

Nobody to Blame

*Attribution, the Hiring Margin, and the Coming Politics
of Artificial Intelligence*

Aaron Long

Draft 4.1 · Full edition · August 2026

*Working paper. Not peer-reviewed. A theory of political attribution with a set of dated,
falsifiable indicators — not a forecast.*

This draft abandons the thesis of the previous one. Section 3 says why.

Abstract

Economic displacement always produces a political response. What determines the *form* of that response is not how large the displacement is but how it can be **attributed**. Di Tella and Rodrik ran the experiment: an identical 900-worker plant closure, with only the stated cause randomised. Demand for import protection ran 0.09 in the control, 0.13 when the cause was a technology shock, 0.23 for outsourcing to France and 0.29 for outsourcing to Cambodia. Trade elicits protectionism by a factor of two to three. And the sign flips on the other margin: trade attribution *reduces* demand for compensation, while technology and bad-management attribution raise it.

The operative variable is therefore not attributability alone but **agency plus out-group membership**. A cause with an identifiable outside agent produces demands to block the outsider. A cause with an identifiable inside agent produces demands to make them share. A cause with no agent at all produces disaffection with no policy vehicle. **Artificial intelligence is the first major displacement in modern economic history with no foreigner attached to it.**

This corrects the paper's own previous draft, which held that no protective countermovement arrived after 1980. One did: 2016–2026 satisfies every Polanyian criterion. It aimed at trade and migration, and the paper's earlier premise that automation destroyed more manufacturing employment than trade does not survive scrutiny either — robots account for 360,000 to 670,000 US jobs over 1990–2007 against roughly 2.0 to 2.4 million for the China shock. Nor did automation escape political consequence: in the one study that horse-races both shocks in a single regression across fourteen countries, a standard-deviation robot shock raises the radical-right vote by 1.8 points while the trade shock is small and insignificant. The countermovement was not absent and was not asleep. It was aimed.

Two mechanisms then determine what will be visible. First, **adjustment runs through hiring rather than separations**. US hires have fallen a full point since January 2022 while layoffs sit 0.2 points above their all-time low; the best-identified firm-level study, covering 65 million résumés across 280,000 firms, finds junior employment falling in AI-adopting firms with senior employment unchanged, "driven primarily by slower hiring rather than increased separations". Second, **junior work is how senior competence is produced**. Training was historically free to the firm because it was a by-product of output the firm wanted anyway; AI supplies that output directly, so training acquires a positive marginal cost for the first time. Because training equilibria come in pairs, that is a tipping point rather than a slope — invisible for years, then discontinuous.

The paper offers five falsifiable predictions and eight dated indicators. Two tests are already running. The Great American AI Act discussion draft of June 2026 would amend the WARN Act to require employers to state whether AI caused a mass layoff — a statutory device for manufacturing an attribution where none exists — and New York's state-level version of the same device, live since March 2025, read **zero** AI attributions across 160+ layoff notices in its first year, exactly as a hiring-margin adjustment predicts. And the geography inverts: 62 of the 100 most AI-exposed US counties voted Democratic in 2024, so the coalition that produced the 2016–2024 countermovement is not the coalition AI displaces. Whether a countermovement can form at all without a foreigner to blame is, on this account, the central political question of the next decade.

JEL: J23, J24, D72, F16, O33. **Keywords:** automation, artificial intelligence, attribution, trade politics, hiring, entry-level employment, human capital, protectionism, Polanyi.

Claims and citations

May be cited for	May NOT be cited for	Open
The reconciliation: attribution determines the FORM of the countermovement (protection versus compensation), not its magnitude.	The claim that technology escapes political blame. It does not — that was this paper's earlier hypothesis and the evidence broke it.	Whether the form distinction holds outside the survey-experimental setting it is measured in.
The hiring-margin mechanism, and the finding that adjustment is running on entry rather than on separations.	A causal estimate of AI's effect on entry-level employment. The leading rival explanation — remote work — accounts for 64 per cent of the rise in young-graduate unemployment through the same mechanism.	How much of the entry-level weakness is AI, how much is remote work, and how much is the rate cycle.
The unbundling argument: junior training was free as a by-product of wanted output, and AI gives it a positive marginal cost.	The claim that the pipeline is collapsing. No time series on training spend per junior hire exists, and no historical case of a pipeline collapsing while demand held up has been documented.	Whether firms are in fact reducing training, and whether the tipping structure is empirically relevant.
The prediction set and the indicator thresholds in Part V.	A forecast of AI capability or timelines. The argument is designed to be informative whatever the rate.	All five predictions. None is yet settled.

Contents

Part I — The mistake

1. What everyone is watching, and why it is the wrong thing
2. Nobody gets fired
3. The countermovement arrived. I said it didn't.

Part II — The theory

4. Attribution determines form, not magnitude
5. Why AI has no foreigner
6. The feedback loop, and its honest caveat

Part III — The mechanism

7. The hiring margin, and what the statistics do and do not see
8. The pipeline: why this is a cliff and not a slope
9. What the historical record does and does not license
10. Who ends up with the surplus

Part IV — The world in five to ten years

11. Rates: what is measured about how fast this moves
12. Four things that follow

Part V — What would settle it

13. Five predictions
14. Eight indicators, with thresholds
15. Conclusion

Disclosure · References

Epistemic status. Part I includes a retraction of the previous draft's central claim. Part II is a theory grounded in one survey experiment and two observational literatures that disagree with each other until the theory reconciles them. Part III is the paper's strongest empirical material on one limb and its weakest on the other, and says which is which. Part IV is conditional reasoning about a decade, not a forecast. Part V exists to make the whole thing killable, and every prediction in it has a stated falsifier.

Part I · The mistake

1. What everyone is watching, and why it is the wrong thing

The public conversation about artificial intelligence and work is organised around a single question — will it take the jobs — and monitored with a single statistic, the unemployment rate. As of mid-2026 that rate is 4.2 per cent, prime-age employment is near multi-decade highs, and this is widely read as evidence that the worry was overblown.

It is the wrong question monitored with the wrong instrument. The right question is not whether people lose jobs but **who never gets one**, and the right instruments are hiring rates, entry composition and the return to experience. This paper argues that the transition now under way is structurally difficult to see, for three separate reasons that compound: the adjustment runs through a margin the headline statistics were not built to monitor; the harm has no agent to which it can be politically attributed; and the most consequential damage is to a stock rather than a flow, so it reports a decade late.

2. Nobody gets fired

Start with the observation that anyone working in an affected industry can make for themselves. Firms are not dismissing their senior engineers. They are not opening the junior requisition.

That is a different economics and it produces different data. In the US Job Openings and Labor Turnover Survey for May 2026, the hires rate stands at 3.3 per cent against 4.3 per cent in January 2022; quits at 1.9 per cent against a peak of 3.0; and layoffs and discharges at 1.1 per cent — just 0.2 points above their all-time low. **Nothing is moving, and what moved most is hiring.** A frozen labour market does not distribute its pain evenly: it concentrates it entirely on the people trying to get in.

The firm-level evidence points the same way and is more direct. Using 65 million résumés across more than 280,000 firms, Hosseini Maasoum and Lichtinger find junior employment falling in AI-adopting firms while senior employment is unchanged — and, in their own words, this is "driven primarily by slower hiring rather than increased separations". That is the mechanism stated by the authors of the study, not inferred by me.

Why this is hard to see, and one correction to a claim I made earlier. I previously wrote that non-hiring is invisible to unemployment statistics. That is wrong, and the correction is more interesting than the error: a new graduate who actively searches is counted, and the BLS publishes a dedicated series for exactly this population. **Unemployed new entrants** — people seeking a first job — rose from 509,000 in June 2019 to **772,000 in June 2026**, up 52 per cent, and their share of total unemployment from 8.6 to 10.9 per cent. It is not invisible. It is visible in a series nobody looks at, while the series everyone looks at reads 4.2 per cent.

The genuinely uncounted margin is different, and larger. The New York Fed puts recent-graduate *underemployment* — working in a job that does not require a degree — at 41.5 per cent, and there is no BLS underutilisation measure that captures it. U-6 counts involuntary part-time work, not skill mismatch. A philosophy graduate working forty hours a week as a barista is fully employed in U-3 and in U-6 alike. Meanwhile the measures that *would* capture discouragement — U-4 and U-5 — are flat year on year, which disposes of the simpler version of this story: these people are not giving up, they are taking worse jobs.

