1 reports the tension and shows the general equilibrium feedback is negligible under either.

Notes: Physical δ_R is the Chávez (2026) minimum-distance estimate, inside the range [0.6, 2.0] implied by Humlum (2022)’s within-firm evidence from Danish manufacturing. Remaining parameters are calibrated from the sources listed. CPS is the Current Population Survey and CEX the Consumer Expenditure Survey. The local incidence relation that maps $(\varepsilon_D, \varepsilon_S^{\text{mig}})$ into the incidence share $\eta_S = \varepsilon_D / (\varepsilon_D + \varepsilon_S^{\text{mig}})$ follows Moretti (2011) and Notowidigdo (2020), whose estimates of differential mobility place non-college workers at the lower migration elasticity. The baseline quantitative model aggregates the two final goods by Cobb–Douglas. The calibrated elasticity $\rho = 0.80$ enters the general-equilibrium channels computed in Section 4.1 and the CES sensitivity in Table G.3. The entry-margin symbols (η, c_0) of equation (10) and the dynamic allocation symbols (ζ_j, ν, α_r) require no calibration: δ_j is estimated directly, and the amplification factor reduces to $(\sigma_c, \rho_S, \rho_H, d_A)$ in Appendix C.

3.5 Why Two Identification Strategies

The asymmetry in identification, with δ_A estimated here from differential local exposure and δ_R recovered from the general-equilibrium level mapping, reflects a substantive economic difference between the two frontiers, not a limitation of the data.

Robot exposure, as measured by the Bartik instrument constructed from European adoption rates, identifies an *industry-level* demand shock. When robots are adopted in the auto industry, all workers in auto-intensive commuting zones experience wage and employment effects, regardless of their individual occupation or task content. The restriction $\psi \equiv 1$, robots hitting college and non-college workers equally through the demand channel, is maintained on these construction grounds (Section 3.2). At the CZ level, the robot Bartik identifies Λ_R , the total industry demand effect inclusive of local multipliers, not the task-level reinstatement rate.

AI exposure, by contrast, is constructed from occupation-level task content and varies across commuting zones with their baseline occupational composition. Its coefficient is estimated separately for the non-college wage, college wage, and college-employment-share outcomes. These three CZ-level moments jointly inform δ_A and φ through the Acemoglu–Restrepo task structure.

That robot effects are identified at the industry level while AI effects are identified at the task level suggests the two technologies operate through different transmission channels, even within the same task-based production function. The macro model uses δ_R and δ_A symmetrically in the production function, but the identification comes from the channel each technology’s data best identify.

Whether the CZ-level data can independently speak to the reinstatement rate $\delta_R = 1.14$, the estimate of [Chávez \(2026\)](#) bracketed by the Humlum firm-level range, is a fair question. The answer requires understanding what identifies each parameter. Λ_R is the local displacement rate: it measures how much *worse off* a high-robot-exposure CZ is relative to a low-exposure CZ, holding national aggregates constant. δ_R is the national reinstatement rate: it governs how many new tasks are created per unit of automation at the *aggregate* level. In the CZ regression, reinstatement is absorbed by the intercept because new tasks do not disproportionately appear in robot-heavy CZs. They are created economy-wide.⁴ This does not contradict the source of the calibration: the

⁴This is a standard feature of shift-share designs applied to automation: the local exposure captures the *differential* displacement across regions, while the aggregate productivity and reinstatement effects are absorbed by time effects

structural estimation of [Chávez \(2026\)](#) recovers δ_R from CZ wage moments through the model’s general-equilibrium mapping, in which the economy-wide reinstatement level is a parameter of the moment function, not from differential local exposure ([Appendix D](#) develops the point). The two parameters therefore load on different margins of the same moment vector: Λ_R on the cross-CZ slope, δ_R on the level that the general-equilibrium mapping implies. This is a separation of margins, not orthogonality, and the two are not identified from independent data. A confound that shifts the level of CZ wage growth enters both moment functions, so the δ_R estimate is not immune to the CZ-level confounds that could bias Λ_R . What disciplines it against them is the external firm-level bracket of [Humlum \(2022\)](#), estimated on different data under a different design, inside which the estimate sits.

Although Λ_R and δ_R cannot be inverted from each other, I verify that they are *jointly consistent* with the macro model. At $\delta_R = 0$ (displacement only, no reinstatement), the macro model implies $\partial \ln W_S / \partial I_R = -0.139$: non-college wages fall by 13.9 percent per unit of frontier advance. This is the model’s prediction for the *local* wage effect in a CZ that absorbs the full frontier shock with no offsetting reinstatement, qualitatively consistent with the estimated $\Lambda_R \cdot \eta_S = 0.0085 \times 0.794 = 0.0067$, or a 0.67 percent non-college wage decline per unit of robot exposure. The comparison is a direction check and no more. $\hat{\Lambda}_R$ carries a bootstrap standard error of 0.019 and is not distinguishable from zero, which is the same weak design-based precision on the robot side documented in [Section 3.2](#). Here $\eta_S \equiv \varepsilon_D / (\varepsilon_D + \varepsilon_S^{\text{mig}}) = 3.0 / (3.0 + 0.78) = 0.794$ is the local incidence share that converts the shock scale into a non-college wage response given local labor demand and migration elasticities. The magnitudes are not directly comparable, because they are measured in different units (frontier position versus robots per thousand workers) and differ by more than an order of magnitude. Both describe wage declines, so the check establishes that the model’s displacement channel runs in the direction the reduced-form evidence indicates, not that it matches it in size.

[Table 4](#) reports how the headline results vary across the range implied by [Humlum \(2022\)](#)’s firm-level estimates, from the narrow bound that counts only robot-complementary tech tasks as reinstatement to the broad bound that attributes the entire net labor-demand expansion of adopting firms to task creation.

or the constant. See [Acemoglu and Restrepo \(2020\)](#) and [Adão et al. \(2019\)](#) for discussion.

Table 4: Sensitivity to δ_R (fixing $\delta_A = 0.99$, $\varphi = 1.99$)

δ_R	π_R	π_A	π_R/π_A	τ_R^*/τ_A^*	Source
0.60	+0.477	+0.481	0.99	1.24	Humlum narrow bound
0.74	+0.517	+0.482	1.07	1.38	Illustrative low value
1.14	+0.622	+0.485	1.28	1.73	Chávez (2026) estimate (baseline)
1.40	+0.682	+0.486	1.40	1.93	Illustrative high value
2.00	+0.804	+0.490	1.64	2.32	Humlum broad bound

Notes: All results computed from the macro model of Section 6, holding δ_A and φ at baseline estimates. π_j denotes the productivity content of frontier j , and τ_R^*/τ_A^* is the optimal subsidy ratio from the Bergson–Samuelson welfare function with $\lambda = 1.5$. The Humlum bounds $[0.60, 2.00]$ are the creation-to-displacement value ratios implied by Humlum (2022)’s firm-level event studies, derived in Appendix D.3. The baseline is the Chávez (2026) minimum-distance estimate $\hat{\delta}_R = 1.14$. Parity occurs at $\delta_R \approx 0.61$, with robot dominance above it, and the subsidy ratio stays above one even at parity because robot-created tasks disproportionately benefit non-college workers, who carry higher welfare weight.

The result $\pi_R > \pi_A$ requires $\delta_R \gtrsim 0.61$, well below the baseline estimate of 1.14 and just above the narrow Humlum bound of 0.60. Across the Humlum range $[0.60, 2.00]$, the productivity ratio runs from 0.99 to 1.64 and the optimal policy ratio from 1.24 to 2.32. Robots dominate AI in productivity content everywhere above the parity threshold, which is to say at the baseline and across most of the range, and fall to approximate parity at the narrow bound, where reinstatement is restricted to robot-complementary tech tasks. Even there the optimal subsidy ratio stays at or above one, carried by the equity motive.

The CZ-level displacement rate in this paper is consistent with prior estimates from different countries and time periods. Acemoglu and Restrepo (2020) find that one additional robot per thousand workers reduces local wages by about 0.42 percent in U.S. CZs over 1990–2007. This paper’s non-college coefficient ($\hat{\beta}_{R,NC} = -0.79$ percent per unit exposure) has the same negative sign despite covering a later period (2005–2019); its magnitude is not directly comparable because the exposure units differ. Dauth et al. (2021) find a small and statistically insignificant overall wage effect in German *Kreise* over 1994–2014, despite Germany’s much higher robot density (7.6 versus 1.6 robots per thousand workers). The contrast is informative: in the model, higher robot

density means more reinstatement has already occurred, so the *net* displacement effect (what the CZ regression captures) is smaller. The pattern, large negative CZ coefficients in the U.S. and small or zero coefficients in Germany, is qualitatively consistent with $\delta_R > 0$. Humlum (2022) directly observes within-firm task reallocation in Denmark, bracketing δ_R between 0.60 and 2.00. The three identification strategies, between-CZ displacement in the U.S., net effects in Germany, and within-firm reinstatement in Denmark, provide convergent evidence that both displacement and reinstatement are economically important, and that their relative magnitudes are consistent with the calibration used in this paper.

The estimated $\hat{\varphi} = 1.99$ captures the sorting disruption induced by AI in the CZ data over 2005–2019. A natural concern is that part of this effect reflects transitional reallocation friction, workers temporarily mismatched to tasks, and if so the long-run φ is smaller than the estimate and the productivity asymmetry narrows. The adjustment half-lives reported below make φ_{LR} near half the estimate the central long-run case, so this section reports the full decomposition, including the rows where the result reverses.

Table 5 reports the headline results under different assumptions about the transitional share of φ , at the point estimate and across the bootstrap set. At a 25 percent transitional share the productivity ratio is 1.16 and 98.0 percent of the bootstrap set remains robot-dominant. At the half-life central case, 50 percent transitional, the point ratio is 1.06 and set-level dominance falls to 73.0 percent: near-parity at the point, with roughly three draws in four still robot-dominant. At 75 percent transitional the point reverses (0.96) and 25.0 percent of the set remains dominant. Parity occurs at $\varphi_{LR} = 0.70$ and robot dominance requires $\varphi_{LR} > 0.70$, so dominance survives transitional shares up to 65 percent. The long-run asymmetry is more sensitive to the transitional decomposition than the contemporaneous asymmetry is to estimation uncertainty, and the geometry explains why: draws with high δ_A carry high φ , which protects them in the contemporaneous comparison and exposes them when φ is scaled down.

Two caveats bound what the exercise can show. First, the set-level rows scale φ uniformly across draws, which is an assumption about which component of the composite is transitional. The scaled draws are themselves outside the bootstrap set. Second, the decomposition treats only φ as transitional when all four structural parameters are medium-run objects. Reinstatement operates at decadal lags in the evidence, so the measured reinstatement rates may understate their long-run

values too, and discounting φ alone is an assumption about which margin adjusts first, evaluated here because it is the margin with direct half-life evidence attached.

Table 5: Sensitivity to Long-Run Sorting Disruption ($\delta_R = 1.14$, $\delta_A = 0.99$)

φ_{LR}	π_R	π_A	π_R/π_A	Set dominant	Interpretation
0.05	+0.626	+0.714	0.88	—	Near-zero sorting disruption
0.50	+0.625	+0.650	0.96	25.0%	75% transitional
0.70	+0.625	+0.625	1.00	—	Parity threshold
0.99	+0.624	+0.591	1.06	73.0%	50% transitional
1.49	+0.623	+0.535	1.16	98.0%	25% transitional
1.99	+0.622	+0.485	1.28	100.0%	Baseline estimate

Notes: φ_{LR} is the permanent component of sorting disruption, $\varphi_{LR} = (1 - q)\hat{\varphi}$ for transitional share q (a symbol distinct from the R&D curvature parameters ξ_j) and $\hat{\varphi} = 1.99$, with other parameters at the baseline estimates. “Set dominant” is the share of the 200 bootstrap draws with $\pi_R/\pi_A > 1$ when each draw’s φ is scaled by $(1 - q)$, under the assumption that transition scales φ uniformly across the set. Parity occurs at $\varphi_{LR} \approx 0.70$, so robot dominance requires φ_{LR} above that value, a maximum transitional share of 65 percent.

Figure 2 plots π_R and π_A continuously over φ_{LR} .

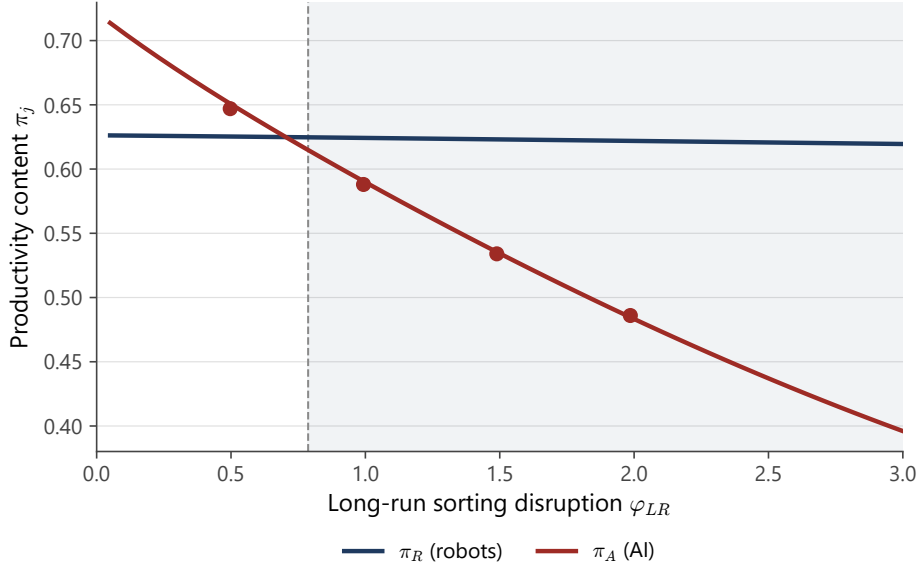


Figure 2: Sensitivity to long-run sorting disruption. π_R (navy) is nearly flat because sorting disruption operates exclusively through the cognitive sector. π_A (red) declines with φ_{LR} . Markers indicate the transitional-share calibrations in Table 5. The shaded region ($\varphi_{LR} > 0.70$) is the parameter space where $\pi_R > \pi_A$.

One feature of this decomposition matters most. φ affects *only* π_A , not π_R : the robot productivity content is nearly flat across the entire range because sorting disruption operates exclusively through the cognitive sector. The asymmetry between frontiers therefore has two independent sources: the reinstatement channel ($\delta_R > \delta_A$, which raises π_R relative to π_A) and the sorting disruption channel ($\varphi > 0$, which lowers π_A). Even without sorting disruption, reinstatement alone holds π_R at 0.626. The result reverses only because π_A *rises* to 0.714 when the sorting penalty is removed.

The speed-of-adjustment literature provides further discipline. [Dix-Carneiro \(2014\)](#) estimates occupational adjustment half-lives of 4–7 years. [Traiberman \(2019\)](#) finds similar magnitudes for Denmark. If AI exposure accumulated gradually over 2010–2019 rather than arriving as a single shock in 2005, workers at the end of the sample have had only 5–9 years of adjustment, closer to one half-life than three. At one half-life, roughly half of the transitional friction remains embedded in the regression estimates, placing φ_{LR} at approximately $0.5\hat{\varphi} \approx 0.99$. Table 5 shows that at $\varphi_{LR} = 0.99$ the ratio is 1.06, above parity at the point but by a thin margin, with set-level dominance at 73.0 percent. The result reverses if more than 65 percent of the estimated sorting

disruption is transitional ($\varphi_{LR} < 0.70$). The contemporaneous asymmetry is what the data identify sharply. The steady-state ranking rests on the transitional decomposition, and the paper reports that dependence in full.

3.6 Mapping CZ-Level Estimates to the National Model

The argument for treating CZ-level estimates as national production-function parameters follows three steps. First, the reinstatement rates enter the production function through the task content of labor $\tilde{\Gamma}_j$, which has the same functional form at both the CZ and national levels. Migration and commuting affect local labor supply but not the task production function. Second, the estimation adjusts for migration on the non-college margin, scaling the local labor supply response by the calibrated elasticity $\varepsilon_S^{\text{mig}} = 0.78$ (Table 3), so the structural δ_A is identified after accounting for the mobility the wage moments would otherwise absorb. Third, the external firm-level bracket for δ_R , 0.60 to 2.00, is consistent with CZ-level patterns across the U.S., Germany, and Spain despite institutional differences, which supports treating reinstatement rates as properties of the technology rather than of local institutions.

The residual concern is spatial spillovers: if robot adoption in one CZ generates task creation in neighboring CZs through supply chain linkages, the CZ-level estimate $\hat{\Lambda}_R$ understates the national effect (or overstates it, depending on the direction of spillovers). This concern primarily affects Λ_R , which enters the macro model only through the sensitivity analysis. The reinstatement rate $\delta_R = 1.14$ is itself a CZ-level estimate, so it carries the same exposure. The firm-level Humlum bracket is an external check that does not, and δ_R sits inside it.

4 Static General Equilibrium

4.1 The General Equilibrium Feedback

A frontier advance moves wages, wages move the automation incentive, and that feedback loop is what could in principle overturn a partial-equilibrium ranking. Write $\mathcal{F}$ for the coefficient of that loop: the response of the automation incentive to the wage change a unit of frontier advance induces. The equilibrium is unique, and the partial-equilibrium ranking survives, precisely when $\mathcal{F}$ is small.

The quantitative core of this paper is two-sector. It contains no service sector, no participation margin and no heterogeneity in marginal propensities to consume, so it produces only the wage-feedback component $\mathcal{F}^{PE}$. At the core’s Cobb–Douglas aggregation the relative-price channel is exactly zero, so the composition channel $\mathcal{F}^{\text{demand}}$ is reported as the increment from evaluating the top-level aggregator at the calibrated $\rho = 0.80$. At the estimated parameters,

$$\mathcal{F}^{PE} = 0.00208, \quad \mathcal{F}^{\text{demand}} = -0.00005. \quad (12)$$

The composition channel is negative. Its sign and magnitude come from the extended model of Appendix B.1, where the skill-specific expenditure shares that drive it are defined: there, non-college households spend more on services and therefore less on physical goods, so the relative-price movement induced by a frontier advance dampens the feedback rather than amplifying it. Equation (12) reports the channel in isolation (-0.00005); Table 6 reports the sequential increment inside the extended model (-0.00004), whose column sums to $\mathcal{F}^{GE}$, and the two differ in the last displayed digit.

The three demand-side blocks of Sections 2.2–2.4 add three further channels, and they are not present in the two-sector core. Appendix B.1 specifies them as an equilibrium, a demand system for services, a reservation-wage distribution delivering the calibrated participation elasticities, and an explicit allocation of automation and reinstatement rents across households, and solves it. With services switched off, the extended model reproduces the two-sector core. Table 6 reports the budget-consistent service calibration; Appendix B.1 also reports variants that separately impose the two Table 3 service anchors, which cannot be matched simultaneously.

Table 6: General equilibrium feedback channels, computed

Channel		Robot frontier	Cognitive frontier
$\mathcal{F}^{PE}$	wage feedback	0.00208	0.00004
$\mathcal{F}^{\text{demand}}$	composition	-0.00004	
$\mathcal{F}^{NT}$	non-tradable services	+0.00012	+0.00033
$\mathcal{F}^m$	MPC concentration	-0.00002	
$\mathcal{F}^{\text{part}}$	participation	-0.00010	
$\mathcal{F}^{GE}$	total	0.00205	0.00033

Notes: Computed in the extended equilibrium of Appendix B.1 at the budget-consistent service calibration. With services switched off the model reproduces the two-sector core to machine precision. The composition, MPC and participation channels are reported on the robot frontier, which carries the headline. Two features are worth stating plainly. The MPC channel is *negative* at the baseline calibration, so it dampens rather than amplifies there. And on the cognitive frontier $\mathcal{F}^{PE}$ nearly cancels, so the non-tradable channel is larger than it, which is a property of the denominator and not a sign that demand matters there: both quantities are more than two orders of magnitude below one. Components are rounded to five decimals and need not sum to the displayed total.

The net demand-side correction on the robot frontier is -1.3 percent of $\mathcal{F}^{PE}$, computed from the unrounded channels (the displayed two-decimal entries imply -1.4). It is negative. On the robot frontier, the demand side, computed rather than bounded, *dampens* the feedback.

Proposition 2 (Supply-Side Dominance). *Across both frontiers and every specification of the demand block considered in Appendix B.1, the total feedback satisfies $|\mathcal{F}^{GE}| \leq 0.0023$. The demand-side channels are more than two orders of magnitude too small to move the automation mix materially, and the distributional consequences of asymmetric automation are therefore supply-side determined.*

The claim is deliberately absolute rather than relative. A ratio of the demand channels to $\mathcal{F}^{PE}$ is not informative, because $\mathcal{F}^{PE}$ itself nearly cancels on the cognitive frontier: the ratio there is large while both quantities are negligible. What matters for the results of this paper is the absolute size of the feedback, and it is negligible on both frontiers.⁵

⁵An earlier version of this section bounded the three channels analytically rather than solving them, using chains

4.2 Contraction and Uniqueness

The feedback loop of Section 4.1 induces a mapping Φ^{GE} on the frontier position: a candidate position x generates wages in the extended equilibrium of Appendix B.1, and the automation incentive at those wages returns the implied position $\Phi^{GE}(x)$. Task shares lie in $[0, \bar{x}]$ with $\bar{x} = 1$.

Proposition 3 (GE Contraction). *The augmented mapping Φ^{GE} satisfies $\sup_x |(\Phi^{GE})'(x)| \leq 0.0023$ on the calibrated domain, and the equilibrium is unique.*

Proof. The derivative of the augmented mapping is the total feedback coefficient, so $|(\Phi^{GE})'| = |\mathcal{F}^{GE}| \leq 0.0023$ by Proposition 2, computed across both frontiers and every specification of the demand block. The bound is two-sided, which is what the contraction argument requires, and it clears the condition $\sup_x |(\Phi^{GE})'(x)| < 1$ by more than two orders of magnitude. Since Φ^{GE} is continuous on the compact interval $[\varepsilon, \bar{x} - \varepsilon]$ for any $\varepsilon > 0$, the Banach fixed-point theorem gives uniqueness. The margin is large because the wage response to a frontier advance is itself small, and every demand-side channel that feeds back through it is smaller still. $\square$

5 Dynamic Extension

The economy is described by automation capital stocks (K_R, K_A) and household assets (A_S, A_H) , with laws of motion driven by R&D investment allocation and heterogeneous savings behavior.