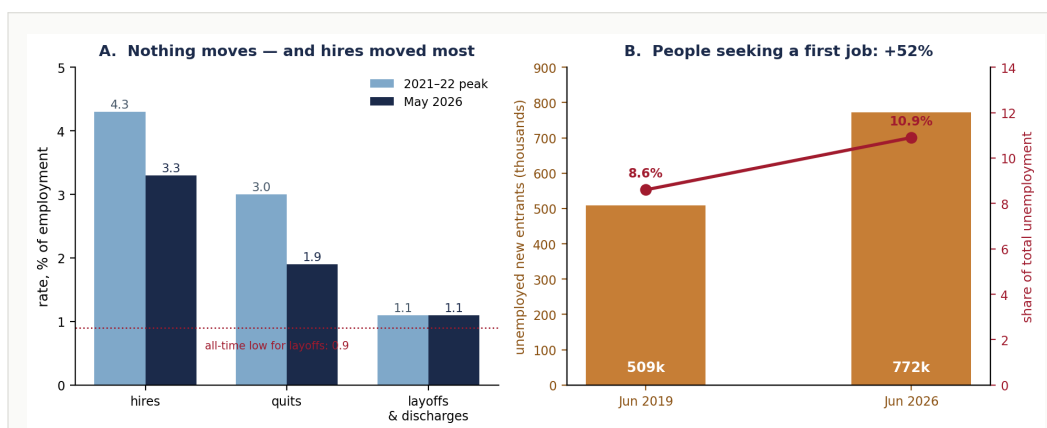


Figure 1. A frozen market, and the series nobody looks at.

What it is: JOLTS rates at their 2021–22 peak against May 2026, and the BLS unemployed-new-entrants series. What it does not prove: none of this is attributed to AI. Hires fell for several reasons, the rate cycle prominent among them, and the new-entrant series responds to anything that slows entry hiring. The claim here is about which margin is moving, not about what is moving it.

US Bureau of Labor Statistics: Job Openings and Labor Turnover Survey, May 2026; and Current Population Survey, unemployed new entrants, June 2019 and June 2026.

3. The countermovement arrived. I said it didn't.

The previous draft of this paper argued that no protective countermovement followed the post-1980 displacement, and offered as evidence that US union density roughly halved over the period. Forty-six years of measured harm, no institutional response: on that reading the Polanyian mechanism had simply failed. **That was wrong, and the way it was wrong is the subject of this paper.**

Consider what actually happened between 2016 and 2026: Brexit; two Trump administrations; Section 301 tariffs from 2018 and their expansion in 2025, with the effective US tariff rate reaching 11.8 per cent in April 2026, the highest since the early 1940s; sustained immigration restriction, with net migration possibly negative for the first time in more than fifty years; the CHIPS Act and the Inflation Reduction Act, and the general return of industrial policy. Measured against Polanyi's own criteria — protective in aim, arising from those most immediately affected, piecemeal, non-doctrinal, not centrally coordinated — that is a countermovement, and a large one. Union density was simply the wrong instrument for detecting it.

Two further premises of the earlier draft also fail, and both are worth stating plainly because they are widely believed.

- **Automation did not destroy more US manufacturing employment than trade did.** Acemoglu and Restrepo's own aggregate for robots over 1990–2007 is 360,000 to 670,000 jobs; the China shock is roughly 2.0 to 2.4 million. Trade is three to seven times larger. The widely circulated claim that automation accounts for around 88 per cent of manufacturing job loss comes from a productivity residual whose own two components sum to 101.4 per cent, and Houseman's critique is decisive on the measurement: computers are 10 to 13 per cent of manufacturing value added, and excluding them, manufacturing real value added was *lower* in 2011 than in 2000.
- **Automation did not escape political consequence either.** In the one study that puts both shocks in a single regression across fourteen countries, Anelli, Colantone and Stanig find a standard-deviation robot shock raising the radical-right vote by 1.8 points on a 5.7 per cent base, while the China shock is positive but small and insignificant. Frey, Berger and Chen find that 95 per cent lower robot exposure flips

Michigan, Pennsylvania and Wisconsin in 2016 — the only "this flipped the Rust Belt" counterfactual in the literature is about robots, not imports.

So the countermovement was neither absent nor asleep, and technology was not politically invisible. What then distinguishes the two shocks? Not their size, and not the political energy they released. **What differs is what that energy asked for.**

Part II · The theory

4. Attribution determines form, not magnitude

Di Tella and Rodrik ran the decisive experiment. Respondents read an identical vignette — a garment plant employing 900 workers closes — with only the stated *cause* randomised across arms. Demand for import protection came out as follows: 0.09 in the control, 0.13 when the cause was given as a technology shock, 0.23 for outsourcing to France, and 0.29 for outsourcing to Cambodia. **Identical harm; the cause alone raises protectionist demand by a factor of two to three.**

The second finding is the one that makes this a theory rather than an observation. The effect on the *other* policy margin runs the opposite way: trade attribution *reduces* demand for compensating the losers, by six to eight points, while technology attribution and bad-management attribution *raise* it. Attribution does not merely scale the response. It redirects it.

And the refinement that identifies the actual variable: the arm in which the closure is attributed to **bad management** — an agent, fully identifiable, but domestic — produces essentially zero protectionism and the largest transfer demand of all. So the operative variable is not attributability alone. It is **agency plus out-group membership**.

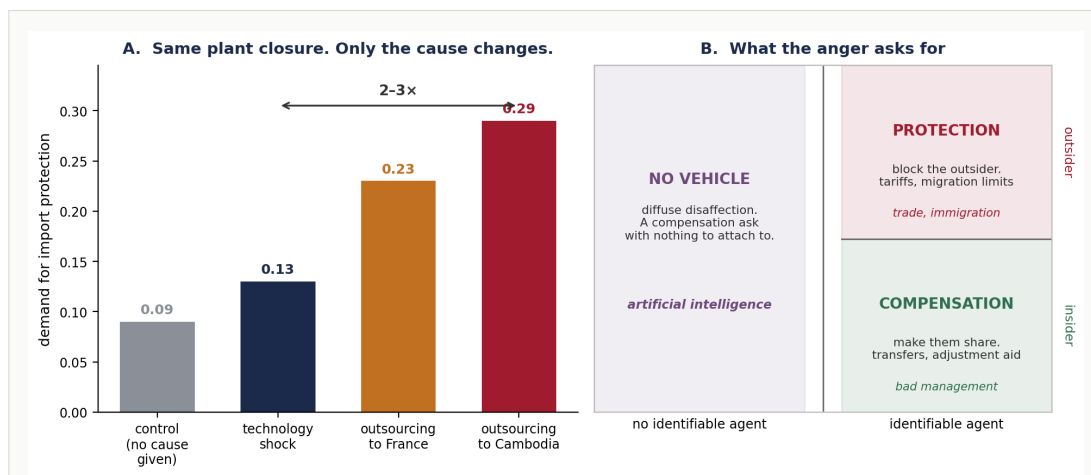


Figure 2. Same harm, different cause, different demand.

What it is: panel A gives Di Tella and Rodrik's experimental result; panel B the classification it implies, with the left column merged because when there is no agent the insider/outsider distinction is undefined. What it does not prove: this is a survey experiment on stated preferences, and the mapping from stated preference to enacted policy is an assumption. The observational literature in Section 3 is consistent with it but was not designed to test it.

Di Tella, R., and Rodrik, D. (2020), *Economic Journal* 130(628): 1008–1030. Classification constructed for this paper.

The reconciliation. Two literatures appear to disagree. One finds trade displacement driving political realignment; the other finds robot displacement doing it at least as strongly. Both are right, and the apparent conflict dissolves once magnitude and form are separated. **Displacement of any origin produces disaffection and radical-right voting at broadly comparable magnitude. Attribution determines whether the resulting demand is for protection or for compensation.** I have not found this reconciliation stated anywhere, and it is the paper's central theoretical claim.

Two authors get close and should be credited. Rodrik observes that "we do not see populists campaign against technology or automation" and that trade is "a convenient scapegoat, since politicians can point to identifiable foreigners... as the source of the problem" — though his emphasis falls on procedural unfairness first and identifiability second. And Mutz documents an eight-to-one ratio of trade to automation coverage in newspapers, and finds that trade attribution raises the belief that "the jobs can come back" — a cleaner mechanism than mine, and one I adopt. What is added here is the separation of magnitude from form, and the observation that this is what reconciles the two voting literatures.