5.1 Setup

Automation capital evolves as $\dot{K}_j = s_j(t) \cdot \mathcal{S}(t) - d_j K_j$, where $\mathcal{S}(t) = Y(t) - C(t)$ is aggregate savings (calligraphic, distinct from the skill subscript S) and s_j is the investment share allocated to frontier j . Frontier thresholds respond to capital: $\dot{I}_j = \zeta_j K_j^{\xi_j}$. Investment shares track relative returns: $s_R/s_A = (\mathcal{R}_R^{\text{auto}}/\mathcal{R}_A^{\text{auto}})^{1/\nu}$.

Households solve standard consumption-savings problems with heterogeneous time preference ($\rho_S > \rho_H$, generating endogenous MPC heterogeneity), yielding Euler equations $\dot{C}_k/C_k = \sigma_c(r(t) - \rho_k)$.

that scale each channel to $\mathcal{F}^{PE}$. That proportionality does not hold: the non-tradable channel's scale is set by the service sector, independently of the wage-composition index, and on the cognitive frontier the analytic bound is violated. The computed channels reported here replace those bounds.

5.2 The Savings Composition Channel

Consider a permanent decline in the user cost of AI capital, r_A . The impact effect is identical to the static counterfactual. The dynamic model generates an additional channel: if non-college workers are hand-to-mouth ($C_S = m_S W_S L_S$, justified by high ρ_S) while college workers are Euler-equation savers, then AI displacement raises college income relative to non-college, increases college savings, and tilts investment toward AI, amplifying the initial shift.

Proposition 4 (Dynamic Amplification). *Under differential time preference ($\rho_S > \rho_H$), a permanent decline in r_A generates:*

- (a) *An impact effect identical to the static counterfactual.*
- (b) *A savings composition effect that tilts investment toward AI over time.*
- (c) *Convergence to a new steady state with lower ω^* and wider skill premium.*

At the calibrated non-college marginal propensity to consume ($m_S = 0.75$, Table 3), the dynamic amplification factor is $\mathcal{A}^{dyn} \approx 1.06$ (Appendix F). The hand-to-mouth limit $m_S = 1$ makes non-college assets $A_S = 0$ and yields the tractable closed form

$$\mathcal{A}^{dyn} \approx 1 + \frac{\sigma_c(\rho_S - \rho_H)}{d_A + \sigma_c \rho_H} = 1.13, \quad (13)$$

a first-order expansion of the exact linearized value $\mathcal{A}^{dyn} = [1 - \sigma_c(\rho_S - \rho_H)/(d_A + \sigma_c \rho_H)]^{-1} = 1.15$ (Appendix C). Both are knife-edge magnitudes, attained only at $m_S = 1$, and the first-order expansion of 1.13 is the more conservative of the two. The economy takes approximately 8 years to close half the gap between the impact response and the new steady state.

Proof. Obtained by log-linearizing the four-dimensional dynamic system around the initial steady state and solving for the ratio of new-steady-state to impact responses. The derivation proceeds in seven steps (state simplification, savings characterization, investment log-linearization, capital dynamics, savings response, steady-state comparison, and reduction to primitives) and is valid for small perturbations under four conditions: local approximation, steady-state comparison, log-linearity (second-order terms are of order $(\Delta r_A)^2$), and separability in the patience gap. The

calibrated $m_S = 0.75$ value of 1.06 is derived in Appendix F. See Appendix C for the full derivation. $\square$

The hand-to-mouth assumption is convenient but loaded: it makes $A_S = 0$, yielding tractable separability. With positive but low non-college savings (a Bewley-style specification), a dampening cross-term proportional to $(1 - m_S) \cdot (\partial W_S / \partial \omega)$ enters the amplification factor, vanishing at the hand-to-mouth limit $m_S = 1$. The qualitative result survives, but the quantitative $\mathcal{A}^{\text{dyn}} = 1.13$ is tied to the knife-edge $m_S = 1$. Appendix F shows that relaxing to $m_S = 0.85$ reduces $\mathcal{A}^{\text{dyn}}$ to approximately 1.09, and at the static calibration's $m_S = 0.75$ (Table 3) it falls to approximately 1.06.

Figure 3 illustrates the transition. Panel (a) shows the automation mix $\omega = I_R / (I_R + I_A)$ following a permanent decline in the user cost of AI capital, r_A . The impact effect shifts the mix toward AI. The savings composition channel then drives a further gradual decline as college workers, who benefit disproportionately from cognitive automation and save at higher rates, tilt aggregate investment toward the cognitive frontier.⁶ Panel (b) shows the corresponding path of the skill premium.

⁶Most of the total movement in ω is the impact effect, and the post-impact transition is small. That is what the amplification factor says: $\mathcal{A}^{\text{dyn}} = 1.15$ means the savings channel adds fifteen percent on top of the impact response, not an amount comparable to it. The mix moves from 0.545 to 0.519 on impact and then only to 0.515 over the transition. The half-life governs how fast that residual gap closes, not how large it is, so a small gap and a slow adjustment are not in tension.

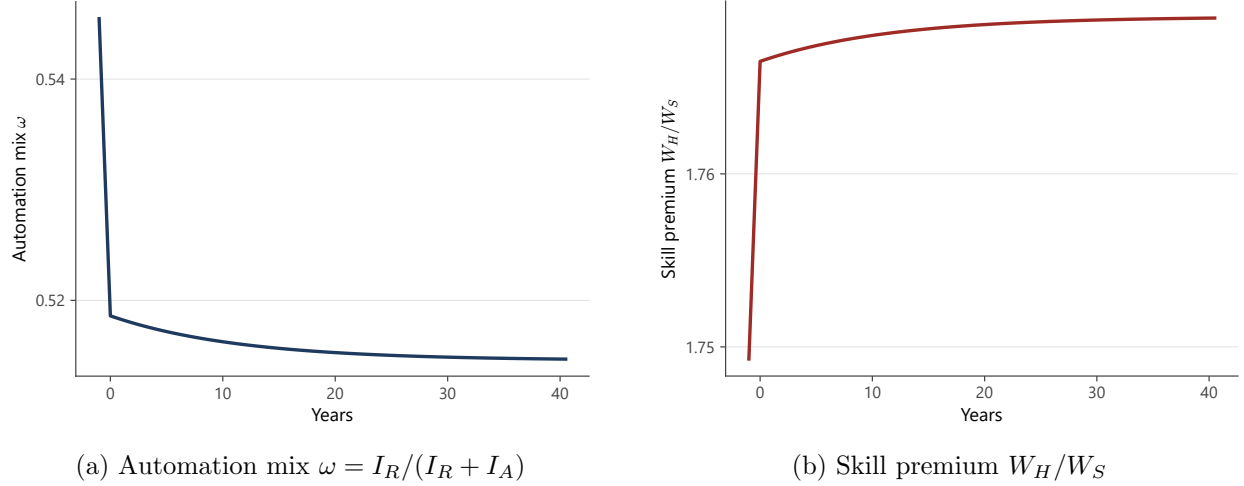


Figure 3: Transition after a permanent decline in the user cost of AI capital, computed at the hand-to-mouth dynamic calibration ($m_S = 1$), with the frontier parameters at their baseline estimates. The automation mix falls from 0.545 to 0.515 and the skill premium rises from 1.75 to 1.77. The impact effect is immediate, and the savings composition channel closes the remaining gap with a half-life of approximately 8 years.

6 Results

6.1 Frontier-Specific Productivity Decomposition

The productivity content of a unit of frontier advance at frontier j decomposes into three components:

$$\pi_j = \underbrace{\pi_j^{\text{direct}}}_{\text{capital replaces labor}} + \underbrace{\pi_j^{\text{reinst}}}_{\text{new complementary tasks}} + \underbrace{\pi_j^{\text{realloc}}}_{\text{worker reallocation}}. \quad (14)$$

Table 7 and Figure 4 report the decomposition at the baseline estimates. The reinstatement-plus-reallocation components are nearly identical: +0.349 for robots against +0.336 for AI. At reinstatement rates of 1.14 and 0.99, the two frontiers rebuild labor’s task content at the same pace once displacement has occurred. If task creation were the whole story, the two technologies would be nearly interchangeable.

The asymmetry lives in the direct component. Shutting down reinstatement ($\delta_j = 0$), a unit of robot advance delivers +0.273 while a unit of AI advance delivers +0.148. The wedge is sorting disruption: at $\varphi = 1.99$, AI’s automation of cognitive tasks degrades the match between skilled workers and the tasks that remain, and this drag depresses the value extracted from every cognitive-

frontier task, automated and reinstated alike. The physical sector faces no comparable friction, which is a modeling assumption (φ enters only the cognitive frontier) rather than something the data impose.

The combined effect is $\pi_R/\pi_A = 1.28$. A marginal unit of robot frontier advance buys about a quarter more output growth than a marginal unit of AI advance, with the sensitivity range reported in Table 8.

The memo rows in Table 7 preview the distributional consequences developed in Section 6.3. Robot advances raise both non-college and college wages (+0.389 and +0.465).⁷ AI advances lower non-college wages ($\partial \ln W_S / \partial I_A = -0.111$) while raising college wages ($\partial \ln W_H / \partial I_A = +0.203$), so AI’s distributional action runs through the skill premium. Sorting disruption concentrates AI’s productivity gains among workers whose skills complement the remaining cognitive tasks. The labor share falls under both frontiers ($\partial s_L / \partial I_A = -0.264$ against -0.119 for robots), and Section 6.3 develops the asymmetry.

⁷These are *national* general-equilibrium semi-elasticities. Advancing the robot frontier raises both wages because reinstatement operates economy-wide. The reduced-form coefficients in Table 1 identify the *local* exposure effect, which is negative, because a commuting zone that absorbs the displacement without the economy-wide reinstatement sees wages fall. Section 3.5 reports the $\delta_R = 0$ local value $\partial \ln W_S / \partial I_R = -0.139$. The positive national sign is a model prediction rather than an estimate, and the local–national gap is the standard Acemoglu and Restrepo (2020) decomposition.

Table 7: Frontier-Specific Productivity Decomposition

Component	Robot (π_R)	AI (π_A)	Ratio
Direct automation	+0.273	+0.148	1.84
Reinstatement + reallocation	+0.349	+0.336	1.04
Total productivity content	+0.622	+0.485	1.28
<i>Memo: distributional effects</i>			
$\partial \ln W_S / \partial I_j$	+0.389	−0.111	
$\partial \ln W_H / \partial I_j$	+0.465	+0.203	
$\partial s_L / \partial I_j$	−0.119	−0.264	
$\partial \ln(W_H/W_S) / \partial I_j$	+0.076	+0.314	

Notes: $\pi_j \equiv d \ln Y / d I_j$ by finite differences ($dI = 0.01$). “Direct” sets $\delta_j = 0$ for the shocked frontier. Components are computed unrounded and displayed to three decimals, so they need not sum exactly to the displayed total (for AI, $0.1484 + 0.3363 = 0.4847$). Calibration at the baseline estimates: $\delta_R = 1.14$, $\delta_A = 0.99$, $\varphi = 1.99$, $I_R = 0.30$, $I_A = 0.25$, Cobb–Douglas final-good aggregation (Table G.3 varies the top-level elasticity, with the calibrated $\rho = 0.80$ moving the ratio to 1.27). $\partial \ln W_k / \partial I_j$ is the semi-elasticity of skill- k wages with respect to frontier j . Robot advances raise both groups’ wages. AI advances raise college wages and lower non-college wages, making the skill premium AI’s dominant distributional margin.

Figure 4 makes the source of the asymmetry visible. The top segments, reinstatement plus reallocation, are nearly identical in height: after displacement, the two frontiers rebuild labor’s task content at the same rate. The bottom segments diverge, because sorting disruption degrades the value the cognitive frontier extracts from its automated and reinstated tasks. At these estimates, the policy-relevant gap between the two technologies lives in this sorting friction.

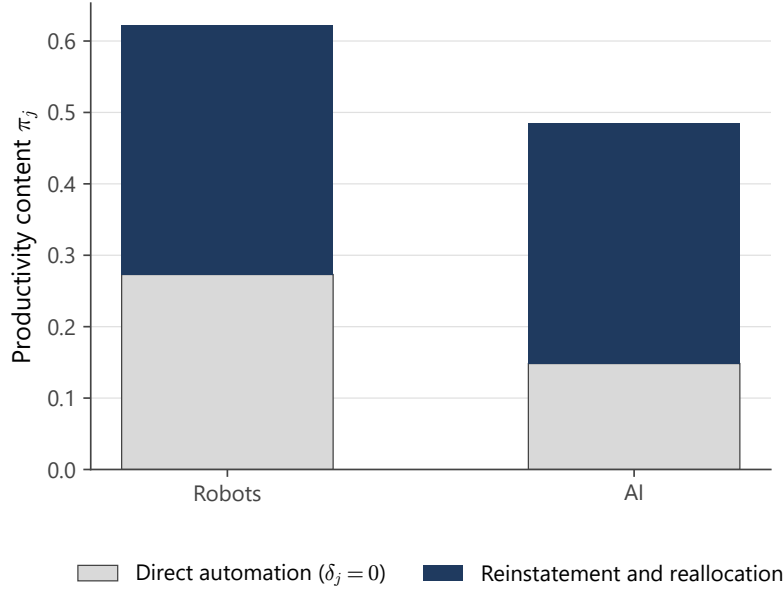


Figure 4: Frontier-specific productivity decomposition at the baseline estimates. Bottom segments: direct automation ($\delta_j = 0$). Top segments: reinstatement and reallocation. The reinstatement-plus-reallocation components are nearly equal across frontiers (+0.35 for robots and +0.34 for AI). The asymmetry is driven by the sorting-disruption drag on the cognitive frontier’s direct component.

The quantitative core aggregates the two final goods by Cobb–Douglas ($\rho = 1$). Adopting the calibrated CES elasticity $\rho = 0.80$ moves the headline ratio by 0.01, from 1.28 to 1.27 (Table G.3), so the choice of aggregator is immaterial to every result and the paper reports the Cobb–Douglas value throughout. The general equilibrium channels of Section 4.1 are computed on the Cobb–Douglas core, with the composition channel reported as the increment from the calibrated $\rho = 0.80$ aggregation, and the productivity decomposition and every headline magnitude likewise use the Cobb–Douglas core.

The productivity ratio is robust to parameter uncertainty.

Proposition 5 (Asymmetric Productivity Content). *The productivity content of robot automation exceeds AI automation by a factor of $\pi_R/\pi_A = 1.28$ at the baseline estimates. Across a joint grid spanning δ_R over its external Humlum bracket $[0.60, 2.00]$ and δ_A over ± 2 bootstrap SE, $\pi_R/\pi_A > 1$ in 20 of 25 cells, with the exceptions confined to the low- δ_R , high- δ_A corner. Over the jointly identified (δ_A, φ) set, robot dominance holds in all 200 bootstrap draws.*

Table 8: Sensitivity of Productivity Ratio π_R/π_A

δ_R	δ_A				
	0.62	0.80	0.99	1.18	1.36
0.60	1.29	1.12	0.99	0.89	0.82
0.74	1.40	1.21	1.07	0.97	0.89
1.14	1.67	1.45	1.28	1.16	1.06
1.40	1.83	1.59	1.40	1.26	1.16
2.00	2.13	1.85	1.64	1.48	1.36

Notes: Each cell reports π_R/π_A by finite differences ($dI = 0.01$) at the indicated (δ_R, δ_A) pair, all other parameters at the baseline estimates. δ_R spans the external Humlum bracket $[0.60, 2.00]$, wider than its bootstrap 95% interval $[0.84, 1.53]$ (SE 0.18), with the baseline 1.14 a node. δ_A spans ± 2 bootstrap SE around its point estimate (0.989 ± 0.187). Bold: baseline. Five of 25 cells fall below parity, all in the low- δ_R , high- δ_A corner.

6.2 What If AI Learns to Reinstate?

The preceding results condition on the current estimated reinstatement rates. But the cognitive frontier is evolving rapidly, and the CZ-level estimates in Table 2 are identified from 2005–2019 data that largely predate generative AI.

Varying δ_A alone while holding $\hat{\varphi} = 1.99$ fixed and solving $\pi_R = \pi_A$ yields:

$$\delta_A^{**} \approx 1.50 > \delta_R = 1.14. \quad (15)$$

AI must reinstate *more* tasks per automation event than robots to achieve equal productivity content. Even at equal reinstatement rates, AI trails robots because sorting disruption ($\varphi = 1.99$) imposes a reallocation cost on the cognitive sector with no physical-sector analogue. The threshold is 52% above the estimate of $\hat{\delta}_A = 0.99$.

But this one-dimensional exercise is incomplete. As AI capability evolves from routine cognitive automation to generative AI, the sorting disruption parameter φ may also change. Whether it rises or falls depends on the *type* of complementary tasks that generative AI creates and how well they match the skills of displaced workers. I extend the analysis to the full (δ_A, φ) plane, examining three

objects: (i) the iso-productivity contour, the set of (δ_A, φ) pairs at which $\pi_R = \pi_A$, (ii) structurally disciplined trajectories through (δ_A, φ) space corresponding to different regimes of AI development, and (iii) bounds from emerging micro evidence on generative AI.

6.2.1 The Generalized Threshold

Define the threshold curve $\varphi^*(\delta_A)$ as the sorting disruption level at which robot and AI productivity equalize:

$$\varphi^*(\delta_A) \equiv \inf \{ \varphi \geq 0 : \pi_R(\delta_R, \delta_A, \varphi) = \pi_A(\delta_R, \delta_A, \varphi) \}, \quad (16)$$

with all remaining parameters held fixed. Both productivity objects carry the full argument list because the equilibrium links the frontiers: π_R moves with δ_A even at fixed (δ_R, φ) . Figure 5 plots π_R/π_A across the (δ_A, φ) plane with the threshold contour. Three features merit comment.

First, at the baseline estimates ($\hat{\delta}_A = 0.99, \hat{\varphi} = 1.99$), robots are 28% more productive per unit of frontier advance ($\pi_R/\pi_A = 1.28$). The estimate lies inside the robot-dominant region, and equalization at this reinstatement rate requires sorting disruption to fall to approximately 0.70, well under the estimate. Sorting is the binding margin.

Second, the threshold curve rises with reinstatement. At $\delta_A = 1.0$, equalization requires $\varphi^* \approx 0.73$. At $\delta_A = 1.2$, $\varphi^* \approx 1.24$. At $\delta_A = 1.5$, $\varphi^* \approx 1.98$. Higher reinstatement buys AI tolerance for more sorting disruption, and near the estimate the locus climbs at roughly 2.5 units of φ per unit of δ_A , the slope that makes the identified ridge run nearly parallel to it.

Third, the 90% region for the estimates ($\delta_A \in [0.71, 1.32], \varphi \in [1.48, 2.49]$) lies almost entirely in the robot-dominant zone, with all 200 bootstrap draws above parity.

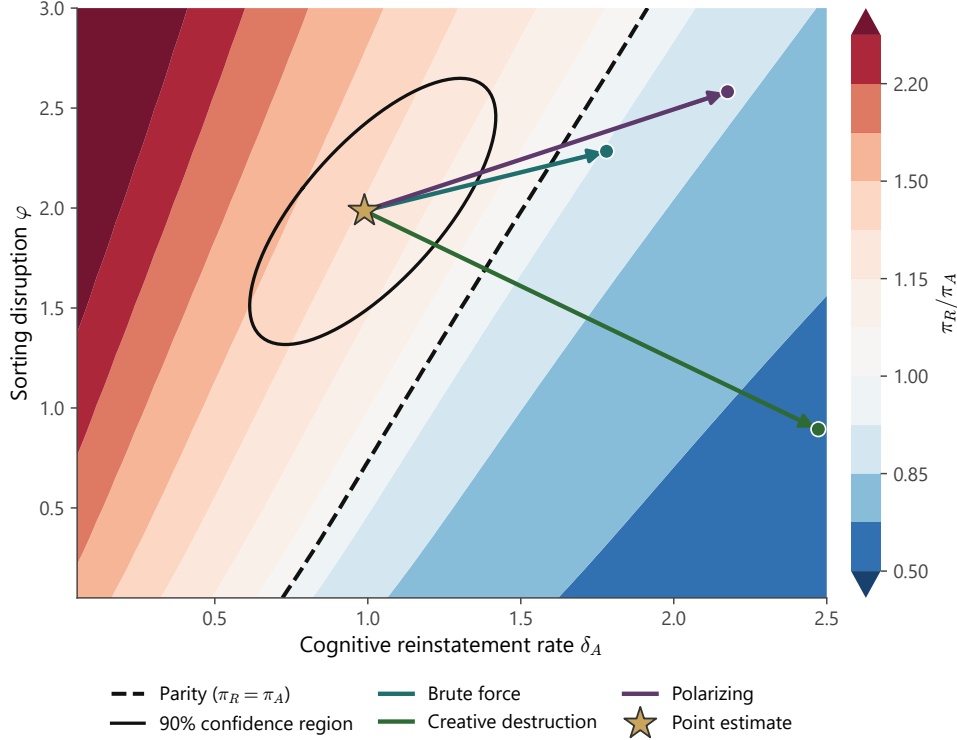


Figure 5: The productivity ratio π_R/π_A over the (δ_A, φ) plane at the baseline estimates. Red shades are robot-dominant, blue shades AI-dominant, and the dashed contour is the parity locus. The star marks the estimate $(\hat{\delta}_A = 0.99, \hat{\varphi} = 1.99)$ and the black ellipse its 90% confidence region. Each colored path runs from the estimate to the endpoint marker for that regime’s full-maturity configuration, with the arrow showing the direction of rising AI capability.

Recent experimental studies on generative AI provide independent information on the plausible region of (δ_A, φ) space. Figure 6 overlays three pieces of micro evidence on the (δ_A, φ) heatmap from Section 6.2.1.

Noy and Zhang (2023) find that access to ChatGPT reduces completion time by 40 percent and raises output quality by 18 percent on professional writing tasks. Two features of their results map directly into the model. First, ChatGPT benefits low-ability workers disproportionately, compressing the productivity distribution. This is a statement about the heterogeneous direct productivity effect of using the tool, and it does not by itself identify the sign of φ as defined here, which governs the degradation of non-college effective productivity as cognitive tasks are *automated* rather than assisted. The two margins can move in opposite directions, and the experiments speak to the first. Second, treated workers restructure their effort toward idea-generation and editing, with ChatGPT handling rough-drafting, a within-job analog of task reinstatement ($\delta_A > 0$). Be-

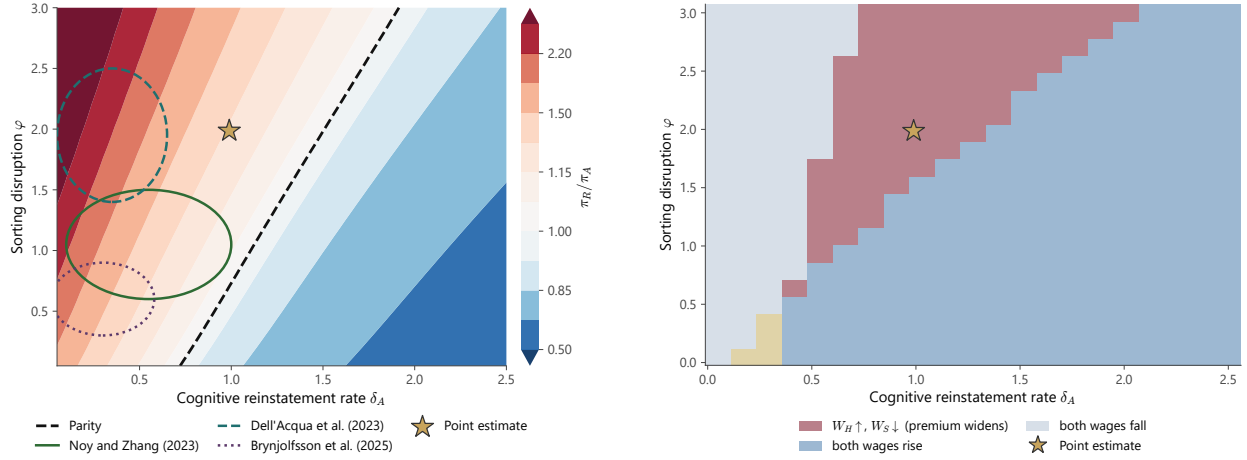
cause the restructuring occurs *within* existing occupations without creating entirely new ones, the evidence suggests moderate δ_A (below 1.0) rather than the large values that would characterize new-occupation creation. The task-restructuring evidence suggests a moderate positive δ_A , placing the study in the left half of Figure 6(a); the experiment does not locate φ . A formal mapping from the 40 percent time reduction to a specific δ_A would require decomposing the productivity gain into a direct automation component (the share of writing time absorbed by ChatGPT) and a reinstatement component (the share redirected to new editing and idea-generation tasks). The experiment does not report this decomposition, so the ellipse in Figure 6(a) should be interpreted as a qualitative consistency region and not a structural identification.

Dell’Acqua et al. (2023) document a “jagged technological frontier” among Boston Consulting Group (BCG) consultants: for tasks within AI’s capability frontier, productivity rises by 40 percent, but for tasks outside it, the likelihood of a correct solution is 19 percentage points *lower* as workers over-rely on AI. The performance degradation outside the frontier resembles a misallocation mechanism potentially related to sorting disruption, in that workers misallocate effort to AI-assisted strategies that fail on tasks requiring human judgment, but it does not identify φ as defined in this model. The evidence suggests limited δ_A ; the study’s marker in the upper-left region of the parameter space is schematic.

Brynjolfsson et al. (2025) find a 14 percent productivity gain in customer service, again concentrated among novice workers. The more modest aggregate gain and the absence of task restructuring (agents continue performing the same tasks, just faster) suggest lower δ_A than the writing-task evidence, consistent with the lower-left region of the heatmap.

Panel (b) of Figure 6 provides a complementary diagnostic: it maps the *sign pattern* of AI’s wage effects across skill groups. The CZ estimates show that AI exposure raises college wages ($\hat{\beta}_{A,C} = +0.52\%$, conventional $t = 2.07$, but exposure-robust $t = 1.46$) and lowers non-college wages ($\hat{\beta}_{A,NC} = -0.69\%$, conventional $t = -2.88$, exposure-robust $t = -2.24$). Taken on its own the college-wage coefficient is not significant under the design-based inference this paper prefers. The diagnostic therefore rests on the skill-premium contrast, which is significant under that inference (exposure-robust $t = 3.60$), and on the sign pattern rather than on the college-wage level effect alone. This wage sign pattern (W_H up, W_S down) constrains the model to the premium-widening region of panel (b), which requires $\varphi \gtrsim 0.5$ at the baseline δ_A . The point estimate

$(\delta_A = 0.99, \varphi = 1.99)$ sits squarely within this region, where the model predicts exactly the sign pattern the completed-specification moments deliver.



(a) Productivity ratio π_R/π_A with micro evidence regions. Ellipses indicate approximate (δ_A, φ) ranges consistent with each study's findings. The dashed line marks parity ($\pi_R = \pi_A$).

(b) Sign pattern of AI wage effects. The red-toned region has W_H rising and W_S falling, matching the CZ estimates. The star marks the point estimate, inside the CZ-consistent region.

Figure 6: Micro evidence anchors on the (δ_A, φ) space. The experimental-study markers are *schematic* placements illustrating the task-creation evidence. Their vertical coordinates are not empirical estimates of φ , and should not be read as such. They do not pin the sorting friction: the heterogeneous productivity effects those experiments measure concern AI used as a tool, not the degradation of skill–task allocation as cognitive work is automated, and the two margins need not move together. It is the CZ wage sign pattern that requires $\varphi \gtrsim 0.5$. The point estimate (star) lies in the region where $\pi_R/\pi_A > 1$ and is consistent with the CZ sign pattern.

Taken together, the experiments suggest positive but moderate task restructuring ($\delta_A \lesssim 1.0$) within existing occupations. Separately, the CZ wage sign pattern requires $\varphi \gtrsim 0.5$ at the baseline δ_A . The baseline estimate lies at the intersection of these two constraints.

6.2.2 Structural Co-Movement of δ_A and φ

As AI capability advances, both δ_A and φ respond. I parameterize AI capability by $\alpha \in [0, 1]$, where $\alpha = 0$ represents the current regime (routine cognitive automation, 2005–2019) and $\alpha = 1$ represents fully mature generative AI. From the free-entry condition for reinstated tasks (10), the reinstatement rate depends on the task complementarity parameter θ_A and the match quality between displaced workers' skills and new tasks. I consider three regimes spanning the structural possibilities.

Regime 1: Brute force. AI automates more cognitive tasks without changing the type of complementary tasks it creates. δ_A rises modestly as more automation generates entry opportunities, but φ stays high or rises because the automated tasks are increasingly non-routine, worsening skill-task matching. Parameterization: $\delta_A(\alpha) = 0.989(1 + 0.80\alpha)$, $\varphi(\alpha) = 1.986(1 + 0.15\alpha)$. Endpoint: $(\delta_A, \varphi) = (1.78, 2.28)$.

Regime 2: Creative destruction. Generative AI creates new task categories (prompt engineering, algorithmic auditing, human-AI collaboration) structurally well-matched to displaced workers’ existing skills. δ_A rises substantially and φ falls. This is the scenario Autor (2024) envisions when arguing that generative AI could “democratize elite expertise.” Parameterization: $\delta_A(\alpha) = 0.989(1 + 1.50\alpha)$, $\varphi(\alpha) = 1.986(1 - 0.55\alpha)$. Endpoint: $(\delta_A, \varphi) = (2.47, 0.89)$.

Regime 3: Polarizing. AI creates high-productivity complementary tasks requiring new skills (ML engineering, AI safety, complex system design) poorly matched to displaced workers (data analysts, junior lawyers, routine knowledge workers). δ_A rises but φ also rises. Parameterization: $\delta_A(\alpha) = 0.989(1 + 1.20\alpha)$, $\varphi(\alpha) = 1.986(1 + 0.30\alpha)$. Endpoint: $(\delta_A, \varphi) = (2.18, 2.58)$.

Table 9: Productivity ratio along regime trajectories at full AI maturity ($\alpha = 1$)

Regime	$\delta_A(1)$	$\varphi(1)$	$\pi_R/\pi_A(1)$	Crosses $\pi_R = \pi_A$?
1: Brute Force	1.78	2.28	0.93	Yes ($\alpha \approx 0.75$)
2: Creative Destruction	2.47	0.89	0.63	Yes ($\alpha \approx 0.25$)
3: Polarizing	2.18	2.58	0.86	Yes ($\alpha \approx 0.55$)

Notes: Each regime parameterizes the evolution of $\delta_A(\alpha)$ and $\varphi(\alpha)$ as AI capability α rises from 0 (current) to 1 (full maturity). $\delta_A(1)$ and $\varphi(1)$ are the terminal values. $\pi_R/\pi_A(1)$ is the productivity ratio at full maturity, computed from the macro model with $\delta_R = 1.14$. “Crosses $\pi_R = \pi_A$ ” reports where the trajectory passes through the parity locus. All three regimes raise δ_A by 80 to 150 percent from the estimate, which is why all three eventually cross. The starting point itself is not marginal: at the estimated sorting disruption, parity would require δ_A to rise by half. Terminal values are computed at the unrounded estimates ($\delta_A = 0.989$, $\varphi = 1.986$), so they can differ in the last digit from arithmetic on the rounded factors. See Figure 7.

Table 9 reports the results. At the baseline estimates all three regimes eventually cross parity. Regime 2 crosses early ($\alpha \approx 0.25$), requiring about a quarter of AI’s hypothetical mature capability. Regime 3 crosses at $\alpha \approx 0.55$ and Regime 1, the slowest, at $\alpha \approx 0.75$. The regimes therefore answer how much of the path to maturity robots dominate, given that all three eventually reach parity. Figure 7 traces the productivity ratio along each trajectory.⁸

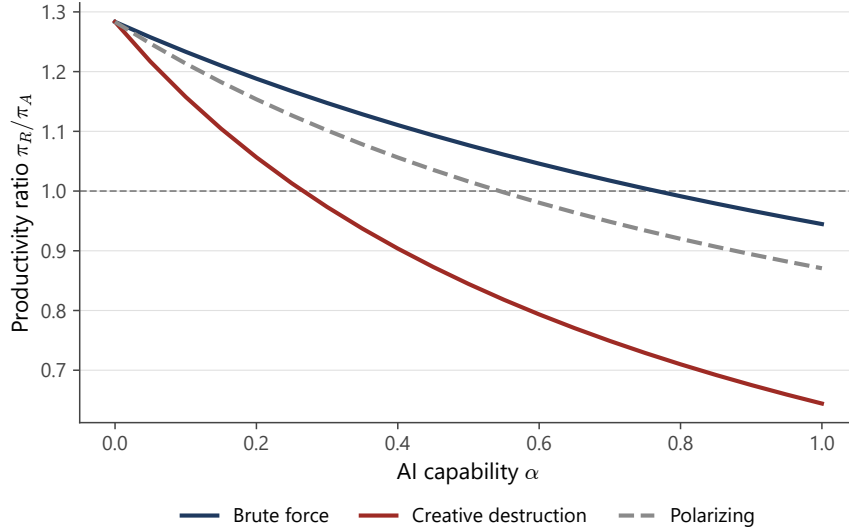


Figure 7: Productivity ratio π_R/π_A along each regime trajectory as AI capability α increases. At the baseline estimates all three trajectories eventually cross parity: Creative Destruction at $\alpha \approx 0.25$, Polarizing at $\alpha \approx 0.55$, and Brute Force at $\alpha \approx 0.75$.

6.2.3 Disciplining the Trajectories with Micro Evidence

Which regime is empirically relevant? Three recent findings provide partial guidance.

Sorting disruption is real and large for LLM automation. Freund and Mann (2026) build a quantitative task-based model with multi-dimensional skill heterogeneity and show that automation driven by large language models (LLMs) generates *larger* shifts in occupational skill composition than past robot automation. Workers specialized in the most LLM-exposed tasks leave their transformed jobs and experience earnings losses. Two of their other findings point the opposite way and should be stated: moderately exposed workers benefit on average, and lower earners gain more than higher earners, a distributional sign opposite to the CZ premium-widening moments

⁸The regime parameterizations are transparent and intentionally stylized. The point is that the three regimes span the structural possibilities, and the medium-run conclusion, robot dominance over most of the capability path under Regimes 1 and 3, does not depend on specific values within each regime.

here. The tension is informative about which margin each object measures. The sorting parameter φ is identified by the reallocation cost at the high-exposure margin, precisely the within-occupation skill-composition shifts their model quantifies, while their average-benefit and compression results describe the direct productivity channel at moderate exposure, which in this model raises π_A without loading on φ . Read at the margin the model uses, their evidence is consistent with $\varphi > 1$ and potentially $\varphi \geq \hat{\varphi}$ for generative AI, though the distributional disagreement at moderate exposure remains an open question a validation panel can settle. This evidence is inconsistent with Regime 2’s assumption that φ falls rapidly, and is most consistent with Regimes 1 or 3.

Aggregate labor market effects remain small. Humlum and Vestergaard (2025) link adoption surveys to Danish administrative records and find precisely estimated null effects on earnings and hours through 2024, despite 50%+ adoption rates in exposed occupations. This is consistent with the model’s prediction that the net labor market effect of cognitive automation, displacement minus reinstatement plus sorting disruption, is small when $\delta_A \approx 0.99$ and $\varphi \approx 1.99$. The null is inconsistent with scenarios in which δ_A has risen substantially without a corresponding decline in φ . The null’s power is directional: at current adoption scale it cannot distinguish $\delta_A = 0.99$ from the near-zero estimate of Chávez (2026), since both configurations predict effects inside the band (Appendix G.8), but a large rise in δ_A at unchanged φ would push predicted effects outside it. The null therefore disciplines upward departures only, which is exactly how this paragraph and the appendix use it.

Task-level productivity gains are concentrated in direct displacement. The experimental literature reports a 14% productivity gain in customer service (Brynjolfsson et al., 2025), a 40% productivity gain for in-frontier consulting tasks (Dell’Acqua et al., 2023), and a 40% reduction in completion time with an 18% quality gain in writing tasks (Noy and Zhang, 2023). These gains operate through the direct automation channel, not through new task creation. The reinstatement rate δ_A governs what happens *after* the task-level gain is realized. The micro evidence shows the direct channel is active but is largely silent on reinstatement.

Taken together, the evidence favors Regime 1 or Regime 3 over Regime 2 for the current generation of AI. Sorting disruption appears persistent and possibly increasing with LLM capability, not declining. Under the slowest-crossing trajectory, robot automation remains more productive per unit of frontier advance through roughly three quarters of the path to full AI maturity, and the

crossing points themselves inherit the transitional-share uncertainty bounded in Table 5.

6.2.4 A Sufficient Condition for Persistent Robot Dominance

The joint analysis yields a sharper version of Proposition 5.

Proposition 6 (Generalized Asymmetric Productivity Content). *The productivity content of robot automation exceeds AI automation ($\pi_R/\pi_A > 1$) for all (δ_A, φ) satisfying $\delta_A < \delta_A^{\min}(\varphi)$, where:*

$$\delta_A^{\min}(\varphi) \equiv \inf \{ \delta : \pi_R = \pi_A \text{ at } (\delta, \varphi) \}. \quad (17)$$

At the calibrated production structure, $\delta_A^{\min}(\varphi) > 0.72$ for all $\varphi \geq 0.05$. Robot dominance is guaranteed whenever $\delta_A < 0.72$, for any sorting disruption above $\varphi = 0.05$.

The 90% interval for $\hat{\delta}_A$ from the bootstrap is $[0.71, 1.32]$. At its upper end with negligible sorting disruption, $\pi_R/\pi_A = 0.71$ and AI dominates, but reaching that point requires $\varphi \approx 0$, which the estimate and all micro evidence reject. At $\delta_A = 1.32$ with $\hat{\varphi} = 1.99$, the ratio is 1.08 and robots still dominate. The one-dimensional threshold $\delta_A^{**} = 1.50$ is valid at $\hat{\varphi}$ but understates the difficulty of equalization: in the joint (δ_A, φ) space, AI must simultaneously achieve high reinstatement *and* low sorting disruption, a combination the micro evidence suggests is structurally difficult for cognitive automation.⁹

Outside the estimated parameters, the productivity ratio's *level* depends on the calibrated skill elasticity of substitution $\varepsilon_w = 1.50$. Re-estimating the frontier parameters at each ε_w across the range the labor literature spans, $[1.40, 2.50]$, π_R/π_A declines from 1.38 to 1.06 but never reverses, and the optimal subsidy ratio stays above 1.3 throughout (Appendix G). The ranking is thus robust to the skill elasticity while its magnitude is not, which is one reason the policy conclusion, anchored by the equity multiplier that holds the subsidy above one even at productivity parity, is the more durable of the two claims.

Re-estimation is what makes that statement true, and the distinction is worth being explicit about, because the alternative exercise gives the opposite answer. Holding the frontier parameters at their baseline values and moving only ε_w , the ratio falls below one at $\varepsilon_w \approx 1.90$ and reaches 0.95 at 2.00. That calculation is not a valid counterfactual: δ_A and φ are estimated conditional

⁹Appendix G, Section G.5 presents a 3×3 joint sensitivity grid of π_R/π_A across (δ_R, φ) pairs.

on ε_w , so at a different value the frozen pair no longer satisfies the moment conditions, and the exercise evaluates the model at parameters the data reject. The re-estimated path is the internally consistent one. But the frozen calculation is reported here because the calibrated $\varepsilon_w = 1.50$ carries no standard error and sits near the bottom of the interval just cited, and a reader is entitled to know how much of the result rests on it.

The re-estimation is not free either. Across the range $\hat{\delta}_A$ falls from 1.19 to 0.33 while δ_R is held at its imported value, so the reinstatement gap widens as ε_w rises. The fit loosens as it does: the Hansen statistic rises from 0.012 at the calibrated value to 1.405 at $\varepsilon_w = 2.50$. On two overidentifying restrictions that is $p = 0.50$, so the restrictions are not rejected anywhere in the range; what deteriorates is the margin by which the model satisfies them, not the model. What that does to the mechanism is worth tracing, because the ranking and the decomposition behind it are not equally robust. Setting $\varphi = 0$ leaves π_R/π_A below one at every point in the range, from 0.75 at $\varepsilon_w = 1.40$ to 0.96 at 2.50: without the sorting friction, robots never dominate, whatever the skill elasticity. But the amount of work the friction does falls steadily. It supplies 0.48 of the ratio at the calibrated $\varepsilon_w = 1.50$ and only 0.10 at 2.50, where the widened reinstatement gap carries most of the distance on its own. At the other end the pattern is starker still: at $\varepsilon_w = 1.40$ the estimated cognitive reinstatement rate *exceeds* the physical one, 1.19 against 1.14, so reinstatement works against the ranking and sorting disruption delivers the entire asymmetry unaided. The ordering of the two frontiers is robust to the skill elasticity. The division of labor between the two channels that produces it is not.

6.3 Labor Share Dynamics

Both frontiers lower the labor share, a generic property of task-based models where capital produces automated tasks, but with asymmetric magnitudes:

$$\frac{\partial s_L}{\partial I_R} = -0.119, \tag{18}$$

$$\frac{\partial s_L}{\partial I_A} = -0.264. \tag{19}$$

Robot automation is approximately 45% as damaging as AI per unit of frontier advance ($|\partial s_L/\partial I_R|/|\partial s_L/\partial I_A| = 0.45$, bootstrap 5th to 95th percentile range 0.43 to 0.48). The intuition runs through sorting. At

the baseline estimates the two frontiers reinstate at comparable rates, and robots at $\delta_R = 1.14$ roughly replace what they automate in labor’s task content. On the cognitive side, reinstatement at $\delta_A = 0.99$ would do the same, but sorting disruption at $\varphi = 1.99$ degrades the value of the reinstated tasks, so labor’s task content recovers less than the raw reinstatement rate suggests.

The asymmetry is visible when viewed from the skill premium. Both frontiers widen the skill premium, but AI is substantially more skill-biased: $\partial \ln(W_H/W_S)/\partial I_R = +0.076$ versus $\partial \ln(W_H/W_S)/\partial I_A = +0.314$. The mechanism is sorting disruption: as cognitive tasks are automated, non-college workers in the cognitive sector lose access to tasks where they had a foothold, while college workers, better matched to the remaining high-complexity tasks, capture the gains from whatever complementary tasks AI does create.

Neither technology raises the labor share. The policy question is not “which technology is good for labor?” It is “which damages labor least?”

7 Welfare and Optimal Policy

7.1 Optimal Frontier-Specific Subsidies

A social planner choosing frontier-specific R&D subsidies maximizes a Bergson–Samuelson welfare function with Atkinson inequality aversion λ :

$$\mathcal{W} = \sum_{k \in \{S, H\}} L_k u(c_k), \quad u(c) = \frac{c^{1-\lambda}}{1-\lambda} \quad (\lambda \neq 1), \quad u(c) = \ln c \quad (\lambda = 1), \quad (20)$$

where non-college workers are hand-to-mouth ($c_S = W_S$) and college workers are Euler savers who also receive the capital income $\Pi \equiv Y - L_S W_S - L_H W_H$. The baseline convention evaluates the marginal utility weights at wage consumption, $W_S^{-\lambda}$ and $W_H^{-\lambda}$ in equation (21) below, while assigning the full capital-income increment $d\Pi/dI_j$ to college households. Section 7.2 reports the consumption-based variants, all of which raise the ratio, so the baseline convention is the conser-

vative one.¹⁰ The social marginal product of advancing frontier j by dI is:

$$SMP_j = W_S^{-\lambda} L_S \frac{dW_S}{dI_j} + W_H^{-\lambda} \left(L_H \frac{dW_H}{dI_j} + \frac{d\Pi}{dI_j} \right). \quad (21)$$

At $\lambda = 0$ (utilitarian, no inequality aversion), $SMP_j = dY/dI_j$, and the optimal ratio reduces to the pure efficiency benchmark π_R/π_A . At $\lambda > 0$, non-college workers receive higher marginal utility weight ($W_S < W_H \Rightarrow W_S^{-\lambda} > W_H^{-\lambda}$), and the ratio tilts further toward the frontier that benefits them more, robots.

The optimal subsidy ratio equates social returns across frontiers:

$$\frac{\tau_R^*}{\tau_A^*} = \frac{SMP_R}{SMP_A} = \underbrace{\frac{\pi_R}{\pi_A}}_{1.28 \text{ (efficiency)}} \times \underbrace{\frac{1 + \mathcal{W}_R^{\text{dist}}}{1 + \mathcal{W}_A^{\text{dist}}}}_{\approx 1.35 \text{ (equity)}} \approx 1.73. \quad (22)$$

Here $\mathcal{W}_j^{\text{dist}}$ is the distributional wedge on frontier j , the share of SMP_j that comes from re-weighting each group's wage gain by its marginal utility rather than by its income, so that $\mathcal{W}_j^{\text{dist}} = 0$ at $\lambda = 0$ and the equity factor collapses to one. It is the object the table below reports as the equity multiplier. The efficiency motive alone justifies a 1.28:1 ratio. The equity motive adds approximately 35 percent because robot reinstatement creates physical tasks that disproportionately benefit non-college workers, who have higher marginal utility under concave u .

7.2 Robustness Across Social Preferences

Table 10 reports the optimal subsidy ratio across social preferences, from utilitarian ($\lambda = 0$) to high inequality aversion ($\lambda = 3$). The Rawlsian limit is not reported as a ratio because it is not defined as one. AI advances lower non-college wages, so a maxmin planner taxes AI research outright rather than granting it a smaller subsidy.