5. Why AI has no foreigner

Every large displacement of the modern era came with a nameable outsider. The China shock had China. Offshoring had Mexico and then Vietnam. Immigration has immigrants. Even the robot, which has no nationality, arrived inside a factory owned by an identifiable firm making an identifiable decision, and the political response could and did attach to the firm.

Artificial intelligence has none of this, and the reason is structural rather than rhetorical. The displacement does not occur at a moment. There is no plant closure, no relocation announcement, no arrival. It occurs through a requisition that is never posted — an absence, distributed across tens of thousands of firms, none of which has done anything identifiable, and most of which would truthfully say they have not reduced headcount. **You cannot hold a rally against a hiring plan that was quietly revised downward.**

On the model of Section 4, that places AI in the cell with no policy vehicle: a cause with no agent generates disaffection and a compensation impulse with nothing to attach to. Which yields the paper's central political prediction, and it is not the one people expect. **The risk is not that AI produces a protectionist backlash. It is that it produces a large quantity of political energy with no available target, and that this energy is therefore either dissipated or attached to a manufactured one.**

Two developments already test this, and both are live.

- **Attribution is being manufactured by statute.** The Great American AI Act discussion draft of June 2026 would amend the WARN Act to require employers to state whether artificial intelligence caused a mass layoff, and would study an AI adjustment assistance programme modelled explicitly on trade adjustment assistance. Read through this paper's model, that is a device for legislating a name for a cause that does not have one — creating the attribution that the compensation demand requires in order to attach to anything. Whether it passes is one of the indicators in Part V.

- **And the first such device has already run, and read zero.** New York added an AI-attribution disclosure to its state WARN system in March 2025: employers filing layoff notices must indicate whether technological innovation or automation contributed, and if so name the technology. In the first year, more than 160 companies filed WARN notices. **None attributed the layoff to AI.** The instrument has real defects — the statutory text was never amended, the key phrase is undefined, and it is a checkbox rather than an audit — but the zero is also exactly what the hiring-margin mechanism predicts: WARN triggers on *separations*, and separations are the margin that is not moving. An attribution device bolted to the layoff process will read zero even while the adjustment is large, because the adjustment is running through the requisition that never opens.

- **The geography inverts.** The China shock and the robot shock hit substantially the same places, which is why they produced one coherent coalition. AI does not: 62 of the 100 most AI-exposed US counties voted Democratic in 2024. The population this displacement threatens is not the population the last one threatened, and the political vehicle built by the last one is therefore unavailable to it. This is falsifiable

within a single electoral cycle.

6. The feedback loop, and its honest caveat

There is an uncomfortable corollary. Protection and immigration restriction both raise the price of labour relative to capital, and labour scarcity induces automation. The clearest natural experiment is Clemens, Lewis and Postel on the 1964 termination of the Bracero programme: excluding roughly half a million Mexican farm workers produced no detectable wage gain for domestic workers, and immediate adoption of the mechanical tomato harvester. Acemoglu and Restrepo find that demographic ageing — scarcity of middle-aged workers — explains around 40 per cent of cross-country variation in robot adoption. **A countermovement aimed at the wrong target does not merely fail. It accelerates the thing it is reacting to.**

The caveat, which cuts against the clean version and which I would rather state than bury: the 2024 Section 301 escalation placed duties of 25 to 50 per cent on semiconductors, batteries and solar — that is, tariffs on *capital goods*, which raise the price of automation rather than lowering it. Flaaen and Pierce found machinery and computers among the sectors most damaged by the 2018 round, and estimate a net effect on manufacturing employment of **−2.7 per cent**: rising input costs and retaliation outweighed the protection. So the enacted policy is ambiguous in sign for the automation loop, even as it is unambiguously negative for the employment it was enacted to protect. The general mechanism is well evidenced; its application to the specific tariff schedule of 2018–2026 is not.

Part III · The mechanism

7. The hiring margin, and what the statistics do and do not see

If the political question is what displacement gets blamed on, the empirical question is where displacement shows up. The answer, and it is the most load-bearing empirical claim in this paper, is that it is not showing up where everyone is looking. **Adjustment is running through hiring, not through separations.** Firms are not dismissing senior staff. They are not opening junior requisitions.

Figure 1 gave the shape of it: a market in which very little happens. That is not the signature of a technology destroying jobs in the way the public argument imagines — a wave of firings, a spike in the unemployment rate. It is the signature of a market adjusting on the only margin that is cheap: the requisition that never opens. Nobody is fired, so nobody appears in the layoff statistics, so nobody appears on the news. The adjustment is real and it is silent, and it falls entirely on people who do not yet have the job.

The best-identified firm-level evidence says the same thing in the authors' own words. Hosseini Maasoum and Lichtinger, matching 65 million résumés to more than 280,000 firms, find that junior employment falls in AI-adopting firms while senior employment is unchanged, and that the fall is "driven primarily by slower hiring rather than increased separations". That is the mechanism stated as a finding rather than as a hypothesis, and it is why this paper puts weight on it.

7.1 A correction, and then the margin that really is invisible

An earlier version of this argument claimed the entry-level adjustment was invisible to official statistics. That is wrong and worth correcting precisely, because the truth is more useful. It is not invisible. It is visible in a series nobody looks at. BLS has published unemployed new entrants for decades; a 52 per cent rise in seven years is a large move in a series with no constituency.

What is invisible is the underemployment margin described in Section 2, and it matters here for a specific reason. If AI closes the graduate-entry rung without reducing labour demand in aggregate, the observable consequence is not unemployment at all. It is a cohort working, earning, counted as employed — and not accumulating the human capital that the job title used to come bundled with. The headline series is structurally incapable of showing that. Which is the subject of the next section.

What cuts against this section. Four things, and all are serious. The recent-graduate and overall unemployment rates crossed over in **February 2019** — before the pandemic and well before any plausible AI effect. Emanuel, Harrington and Pallais find that **remote work explains 64 per cent** of the rise in young-graduate unemployment from 2017–19 to 2022–24, through the same lost-mentorship mechanism this paper describes — that paper concedes the mechanism and takes the attribution. Fast and Gould at EPI note that young *noncollege* unemployment rose by a similar amount over the same window, to 7.1 per cent, so nothing college-specific is happening — their preferred story is credential erosion, with college attainment among young adults up from 18.0 per cent in 1979 to 31.6 per cent now. And Audoly, Guerin and Topa at the New York Fed find the relative posting decline in AI-exposed occupations **began before ChatGPT**. One reply is available to the third point: the hiring-margin mechanism does not predict a college-specific effect — entry rungs close for noncollege juniors too. But anyone using the entry-level evidence here as an AI result is overreading it: what the evidence supports is that the entry rung is under strain and that the mechanism is mentorship, not that AI is the cause.

7.2 The randomised trials look like they refute this. They do not.

There is an apparent contradiction that should be faced directly. Every well-identified trial of AI assistance finds gains that are *junior*-biased. Cui and co-authors, pooling three randomised controlled trials across 4,867 developers, find juniors gaining 21 to 40 per cent and seniors 7 to 16. Brynjolfsson, Li and Raymond find novice customer-support agents gaining 34 per cent and experienced agents approximately nothing. If AI helps juniors most, how can it be displacing them?

The mechanism does not require seniors to become more productive. It requires **senior + AI** to substitute for **senior + junior**. That condition holds precisely when AI is most productive at junior-type tasks — which is exactly what the trials find. Read correctly, the junior-biased gradient is evidence *for* the mechanism, not against it.

One citation hazard belongs here because it is circulating. METR's February 2026 follow-up found **−18 per cent** for the original cohort (confidence interval −38 to +9) and −4 per cent for new recruits, and the team is redesigning the study because selection broke it. A widely shared 2026 summary reports this as a "+18 per cent speedup" — a sign error on a headline result. The only sources claiming senior-biased gains are vendor telemetry with no stated methodology. The honest position is that the trial literature supports junior-biased task gains and says nothing yet about the substitution the mechanism needs.

8. The pipeline: why this is a cliff and not a slope

The hiring margin explains why displacement is quiet. The pipeline explains why quiet is dangerous. The argument is narrow and it is the paper's own, so it should be stated in one paragraph and then defended.