¹⁰ AI advances lower non-college wages ($\partial \ln W_S / \partial I_A = -0.111$) and the labor share ($\partial s_L / \partial I_A = -0.264$). The CZ estimate $\hat{\beta}_{A,NC} = -0.69\%$ is the local partial-equilibrium counterpart, and the general-equilibrium and local estimates agree in sign. The welfare calculations throughout this section use the general-equilibrium derivatives.

Table 10: Optimal subsidy ratio τ_R^*/τ_A^* across social preferences

λ	Social preference	τ_R^*/τ_A^*	Equity multiplier	Favors robots?
0.0	Utilitarian (efficiency only)	1.28	1.00	Yes
0.5	Mild inequality aversion	1.39	1.08	Yes
1.0	Log utility	1.53	1.19	Yes
1.5	Baseline	1.73	1.35	Yes
2.0	Moderate inequality aversion	2.03	1.58	Yes
3.0	High inequality aversion	3.18	2.48	Yes

Notes: The planner maximizes a Bergson–Samuelson SWF $\mathcal{W} = \sum_k L_k u(c_k)$ with $u(c) = c^{1-\lambda}/(1-\lambda)$

for $\lambda \neq 1$ and $u(c) = \ln c$ at $\lambda = 1$. The optimal ratio τ_R^*/τ_A^* equates social marginal products across frontiers. The equity multiplier is $(\tau_R^*/\tau_A^*)/(\pi_R/\pi_A)$, i.e. the factor by which distributional concerns amplify the efficiency case. At $\lambda = 0$ it equals 1 by construction. At the estimated parameters the Rawlsian limit is not well defined as a subsidy ratio, because AI advances lower non-college wages ($\partial \ln W_S / \partial I_A = -0.111$): a maxmin planner assigns AI research a tax, not a smaller subsidy. The table therefore reports finite values of λ only. All parameters at baseline: $\delta_R = 1.14$, $\delta_A = 0.99$, $\varphi = 1.99$.

Two features of Table 10 deserve emphasis. First, the result $\tau_R^* > \tau_A^*$ holds at *every* value of λ in the table, from pure efficiency ($\lambda = 0$, $\tau_R^*/\tau_A^* = 1.28$) to high inequality aversion, and the maxmin limit preserves the ranking in a stronger form, with AI taxed outright while robots remain subsidized. The finding is not an artifact of a particular ethical stance. Second, the equity multiplier rises monotonically with λ : distributional concerns *amplify* the efficiency case for robot subsidies instead of offsetting it. At baseline ($\lambda = 1.5$), the equity channel adds 35 percent on top of the efficiency ratio.

At today’s estimated parameters, the planner subsidizes robots more than AI across every distributional and preference specification. The magnitude (1.73 : 1 at the baseline λ) carries the uncertainty of the underlying $(\delta_R, \delta_A, \varphi)$ estimates, but the direction is highly robust. This robustness is to the estimation uncertainty and to the transition horizon over which the subsidy is computed at today’s (δ_A, φ) . It is not unconditional in AI’s development path. Under the regimes the current micro evidence favors, the brute-force and polarizing regimes of Section 6.2, the productivity ratio holds above parity over most of the capability path, crossing only at $\alpha \approx 0.75$

and $\alpha \approx 0.55$. Under the creative-destruction regime, the Autor (2024) scenario in which generative AI democratizes elite expertise, it crosses far earlier, after about a quarter of the path to AI maturity (Table 9). The subsidy tilt outlives productivity parity, because the equity motive holds the ratio above one even where the two frontiers are equally productive, but on a creative-destruction path it too eventually reverses. The frontier-specific subsidy is therefore a bet on which development path obtains, one that post-deployment data on generative AI’s reinstatement and sorting will settle.

The welfare block also embeds two knife-edge distributional assumptions. The first is that all capital income, including robot reinstatement rents, accrues to college households at the margin. The second is that marginal utilities are evaluated at wage consumption ($c_S = W_S$), which implicitly imposes $m_S = 1$ and so sits in tension with the static calibration’s own $m_S = 0.75$. Relaxing either one moves the ratio up rather than down over the range examined. Evaluating marginal utility at full per-capita consumption, with college households receiving all capital income, raises the ratio to 4.61, because capital income lowers the college welfare weight. Letting non-college households hold 10 or 20 percent of capital income gives 2.37 and 1.81, so the variants approach the baseline from above as the non-college capital share rises, and the claim is bounded by that share rather than global. The baseline 1.73 is the most conservative of these variants. The hand-to-mouth structure therefore works against the robot preference rather than manufacturing it. Proposition 7 below collects the full robustness envelope, incorporating sensitivity to λ , δ_R and ε_H .

7.3 Second-Best Considerations

The preceding analysis assumes the planner can set frontier-specific subsidies (τ_R, τ_A) . In practice, industrial policy instruments are blunt: governments may face political or informational constraints that limit them to a uniform technology subsidy τ or, at best, broad sectoral support. Three second-best considerations are relevant.

If lump-sum redistribution is available, the planner needs only to equate private marginal products across frontiers, yielding $\tau_R^*/\tau_A^* = \pi_R/\pi_A = 1.28$. The gap between the optimal ratio (1.73 at the point estimate, with a median of 1.68 across bootstrap draws) and the productivity ratio (1.28) reflects an Atkinson wedge: without perfect transfers, technology policy must do double duty as both an efficiency and a redistribution instrument.

The decomposition of social marginal products reveals why. The wage-bill share of a marginal

output gain is $s_L + \partial s_L / \partial \ln Y$, where $s_L = 0.618$ is the model’s baseline labor share and $\partial s_L / \partial \ln Y = (\partial s_L / \partial I_j) / \pi_j$ from the labor-share derivatives in Table 7. For robots, 42.7 percent of the output gain from a frontier advance accrues as wage income and 57.3 percent as capital income. For AI only 7.4 percent reaches wages, with 92.6 percent accruing as capital income, and even that sliver nets a negative non-college contribution against a positive college one. The part of AI’s return that reaches workers is small and concentrated at the top of the skill distribution. A planner who cannot redistribute capital income therefore has a strong incentive to tilt technology policy toward the wage-intensive frontier.

The second consideration reinforces this conclusion. The productivity ratio π_R / π_A measures output elasticity with respect to frontier *position* and is invariant to capital depreciation rates. But the cost of maintaining a given frontier level scales with d_j , so the cost-adjusted subsidy ratio is $(\pi_R / \pi_A) \cdot (d_A / d_R)$. If AI capital depreciates faster than physical robots, the cost-adjusted ratio exceeds the output-based ratio. This is plausible given rapid model obsolescence, with frontier LLMs depreciating on 1–2 year cycles versus 10–15 year robot lifespans. Taken literally, the LLM cycle implies $d_A \approx 0.5$ or higher and multiplies the cost-adjusted ratio five-fold or more, further strengthening the case for robot subsidies (Appendix G).

If the planner is restricted to a uniform subsidy $\tau_R = \tau_A = \tau$, the welfare cost depends on the gap in social marginal products at the undifferentiated allocation. At the baseline estimates, SMP_R exceeds SMP_A by 73 percent, a substantial misallocation. The uniform subsidy over-invests in the AI frontier (relative to the social optimum) and under-invests in robot adoption.

Under the quadratic approximation to the social welfare function, the optimal reallocation shifts a modest share of the AI R&D budget toward robots, and even a partial tilt captures most of the attainable welfare gain, because the social return function is locally concave and welfare losses from overshooting are second-order. Frontier-specific subsidies need not be precisely calibrated to deliver most of their benefit. The first-order gain from *any* differentiation dominates the second-order loss from imprecision.

Proposition 7 (Policy Robustness). *Across the bootstrap draws of (δ_A, φ) , the transition-horizon subsidy ratio τ_R^* / τ_A^* lies between 1.41 and 2.14 (5th to 95th percentile of those draws, computed under the baseline wage-weight convention of Section 7, the most conservative of the distributional*

variants), and it exceeds one for every finite value of the inequality-aversion parameter λ reported in Table 10 and across the bootstrap 95 percent interval for δ_R . In the maxmin limit the ratio is not defined because the planner taxes AI research rather than granting it a positive subsidy. It falls below one only in the joint low- δ_R , high- δ_A corner of the wider Humlum grid, holding φ at its estimate, the same corner that reverses the productivity ranking. Across the bootstrap draws it exceeds one everywhere, and the equity motive would hold it above one even at productivity parity. Because the optimal frontier-specific subsidies are positive for both frontiers at every finite λ , and the social return function is locally concave, a uniform automation tax moves both instruments away from their optima and is locally welfare-reducing. The statement concerns marginal uniform taxation near the optimum, not taxes of arbitrary size.

I leave the formal dynamic optimal policy problem, where $\delta_A(\alpha)$ and $\varphi(\alpha)$ evolve endogenously with the AI capability frontier, to future work.

8 Conclusion

In a two-frontier growth model estimated on U.S. commuting-zone data, robot automation generates about a quarter more aggregate productivity per unit of frontier advance than AI automation, a ranking that holds under both this paper’s cognitive estimate and a near-zero one. What carries the asymmetry depends on the estimate. Where the two rebuild at comparable rates, a cognitive-specific sorting friction accounts for it, as AI degrades the matching of skilled workers to the cognitive tasks that remain. Where AI rebuilds little, the reinstatement gap accounts for it. Either way the drag falls on the cognitive frontier, with no counterpart on the physical side. The steady-state ranking depends on how much of the measured sorting friction is transitional, and the paper bounds that dependence and leaves its resolution to post-deployment data.

Three macro results followed from a single growth decomposition. Demand channels are second-order: the distributional consequences of asymmetric automation are set on the supply side. Both frontiers erode the labor share, but robots do roughly half the damage per unit of advance, the most tightly pinned down of the paper’s headline ratios. And over the transition horizon a planner subsidizes robot research at between 1.4 and 2.1 times the rate of AI research across the central ninety percent of the estimated parameter set, and at between 1.2 and 2.6 times across all of

it, a ranking robust to the estimation uncertainty and holding over most of the capability path under the development regimes the current evidence favors, though it erodes as AI matures, fastest along a creative-destruction path that democratizes elite expertise. The sharpest empirical result is distributional: AI raises the skill premium, significantly under exposure-robust inference and the within-skill-bin permutation null, so AI’s productivity gains concentrate among the workers already earning the most. The fully dynamic optimal-policy problem, in which the planner sets frontier-specific subsidies while δ_A and φ themselves evolve along the AI development path, is left for future work.

A limitation of the empirical framework is that the reduced-form moments are estimated from 2005–2019 commuting zone data, with AI exposure measured prior to the deployment of large language models. The structural estimates therefore reflect the reinstatement and sorting characteristics of narrow, task-specific AI, not frontier generative models. Two observations temper this concern. First, the direction favored by the current micro evidence, persistent or increasing sorting disruption, delays the reversal: generative AI plausibly intensifies sorting disruption, since language models can perform cognitive tasks that previously required college-level judgment, while the effect on reinstatement is ambiguous. Creative-destruction paths in which sorting disruption falls instead weaken the result and produce an earlier crossing (Section 6.2). The identification analysis shows that the headline finding survives across the entire range of parameter combinations consistent with the data, including values well beyond the point estimates. Second, the policy comparison is a relative statement about marginal returns to two types of automation. Reversing it would require either cognitive automation to generate complementary tasks at rates well above those of industrial robots, or the sorting friction to be substantially smaller than estimated. The second route is the live one: at the estimated reinstatement rates, which are nearly symmetric, it is the sorting friction alone that puts robots ahead, and a sufficiently low value of it reverses the ranking without any change in the rate at which the two frontiers create tasks. Updating this ranking as post-GPT data accumulate is a priority for future work.

The qualitative message is robust: automation technologies that create complementary human tasks without degrading the skill–task allocation are more productive than technologies whose advance imposes larger sorting losses, and the distinction matters quantitatively for growth, distribution, and policy.

References

- Acemoglu, D. (2024). The simple macroeconomics of AI. NBER Working Paper No. 32487.
- Acemoglu, D. and Restrepo, P. (2018). The race between man and machine: Implications of technology for growth, factor shares, and employment. *American Economic Review*, 108(6):1488–1542.
- Acemoglu, D. and Restrepo, P. (2020). Robots and jobs: Evidence from US labor markets. *Journal of Political Economy*, 128(6):2188–2244.
- Acemoglu, D. and Restrepo, P. (2022). Tasks, automation, and the rise in US wage inequality. *Econometrica*, 90(5):1973–2016.
- Adão, R., Kolesár, M., and Morales, E. (2019). Shift-share designs: Theory and inference. *Quarterly Journal of Economics*, 134(4):1949–2010.
- Angrist, J. D. and Pischke, J.-S. (2009). *Mostly Harmless Econometrics*. Princeton University Press.
- Autor, D. (2024). Applying AI to rebuild middle class jobs. NBER Working Paper No. 32140.
- Autor, D. H., Dorn, D., and Hanson, G. H. (2013). The china syndrome: Local labor market effects of import competition in the United States. *American Economic Review*, 103(6):2121–2168.
- Babina, T., Fedyk, A., He, A., and Hodson, J. (2024). Artificial intelligence, firm growth, and product innovation. *Journal of Financial Economics*, 151:103745.
- Baumol, W. J. (1967). Macroeconomics of unbalanced growth: The anatomy of urban crisis. *American Economic Review*, 57(3):415–426.
- Borusyak, K., Hull, P., and Jaravel, X. (2022). Quasi-experimental shift-share research designs. *Review of Economic Studies*, 89(1):181–213.
- Brynjolfsson, E., Li, D., and Raymond, L. R. (2025). Generative AI at work. *Quarterly Journal of Economics*, 140(2):889–942.
- Card, D. (2009). Immigration and inequality. *American Economic Review*, 99(2):1–21.

- Chávez, C. (2026). Automation is not a sufficient statistic: Robots, ai, and inequality. Working paper, University of Chicago.
- Chetty, R., Guren, A., Manoli, D., and Weber, A. (2013). Are micro and macro labor supply elasticities consistent? A review of evidence on the intensive and extensive margins. *American Economic Review: Papers & Proceedings*, 103(3):471–475.
- Dauth, W., Findeisen, S., Suedekum, J., and Woessner, N. (2021). The adjustment of labor markets to robots. *Journal of the European Economic Association*, 19(6):3104–3153.
- Dell’Acqua, F., McFowland, E., Mollick, E. R., Lifshitz-Assaf, H., Kellogg, K., Rajendran, S., Kraymer, L., Candelon, F., and Lakhani, K. R. (2023). Navigating the jagged technological frontier: Field experimental evidence of the effects of AI on knowledge worker productivity and quality. Harvard Business School Working Paper No. 24-013.
- Dix-Carneiro, R. (2014). Trade liberalization and labor market dynamics. *Econometrica*, 82(3):825–885.
- Dynan, K. E., Skinner, J., and Zeldes, S. P. (2004). Do the rich save more? *Journal of Political Economy*, 112(2):397–444.
- Eloundou, T., Manning, S., Mishkin, P., and Rock, D. (2024). GPTs are GPTs: An early look at the labor market impact potential of large language models. *Science*, 384(6702):eadj2942.
- Freund, L. B. and Mann, L. F. (2026). Job transformation, specialization, and the labor market effects of AI. CESifo Working Paper No. 12072.
- Goldsmith-Pinkham, P., Sorkin, I., and Swift, H. (2020). Bartik instruments: What, when, why, and how. *American Economic Review*, 110(8):2586–2624.
- Graetz, G. and Michaels, G. (2018). Robots at work. *Review of Economics and Statistics*, 100(5):753–768.
- Hémous, D. and Olsen, M. (2022). The rise of the machines: Automation, horizontal innovation, and income inequality. *American Economic Journal: Macroeconomics*, 14(1):179–223.

- Herrendorf, B., Rogerson, R., and Valentinyi, Á. (2014). Growth and structural transformation. In Aghion, P. and Durlauf, S. N., editors, *Handbook of Economic Growth*, volume 2, pages 855–941. Elsevier.
- Humlum, A. (2022). Robot adoption and labor market dynamics. Working paper, ROCKWOOL Foundation Study Paper No. 175.
- Humlum, A. and Vestergaard, E. (2025). Large language models, small labor market effects. NBER Working Paper No. 33777.
- Karabarbounis, L. and Neiman, B. (2014). The global decline of the labor share. *Quarterly Journal of Economics*, 129(1):61–103.
- Katz, L. F. and Murphy, K. M. (1992). Changes in relative wages, 1963–1987: Supply and demand factors. *Quarterly Journal of Economics*, 107(1):35–78.
- Koch, M., Manuylov, I., and Smolka, M. (2021). Robots and firms. *Economic Journal*, 131(638):2553–2584.
- Moretti, E. (2011). Local labor markets. *Handbook of Labor Economics*, 4:1237–1313.
- Notowidigdo, M. J. (2020). The incidence of local labor demand shocks. *Journal of Labor Economics*, 38(3):687–725.
- Noy, S. and Zhang, W. (2023). Experimental evidence on the productivity effects of generative artificial intelligence. *Science*, 381(6654):187–192.
- Piketty, T. (2014). *Capital in the Twenty-First Century*. Harvard University Press, Cambridge, MA.
- Solon, G., Haider, S. J., and Wooldridge, J. M. (2015). What are we weighting for? *Journal of Human Resources*, 50(2):301–316.
- Sun, Y. (2006). The exact law of large numbers via Fubini extension and characterization of insurable risks. *Journal of Economic Theory*, 126(1):31–69.
- Traiberman, S. (2019). Occupations and import competition: Evidence from Denmark. *American Economic Review*, 109(12):4260–4301.

Webb, M. (2020). The impact of artificial intelligence on the labor market. *Working paper, Stanford University*.

A Production Function Derivation

This appendix derives Lemma 1 and Proposition 1 from the disaggregated task production function (2). The derivation proceeds in five steps. Throughout, $\varsigma \equiv (\sigma - 1)/\sigma$ with $\sigma > 1$.

Assumption A1 (Uniform productivity within segments). *All surviving labor tasks in sector j share a common productivity: $g_{Lj}(z) = \bar{g}_{Lj}$ for all $z \in (I_j, 1]$. All reinstated tasks share a common productivity: $g_{Rj}(z) = \gamma_j B_j$ for all $z \in (1, 1 + \delta_j I_j]$.*

This is the standard assumption in Acemoglu and Restrepo (2020). We derive all results under A1, then state in Lemma 2 the exact aggregation result that allows $\gamma_j B_j$ to be interpreted as a CES power mean when reinstated-task productivities are heterogeneous.

Step 1: Optimal Capital Allocation Across Automated Tasks

The firm allocates capital across automated tasks $z \in [0, I_j]$. At task z , output is $y_{Kj}(z) = g_{Kj}(z) k_j(z)$. The firm solves:

$$\max_{\{k_j(z)\}} \int_0^{I_j} (g_{Kj}(z) k_j(z))^\varsigma dz \quad \text{s.t.} \quad \int_0^{I_j} k_j(z) dz = K_j. \quad (\text{A.1})$$

First-order condition:

$$g_{Kj}(z)^\varsigma k_j(z)^{\varsigma-1} = \lambda_K \quad \forall z \in [0, I_j]. \quad (\text{A.2})$$

Since $\varsigma - 1 < 0$, solving for $k_j(z)$:

$$\bar{k}_j(z) = \frac{g_{Kj}(z)^{\varsigma/(1-\varsigma)}}{\int_0^{I_j} g_{Kj}(z')^{\varsigma/(1-\varsigma)} dz'} \cdot K_j = \frac{g_{Kj}(z)^{\sigma-1}}{\int_0^{I_j} g_{Kj}(z')^{\sigma-1} dz'} \cdot K_j, \quad (\text{A.3})$$

where we used $\varsigma/(1-\varsigma) = (\sigma-1)/\sigma \cdot \sigma = \sigma-1$. Capital is allocated proportionally to $g_{Kj}(z)^{\sigma-1}$.

Substituting back:

$$\begin{aligned} Z_j &\equiv \int_0^{I_j} g_{Kj}(z)^\varsigma \bar{k}_j(z)^\varsigma dz \\ &= \frac{K_j^\varsigma}{[\mathcal{G}_{Kj}]^\varsigma} \int_0^{I_j} g_{Kj}(z)^{\varsigma+(\sigma-1)\varsigma} dz \\ &= \frac{K_j^\varsigma}{\mathcal{G}_{Kj}^\varsigma} \cdot \mathcal{G}_{Kj} = \mathcal{G}_{Kj}^{1-\varsigma} \cdot K_j^\varsigma = \mathcal{G}_{Kj}^{1/\sigma} \cdot K_j^\varsigma, \end{aligned} \quad (\text{A.4})$$

where $\mathcal{G}_{Kj} \equiv \int_0^{I_j} g_{Kj}(z)^{\sigma-1} dz$ and we used $\varsigma + (\sigma-1)\varsigma = \sigma\varsigma = \sigma-1$ and $1-\varsigma = 1/\sigma$.

Units: $[g_{Kj}] = [\text{output/capital}]$, so $[Z_j] = [(\text{output/capital})^{\sigma-1}]^{1/\sigma} \cdot [\text{capital}]^\varsigma = [\text{output}]^\varsigma$. $\square$

Step 2: Optimal Labor Allocation Across Labor Tasks

The firm allocates effective labor L_j^{eff} across surviving tasks $z \in (I_j, 1]$ and reinstated tasks $z \in (1, 1 + \delta_j I_j]$. Under Assumption A1, the productivities are $\bar{g}_{Lj}$ and $\gamma_j B_j$ respectively. The firm solves:

$$\max_{\{\ell_j(z)\}} \int_{I_j}^{1+\delta_j I_j} (g_j(z) \ell_j(z))^\varsigma dz \quad \text{s.t.} \quad \int_{I_j}^{1+\delta_j I_j} \ell_j(z) dz = L_j^{\text{eff}}. \quad (\text{A.5})$$

The first-order condition gives optimal allocation proportional to $g_j(z)^{\sigma-1}$:

$$\bar{\ell}_j(z) = \frac{g_j(z)^{\sigma-1}}{\mathcal{G}_{Lj}} \cdot L_j^{\text{eff}}, \quad \mathcal{G}_{Lj} \equiv \int_{I_j}^{1+\delta_j I_j} g_j(z)^{\sigma-1} dz. \quad (\text{A.6})$$

Substituting back (identical algebra to Step 1):

$$\int_{I_j}^{1+\delta_j I_j} g_j(z)^\varsigma \bar{\ell}_j(z)^\varsigma dz = \mathcal{G}_{Lj}^{1/\sigma} \cdot (L_j^{\text{eff}})^\varsigma. \quad (\text{A.7})$$

Under Assumption A1, separate $\mathcal{G}_{Lj}$ into segments:

$$\mathcal{G}_{Lj} = (1 - I_j) \bar{g}_{Lj}^{\sigma-1} + \delta_j I_j (\gamma_j B_j)^{\sigma-1}. \quad (\text{A.8})$$

Define:

$$\tilde{\Gamma}_j \equiv \mathcal{G}_{Lj}^{1/\sigma} = \left[(1 - I_j) \bar{g}_{Lj}^{\sigma-1} + \delta_j I_j (\gamma_j B_j)^{\sigma-1} \right]^{1/\sigma}. \quad (\text{A.9})$$

This is the exact expression for $\tilde{\Gamma}_j$: it is a CES aggregate of the surviving-task and reinstated-task content, raised to $1/\sigma$.

Relationship to the main text notation. Equation (7) in the main text is exactly (A.9): under Assumption A1, $\int_{I_j}^1 g_{Lj}(z)^{\sigma-1} dz = (1 - I_j) \bar{g}_{Lj}^{\sigma-1}$, so the two displays coincide, and without A1 the integral form of (7) is the general statement of Step 2. Proposition 1 is therefore exact.

Some expository passages use the additive shorthand for the two components of labor's task content:

$$\underbrace{\left[(1 - I_j) \bar{g}_{Lj}^{\sigma-1} + \delta_j I_j (\gamma_j B_j)^{\sigma-1} \right]^{1/\sigma}}_{\text{exact CES aggregate, eq. (7)}} \quad \text{vs.} \quad \underbrace{(1 - I_j) \bar{g}_{Lj}^\varsigma + \delta_j I_j (\gamma_j B_j)^\varsigma}_{\text{additive shorthand}}. \quad (\text{A.10})$$

The two are equal only in special cases, including when a single segment is active, because $[a^{\sigma-1} + b^{\sigma-1}]^{1/\sigma} \neq a^\varsigma + b^\varsigma$ generically for $\sigma \neq 1$. They differ by the concavity of $x \mapsto x^{1/\sigma}$. At the calibrated parameters ($\sigma = 2$, $I_R = 0.30$, $I_A = 0.25$, task productivities close across segments) the discrepancy is small, and the shorthand follows the [Acemoglu and Restrepo \(2018\)](#) efficiency-unit convention in which g^ς names the reduced-form task-content coefficient. All quantitative results in the paper are computed numerically from the full CES equilibrium, never from the shorthand.

Units: $[\tilde{\Gamma}_j] = [(\text{output}/\text{labor})^{\sigma-1}]^{1/\sigma} = [\text{output}/\text{labor}]^\varsigma$, so $[\tilde{\Gamma}_j \cdot (L_j^{\text{eff}})^\varsigma] = [\text{output}]^\varsigma = [Z_j]$. $\square$

Step 3: CES Aggregation of Heterogeneous Reinstated Tasks

Step 2 assumed all reinstated tasks share productivity $\gamma_j B_j$ (Assumption A1). We now relax this for reinstated tasks and show that $\gamma_j B_j$ can be interpreted as a power mean.

Lemma 2 (Heterogeneous Reinstatement Aggregation). *Let reinstated tasks at frontier j have productivities $g_{Rj}(z) = \gamma(z) \cdot B_j$, where $\gamma(z)$ are i.i.d. draws from G_j^* . Under optimal labor allocation and the exact law of large numbers for continuum economies ([Sun, 2006](#)):*

$$\int_1^{1+\delta_j I_j} (\gamma(z) B_j)^{\sigma-1} dz = \delta_j I_j \cdot B_j^{\sigma-1} \cdot \mathbb{E}_{G_j^*}[\gamma^{\sigma-1}] \quad a.s. \quad (\text{A.11})$$

Define the $(\sigma - 1)$ -power mean:

$$\gamma_j \equiv \left(\mathbb{E}_{G_j^*}[\gamma^{\sigma-1}] \right)^{1/(\sigma-1)}. \quad (\text{A.12})$$

Then the reinstated-task component of $\mathcal{G}_{Lj}$ is $\delta_j I_j \cdot (\gamma_j B_j)^{\sigma-1}$, and all results from Step 2 hold with $\gamma_j B_j$ replaced by the power mean (A.12). When G_j^* is degenerate (homogeneous reinstated tasks), the power mean reduces to the common productivity.

Proof. Equation (A.11) follows from the exact LLN applied to the i.i.d. continuum $\{(\gamma(z)B_j)^{\sigma-1}\}_{z \in (1, 1+\delta_j I_j]}$. Substituting $\mathbb{E}[\gamma^{\sigma-1}] = \gamma_j^{\sigma-1}$ into (A.8) gives $\mathcal{G}_{Lj} = (1 - I_j) \bar{g}_{Lj}^{\sigma-1} + \delta_j I_j (\gamma_j B_j)^{\sigma-1}$, identical in form to the homogeneous case. The remainder of Step 2 is unchanged. $\square$

Remark on the power mean exponent. The object that emerges from the CES optimal allocation is the $(\sigma - 1)$ -power mean (A.12), not the ς -power mean. At $\sigma = 2$: $\gamma_j = \mathbb{E}[\gamma]$ (arithmetic mean). At general σ : $\gamma_j = (\mathbb{E}[\gamma^{\sigma-1}])^{1/(\sigma-1)}$. Under the AR efficiency-unit convention in the main text, where $(\gamma_j B_j)^\varsigma$ denotes the reduced-form task content, the distinction is absorbed into the definition of γ_j . When G_j^* is degenerate, all power means coincide and the exponent is irrelevant.

Step 4: Assembly (Proof of Proposition 1)

Proof. The firm's problem separates into capital allocation (automated tasks) and labor allocation (surviving + reinstated tasks) because the two sets use different factors and enter the CES aggregator additively.

Capital side (Step 1):

$$\int_0^{I_j} y_{Kj}(z)^\varsigma dz = Z_j(I_j),$$

with $[Z_j] = [\text{output}]^\varsigma$.

Labor side (Steps 2–3):

$$\int_{I_j}^{1+\delta_j I_j} y_{Lj/Rj}(z)^\varsigma dz = \tilde{\Gamma}_j(I_j, \delta_j) \cdot (L_j^{\text{eff}})^\varsigma,$$

with $[\tilde{\Gamma}_j \cdot (L_j^{\text{eff}})^\varsigma] = [\text{output}]^\varsigma$.

Substituting into (2):

$$Y_j = \left[Z_j(I_j) + \tilde{\Gamma}_j(I_j, \delta_j) \cdot (L_j^{\text{eff}})^\varsigma \right]^{1/\varsigma}.$$

Both terms inside the bracket have units $[\text{output}]^\varsigma$. Y_j has units $[\text{output}]$. $\square$

Step 5: Reinstatement Rate and Frontier Position

Equation (10) expresses δ_j as a function of $(I_j, \theta_j, v_j, W_{k(j)}^d, \eta, c_0)$. The total measure of reinstated tasks is:

$$\delta_j \cdot I_j = \left[\frac{\theta_j}{c_0(1+\eta)} (v_j - W_{k(j)}^d) \right]^{1/\eta}, \quad (\text{A.13})$$

which does not depend directly on I_j . The equilibrium objects $(v_j, W_{k(j)}^d)$ generically do depend on I_j , so the total measure is not invariant to the frontier position. The estimates $\delta_R = 1.14$ and $\delta_A = 0.99$ are evaluated at $(I_R, I_A) = (0.30, 0.25)$. For comparative statics via finite differences ($dI = 0.01$), the equilibrium dependence through $(v_j, W_{k(j)}^d)$ is handled numerically rather than suppressed.

B Proofs and Derivations

B.1 The Demand-Side Equilibrium

The quantitative core of Section 6 is two-sector. This appendix specifies the three demand-side blocks of Sections 2.2–2.4 as an equilibrium and solves it, which is what produces the computed channels of Table 6.

Services. A non-tradable sector produces $Y_N = A_N L_N$ with labor only and is not automated. Households of skill k spend a fixed share α_k^N of income on services, with $\alpha_S^N = 0.40 > \alpha_H^N = 0.25$ (Table 3), so non-college households are the more service-intensive. Service-market clearing sets $p_N Y_N$ equal to aggregate service expenditure, and the labor released from services is absorbed by the two tradable sectors at the prevailing skill wages.

Participation. Reservation wages are drawn from a distribution with constant elasticity, so that labor supply is $L_k = \bar{L}_k (W_k / \bar{W}_k)^{\varepsilon_k^{LS}}$ with $\varepsilon_S^{LS} = 0.25$ and $\varepsilon_H^{LS} = 0.10$ (Table 3). This is the simplest distribution consistent with the calibrated elasticities, which are all the model of Section 2.3 pins down.

Capital income. Automation and reinstatement rents accrue as the residual $\Pi = Y - W_S L_S - W_H L_H$ and are distributed across households, who consume out of them at $m_K = 0.20$ against $m_S = 0.75$ and $m_H = 0.45$ (Table 3). Unconsumed income is saved and spent on tradables.

Two checks discipline the exercise, and one tension must be reported. With services switched off,

the extended model reproduces the two-sector core of Section 6 to machine precision. The service block is then calibrated to a single target, the share of service-sector workers who are non-college, 0.68. It is not validated against any moment it was not fitted to. The effective service expenditure share of 0.32 and the local demand multiplier of 1.32 that the block reports are not independent checks on it: the multiplier is an algebraic transform of the expenditure share, which is itself an income-weighted average of the two calibrated expenditure shares and is therefore insensitive to the scale of the service sector.

The tension is that the resulting service sector is smaller than Table 3 assumes. The equilibrium delivers a service share of value added of 0.15 and a share of non-college workers employed in services of 0.27, against the $\mu_N = 0.35$ service-sector aggregate-share anchor and the $s_N^S = 0.45$ share of non-college workers employed in services. The latter uses the non-college workforce as its denominator, whereas the 0.68 target above is the share of service-sector workers who are non-college. No service-sector scale satisfies both of those anchors at once. The channels below are therefore reported at the budget-consistent calibration, and also at two variants that force each of the table’s anchors in turn. Forcing $\mu_N = 0.35$ raises $\mathcal{F}^{NT}$ to +0.00031 and leaves every conclusion of this section intact. That is why the claim made in the text is an absolute bound holding across all three calibrations, rather than a point estimate at any one of them.

The computed channels are in Table 6. Three features of them are worth stating, because none was visible in the analytic bounds this appendix replaces.

First, the MPC channel is *negative*. A shift toward AI raises the capital income share and so lowers the aggregate propensity to consume, which contracts service demand rather than expanding it. Its sign depends on the service expenditure share out of capital income, which the calibration does not pin down; across the range $[0, 0.40]$ the channel runs from +0.00002 to -0.00004 , and it is negligible throughout.

Second, the participation margin enters negatively, at -0.00010 : the wage dampening exceeds the income amplification, so participation reduces the total feedback, as Section 2.3 conjectured.

Third, on the cognitive frontier the wage-feedback component $\mathcal{F}^{PE}$ nearly cancels, at 0.00004, and the non-tradable channel is larger than it. That is a fact about the denominator. The non-tradable channel has its own scale, set by the size of the service sector rather than by the wage-composition index, so any statement of the form “the demand channels are a fraction of the partial-

equilibrium feedback” fails where that feedback is small. The claim this paper makes is absolute instead: $|\mathcal{F}^{GE}| \leq 0.0023$ on both frontiers, more than two orders of magnitude below the threshold that would matter, and that claim is robust to every specification choice above.

C Dynamic Model Derivation

This appendix provides the full derivation of the dynamic amplification factor $\mathcal{A}^{\text{dyn}}$ stated in Proposition 4. The derivation proceeds in seven steps.

C.1 State Variables and Laws of Motion

The economy is characterized by the vector (K_R, K_A, A_S, A_H) with laws of motion:

$$\dot{K}_R = s_R(K_R, K_A, A_S, A_H) \cdot S - d_R K_R, \quad (\text{C.1})$$

$$\dot{K}_A = s_A(K_R, K_A, A_S, A_H) \cdot S - d_A K_A, \quad (\text{C.2})$$

$$\dot{A}_S = r A_S + W_S L_S - C_S, \quad (\text{C.3})$$

$$\dot{A}_H = r A_H + W_H L_H - C_H, \quad (\text{C.4})$$

where $\mathcal{S}(t) = Y(t) - C(t)$ is aggregate savings (calligraphic, distinct from the skill subscript S), $s_j \in [0, 1]$ is the share allocated to frontier j (with $s_R + s_A = 1$), and d_j is depreciation. Frontier thresholds evolve as $\dot{I}_j = \zeta_j K_j^{\xi_j}$, and the investment allocation tracks relative returns: $s_R/s_A = (\mathcal{R}_R^{\text{auto}}/\mathcal{R}_A^{\text{auto}})^{1/\nu}$, where $\nu > 0$ governs responsiveness. Household type k has Euler equation $\dot{C}_k/C_k = \sigma_c(r(t) - \rho_k)$.

C.2 Simplify to the Relevant Margin

The static equilibrium conditions (labor markets, automation FOCs, reinstatement free entry) determine ω as a function of the state at each instant: $\omega(t) = \tilde{\omega}(K_R(t), K_A(t), A_S(t), A_H(t))$. The impact response to a decline in r_A depends only on the direct effect on the automation FOC (no state variable moves at $t = 0$):

$$\Delta\omega^{\text{impact}} = \frac{\partial \tilde{\omega}}{\partial r_A} \Delta r_A. \quad (\text{C.5})$$

C.3 Characterize the Savings Channel

With heterogeneous ρ_k , I adopt a tractable two-asset structure: non-college workers are hand-to-mouth ($C_S = m_S W_S L_S$ with $m_S = 1$) and college workers are Euler-equation savers. Aggregate savings are:

$$S = W_H L_H - C_H + \Pi, \quad (\text{C.6})$$

where C_H satisfies $\dot{C}_H/C_H = \sigma_c(r - \rho_H)$. In steady state: $\mathcal{S}^* = d_R K_R^* + d_A K_A^*$ and $r^* = \rho_H$.

C.4 Log-Linearize the Investment Allocation

From the investment share equation, log-linearizing around steady state:

$$\hat{s}_A \approx \frac{s_R^*}{\nu} (\hat{\mathcal{R}}_A - \hat{\mathcal{R}}_R),$$

where hats denote log-deviations. Since returns depend on ω :

$$\hat{s}_A = \frac{s_R^*}{\nu} \frac{\partial \ln(\mathcal{R}_A/\mathcal{R}_R)}{\partial \omega} \hat{\omega} \equiv \eta_{\text{sav}} \hat{\omega},$$

where η_{sav} is the savings-reallocation elasticity (a symbol distinct from the incidence share η_S of Section 3.6).

C.5 Log-Linearize the Capital Law of Motion

From (C.2): $\dot{K}_A = s_A \mathcal{S} - d_A K_A$. Using $s_A^* \mathcal{S}^*/K_A^* = d_A$ in steady state:

$$\dot{\hat{K}}_A = d_A (\hat{s}_A + \hat{\mathcal{S}} - \hat{K}_A). \quad (\text{C.7})$$

C.6 Close with the Savings Response

The permanent decline in r_A raises AI returns, so $\hat{s}_A > 0$. College income $W_H L_H$ and profits Π also respond. Under the hand-to-mouth assumption, aggregate savings respond only through college income and profits. Log-linearizing:

$$\hat{\mathcal{S}} = \eta_W \hat{\omega},$$

where $\eta_W \equiv (W_H^* L_H^* / \mathcal{S}^*)(\partial \ln W_H / \partial \omega) + (\Pi^* / \mathcal{S}^*)(\partial \ln \Pi / \partial \omega)$ captures how the automation mix affects college savings through wages and profits.

C.7 Solve for the New Steady State

In the new steady state, $\dot{K}_A = 0$, so $\hat{K}_A^{ss} = \hat{s}_A^{ss} + \hat{\mathcal{S}}^{ss}$. The automation mix responds to capital: $\hat{\omega}^{ss} = \hat{\omega}^{\text{impact}} + \eta_K \hat{K}_A^{ss}$, where $\eta_K \equiv (\partial \omega / \partial K_A)(K_A^* / \omega^*)$. Substituting:

$$\hat{\omega}^{ss} = \hat{\omega}^{\text{impact}} + \eta_K(\eta_{\text{sav}} + \eta_W)\hat{\omega}^{ss}.$$

Solving:

$$\hat{\omega}^{ss} = \frac{\hat{\omega}^{\text{impact}}}{1 - \eta_K(\eta_{\text{sav}} + \eta_W)}. \quad (\text{C.8})$$

The amplification factor is:

$$\mathcal{A}^{\text{dyn}} = \frac{1}{1 - \eta_K(\eta_{\text{sav}} + \eta_W)}. \quad (\text{C.9})$$

C.8 Express in Terms of Primitives

The composite elasticity $\eta_K(\eta_{\text{sav}} + \eta_W)$ depends on three primitive objects:

- *College savings response*: proportional to σ_c (a higher intertemporal elasticity of substitution, IES, means college workers smooth less, save more in response to income gains).
- *Patience gap*: $\rho_S - \rho_H$ determines how fast wealth inequality grows along the transition, which determines how strongly investment tilts toward AI.
- *Depreciation*: d_A determines how quickly the capital stock adjusts.

Under the simplifying assumptions that (i) the return differential is approximately proportional to $\omega - \omega^*$ with coefficient α_r , (ii) savings respond to college income with elasticity σ_c , and (iii) the patience gap drives the wealth accumulation rate, the composite elasticity simplifies to:

$$\eta_K(\eta_{\text{sav}} + \eta_W) = \frac{\sigma_c(\rho_S - \rho_H)}{d_A + \sigma_c \rho_H}. \quad (\text{C.10})$$

Substituting yields:

$$\mathcal{A}^{\text{dyn}} = \frac{1}{1 - \sigma_c(\rho_S - \rho_H)/(d_A + \sigma_c \rho_H)}. \quad (\text{C.11})$$

For small numerator (which holds empirically: $\sigma_c(\rho_S - \rho_H) = 0.015 \ll d_A = 0.10$), this is approximately $1 + \sigma_c(\rho_S - \rho_H)/(d_A + \sigma_c\rho_H)$.

At calibrated values ($\sigma_c = 0.5$, $\rho_S = 0.06$, $\rho_H = 0.03$, $d_A = 0.10$), the exact formula gives:

$$\mathcal{A}^{\text{dyn}} = \frac{1}{1 - 0.015/0.115} = \frac{1}{1 - 0.1304} = 1.15,$$

and the first-order expansion gives $1 + 0.015/0.115 = 1.13$, the value reported in equation (13).

The linearized system has a stable eigenvalue $\lambda_{\text{dyn}} = -d_A(1 - \eta_K(\eta_{\text{sav}} + \eta_W)) = -0.10(1 - 0.13) = -0.087$ (subscripted to distinguish it from the inequality-aversion parameter λ of Section 7). The half-life is $t_{1/2} = \ln 2/0.087 \approx 8.0$ years.

Along the transition:

$$\omega(t) = \omega^{*,\text{new}} + (\omega^{\text{impact}} - \omega^{*,\text{new}})e^{\lambda_{\text{dyn}} t}. \quad (\text{C.12})$$

C.9 Validity Conditions

The formula (C.11) is:

- (i) *Local*: valid for small perturbations Δr_A around the initial steady state.
- (ii) *Steady-state comparison*: $\mathcal{A}^{\text{dyn}}$ compares the new steady state to the impact response, not the transitional path.
- (iii) *Log-linear*: obtained by first-order Taylor expansion. Second-order terms are $O((\Delta r_A)^2)$.
- (iv) *Separable* in the patience gap because the hand-to-mouth assumption makes non-college savings identically zero. Without this assumption, a cross-term proportional to $(1 - m_S) \cdot (\partial W_S / \partial \omega)$ would enter (see Appendix F).

The approximation is local. For shocks several times the size of the baseline move in ω , which is 0.031 at the estimated parameters (the two-decimal endpoints displayed in Figure 3 imply 0.030), the log-linear characterisation should not be relied on.

D Aggregation from CZ-Level to National Parameters

D.1 The Aggregation Problem

This paper uses two reinstatement rates as production-function parameters in a national macro model: $\delta_R = 1.14$ from the CZ-level estimation in [Chávez \(2026\)](#), bracketed by the firm-level Danish evidence of [Humlum \(2022\)](#), and $\delta_A = 0.99$ from this paper’s own CZ-level moments. The aggregation question is whether these estimates, identified at different levels of analysis and in different institutional environments, can be treated as national production-function parameters.

The MDE in Section 3.3 estimates four structural parameters $(\Lambda_R, \delta_A, \varphi, \varepsilon_H)$ from six reduced-form moments, four wage moments and two college-employment-share moments. Each skill-group wage equation is a weighted cross-section over the 722 commuting zones,

$$\Delta \ln W_{k,c} = \beta_k^R \cdot \text{Robot}_c + \beta_k^A \cdot \text{AI}_c + X_c' \Gamma_k + \iota_{r(c)} + \varepsilon_{k,c},$$

where c indexes commuting zones, $k \in \{S, H\}$ indexes skill groups, X_c are the completed set of pre-determined composition shares (Section 3.1), and $\iota_{r(c)}$ are census-region fixed effects (a symbol distinct from the reinstatement rates δ_j). With one observation per zone per equation, the controls enter as $X_c' \Gamma_k$ rather than a commuting-zone fixed effect, which would be collinear. The skill-premium moment stacks the two equations and takes the contrast,

$$\Delta \ln(W_H/W_S)_c = (\beta_H^R - \beta_S^R) \text{Robot}_c + (\beta_H^A - \beta_S^A) \text{AI}_c + \tilde{\varepsilon}_c,$$

in which any commuting-zone-level shock common to both skill groups differences out. This is why the premium is estimated more precisely than either wage level. The robot Bartik identifies a composite industry demand parameter $\Lambda_R = 0.0085$, not a task-level reinstatement rate. The AI moments identify δ_A and φ directly through the Acemoglu–Restrepo task structure. The physical reinstatement rate δ_R is not identified by this paper’s six-moment differential-exposure system; it is imported from the distinct general-equilibrium moment system of [Chávez \(2026\)](#), estimated on the same underlying U.S. data.

D.2 Aggregation of δ_A : CZ-Level to National

The cognitive reinstatement rate $\hat{\delta}_A = 0.99$ is identified from cross-sectional CZ-level variation in AI exposure. The question is: when does the CZ-level estimate equal the national parameter δ_A^{nat} ?

The two objects coincide when three conditions hold:

- (i) *Homogeneous technology*: the reinstatement rate is a property of the automation technology, not of the local labor market. The same AI system generates the same number of complementary tasks regardless of the CZ in which it is deployed.
- (ii) *No spatial spillovers*: AI adoption in CZ c does not affect task creation in neighboring CZ c' .
- (iii) *Correct spatial equilibrium*: the CZ-level estimation accounts for migration and commuting responses, so that the estimated δ_A reflects task creation after netting out worker reallocation.