The unbundling argument. Junior training was historically free to the firm, because it was a by-product of output the firm wanted anyway: the associate's document review, the resident's rounds, the graduate engineer's tickets. The firm bought the output and got the training as change. AI supplies that output directly. Training therefore acquires a positive marginal cost **for the first time** — and a cost that must be paid deliberately, by a firm that does not capture its full return, because trained workers leave. That is a textbook underprovision problem newly created by a technology, not an old one made worse.

The price of the unbundled half is measurable. Emanuel, Harrington and Pallais put a number on mentorship: across 1,055 engineers, senior engineers write **0.76 fewer programs per month** while mentoring. That cost was always there. What has changed is that the offsetting benefit — the junior's output — can now be bought elsewhere. When the two were bundled, the mentorship cost was paid without ever being priced. Unbundled, it is a line item, and line items get cut.

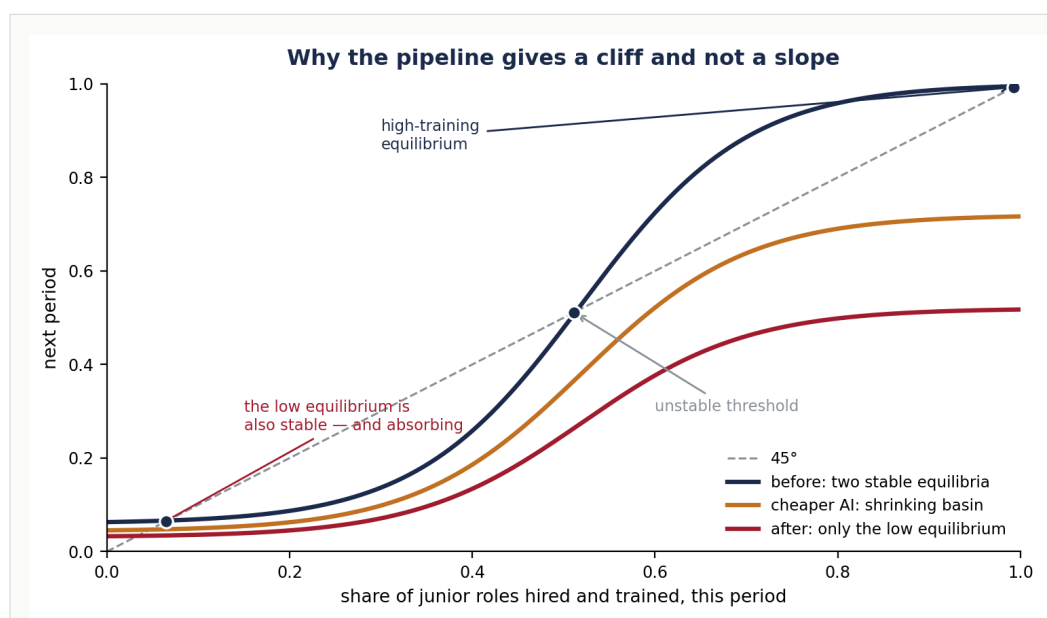


Figure 3. Training equilibria come in pairs, which is where the discontinuity comes from.

Stylised. The share of junior roles trained next period as a function of the share trained this period, against the 45-degree line. Firms train when other firms train, because a thick junior market lowers the cost of replacing leavers — so the map is S-shaped and crosses the diagonal three times: a high stable equilibrium, an unstable threshold, and a low stable one. Cheaper AI shifts the map down. Past a point the high equilibrium does not shrink, it *ceases to exist*, and the system falls to the absorbing state. The curves are illustrative; the pair structure is a theorem.

Structure from Acemoglu & Pischke (1998) and Afrouzi, Blanco, Drenik & Hurst (2026). Figure generated by *attribution_figures.py*.

Afrouzi, Blanco, Drenik and Hurst formalise entry tasks as "the curriculum through which workers accumulate human capital" and prove that equilibria always exist in *pairs* — a high-learning one and a low-learning one — with cheaper technology raising learning in the first and driving it to zero in the second.

Acemoglu and Pischke had the same structure in 1998 for general training under labour-market frictions. This is where the discontinuity comes from, and it is why the right mental image is not the post-2008 cohort-scarring literature. Scarring is a slope: a bad cohort earns less for a decade and the gap closes. A training equilibrium is a tipping point between two stable states, and the transition is invisible while the existing stock of seniors is still working.

That last clause is the whole forecasting problem. The stock of senior competence in any profession was produced by a pipeline that ran ten to twenty years ago. If the pipeline stops today, output is unaffected today, and is unaffected for as long as the existing stock lasts. The firm that stops training sees no cost in its own accounts within any horizon over which its managers are evaluated. The consequence arrives as a step, a decade later, in a labour market that has no way of pricing it in advance.

8.1 The reverse experiment: PATCO, 1981–1992

No historical case of a pipeline collapsing while demand held up has been found — the previous draft said so and it remains true. But the *reverse* experiment exists, and it puts a number on the one thing the argument needs a number for. On 5 August 1981 the federal government terminated roughly 10,400 of the FAA’s 17,000 air traffic controllers, leaving 4,669 operational controllers. Demand did not fall — air traffic kept growing — and the training pipeline was not merely intact but run at the greatest throughput in its history: the FAA administered its aptitude test over **400,000** times between 1981 and 1992, nearly 50,000 in the peak year, and sent 25,277 selectees to its Academy. The agency’s own retrospective declares workforce recovery complete in **mid-1992** — eleven years — and on the stricter headcount reading the workforce reached about 14,400, approaching pre-strike levels, only by **mid-1995**.

Read carefully, this is a measurement of the pipeline’s **time constant**: with demand guaranteed, funding guaranteed, and selection run at maximum scale, converting novices into a stock of senior competence took **eleven to fourteen years**. The bound runs in one direction only. A system trying its hardest took that long; a system that has quietly stopped trying will not be faster. It also prices the asymmetry the forecasting problem turns on: destroying the stock took one afternoon.

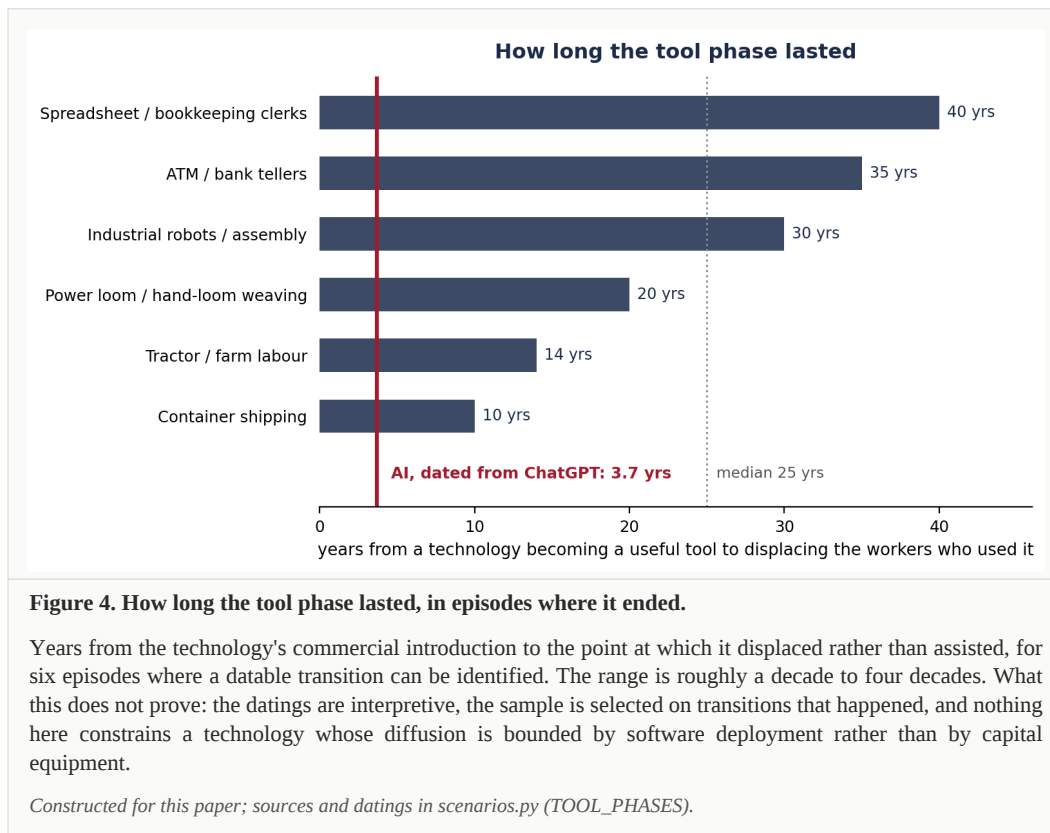
The honest weakness, restated. The *forward* case — a pipeline collapsing while demand held up, with clean numbers and a decade-lagged consequence — remains undocumented. UK apprenticeships roughly halved between 1979 and 1995, but the industries died first, so the causation runs the wrong way. PATCO bounds the recovery time, not the probability of collapse. The pipeline argument is therefore best read as: theoretically well-founded, unprecedented in its forward form, and carrying a measured eleven-to-fourteen-year repair time if it fires. The indicators in Part V are what would tell us.

9. What the historical record does and does not license

Both sides of the public argument reach for history, and both reach for the wrong level of it. The optimist says every previous wave created more jobs than it destroyed; the pessimist says this time the machines take everything. The record answers three different questions with three different answers, and almost all confusion in this area is a matter of answering one while being asked another.

Question	Answer in the record	What it does not license
Do displaced occupations recover?	No. Essentially never. Hand-loom weavers 240,000 (1820) to 10,000 (1860). UK coal 1,191,000 (1920) to 360 (2022). US coal 863,000 (1923) to 41,600 (2020) — while output rose to an all-time peak in 2008. US agriculture 41 per cent of employment (1900) to 1.9 (2000).	Almost nothing on its own. This is selection on the dependent variable: occupations that were not destroyed are not in the sample.
Do the specific displaced people recover?	Mostly not, within their own working lives — roughly one in seven. The China-shock literature finds 86 per cent of net job loss absorbed by a falling employment rate rather than by reallocation, with effects still measurable in 2019. Displaced workers lose about 1.4 years of prior earnings in expansions and 2.8 in recessions.	This is the level the public argument is actually about, it is the level at which the record is most pessimistic, and it is the level least often cited. It licenses no claim about aggregates.
Does labour in aggregate recover?	Always, so far. Roughly one for one. No episode in the record reduced aggregate employment over the medium run; Bessen, across 317 occupations and 243 industries from 1980 to 2013, puts the net effect of computer use on total employment at -0.07 per cent a year. Composition, not level.	Every recovery ran through generational replacement over fifty to a hundred years, not through retraining. Aggregate recovery is compatible with total individual ruin, and historically has been.

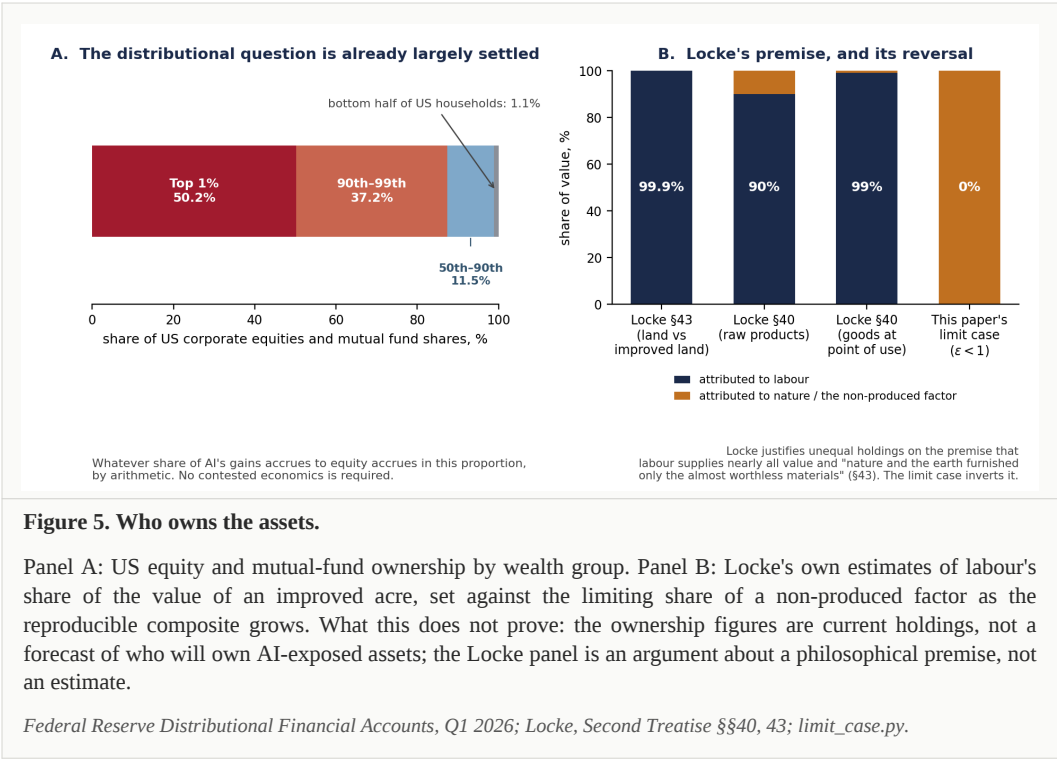
The three answers are not in tension. They describe a process in which the aggregate is repaired by the arrival of people who were never in the destroyed occupation, on a timescale longer than a career. That is the relevant constraint on any AI transition, and it is a constraint on *speed* rather than on *level*. It is also why the political question in this paper is the operative one: over the interval in which the aggregate repairs itself, the people bearing the cost will be voting.



The tool-versus-replacement distinction is worth taking seriously precisely because so many people report the former. Nearly every senior software engineer describes current AI as a large productivity gain and not as a substitute, and they are reporting accurately. The historical point is only that this is what the interior of a transition feels like. Power looms assisted before they replaced; the tractor assisted for fourteen years before the 1924 Farmall closed the remaining task set. The phase is real, its duration has ranged from about ten to about forty years, and the phase ending has never been announced in advance.

10. Who ends up with the surplus

Two questions are routinely fused in public argument and should be separated: how much surplus AI produces, and who receives it. The first is genuinely uncertain. The second is close to arithmetic.



The top one per cent of US households hold **50.2 per cent** of corporate equities and mutual-fund shares; the ninetieth to ninety-ninth percentiles hold a further 37.2, so the top decile holds about **87 per cent**. The bottom half holds about 1 per cent. Whatever share of AI's gains accrues to equity accrues roughly half to the top one per cent and roughly seven-eighths to the top decile. One need believe nothing about elasticities or timelines to reach that; one needs a share register. It also means the distributional worry does not wait on the technological one. A weaker AI produces a smaller surplus distributed in the same proportions. **The proportion is the settled part.**

10.1 The limiting case, and why "only capital survives" is wrong

The natural fear — if AI drives labour's marginal product to nothing, capital owns everything — does not survive the model that generates it. Capital is *reproducible*. In a task framework where a growing share of tasks becomes automatable, labour's share does fall toward zero, but the returns do not accumulate to the owners of machines, because machines can be built. They accumulate to whatever cannot be reproduced: land, energy, spectrum, compute at a binding physical constraint, regulatory position, data with a legal moat. Whether that is a large share or a trivial one turns on a single elasticity, ϵ , between the reproducible composite and the non-produced factor.

ϵ	Terminal share of the non-produced factor	Reading
0.5 (complements)	96.5%	Rentier terminus. Nearly the whole product accrues to whoever owns the thing that cannot be built.