Condition (i) is plausible for cognitive automation: AI tools are digital, scalable, and deployed uniformly across locations. Unlike robots, whose complementary tasks may depend on local supply chains, AI-generated tasks (e.g., prompt engineering, algorithmic auditing, AI-assisted analysis) are inherently non-spatial.

Condition (ii) is less restrictive for AI than for robots because cognitive tasks are not geographically tied to the adopting firm's CZ. Supply chain spillovers, the primary concern for physical automation, are largely absent for cognitive automation.

Condition (iii) is addressed by the migration adjustment, which scales the non-college labor supply response by the calibrated elasticity $\varepsilon_S^{\text{mig}} = 0.78$, so the structural δ_A is identified after accounting for worker mobility.

D.3 External Validation of δ_R : The Humlum Bracket

The physical reinstatement rate $\delta_R = 1.14$ enters this paper as a calibrated input: the minimum-distance estimate of [Chávez \(2026\)](#) from U.S. commuting-zone robot wage moments, with bootstrap standard error 0.18. It is a task-level parameter, the mass of new labor tasks created per unit of physical-frontier advance, and so transports across models that share the task structure. This paper's final-good aggregation is more general than the specialization estimated there ($\rho = \sigma$), so δ_R is imported rather than re-estimated here. Re-identifying it under the more general aggregation would not sharpen that estimate and could conflate it with the outer elasticity. What this subsection

makes self-contained is the external validation, deriving from Humlum (2022)’s firm-level evidence the model-free range of δ_R that evidence implies.

Humlum (2022) links workers, firms, and robots in Danish administrative data and estimates firm-level responses to robot adoption with a matched-control event-study design. His paper does not report a reinstatement rate directly, so I translate his estimates into the units of δ_R , the value of new labor tasks created per unit of task value displaced at the physical frontier.

The inputs are Humlum’s four-year event-study responses, the moments his own structural estimation targets. Adopting firms expand sales by 20 percent, cut the production-worker wage bill by 20 percent, raise the tech-worker wage bill by 30 percent, and raise the total wage bill by 8 percent. In the year before adoption, production workers account for 39.9 percent of the adopters’ wage bill and tech workers for 16.0 percent.

In value terms within adopting firms, displaced production-task value is $0.20 \times 39.9 = 8.0$ percent of the firm wage bill. Created robot-complementary task value, the tech department that operates and maintains the robots, is $0.30 \times 16.0 = 4.8$ percent. The remaining occupations expand with the output scale-up, bringing the total wage bill to +8 percent.

Two bounds follow. Counting only robot-complementary tech tasks as reinstatement gives $\delta_R \approx 4.8/8.0 = 0.60$, a narrow lower bound that excludes new task content in other occupations. Attributing the entire net labor-demand expansion to task creation gives $\delta_R \approx (8.0+8.0)/8.0 = 2.00$, an upper bound that conflates reinstatement with scale. The implied range is $[0.60, 2.00]$, and the Chávez (2026) estimate $\delta_R = 1.14$ lies inside it. Humlum’s central qualitative finding is sharper than either bound: total labor demand rises as production tasks fall, which in the model is the reinstatement-dominant regime, the one in which task creation raises total labor demand despite declining production-task demand. His evidence supports the regime the estimate implies and brackets its magnitude. It does not supply a point estimate, and no standard error or confidence interval for δ_R should be attributed to it. The 0.18 reported above is the minimum-distance bootstrap standard error of Chávez (2026), not Humlum’s.

D.4 Bias Bounds for the Reinstatement Calibration

The CZ-level estimate is the source of δ_R : the structural estimation in Chávez (2026) recovers δ_R from CZ wage moments through the model’s general-equilibrium mapping, not from differential

local exposure. The firm-level Danish bracket $[0.60, 2.00]$ (Humlum, 2022) serves as the external check, and the bounds below ask how far either could be off before the headline moves.

The primary bias concern runs in opposite directions for the two estimates. For δ_A , spatial spillovers are minimal (cognitive tasks are non-spatial), so $\hat{\delta}_A^{CZ} \approx \delta_A^{\text{nat}}$. For δ_R , positive between-firm spillovers make the bracket’s lower bound conservative: if adopters’ task creation spills to suppliers (robot adoption at an auto manufacturer generating quality-control tasks there), the reinstatement realized outside the adopting firm is missed by the firm-level bracket, so the anchoring claim only strengthens. Under a simple spillover structure where a fraction χ of reinstatement occurs at other firms, the national parameter is $(1 + \chi)$ times the measured value. At $\chi = 0.15$ (15% of reinstatement occurs at supplier firms), $\delta_R^{\text{nat}} = 1.31$, which would increase π_R/π_A from 1.28 to 1.36 (direct solve at the baseline estimates). At $\chi = 0.30$, $\delta_R^{\text{nat}} = 1.48$ and $\pi_R/\pi_A = 1.44$. The qualitative conclusion, that robot automation is more productive per unit of frontier advance, is strengthened by positive spillovers.

Conversely, if the CZ estimate overstates the national parameter, a 30% downward adjustment gives $\delta_R = 0.80$, inside the firm-level bracket’s lower half. At this value, $\pi_R/\pi_A = 1.11$, still above one. All qualitative results are preserved: robot automation remains more productive, both frontiers lower the labor share, and the optimal policy favors robots. The robustness exercises cut in both directions: spillovers ($\chi > 0$), faster AI depreciation ($d_A > d_R$), and the consumption-based welfare conventions of Section 7.2 strengthen robot dominance, while a CZ estimate that overstates the national rate, a largely transitional sorting friction (Table 5), and a higher skill elasticity all weaken it. The ranking survives the adjustments that push against it, which is the relevant test.

E Service Sector Automation Extension

The baseline model assumes the non-tradable sector is not automated ($I_N = 0$). This appendix evaluates a parametric extension in which AI partially automates service tasks.

E.1 Setup

Replace $Y_N = A_N L_N^{\text{eff}}$ with:

$$Y_N = \left[Z_N(I_N^A) + \tilde{\Gamma}_N(I_N^A, \delta_N^A) \cdot (L_N^{\text{eff}})^\varsigma \right]^{1/\varsigma}, \quad (\text{E.1})$$

where $I_N^A \in [0, I_A]$ is the share of service tasks automated by AI (service automation can only use AI, not robots, by assumption) and $\delta_N^A \leq \delta_A$ is the reinstatement rate in services. This nests the baseline ($I_N^A = 0$) and a partially automated service sector ($I_N^A > 0$).

E.2 Effects on the Feedback Channels

Service automation affects two channels:

1. *Non-tradable multiplier*: When services are partially automated, service employment L_N declines for a given level of demand, weakening the labor reallocation mechanism. $\mathcal{F}^{NT}$ shrinks in proportion to the automated share I_N^A .
2. *Direct service displacement*: AI automation in services directly displaces non-college workers (who are 68% of service employment), adding service-sector exposure to the PE wage channel.

The two effects move $\mathcal{F}^{GE}$ in opposite directions, and both are second order. The reallocation channel is the computed cognitive-frontier $\mathcal{F}^{NT} = +0.00033$ of Table 6 before any of it is automated away, so that is the largest possible weakening. The direct displacement effect adds to the wage-feedback channel a term scaled by the automated service share, which at $I_N^A = 0.30$ is of the same order. Both are more than two orders of magnitude below one, so $|\mathcal{F}^{GE}|$ remains far inside the bound of Proposition 2, the contraction condition of Proposition 3 continues to hold, and the supply-side ranking is unaffected.

Supply-side dominance survives because both channels are absolutely negligible: each is more than two orders of magnitude below the level at which the feedback would matter. The qualitative result, that demand-side channels are second-order, is robust to AI automation of services.

F Bewley Extension

F.1 Relaxing the Hand-to-Mouth Assumption

The dynamic model of Section 5 assumes non-college workers consume all income ($m_S = 1$, $A_S = 0$), which is a knife-edge rather than the calibrated case. This appendix relaxes to $m_S \in (0, 1)$, allowing non-college workers to hold positive assets and smooth consumption.

With $m_S < 1$, the aggregate savings equation becomes:

$$S = (1 - m_S)W_S L_S + W_H L_H - C_H + \Pi,$$

and the log-linearized savings response gains a cross-term:

$$\hat{S} = \eta_W \hat{\omega} + \underbrace{(1 - m_S) \frac{W_S^* L_S^*}{S^*} \frac{\partial \ln W_S}{\partial \omega}}_{\text{non-college savings response}} \hat{\omega}.$$

Since $\partial \ln W_S / \partial \omega > 0$ (robot-dominated automation raises non-college wages), the non-college savings response is positive when ω rises and negative when ω falls. After an AI cost decline ($\hat{\omega} < 0$), non-college savings fall, *partially offsetting* the college savings channel that amplifies the shift toward AI.

F.2 Modified Amplification Factor

The composite elasticity becomes:

$$\eta_K(\eta_{\text{sav}} + \eta_W) = \frac{\sigma_c(\rho_S - \rho_H)}{d_A + \sigma_c \rho_H} - (1 - m_S) \cdot \frac{W_S^* L_S^*}{S^*} \cdot \frac{\partial \ln W_S}{\partial \omega} \cdot \frac{\eta_K}{d_A + \sigma_c \rho_H}.$$

The second term is negative (it reduces the amplification) and proportional to $(1 - m_S)$: the more non-college workers save, the stronger the dampening, and at the hand-to-mouth limit $m_S = 1$ it vanishes, recovering the baseline. Substituting the modified composite elasticity into the amplification formula, at $m_S = 0.85$ (non-college workers save 15% of income):

$$\mathcal{A}^{\text{dyn}} = \frac{1}{1 - \eta_K(\eta_{\text{sav}} + \eta_W)} \Big|_{m_S=0.85} \approx 1.09,$$

At $m_S = 0.75$ (25% savings rate, the highest plausible value for non-college workers):

$$\mathcal{A}^{\text{dyn}} \approx 1.06.$$

Table F.1: Dynamic Amplification Under Alternative Savings Behavior

m_S	$\mathcal{A}^{\text{dyn}}$	Interpretation
1.00 (hand-to-mouth)	1.13	Knife-edge
0.90	1.11	Low NC savings
0.85	1.09	Moderate NC savings
0.75	1.06	High NC savings
0.50	1.03	Near-equal savings rates

Notes: Each row reports the first-order expansion of the amplification factor under the stated m_S , holding all other parameters at baseline values. This is the object reported in equation (13), and it is the more conservative of the two: at the hand-to-mouth limit the exact log-linear value is 1.15 against the 1.13 shown here.

The dynamic amplification ranges from 1.03 to 1.13 across the empirically relevant range of m_S . The qualitative result, the savings composition channel amplifies the shift toward AI, survives at all values, but the magnitude is sensitive to the savings assumption. The knife-edge $\mathcal{A}^{\text{dyn}} = 1.13$ is the value at $m_S = 1$, and it is an upper bound with respect to the savings margin m_S , not with respect to the linearization.

Online Appendix

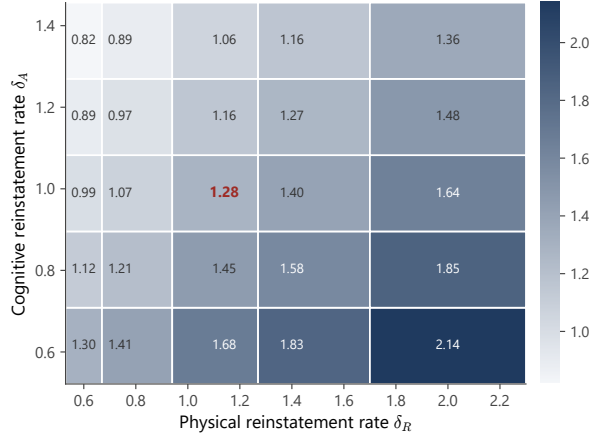
The Macroeconomics of Automation: Sorting Frictions and the Asymmetry Between Robots and AI

For online publication only

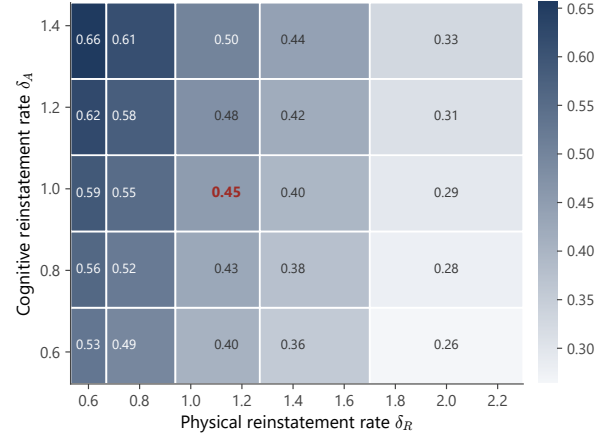
G Extended Robustness

This appendix evaluates the sensitivity of the productivity content ratio π_R/π_A , the feedback coefficient $\mathcal{F}$, and the labor share response to parameters not varied in the main text. All grids are computed at the baseline estimates $(\delta_R, \delta_A, \varphi) = (1.14, 0.99, 1.99)$ with the same solver that produces the main-text results (replication file `appendix_h_sensitivity.py`, whose name reflects an earlier appendix lettering). The baseline rows reproduce the headline decomposition exactly: $\pi_R/\pi_A = 1.28$, direct components $(+0.273, +0.148)$, and reinstatement-plus-reallocation components $(+0.349, +0.336)$.

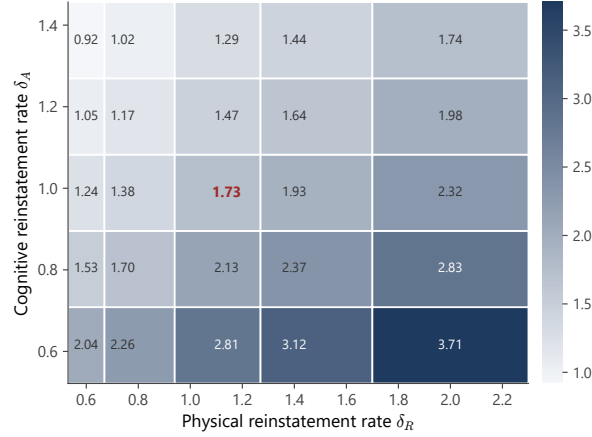
Figure G.1 provides a visual summary of the robustness analysis. All three panels grid δ_R over the Humlum bracket $[0.60, 2.00]$ and δ_A over ± 2 bootstrap SE, centered on the point estimates. Panel (a) shows $\pi_R/\pi_A > 1$ everywhere except the low- δ_R or high- δ_A edges. Panel (b) shows that robots are always less damaging to the labor share. Panel (c) shows that the optimal policy ratio exceeds one outside the same edges. Reversals occur only at the low- δ_R or high- δ_A edges of the displayed grid; away from those edges, the qualitative conclusions hold.



(a) Productivity ratio π_R/π_A



(b) Labor share damage ratio



(c) Optimal policy ratio τ_R^*/τ_A^*

Figure G.1: Sensitivity of main results, gridding δ_R over the Humlum bracket $[0.60, 2.00]$ and δ_A over ± 2 bootstrap SE, centered on the point estimates ($\delta_R = 1.14$, $\delta_A = 0.99$) with φ at its estimate. In panel (b), the labor-share damage ratio, values below one mean robots damage the labor share *less* than AI, so lower is more favorable to the robot frontier.

G.1 Sensitivity to Reinstatement Task Productivity ($\gamma_j B_j$)

The reinstatement productivity parameter $\gamma_j B_j$ governs the quality of newly created tasks. If new AI tasks are more productive than new robot tasks ($\gamma_A B_A > \gamma_R B_R$), this partially offsets the δ asymmetry. At calibrated values, $\gamma_R B_R \approx \gamma_A B_A$ (from the free-entry condition).

Table G.1: Sensitivity to Reinstatement Task Quality

$\gamma_A B_A / \gamma_R B_R$	π_R	π_A	π_R / π_A	$\partial s_L / \partial I_A$
0.50	+0.621	+0.330	1.88	-0.303
0.75	+0.621	+0.410	1.51	-0.283
1.00 (baseline)	+0.622	+0.485	1.28	-0.264
1.50	+0.623	+0.619	1.01	-0.231
2.00	+0.624	+0.736	0.85	-0.203

Notes: Each row recomputes all objects under the stated quality ratio at the baseline estimates $(\delta_R, \delta_A, \varphi) = (1.14, 0.99, 1.99)$. The quality ratio multiplies the productivity of reinstated cognitive tasks relative to reinstated physical tasks, holding surviving-task productivities fixed.

The ratio is monotone in the direction the mechanism implies: more productive AI tasks offset more of the δ asymmetry. Robot dominance survives up to a quality ratio of about 1.52, and at a quality ratio of 1.5 the productivity ratio still sits above parity (1.01). Beyond the crossing it falls below one, reaching 0.85 at a quality ratio of 2.00. Equalizing productivity content therefore requires AI's newly created tasks to be more than half again as productive as robots' newly created tasks, while the free-entry condition pins the calibrated ratio near one.

G.2 Sensitivity to Comparative Advantage (κ)

Higher κ means skill groups are more specialized, increasing the reallocation cost of displacement. AI displacement generates larger skill mismatch under stronger comparative advantage because non-college workers displaced from cognitive tasks have less transferable human capital. The table shows that the opposite force dominates: the ratio narrows in κ because π_A rises faster, the mechanism the paragraph after the table explains.

Table G.2: Sensitivity to Comparative Advantage

κ	π_R	π_A	π_R/π_A	π_R^{r+r}	π_A^{r+r}
1.0 (no CA)	+0.590	+0.361	1.64	+0.373	+0.355
1.2	+0.600	+0.401	1.50	+0.365	+0.349
1.4	+0.609	+0.434	1.40	+0.359	+0.344
1.6	+0.616	+0.461	1.34	+0.353	+0.340
1.8 (baseline)	+0.622	+0.485	1.28	+0.349	+0.336
2.0	+0.627	+0.505	1.24	+0.345	+0.333
2.5	+0.638	+0.546	1.17	+0.337	+0.328
3.0	+0.646	+0.578	1.12	+0.331	+0.323

Notes: Each row recomputes all objects at the baseline estimates $(\delta_R, \delta_A, \varphi) = (1.14, 0.99, 1.99)$. “r+r” denotes the reinstatement-plus-reallocation component. κ governs skill group specialization: higher κ means stronger comparative advantage.

The productivity ratio π_R/π_A is *decreasing* in κ : stronger comparative advantage narrows the gap between frontiers. This is because higher κ raises both π_R and π_A (more specialized workers are more productive in their comparative advantage sector), but raises π_A proportionally more. In all cases $\pi_R/\pi_A > 1$.

G.3 Sensitivity to Top-Level Elasticity (ρ)

Under general CES aggregation ($\rho \neq 1$), a price channel amplifies or dampens $\mathcal{F}$. The realized channel is small and rises with ρ , from -0.0001 at $\rho = 0.50$ to $+0.0001$ at $\rho = 1.50$, so it dampens the feedback below Cobb–Douglas and amplifies it above. The productivity ratio moves with it, from 1.24 to 1.32 across the same range.

Table G.3: Sensitivity to Top-Level Substitution Elasticity

ρ	π_R/π_A	$\mathcal{F}^{PE}$	$\mathcal{F}^{\text{price}}(\rho)$	Uniqueness
0.50	1.24	0.0020	−0.0001	Yes
0.60	1.25	0.0020	−0.0001	Yes
0.80 (calibrated)	1.27	0.0020	−0.0000	Yes
1.00 (Cobb–Douglas)	1.28	0.0021	0.0000	Yes
1.20	1.30	0.0021	+0.0000	Yes
1.50	1.32	0.0022	+0.0001	Yes

Notes: Computed at the baseline estimates $(\delta_R, \delta_A, \varphi) = (1.14, 0.99, 1.99)$ with the top-level aggregator set to the stated ρ . $\mathcal{F}^{PE}$ is the robot-frontier wage-feedback channel computed directly from the equilibrium wage responses, the larger of the two frontiers and the headline value of Section 4.1. $\mathcal{F}^{\text{price}}(\rho)$ is its increment relative to Cobb–Douglas. Cobb–Douglas is the aggregation of the baseline model, and $\rho = 0.80$ is the calibrated top-level elasticity.

The headline ratio moves by at most 0.04 in either direction across the range, the wage-feedback channel is flat, and uniqueness holds throughout.

G.4 Sensitivity to R&D Curvature (ξ_j)

The R&D curvature parameters govern how responsive the automation mix is to wage changes. Two properties pin down their role without a numerical grid. First, the productivity ratio π_R/π_A does not depend on (ξ_R, ξ_A) : the ratio is computed at fixed frontier positions, and the curvature parameters enter only the R&D response that moves those positions. Second, the feedback scales with the responsiveness of the automation mix to wages, which the R&D cost curvature governs: steeper diminishing returns (lower ξ) weaken the feedback and flatter returns strengthen it, while the realized wage-feedback at the calibrated curvature sits far below unity, so no admissible curvature configuration threatens uniqueness.

G.5 Joint Sensitivity to δ_R and φ

Because δ_R and φ enter different channels of the productivity decomposition, δ_R operates through the robot frontier while φ operates through the cognitive sector, their effects on the productivity

ratio are close to additive, with a modest positive interaction. Table G.4 illustrates this with a 3×3 grid computed at the baseline estimate $\delta_A = 0.99$.

Table G.4: Joint sensitivity to δ_R and φ (fixing $\delta_A = 0.99$)

		φ		
		0.50	1.99	3.00
		(75% trans.)	(baseline)	(high)
	0.60 (bracket lower)	0.74	0.99	1.21
δ_R	1.14 (Chávez 2026)	0.96	1.28	1.56
	2.00 (bracket upper)	1.23	1.64	1.99

Notes: Each cell reports π_R/π_A at the estimated $\delta_A = 0.99$. Row anchors are the firm-level bracket endpoints and the CZ estimate of Chávez (2026). The 75%-transitional column sets $\varphi = 0.25 \times 1.99$.

The effect of δ_R is amplified at higher φ (moving from $\delta_R = 0.60$ to 2.00 adds +0.49 when $\varphi = 0.50$ but +0.78 when $\varphi = 3.00$), indicating positive complementarity: sorting disruption and robot reinstatement reinforce each other in generating the asymmetry. Three cells sit at or below parity: the bracket’s lower bound at 75%-transitional sorting (0.74) and at baseline sorting (0.99), and the Chávez (2026) estimate at 75%-transitional sorting (0.96). The last cell restates the transitional-share result of Table 5, which reports the same 0.96 point ratio at a 75 percent transitional share. Away from these combinations, robot dominance holds.