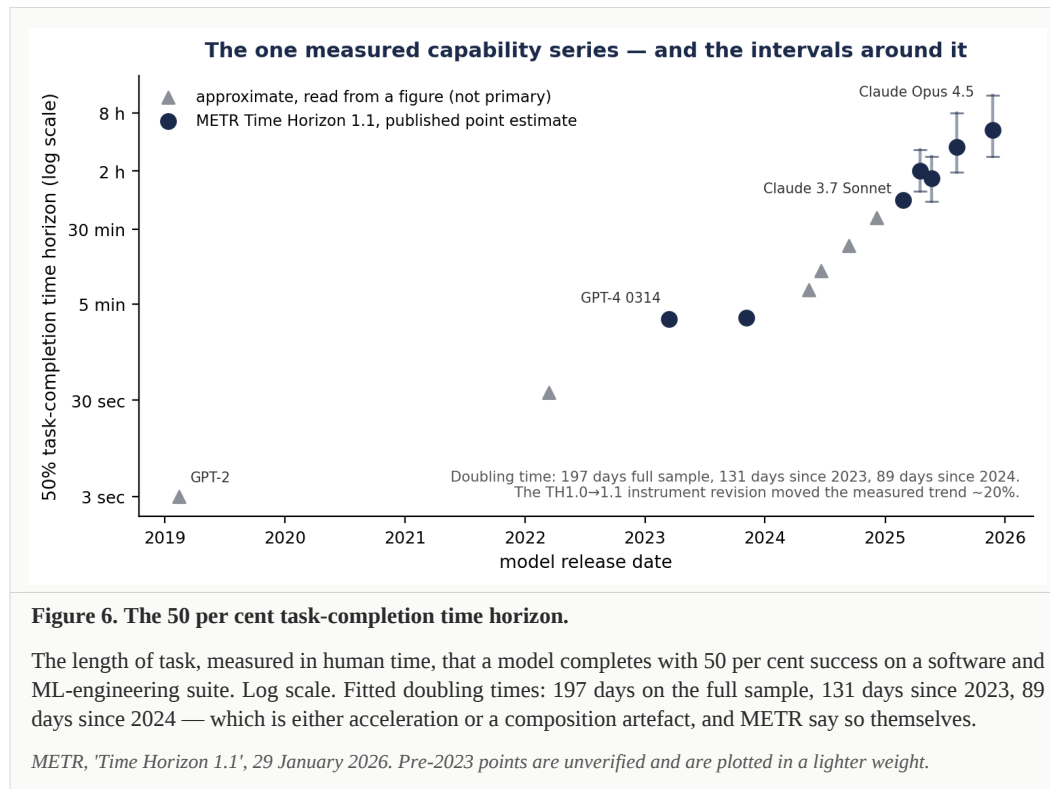
ε	Terminal share of the non-produced factor	Reading
1.0 (Cobb–Douglas)	5.0%	Shares are constant by construction. Nothing happens.
1.5 (substitutes)	0.65%	The constraint is engineered around and the non-produced factor becomes irrelevant.

Three orders of magnitude between the cases, and ε **is not identified**. That is the honest state of the question, and it is why this paper does not forecast a distributional terminus. What the exercise does establish is the direction of the correction: the intuition that "only capital survives" points at the wrong asset. Reproducible things do not stay scarce. The candidate for terminal concentration is not capital in general but the non-produced residual — which is a claim about land, energy and law, and which implies a completely different policy conversation from the one currently being had about robots.

Part IV · The world in five to ten years

11. Rates: what is actually measured about how fast this moves

"Things are exponential" is the most common premise in this discussion and the least often given a number. There is one series that measures capability on a time scale rather than on a benchmark score, and it is worth using carefully because it is also the series most often abused.



Four caveats, all from METR. Confidence intervals are wide — upper bounds run 1.66× to 2.28× the point estimates. Only 5 of 31 tasks over eight hours have measured human baselines. The instrument revision from TH1.0 to TH1.1 moved the measured trend by about 20 per cent, with old models falling 35–57 per cent and new models rising 11–55 — so part of the apparent acceleration is an artefact of the instrument. And the suite is software and ML engineering only. Nothing here licenses extrapolation to legal, medical, managerial or physical work.

Taken at face value, a doubling time of three to six months on task length is a fast rate by any historical standard for a general-purpose technology. But the rate that matters for this paper is not the capability rate. It is the **deployment** rate, and the historical record on that is unambiguous and much slower: electrification took roughly four decades from the first central station to the productivity surge, because the complementary reorganisation — unit drive, factory layout, the whole management system — had to be invented separately. The three lags that separate a capability from an economic effect are invention, diffusion, and reorganisation, and only the first is on the METR curve.

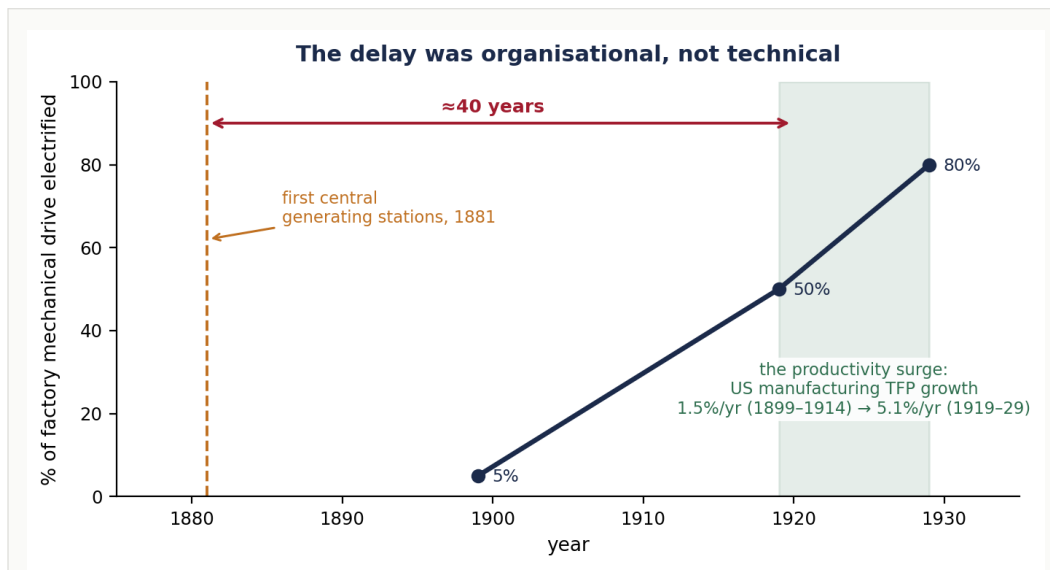


Figure 7. The delay was organisational, not technical.

Share of US factory mechanical drive electrified, against the dating of the first central generating stations. Roughly forty years separate the availability of the technology from the productivity surge, and the binding constraint was the reorganisation — unit drive, factory layout, the whole management system — rather than the electricity. What this does not prove: software diffuses genuinely faster than capital equipment, so the forty-year figure is an upper bound and not a forecast. The point is that the reorganisation lag has no particular reason to have shortened, because it is bounded by how fast organisations and the people in them can change.

Devine (1983) and the electrification series in data.py; TFP figures from Kendrick via David (1990). See also Crafts (2022).

This is the substantive reason the paper does not forecast. If the deployment rate is fast, the mechanisms in Part III fire within the decade and the political question in Part II becomes urgent. If it is slow, they fire later and more of the adjustment runs through ordinary generational replacement. **The argument is designed to be informative either way**, because attribution structure and the hiring margin do not depend on the rate. What depends on the rate is only when the indicators in Part V move.

12. Four things that follow

Granting the theory in Part II and the mechanisms in Part III, four consequences follow for the next five to ten years. Each is conditional, and each is stated so that it could be wrong in a recognisable way.

12.1 The politics gets angrier and less coherent at the same time

This is the direct implication of the attribution result. Displacement produces disaffection at broadly the same magnitude whatever causes it, so the anger arrives on schedule. But a compensation demand needs an attributable cause to attach to, and AI does not supply one — there is no foreigner, no plant relocation, no identifiable decision-maker who chose this. The predicted result is high and rising disaffection with no stable policy vehicle: volatile support, rapid switching between candidates and positions, and a politics that looks irrational from the outside because the demand has no natural object. Anti-incumbency without a programme is what this model predicts, and it is roughly what 2024–2026 has looked like across the OECD.

12.2 Someone will manufacture a target

A cause with no agent is politically unusable, and unusable political demand does not evaporate — it gets supplied with an object. There are three candidates already visible. **Legislative attribution:** the Great American AI Act discussion draft would amend the WARN Act to make employers state whether AI caused a mass layoff, which is precisely a device for creating an attribution where none exists. **Corporate attribution:** the frontier labs are a small, named, wealthy, geographically concentrated set of firms, and they are the most available agent in the picture. **Foreign attribution:** reframing AI as a Chinese technology race converts an agentless domestic cause into an out-group one, and would — on this model — flip the demand from compensation back to protection. The third is the one worth watching, because it is the only route by which AI displacement produces classical protectionism.

12.3 The damage lands on entrants, and lands late

The hiring margin means the visible cost falls on people without a job rather than people losing one, and that population is politically weak: young, unorganised, dispersed across employers, and not represented by any institution built for the purpose. Unions represent incumbents; the displaced-worker programmes are triggered by separations, which is exactly the event that is not happening. The pipeline argument then adds that the aggregate consequence arrives roughly a decade after the cause, when the existing stock of senior competence begins to retire without replacement. **A cost borne by the politically weakest group, arriving on a lag longer than an electoral cycle, attributed to nothing in particular** is close to a worst case for institutional response.

12.4 The compensation question arrives before the automation question

Public argument treats universal basic income as the policy for a world without jobs. On this account the sequence is the other way round. Long before any aggregate employment effect is measurable, the distributional question is live — because whatever surplus exists accrues to equity, and equity ownership is already what it is. The likely political form is therefore not UBI in the technological-unemployment sense but ordinary distributional conflict over an unusually concentrated gain: taxation of returns, sectoral bargaining, adjustment assistance, licensure and professional-scope fights. The Hobbesian version of the worry — that people who contribute nothing to production have no leverage and therefore no claim — is the thing to watch for in the rhetoric, because contractarian justifications of the social order have an exclusion problem exactly when a group's cooperation stops being needed. That is a philosophical vulnerability with a political expression, and it is the reason Locke rather than Hobbes is the better foundation for whatever gets built.

Part V · What would settle it

13. Five predictions

A theory that cannot be killed is not worth having. Each of the following carries an explicit falsifier, and none of them is settled today.

	Prediction	Falsified if
P1	Displacement shows up in hiring rates and in new-entrant unemployment before it shows up in the headline unemployment rate — and possibly without ever showing up there.	The unemployment rate rises materially while the hires rate is stable.
P2	The junior-to-senior hiring ratio falls in AI-exposed occupations, and the fall is visible in the task content of entry-labelled postings before it is visible in their count.	Task-text seniority of entry roles is flat while counts fall — that would be an ordinary demand story, not this one.
P3	Political response to AI displacement takes the form of compensation demands rather than protection demands, because there is no foreigner to exclude — unless a target is manufactured.	AI-driven displacement produces trade or immigration restriction aimed at AI specifically; or produces no policy demand at all.
P4	The AI countermovement does not inherit the trade countermovement's coalition, because it does not hit the same places.	AI-exposed districts move the same way manufacturing-exposed districts did after 2000.
P5	Attempts are made to give AI displacement a legal name , because a compensation demand needs an attributable cause to attach to.	Such attempts stop appearing. Partially confirmed already by the June 2026 WARN Act discussion draft.

P3 and P4 are the ones I would bet against myself on. P3 fails if the manufactured-target route in the previous section works, and there is no strong reason to think it will not. P4 fails if AI exposure at the county level turns out to be a poor proxy for who actually loses work, which is entirely possible given how those exposure measures are built.

14. Eight indicators, with thresholds

These are the series to watch, with what each reads now, what would count as the mechanism firing, and what would kill it. Three of them do not exist yet and someone should build them.

Indicator	Where it stands now	Fires if / Dies if
1. Unemployed new entrants (BLS) <i>watches the hiring margin directly</i>	772,000 in June 2026, up 52 per cent from 509,000 in June 2019; 10.9 per cent of total unemployment against 8.6.	Fires: the share keeps rising while headline unemployment is flat or falling. Dies: it reverts toward 8–9 per cent while AI adoption continues.
2. Recent-graduate underemployment (NY Fed) <i>the genuinely uncounted margin</i>	41.5 per cent, and no BLS underutilisation measure captures it.	Fires: it rises while U-3 and U-6 are stable. Dies: it falls back toward its pre-2019 level.
3. JOLTS hires rate <i>whether adjustment runs on entry rather than separation</i>	3.3 per cent in May 2026 against 4.3 in January 2022, with layoffs at 1.1.	Fires: hires stay depressed with layoffs at record lows — a frozen market. Dies: hires recover without any rise in layoffs.
4. Junior-to-senior posting ratio, measured on task text rather than title <i>the sharpest single test of the pipeline</i>	Contested. Indeed reports entry postings –7.5 per cent year on year against senior +15; two Lightcast-based analyses — including Audoly, Guerin and Topa at the New York Fed — find junior and senior demand moving broadly in parallel. But PwC's June 2026 barometer — one billion postings, 27 countries, 2.4 million US entry-level jobs — finds AI-exposed entry roles seven times more likely to demand traditionally senior skills, with seniorised entry roles up 35 per cent since 2019 while conventional entry roles fell 10; and Indeed finds 30 per cent of entry-level applications now come from people with ten or more years' experience. The bin is stable; its contents are not.	Fires: task-text seniority of entry-labelled roles keeps rising. Dies: entry roles retain entry-level task content. Measuring on title will miss this entirely.
5. Within-occupation return to experience, by AI exposure <i>whether expertise rent is being extinguished</i>	The within-occupation test is still unrun, but a between-occupation cousin now exists: Davis (Dallas Fed, Feb 2026) finds the experience premium — median 40 per cent across 200+ occupations — positively correlated with AI exposure, with exposure reducing wage growth 0.28pp in zero-premium occupations, 0.05pp at the median, and <i>raising</i> it 0.2pp at the 90th-percentile premium; exposed-sector employment losses fall disproportionately on workers under 25. Direction: seniority protects. The compression test proper remains a day's work on CPS or ACS (companion script: <code>indicator5_experience_gradient.py</code>).	Fires: the gradient compresses in exposed occupations and not elsewhere, concentrated in the first decade of tenure. Dies: no differential compression.
6. Firm training and mentorship spend per junior hire <i>whether the unbundling is real</i>	No time series exists. This is the single most important missing measurement for the argument in Section 8.	Fires: training spend falls while senior productivity rises. Dies: training holds up in AI-adopting firms.

Indicator	Where it stands now	Fires if / Dies if
7. Whether AI acquires a legal name as a cause of layoffs <i>watches attribution directly</i>	Partially fired — and it read zero. New York added AI-attribution disclosure to its state WARN system in March 2025; in the first year, 160+ notices were filed and none cited AI or automation. The federal Great American AI Act discussion draft (June 2026) would do the same nationally and study AI adjustment assistance modelled on TAA.	Fires: the federal draft passes, or further jurisdictions mandate attribution — and, per this paper, separation-triggered instruments keep reading zero while hiring-margin indicators move. Dies: attribution mandates lapse and nothing replaces them; or WARN-type instruments start reading materially nonzero, which would mean adjustment has moved to separations and Prediction 1 is in trouble.
8. The political geography of AI exposure <i>whether a coalition can form at all</i>	62 of the 100 most AI-exposed US counties voted Democratic in 2024, inverting the China-shock and robot-shock geography.	Fires: AI-exposed districts shift measurably between 2026 and 2028. Dies: no realignment appears by 2028.

On timing: indicators 1, 2 and 3 are monthly and would give a readable signal within twelve to eighteen months. Indicator 4 requires someone to re-run the posting analyses on task text rather than titles, which is a matter of months and not of data availability. Indicators 5 and 6 need studies that nobody has yet run. Indicator 7 resolves on a legislative calendar. Indicator 8 resolves in November 2026 and more decisively in November 2028. **If four or more of the eight have not moved in the predicted direction by the end of 2028, the paper is wrong and should be discarded rather than patched.**

15. Conclusion

The question the public asks about AI and work is whether the machines will take the jobs, and it is monitored with the unemployment rate. Both the question and the instrument are wrong, in a way that has a specific and checkable structure.

The instrument is wrong because a frozen labour market adjusts on the margin that is cheapest, and the cheapest margin is the requisition that never opens. Nobody is fired. Hires have fallen a full point since 2022 while layoffs sit near an all-time low, junior employment falls in AI-adopting firms while senior employment does not, and the number of people looking for a first job is up by half in seven years. None of that appears in the headline rate, and some of it never will, because the terminal state of a closed entry rung is not unemployment but underemployment — which no official series measures.

The question is wrong because it treats the aggregate as the thing at stake. The historical record is clear that aggregate employment recovers and equally clear that the specific displaced people mostly do not, within their own working lives — and that every recovery in the record ran through generational replacement rather than retraining. What is genuinely new is narrower than either side of the public argument suggests: junior work was the mechanism by which senior competence got produced, and it was free because it was a by-product of output the firm wanted anyway. AI supplies that output directly. Training acquires a positive marginal cost for the first time, in a market that underprovides it, and training equilibria come in pairs — so the failure mode is a tipping point, invisible while the existing stock of seniors is still working, and arriving as a step a decade later. That argument is theoretically well-founded and historically unevicenced, and it is stated here as a mechanism to watch rather than a result.