G.6 Additional Parameter Sensitivity

Four parameters enter the productivity decomposition through channels that are either weakly identified or normalized in the baseline specification. Table G.5 reports sensitivity to each.

Table G.5: Sensitivity to ancillary parameters

Parameter	Range	Channel	π_R/π_A range	τ_R^*/τ_A^* range	Reversal?
<i>Panel A: Parameters affecting robot channel</i>					
ψ (demand symmetry)	[0.70, 1.30]	Effective δ_R	[1.19, 1.36]	[1.58, 1.86]	No
<i>Panel B: Parameters affecting AI channel</i>					
ε_H (college LS elast.)	[0, 6.26]	Implied (δ_A, φ)	[1.28, 1.40]	[1.71, 1.92]	No
d_A (AI depreciation)	[0.05, 0.50]	Cost adjustment	invariant	cost-scaled	Yes, cost-adjusted [‡]
<i>Panel C: Production elasticity</i>					
ε_w (skill elast. subst.)	[1.40, 2.50]	Re-fit (δ_A, φ)	[1.06, 1.38]	[1.32, 1.93]	No

Notes: ψ governs the asymmetry of robot displacement across skill groups, restricted to $\psi = 1$ on the construction grounds of Section 3.2. Freeing ψ changes effective δ_R but leaves (δ_A, φ) unaffected. ε_H is the college local labor-supply (incidence) elasticity, distinct from the participation elasticity ε_H^{LS} (point estimate 1.38, bootstrap SE 2.91, 90% percentile interval $[-0.60, 6.26]$ with the negative reach truncated at zero). Its row re-minimizes the minimum-distance objective with ε_H fixed at each of eight grid points spanning the interval, letting $(\Lambda_R, \delta_A, \varphi)$ adjust. The range minima sit near the point estimate. d_A is the AI capital depreciation rate. π_R/π_A is invariant to d_A because it measures output elasticity with respect to frontier *position*, but the cost of maintaining that position scales with d_j . [‡]For this row, the reversal entry refers to the cost-adjusted policy ranking, not to π_R/π_A or the uncorrected subsidy ratio: the cost-adjusted policy ranking *does* reverse inside the displayed range. The cost-adjusted ratio $(\pi_R/\pi_A)(d_A/d_R)$ falls below one for $d_A/d_R < 0.78$, and if $d_A > d_R$ (AI depreciates faster than physical robots) it *increases*, strengthening the robot case. ε_w is the elasticity of substitution between college and non-college labor, calibrated to 1.50 from commuting-zone wage differentials. Because it is calibrated rather than estimated, its row re-fits the minimum-distance parameters $(\Lambda_R, \delta_A, \varphi, \varepsilon_H)$ to the same six moments at each ε_w : a higher ε_w dampens the relative-wage response to a given skill-biased shock, so the estimation attributes more of the observed premium widening to sorting, and $\hat{\delta}_A$ falls while $\hat{\varphi}$ stays high. The productivity ratio declines monotonically toward parity as ε_w rises but does not cross it: $\pi_R/\pi_A \in [1.06, 1.38]$ and $\tau_R^*/\tau_A^* \in [1.32, 1.93]$ across the [1.40, 2.50] range spanned by the skill-substitution literature (Katz and Murphy, 1992; Card, 2009). The fit is tightest at the calibrated value (Hansen $J = 0.01$ through $\varepsilon_w = 2.2$) and degrades at the top of the range ($J = 1.4$ at $\varepsilon_w = 2.5$, where $\hat{\delta}_A$ approaches zero and $\hat{\varphi}$ leaves the bootstrap 90% region), so the upper endpoint is a different estimation regime rather than a smooth continuation. Other parameters held at baseline except where re-fit.

ψ (robot demand symmetry). The restriction $\psi = 1$ (robots displace non-college and college

labor symmetrically) is maintained on the construction grounds of Section 3.2. To bound the effect, I scale the effective robot reinstatement rate as $\delta_R(\psi) \approx \delta_R \cdot \psi^{1/2}$, reflecting that asymmetric displacement requires different reinstatement to match the same CZ wage moments. Over $\psi \in [0.70, 1.30]$, the productivity ratio ranges from 1.19 to 1.36, modestly affecting the *level* of robot dominance but never approaching parity.

ε_H (*college local labor-supply elasticity*). The point estimate $\hat{\varepsilon}_H = 1.38$ is weakly identified, with a 90 percent bootstrap percentile interval of $[-0.60, 6.26]$. The interval’s negative reach reflects the weak identification, not an economically meaningful labor-supply elasticity below zero. To bound the influence of this parameter, I re-minimize the minimum-distance objective with ε_H fixed at the interval’s economically meaningful endpoints, letting $(\Lambda_R, \delta_A, \varphi)$ adjust to fit the moments. At the 95th percentile ($\varepsilon_H = 6.26$) the implied estimates shift to $\delta_A = 0.77$ and $\varphi = 1.71$, and the productivity ratio rises to 1.40. At the economic lower bound ($\varepsilon_H = 0$) the implied $\delta_A = 0.93$ and the ratio is 1.32. Profiling over an eight-point interior grid confirms the minimum of the profiled ratio is 1.28, attained near the point estimate, so the ratio ranges over $[1.28, 1.40]$ across the interval. The ratio exceeds its point value at both endpoints, so the weak identification of ε_H does not push the model toward parity. The worst case for robot dominance therefore requires δ_R at the bracket’s lower bound and φ mostly transitional simultaneously, and neither margin is reinforced by the uncertainty in ε_H .

d_A (*AI depreciation*). The baseline assumes symmetric depreciation ($d_A = d_R = 0.10$). The productivity ratio π_R/π_A is invariant to d_A because it measures the output elasticity with respect to frontier *position*, not investment flow. However, the annual cost of maintaining a given frontier level is $d_j \cdot K_j$, so the cost-adjusted subsidy ratio scales as $(\pi_R/\pi_A) \cdot (d_A/d_R)$. If AI capital depreciates faster than physical robots ($d_A > d_R$), plausible given rapid model obsolescence, the cost-adjusted ratio *increases*, strengthening the case for robot subsidies. Taking the frontier-LLM release cycle literally gives $d_A \approx 0.5$ or higher (a one-to-two-year effective life) against $d_R \approx 0.07$ – 0.10 for ten-to-fifteen-year robot lifespans, so d_A/d_R is approximately five- to seven-fold at $d_A = 0.50$, can approach fifteen-fold under a one-year AI lifespan, and multiplies the cost-adjusted ratio by the same factor relative to the symmetric baseline. The headline results and Proposition 7 hold $d_A = d_R = 0.10$ fixed, as does the dynamic block, so this cost adjustment is a further tilt in robots’ favor that the reported subsidy ratios do not price in. It does not widen their stated range. The table’s range is

capped at $d_A = 0.50$ because beyond that point the conclusion is monotone in d_A/d_R and no new information arrives. Only if AI infrastructure is substantially more durable than physical robots ($d_A/d_R < 0.78$) would cost adjustment bring the ratio below unity.

ε_w (*skill elasticity of substitution*). The baseline sets $\varepsilon_w = 1.50$ from commuting-zone wage differentials. This parameter cannot be perturbed while holding the structural estimates fixed, because it governs how the wage moments map into skill-specific task content: each value of ε_w requires re-fitting (δ_A, φ) to the same six moments. Across $\varepsilon_w \in [1.40, 2.50]$ the productivity ratio stays within $[1.06, 1.38]$ and the subsidy ratio within $[1.32, 1.93]$, so the ranking never reverses. The magnitude does compress at the top of the range. The skill elasticity therefore bounds how far the headline can be pushed quantitatively, not whether it holds.

G.7 Identification of δ_A and φ

The restricted model estimates four parameters $\theta = (\Lambda_R, \delta_A, \varphi, \varepsilon_H)$ from six moments, three from robot exposure and three from AI exposure, with two overidentifying restrictions. The robot moments identify Λ_R ; ε_H is a cross-block parameter informed by both the robot and AI wage and share moments. The AI moments identify δ_A (AI reinstatement) and φ (cognitive sorting disruption). Because both δ_A and φ enter through the same equilibrium solver, they could in principle move the AI moments in correlated directions, leaving only a composite well-identified. This section shows that (i) the parameters are moderately but not sharply separated by the data, and (ii) the headline result $\pi_R/\pi_A > 1$ is robust to this identification uncertainty.

Table G.6 reports the numerical Jacobian $J_{ij} = \partial m_i(\hat{\theta})/\partial \theta_j$ evaluated at the MDE estimates by centered finite differences. The frontier-specific zero restrictions are exact: the robot moments (rows 1–3) have zero sensitivity to δ_A and φ , and the AI moments (rows 4–6) have zero sensitivity to Λ_R . The matrix is not block-diagonal, because ε_H is shared. The only cross-block channel is ε_H , which enters the AI moments through the demand elasticity $\eta_H = \varepsilon_D/(\varepsilon_D + \varepsilon_H)$. Identification of (δ_A, φ) therefore rests entirely on the three AI wage and skill-share moments.

Table G.6: Jacobian of Moment Conditions at MDE Estimates

	Λ_R	δ_A	φ	ε_H
Robot $\rightarrow W_{NC}$	-78.83	0	0	0
Robot $\rightarrow W_C$	-68.11	0	0	+0.13
Robot $\rightarrow \Delta s_C$	-7.82	0	0	-0.10
AI $\rightarrow W_{NC}$	0	+2.48	-1.58	0
AI $\rightarrow W_C$	0	+2.53	-0.34	-0.20
AI $\rightarrow \Delta s_C$	0	+0.36	+0.19	+0.15

Notes: Each entry is $\partial m_i / \partial \theta_j$ evaluated at the baseline estimates $\hat{\theta} = (0.0085, 0.989, 1.986, 1.379)$ via centered finite differences with step $h = 10^{-4}|\hat{\theta}_j|$. Units are percentage points of moment value per unit change in parameter. The zero entries are exact (structural, not numerical).

The columns for δ_A and φ in the AI block reveal the identification logic. Reinstatement (δ_A) creates new cognitive tasks that raise the marginal product of both skill groups, producing positive entries for all three AI moments. Sorting disruption (φ) degrades non-college workers' effective productivity in the cognitive sector: the skill weight a_S^C in the labor aggregator (8) becomes $a_S^C(1 - I_A)^\varphi$, which falls steeply in φ while the college weight a_H^C is largely unaffected. The φ column is therefore strongly negative on AI $\rightarrow W_{NC}$ (-1.58) but only weakly negative on AI $\rightarrow W_C$ (-0.34).

The two Jacobian columns are nearly anti-parallel (an angle of about 144°), deviating from exact collinearity by 35.7° , with a two-parameter condition number of 4.3, evaluated at the point estimate. This indicates moderate separation: the parameters move the AI moments in meaningfully distinct directions, but are not orthogonal, and the near anti-parallelism is what generates the positive bootstrap correlation below. The primary identifying moment is AI $\rightarrow W_{NC}$, where δ_A and φ pull in opposite directions (+2.48 versus -1.58). AI $\rightarrow W_C$ provides secondary identification, and AI $\rightarrow \Delta s_C$ contributes mainly as an overidentifying restriction.

The contrast between the two AI wage moments is what separates the two parameters, and it is worth making the point explicitly, because it is the whole identification of the mechanism. Reinstatement is a *level* effect: it raises the marginal product of both skill groups by almost the same amount, entering the two wage moments at +2.48 and +2.53. It therefore very nearly cancels

in their difference, which carries a Jacobian of only $+0.05$ in δ_A . Sorting disruption is a *differential* effect: it enters at -1.58 and -0.34 , so the difference retains $+1.24$ of it. The skill-premium response is thus an almost pure reading of φ , with a ratio of φ to δ_A sensitivity at least roughly forty times that of either wage level taken on its own. It is not a seventh moment, being a linear combination of two the estimation already uses, but it is where the identifying content for the sorting friction sits. This is why the premium is the sharpest moment in the paper, and it is also the source of the ridge in Table G.7: the two wage *levels* are nearly collinear in δ_A , and only their contrast pulls the parameters apart.

The same logic says what an additional moment would have to look like to sharpen φ . It would have to difference out a level while preserving a differential. The within-cell dispersion of residual wages is the natural candidate: task creation raises everyone’s wage inside a skill-occupation cell and so leaves its variance unchanged, while sorting disruption degrades match quality and therefore raises it. A displacement moment would be sharper still, since reinstatement absorbs displaced workers while sorting disruption strands them, so the two parameters would enter with opposite signs rather than merely at different magnitudes. Neither is available in the repeated cross-section used here, and both are natural targets for the post-deployment panel that Section 8 calls for.

To characterize the joint sampling uncertainty directly, I examine the joint distribution of $(\hat{\delta}_A, \hat{\varphi})$ across bootstrap replications. Table G.7 reports results for four specifications: the baseline restricted model and three robustness variants.

Table G.7: Bootstrap Identification Diagnostics for (δ_A, φ)

	n	$\bar{\delta}_A$	$\bar{\varphi}$	$\text{Corr}(\delta_A, \varphi)$	a/b	slope
Main (two-way VCV)	200	1.016	1.983	+0.71	2.86	2.0
<i>Alternative-specification diagnostics (correlation structure only)</i>						
Baseline	50	0.815	1.477	+0.66	5.9	—
Spec B	50	0.799	1.377	+0.76	6.4	—
Spec D	50	0.747	1.299	+0.74	5.9	—

Notes: Corr is the bootstrap correlation between $\hat{\delta}_A$ and $\hat{\varphi}$. $a/b = \sqrt{\lambda_{\max}/\lambda_{\min}}$ is the principal-axis ratio of the confidence ellipse; the condition number of the 2×2 bootstrap covariance matrix is $\lambda_{\max}/\lambda_{\min}$. “Slope” is the principal eigenvector slope $\Delta\varphi/\Delta\delta_A$. The main row is the two-way-VCV bootstrap that produces every set statement in the paper. The remaining rows are smaller diagnostic bootstraps under alternative specifications, shown only to establish that the positive (δ_A, φ) correlation is stable, and their means and scales are not comparable to the main row. Spec B is a legacy diagnostic variant with a fuzzy frontier boundary (a task-overlap zone between the two frontiers). Spec D re-estimates with sorting disruption entering only the non-college skill weight, φ_{NC} , holding the college cognitive weight at its base value.

The positive correlation $\text{Corr}(\hat{\delta}_A, \hat{\varphi}) = 0.71$ across the bootstrap draws shows that the data identify a composite of δ_A and φ : bootstrap draws with higher reinstatement also require more sorting disruption to match the observed AI wage moments. The eigenvalue ratio of 8.2 indicates that most of the bootstrap variance lies along a single direction in (δ_A, φ) space, with principal-axis slope $\Delta\varphi/\Delta\delta_A \approx 2.0$: increasing δ_A by 0.1 requires φ to increase by approximately 0.2 to maintain moment fit.

Figure G.2 plots the main two-way-VCV bootstrap draws and their 95% confidence ellipse.

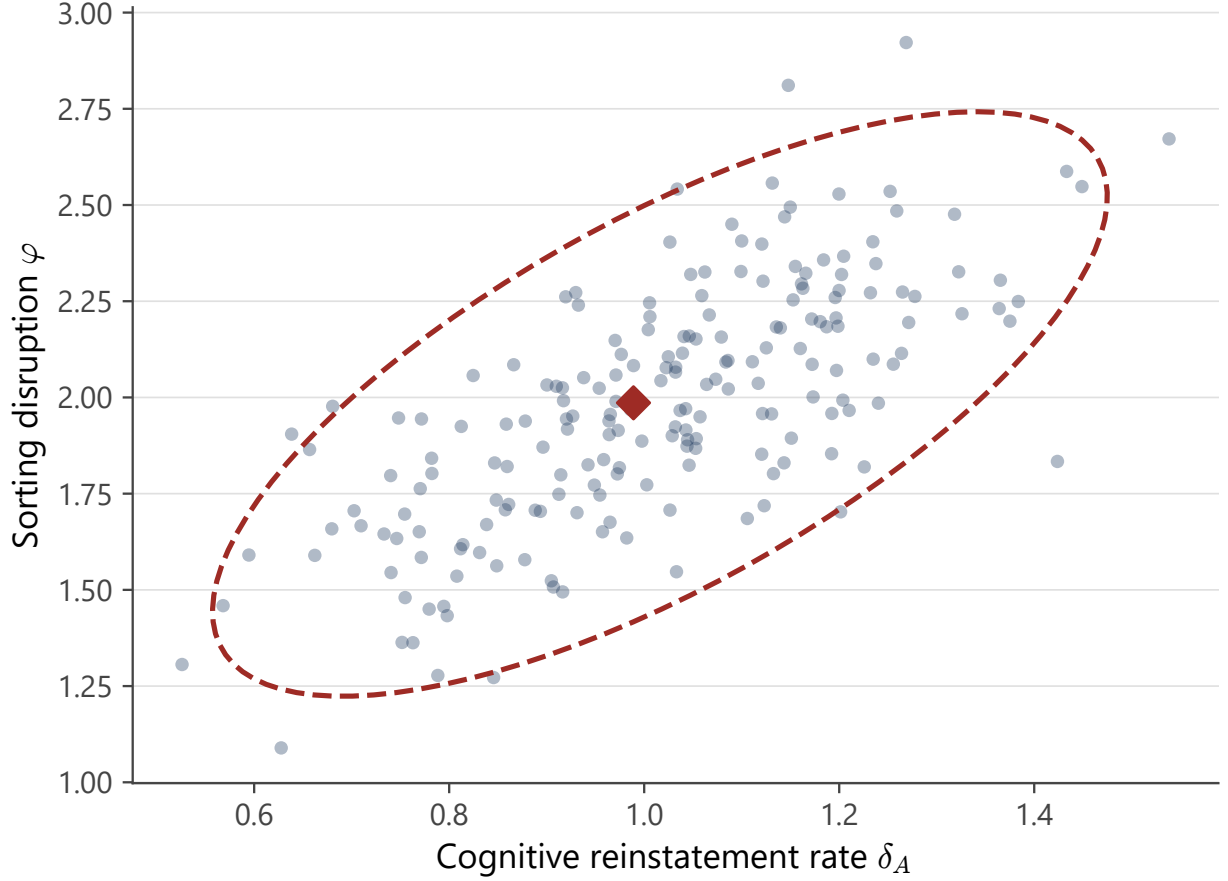


Figure G.2: Bootstrap draws of the cognitive reinstatement rate and sorting disruption parameter.

Notes: Each point is a bootstrap MDE estimate under the two-way VCV. The dashed ellipse is the 95% confidence region from the bootstrap covariance matrix. The diamond marks the point estimate $(\hat{\delta}_A, \hat{\varphi}) = (0.989, 1.986)$. The positive correlation reflects the compensating relationship between reinstatement and sorting disruption in matching the AI wage moments.

The critical question is whether the headline result $\pi_R/\pi_A > 1$ is robust to identification uncertainty. The bootstrap draws constitute the empirical uncertainty set, an approximation to the sampling distribution of the estimates: if $\pi_R/\pi_A > 1$ at every draw, the policy conclusion does not depend on the precise decomposition of δ_A and φ .

Table G.8 reports π_R/π_A and the optimal subsidy ratio evaluated at every bootstrap draw, holding $\delta_R = 1.14$ at the [Chávez \(2026\)](#) estimate.

Table G.8: Macro Stability: π_R/π_A Over Bootstrap Draws

	n	Min	p_5	Median	p_{95}	Max	$\Pr(> 1)$
π_R/π_A	200	1.01	1.12	1.27	1.46	1.61	100%
τ_R^*/τ_A^*	200	1.18	1.41	1.68	2.14	2.56	100%

Notes: Each row evaluates the indicated ratio at all 200 bootstrap draws of $(\hat{\delta}_A, \hat{\varphi})$ under the two-way VCV, holding $\delta_R = 1.14$ at the [Chávez \(2026\)](#) estimate. The last column reports the fraction of draws above one. Every draw exceeds parity, with a minimum ratio of 1.01 (Figure [G.3](#)). τ_R^*/τ_A^* is the equity-adjusted Pigouvian ratio with inequality aversion $\lambda = 1.5$.

π_R/π_A exceeds unity in all 200 bootstrap draws. The minimum value is 1.007 and the lower 5th percentile is 1.121. The policy-relevant ratio τ_R^*/τ_A^* exceeds unity at every draw, with a minimum of 1.182.

These draws hold δ_R at its point estimate, because it is identified on a separate moment block and enters the ratio only after the sorting parameters are recovered. That understates the sampling uncertainty, since δ_R carries a bootstrap standard error of 0.18 of its own. Drawing it from $N(1.14, 0.18^2)$ alongside each (δ_A, φ) replication, robot dominance survives in 98.5 percent of the joint draws, with a fifth percentile of 1.06. The ranking is therefore not an artefact of treating the imported rate as known: it survives that rate’s own sampling error, though not in literally every draw, and the honest statement of the result is the one given here rather than a claim of universality.

The robustness of this result follows from the geometry of the identification problem, and the geometry is sharper than an orthogonality argument. The $\pi_R = \pi_A$ parity contour in (δ_A, φ) space rises with slope $\Delta\varphi/\Delta\delta_A \approx 2.5$ (Figure [5](#)), and the identified ridge rises with slope ≈ 2.0 : the two are nearly parallel, separated by about five degrees. The composite the data cannot pin down therefore moves draws along the parity boundary, not across it. Draws with high δ_A carry high φ , and the higher sorting disruption raises the reinstatement rate AI needs for parity by almost exactly the amount the draw adds, which is why all 200 draws remain robot-dominant despite the partial identification. The parallel configuration protects against sampling uncertainty, not against bias: a systematic error perpendicular to the ridge, from a misspecified moment for example, would move every draw across parity together. That perpendicular direction is disciplined not by the bootstrap

but by the micro evidence anchors of Section 6.2, which place (δ_A, φ) on the dominant side from data the estimation never used.

Figure G.3 visualizes this result.

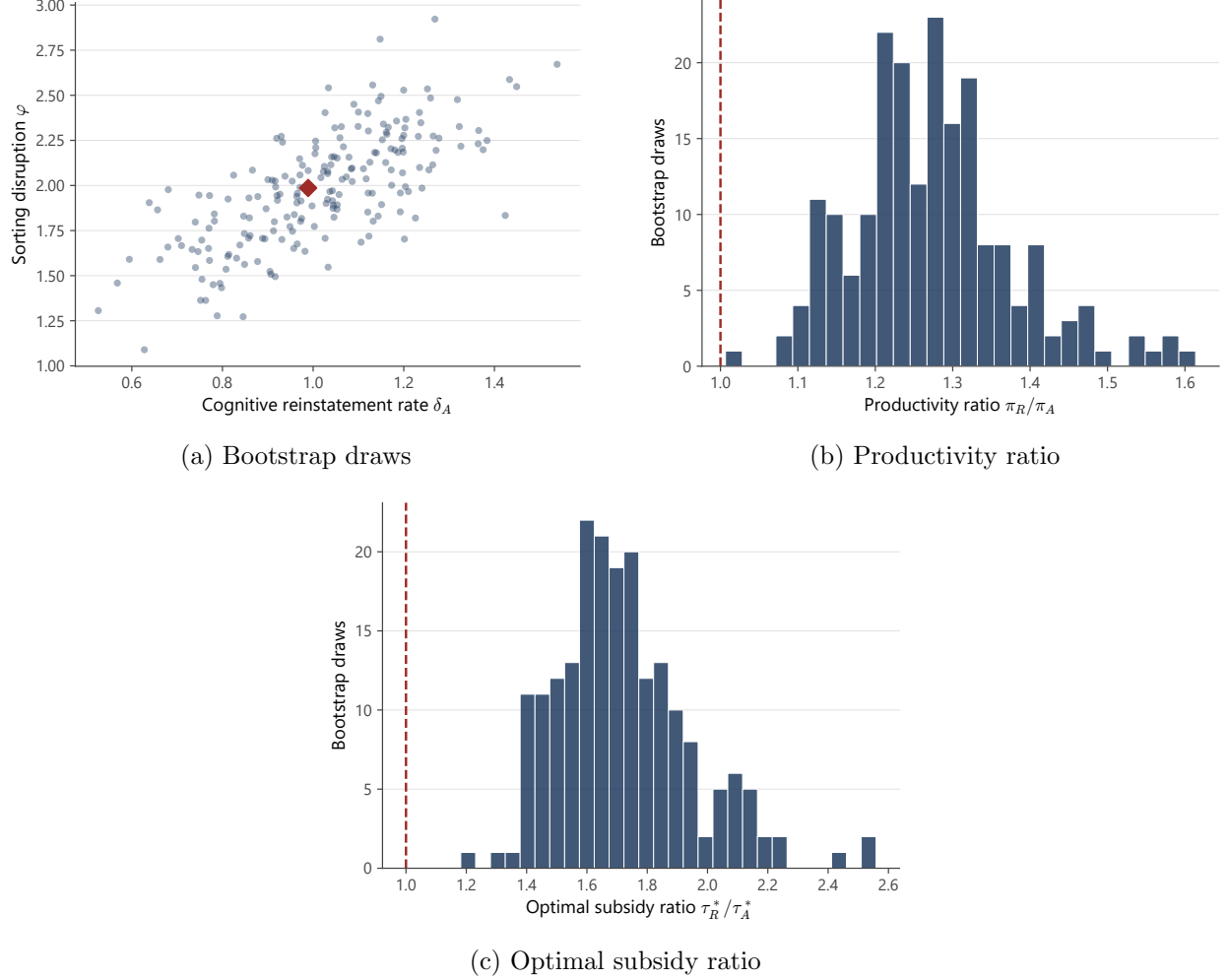


Figure G.3: Macro stability over the bootstrap draws.

Notes: Panel (a) plots each bootstrap draw, with the diamond at the point estimate. Panel (b) shows the distribution of π_R/π_A across draws. All 200 values exceed parity (dashed line). Panel (c) shows the distribution of the equity-adjusted optimal subsidy ratio τ_R^*/τ_A^* .

Although δ_A and φ are not individually sharply identified, reflecting an economic complementarity between reinstatement and sorting disruption in matching the AI wage moments, the data identify a sufficient statistic for the policy comparison. Across the bootstrap draws, π_R/π_A spans $[1.12, 1.46]$ and τ_R^*/τ_A^* spans $[1.41, 2.14]$ at the 5th to 95th percentiles, with all 200 draws above productivity parity. The conclusion that robot automation has higher marginal productivity gains,

and warrants relatively stronger policy support, does not depend on the precise decomposition of AI’s labor market effects into reinstatement and sorting channels.

G.8 Robustness to the Joint Estimates of Chávez (2026)

The baseline macro results use the estimates from Table 2 ($\delta_R = 1.14$, $\delta_A = 0.99$, $\varphi = 1.99$), where δ_A and φ come from this paper’s own minimum-distance estimation. Chávez (2026) jointly estimates all reinstatement and sorting parameters from wage moments in a different model and arrives at a much lower cognitive reinstatement rate: $\hat{\delta}_A = 0.03$, statistically indistinguishable from zero, with a bootstrap-percentile 95 percent interval of $[0.00, 0.35]$, asymmetric because it is bounded below at zero by the estimator’s constraint (draws pile at the bound), and $\hat{\varphi} = 1.11$.¹¹ Table G.9 summarizes both sets of estimates.

Table G.9: Parameter Estimates: Baseline vs. the Joint Estimation of Chávez (2026)

Parameter	Baseline		Chávez (2026) joint	
	Estimate	SE	Estimate	SE
δ_R (robot reinstatement)	1.14 [†]	(0.18)	1.14	(0.18)
δ_A (AI reinstatement)	0.99	(0.19)	0.03	(0.11)
φ (sorting disruption)	1.99	(0.31)	1.11	(0.82)

Notes: [†]The baseline specification takes δ_R from the Chávez (2026) estimate and estimates (δ_A , φ , ε_H) from six CZ-level moments in this paper’s macro model. The Chávez (2026) joint column reports that paper’s five-parameter minimum-distance estimates from wage moments, with bootstrap standard errors from 500 replications. The 0.35 used in the sensitivity tables is the bootstrap 95th percentile; the interval is bounded below at zero by the estimator’s constraint (draws pile at the bound), not a normal approximation. δ_R is point-identified in that paper, with bootstrap SE 0.18 and 95% interval $[0.84, 1.53]$ from the Robot moment block. The Humlum bracket $[0.60, 2.00]$ is a wider external-validation range, and 1.14 is the point estimate.

Two features stand out. First, the physical frontier is common to both by construction: the baseline imports $\hat{\delta}_R = 1.14$ from Chávez (2026), so the comparison is confined to the cognitive

¹¹The specification in Chávez (2026) differs in the moment set (wage levels and the skill premium in place of the college employment share), the sample period, and the exposure measures, so the two estimates of δ_A are not directly comparable. The exercise here asks what the macro conclusions look like if that paper’s cognitive-reinstatement finding is taken at face value.

frontier. Second, its wage system places AI-driven task reinstatement near zero ($\hat{\delta}_A \approx 0$), well below this paper’s joint wage-and-employment-share estimate of 0.99.

Which estimate to believe is an open question, and it is narrower than the headline gap suggests. The two papers agree on the sorting story: both estimate a positive, economically large cognitive sorting friction, at 1.11 (imprecisely) and 1.99, and both place the physical frontier at the same reinstatement rate. The numerical disagreement is concentrated in one reported parameter, but it cannot be attributed to one moment or exposure measure. The papers differ jointly in cognitive exposure, endpoints, controls, outcomes, covariance weighting, and structural mapping. Controlled bridges make this non-portability explicit. Substituting LMOE moments into this paper’s exact-control mapping gives $\hat{\delta}_A = 1.89$ (1.42 in the wage-only system), moving away from rather than toward 0.03. Substituting Webb moments into the Chávez mapping pushes parameters to their bounds and leaves a large objective ($J = 45.1$). The defensible interpretation is therefore design-specific: each δ_A belongs to its complete moment-and-model system, not to LMOE or Webb in isolation. External evidence does not yet arbitrate. At early-LLM-era frontier advances both estimates predict earnings effects inside the precision of the Danish adoption-study nulls: per $\Delta I_A = 0.01$ of frontier advance, this paper’s estimate predicts wage changes of -0.111 percent (non-college) and $+0.203$ percent (college), and the near-zero alternative -0.327 and -0.251 percent, all inside the ± 1 percent band this appendix adopts as a benchmark for those precisely estimated nulls. The predictions first leave the band at cumulative advances of $\Delta I_A = 0.049$ (this paper) and 0.031 (the alternative), each computed as the band divided by that configuration’s largest-magnitude wage derivative, so the null is consistent with this paper’s estimate without discriminating against the near-zero alternative. What does differ is the predicted sign pattern: at $\delta_A = 0.99$ an AI advance raises college wages and lowers non-college wages, widening the premium, while at $\delta_A \approx 0$ both wages fall. A post-deployment validation panel measuring LLM-era exposure against skill-group wage growth is the designed test. Until it runs, the macro results are reported at this paper’s estimate, with that paper’s configuration carried throughout this appendix as the alternative.

To assess robustness, I feed the [Chávez \(2026\)](#) estimates directly into the baseline macro model from Sections 6 and 7. Table G.10 and Figure G.4 compare the results at three calibrations: the baseline, that paper’s 95th percentile for δ_A , and its point estimate.

Table G.10: Macro Results: Baseline and the Joint Estimates of Chávez (2026)

	δ_R	δ_A	π_R	π_A	π_R/π_A	τ_R^*/τ_A^*
Baseline	1.14	0.99	+0.622	+0.485	1.28	1.73
Chávez (2026), 95th pctl	1.14	0.35	+0.622	+0.364	1.71	2.45
Chávez (2026), point est.	1.14	0.03	+0.621	+0.239	2.60	5.88

Notes: Each row feeds the indicated $(\delta_R, \delta_A, \varphi)$ into the baseline macro model from Sections 6 and 7. π_j is the productivity content of frontier j (output elasticity with respect to I_j). τ_R^*/τ_A^* is the optimal subsidy ratio from the Bergson–Samuelson SWF with $\lambda = 1.5$. Both Chávez (2026) rows use the joint estimates of Chávez (2026): δ_A at its point estimate and 95th percentile, each evaluated at that paper’s sorting estimate $\varphi = 1.11$. The baseline row uses this paper’s $\varphi = 1.99$. All hold $\sigma = 2$, $\mu = 0.45$, and the remaining parameters at baseline. Cells are evaluated at the displayed two-decimal δ_A ; at the companion’s unrounded point estimate 0.0265 the productivity cells read +0.238 and 2.61, one digit in the last place.

The bar decomposition makes the same point graphically: the productivity ranking is stable across the calibrations even as the level of each frontier’s contribution moves.

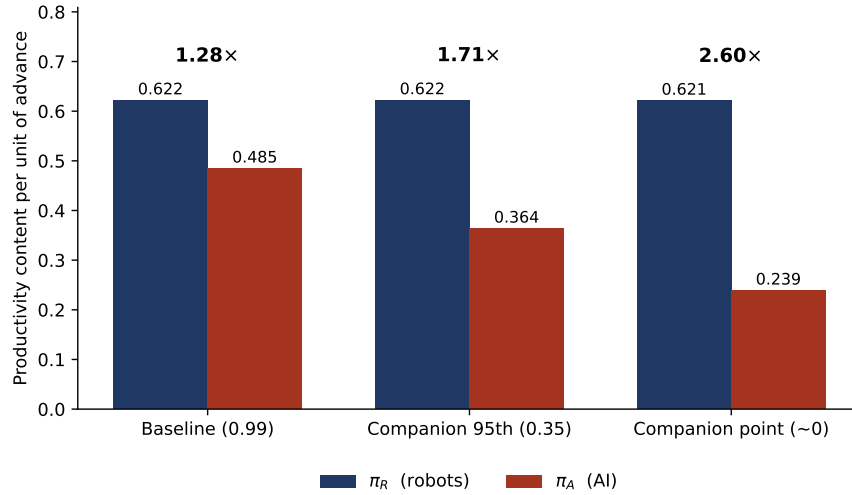


Figure G.4: Productivity content across calibrations. Moving from this paper’s estimate ($\delta_A = 0.99$, $\varphi = 1.99$) to the joint estimate of Chávez (2026) ($\delta_A \approx 0$, $\varphi = 1.11$), π_A declines while π_R holds, widening the asymmetry from $1.28\times$ to $2.60\times$. That paper’s lower φ partly offsets the δ_A -driven decline.

Across all three calibrations, $\pi_R > \pi_A > 0$ and $\tau_R^*/\tau_A^* > 1$. As Figure G.4 makes clear, the asymmetry runs entirely through the cognitive channel: δ_R is common to all three rows, so π_R is

nearly constant at 0.62, while the lower δ_A of [Chávez \(2026\)](#) reduces π_A from 0.485 to 0.239 by all but eliminating the reinstatement channel for AI, its lower φ partly offsetting the fall. At those point estimates, robots are over two and a half times more productive than AI and the optimal subsidy ratio rises to 5.88. π_A remains strictly positive even at $\delta_A \approx 0$: the direct capital-deepening effect ensures that advancing the AI frontier always raises output, even without reinstatement.

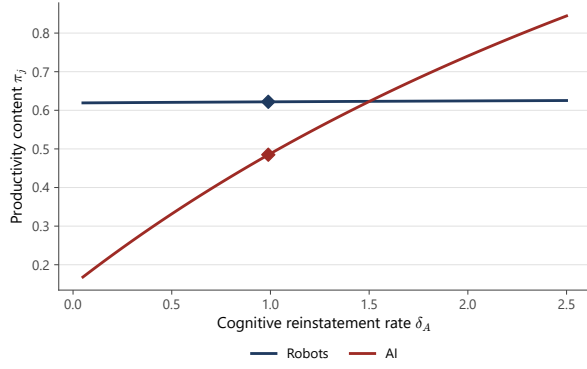
The key partially identified parameter is the AI reinstatement rate δ_A .

Table G.11: Sensitivity to δ_A ($\delta_R = 1.14$, $\varphi = 1.11$)

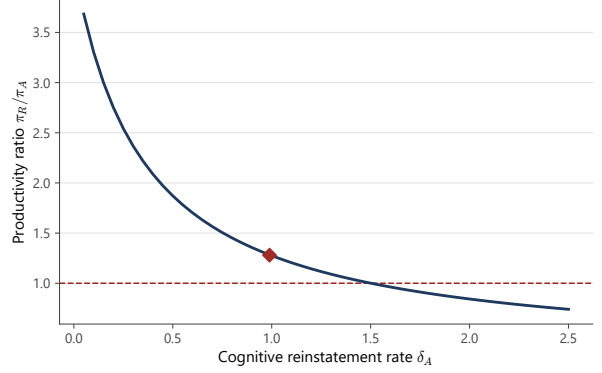
δ_A	π_R	π_A	π_R/π_A	τ_R^*/τ_A^*
0.00	+0.621	+0.227	2.74	6.84
0.35	+0.622	+0.364	1.71	2.45
0.60	+0.623	+0.452	1.38	1.73
0.99	+0.624	+0.577	1.08	1.23
1.20	+0.625	+0.639	0.98	1.07
1.40	+0.625	+0.694	0.90	0.97

Notes: δ_R and φ held at the joint point estimates of [Chávez \(2026\)](#), so the $\delta_A = 0.99$ row differs from the baseline in Table [G.10](#), which uses this paper’s own sorting estimate $\varphi = 1.99$. That paper’s specification places $\delta_A \in [0, 0.35]$ (point estimate to bootstrap 95th percentile), corresponding to $\pi_R/\pi_A \in [1.7, 2.7]$. The $\delta_A = 0.00$ row uses exactly zero. Its *point* estimate is $\hat{\delta}_A = 0.03$, which is why Table [G.10](#)’s “Chávez (2026), point est.” row reads $\pi_A = 0.239$ and $\pi_R/\pi_A = 2.60$ rather than the 0.227 and 2.74 here. Parity ($\pi_R = \pi_A$) occurs at $\delta_A \approx 1.15$, and above it the ranking reverses. See Figure [G.5](#) for the continuous profile.

To trace out how the results vary continuously with this parameter, Figure [G.5](#) plots π_j and the productivity ratio as functions of δ_A , holding $\delta_R = 1.14$ fixed. Table [G.11](#) reports the profile at the [Chávez \(2026\)](#) sorting value $\varphi = 1.11$, while the figure holds φ at this paper’s estimate, so the crossing points differ (1.15 against 1.50).



(a) Productivity content π_R and π_A . Diamonds mark the estimate ($\delta_A = 0.99$).



(b) Productivity ratio π_R/π_A relative to parity (dashed line).

Figure G.5: Sensitivity to the AI reinstatement rate δ_A , holding $\delta_R = 1.14$ and $\varphi = 1.99$ at the baseline estimates. This figure fixes φ at this paper's estimate 1.99, so its parity crossing ($\delta_A \approx 1.50$) differs from the $\delta_A \approx 1.15$ crossing of Table G.11, which fixes φ at the Chávez (2026) value 1.11. The figure and the table answer the same question at different sorting values. That paper's 95% interval $\delta_A \in [0, 0.35]$ lies in the far robot-dominant region, where both ratios are well above unity. In the figure the productivity ratio crosses one at $\delta_A \approx 1.50$, while the subsidy ranking survives past that crossing, so the planner's preference for robots holds throughout the data-consistent region and reverses only if AI reinstatement substantially exceeds robot reinstatement.

Panel (a) of Figure G.5 shows that π_R is nearly flat across the range, robot productivity depends on δ_R , not δ_A , while π_A rises steeply with the AI reinstatement rate. Panel (b) translates this into ratios at this paper's sorting value, crossing parity at $\delta_A \approx 1.50$. At the Chávez (2026) sorting value, Table G.11 shows the ratio falling from 2.74 at $\delta_A = 0$ to parity at $\delta_A \approx 1.15$ and below one beyond it. That paper's specification places $\delta_A \in [0, 0.35]$ (point estimate to 95th percentile), corresponding to $\pi_R/\pi_A \in [1.7, 2.7]$. At this paper's own estimate $\delta_A = 0.99$ with that paper's sorting value, the ratio is 1.08.

The qualitative policy conclusion, that robot automation deserves at least as favorable treatment as AI automation, holds throughout the data-consistent range of δ_A . The productivity ranking crosses parity at $\delta_A \approx 1.15$ under the Chávez (2026) sorting value and at ≈ 1.50 at this paper's estimate, but the subsidy ranking survives past those points and reverses only at a higher δ_A : at $\delta_A = 1.20$ the productivity ratio has already fallen to 0.98 while the subsidy ratio is still 1.07. The equity multiplier is what holds the tilt in place, which is why the policy conclusion is the more durable of the two. The quantitative magnitude depends primarily on δ_A : under that paper's sorting value, the ratio at this paper's $\delta_A = 0.99$ is 1.08, and at its point estimate ($\delta_A \approx 0$) robots

are over two and a half times more productive. The productivity ratio lies in the range $[1.08, 2.7]$ depending on which estimate of δ_A one adopts.

H Robustness of Reduced-Form Moments

Table H.1 reports three robustness checks on the reduced-form moments from Table 1.

Panels A and B are estimated on the completed control set of Section 3.1, so those comparisons are within-specification against Table 1. Panel C instead drops the farm-occupation share, and is the specification discussed in Section 3.1.

Panel A adds the Autor et al. (2013) China import shock as a third regressor. The robot and AI coefficients are virtually identical to the baseline specification: robot exposure on non-college wages remains -0.0079 ($t = -4.26$), and AI exposure on college wages remains $+0.0053$ (conventional $t = +2.10$, and, as in the baseline, not significant under exposure-robust inference). The China shock itself is individually insignificant on both wage outcomes ($\hat{\beta}_{\text{China}} = -0.0006$ on non-college wages, $p = 0.70$, and -0.0021 on college wages, $p = 0.35$). It enters marginally for college share ($+0.0017$, $p = 0.097$), but this does not affect the wage moments that drive identification of Λ_R and the cognitive parameters δ_A and φ . The R^2 values are nearly unchanged, showing that the estimated robot and AI gradients change little when the China import shock is added.

Panel B adds log population and baseline tech employment share to the horse-race specification. The robot coefficient on non-college wages is stable (-0.0075 , $p < 0.001$), and the AI coefficient on non-college wages strengthens slightly and remains significant (-0.0079 , $p = 0.003$), but the AI coefficient on college wages attenuates from $+0.0052$ to $+0.0022$ ($p = 0.44$). The attenuation is therefore specific to the college moment, where the tech-hub geography lives, and it is expected: log population and tech employment share are plausibly endogenous to AI exposure. Commuting zones with high AI-task intensity are disproportionately large metropolitan areas (New York, San Francisco, Boston) with high baseline tech employment. Including these variables absorbs precisely the cross-sectional variation that identifies the AI channel. Such regressors absorb the variation that identifies the channel of interest rather than removing confounds (Angrist and Pischke, 2009). The baseline specification in Table 1, which conditions on pre-determined demographic composition and region but not on endogenous urban characteristics, is therefore the preferred source of moments

for the MDE.

Table H.1: Robustness of Reduced-Form Moments

	$\Delta \ln W_{NC}$	$\Delta \ln W_C$	Δs_C
	(1)	(2)	(3)
<i>Panel A: Horse race with China import shock</i>			
Robot exposure (EU5)	−0.0079*** (0.0019)	−0.0056** (0.0028)	−0.0015 (0.0011)
AI exposure (Webb pctl)	−0.0069*** (0.0024)	0.0053** (0.0025)	0.0029** (0.0013)
China import shock	−0.0006 (0.0014)	−0.0021 (0.0023)	0.0017* (0.0010)
R^2	0.660	0.514	0.418
<i>Panel B: Augmented controls</i>			
Robot exposure (EU5)	−0.0075*** (0.0017)	−0.0045* (0.0026)	−0.0013 (0.0011)
AI exposure (Webb pctl)	−0.0079*** (0.0026)	0.0022 (0.0029)	0.0022 (0.0014)
R^2	0.677	0.521	0.420
<i>Panel C: Dropping the farm-occupation share</i>			
Robot exposure (EU5)	−0.0077*** (0.0019)	−0.0055** (0.0028)	−0.0016 (0.0011)
AI exposure (Webb pctl)	−0.0027 (0.0026)	0.0056*** (0.0021)	0.0017 (0.0012)
R^2	0.636	0.513	0.401
Region FEs	Yes	Yes	Yes
Observations (A / B / C)	722 / 718 / 722	722 / 718 / 722	722 / 718 / 722

Notes: Panels A and B use the completed control specification of Table 1, so those comparisons are within-specification. Panel A adds the [Autor et al. \(2013\)](#) China import shock. Robot and AI coefficients are virtually unchanged, and the China shock is individually insignificant on wage outcomes. Panel B further adds baseline log population and tech employment share. The AI coefficient on college wages attenuates from 0.0052 ($p = 0.039$) to 0.0022 ($p = 0.44$), reflecting controls plausibly endogenous to AI exposure: commuting zones with high AI-task intensity are disproportionately large metropolitan areas with high baseline tech employment. Panel B observations fall to 718 because the augmented controls have missing values. Panel C instead drops the baseline farm-occupation share. The AI college-wage coefficient is essentially unchanged at 0.0056 conventionally, but the non-college wage and college-share moments lose significance, which is why the farm control is retained. Regressions are weighted by baseline CZ employment, robust standard errors in parentheses. * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$.