And the politics is the part this paper started out wrong about. I argued in the previous draft that no protective countermovement followed the post-1980 displacement. One did — 2016 to 2026 satisfies every Polanyian criterion — and it aimed at trade and migration rather than at machines. The reason is not that technology escapes blame. Robot displacement moves the radical-right vote at least as strongly as trade displacement does. The reason is that attribution determines the *form* of the demand and not its magnitude: a cause with an identifiable outside agent produces demands to block the outsider, a cause with an identifiable inside agent produces demands to make them share, and a cause with no agent at all produces disaffection with nowhere to go.

Artificial intelligence is the first major displacement in modern economic history with no foreigner attached to it. The anger will arrive on schedule. What it asks for is undetermined — and, on the evidence here, that is being decided right now, by whoever succeeds in giving the thing a name.

This is not a forecast, and the argument is deliberately built so that the rate of AI progress does not settle it. Five predictions and eight indicators are given above with explicit thresholds, three of which require measurements nobody has yet made. The most useful thing a reader can do with this paper is run indicator 5 — the within-occupation return to experience by AI exposure — which is a day's work on public microdata and which would either strengthen the mechanism considerably or kill it outright.

Disclosure and epistemic status

This is an independent working paper. It has not been peer-reviewed, it is not affiliated with any institution, and it was written with substantial AI assistance: source location, citation verification, adversarial critique and drafting support. Every claim that remains is one I am willing to defend; every claim that was cut was cut because it did not survive scrutiny. Errors are mine.

This draft retracts the central claim of the previous one. Draft 3 argued that no protective countermovement followed the post-1980 displacement, reading collapsing union density as the absence of an institutional response. Section 3 explains why that was wrong. Two of the premises underneath it were also wrong: the widely circulated figure attributing 88 per cent of manufacturing job loss to automation is a residual from an accounting decomposition rather than an estimate, and the claim that technology displacement escapes political consequence is contradicted by the one study that horse-races both shocks in a single specification. Readers of the earlier drafts should treat those sections as superseded.

The manuscript has been through three independent verification passes with different framings. The second found twenty-four items the first missed, including one figure the first had corrected in the *wrong direction*. That is the argument for doing it more than once, and readers should still check any number they intend to use. Specific hazards worth repeating: Acemoglu and Restrepo’s automation TFP figure is 3.4 per cent in the published version and 3.8 in the working paper; Brynjolfsson, Li and Raymond’s published figures are 15 and 30 per cent, not the widely quoted 14 and 34; METR’s February 2026 follow-up is **−18 per cent** and is circulating as “+18 per cent”; and US union density series are not splice-compatible across the 1970s denominator change.

The weakest material in this paper, in order. The pipeline argument in Section 8: its forward form has no documented historical case behind it, and the PATCO episode bounds only the repair time, not the probability of collapse. The attribution of entry-level weakness to AI, where the leading rival explanation — remote work — accounts for 64 per cent of the relevant rise and the graduate/overall unemployment crossover predates any plausible AI effect. The elasticity ϵ in Section 10, which is unidentified and on which the rentier case entirely depends. The AI exposure measures throughout, which are vendor telemetry rather than representative instruments. And the extension of a survey-experimental result to enacted policy, which is the joint on which the whole theory turns and which the experiment was not designed to bear.

Reproduction. The attribution model, the indicator set, the derivations, the figures and the verified data series are in `attribution.py`, `attribution_figures.py`, `scenarios.py`, `limit_case.py`, `toy_model.py`, `generate_figures.py` and `data.py`. Simulations use fixed seeds and common random numbers across parameter sweeps. Every series in the data file carries a source string and a verification tag, and anything tagged unverified is not used in the text. Earlier drafts of this project, which developed the Polanyi and Marx material at greater length, are retained as *The Third Disembedding* (drafts 1.0 and 2.0) and *The Slowness Condition* (draft 3.0).

References

- Acemoglu, D. (2025). The simple macroeconomics of AI. *Economic Policy*, 40(121), 13–58.
- Acemoglu, D., & Johnson, S. (2023). *Power and Progress*. PublicAffairs.
- Acemoglu, D., & Johnson, S. (2024). Learning from Ricardo and Thompson. *Annual Review of Economics*, 16, 597–621.
- Acemoglu, D., & Pischke, J.-S. (1998). Why do firms train? Theory and evidence. *Quarterly Journal of Economics*, 113(1), 79–119.
- Acemoglu, D., & Restrepo, P. (2019). Automation and new tasks. *Journal of Economic Perspectives*, 33(2), 3–30.

- Acemoglu, D., & Restrepo, P. (2022). Tasks, automation, and the rise in U.S. wage inequality. *Econometrica*, 90(5), 1973–2016.
- Acemoglu, D., & Robinson, J. A. (2015). The rise and decline of general laws of capitalism. *Journal of Economic Perspectives*, 29(1), 3–28.
- Acemoglu, D., Autor, D., & Johnson, S. (2026). *Building pro-worker AI*. NBER Working Paper 34854; The Hamilton Project.
- Afrouzi, H., Blanco, A., Drenik, A., & Hurst, E. (2026). Changing entry tasks and human capital accumulation. Working paper.
- Aghion, P., Jones, B. F., & Jones, C. I. (2019). Artificial intelligence and economic growth. In Agrawal, Gans & Goldfarb (eds.), *The Economics of Artificial Intelligence*. University of Chicago Press.
- Alfani, G. (2021–). Work on epidemics and inequality; see *Journal of Economic Literature* and related surveys. [Cited for the finding that seventeenth-century Italian plagues produced no levelling.]
- Allen, R. C. (2009). Engels' pause. *Explorations in Economic History*, 46(4), 418–435.
- Allen, R. C. (2017). *The hand-loom weaver and the power loom*. NYU Abu Dhabi Working Paper 0004.
- Anelli, M., Colantone, I., & Stanig, P. (2021). Individual vulnerability to industrial robot adoption increases support for the radical right. *PNAS*, 118(47).
- Arendt, H. (1958). *The Human Condition*. University of Chicago Press. [Prologue, p. 4 in the 2013 Chicago edition; pagination is edition-dependent.]
- Audoly, R., Guerin, M., & Topa, G. (2026). Do job postings show early labor-market effects of AI? *Liberty Street Economics*, Federal Reserve Bank of New York, 14 May.
- Autor, D. H. (2015). Why are there still so many jobs? *Journal of Economic Perspectives*, 29(3), 3–30.
- Autor, D. H., & Thompson, N. (2025). *Expertise*. NBER Working Paper 33941.
- Autor, D. H., Dorn, D., & Hanson, G. H. (2013 and later). The China syndrome and the China shock literature.
- Becker, L. C. (2005). Reciprocity, justice, and disability. *Ethics*, 116(1).
- Benanav, A. (2020). *Automation and the Future of Work*. Verso.
- Bessen, J. (2016). *How computer automation affects occupations*. Boston University School of Law Working Paper 15-49.
- Bidadanure, J. U. (2019). The political theory of universal basic income. *Annual Review of Political Science*, 22, 481–501.
- Broach, D. (1998). *Recovery of the FAA Air Traffic Control Specialist Workforce, 1981–1992*. FAA Civil Aerospace Medical Institute, DOT/FAA/AM-98/23.
- Brynjolfsson, E., Chandar, B., & Chen, R. (2025). *Canaries in the coal mine?* Stanford Digital Economy Lab, November revision.
- Brynjolfsson, E., Li, D., & Raymond, L. R. (2025). Generative AI at work. *Quarterly Journal of Economics*, 140(2), 889–942.
- Brynjolfsson, E., Rock, D., & Syverson, C. (2021). The productivity J-curve. *American Economic Journal: Macroeconomics*, 13(1), 333–372.
- Clark, G. (2005, 2007). English real wage and factor share series, 1209–2004.
- Clemens, M., Lewis, E., & Postel, H. (2018). Immigration restrictions as active labor market policy: evidence from the Mexican Bracero exclusion. *American Economic Review*, 108(6), 1468–1487.
- Costinot, A., & Werning, I. Robots, trade, and Luddism. *Review of Economic Studies*.
- Crafts, N. (2022). Slow real wage growth during the Industrial Revolution. *Oxford Economic Papers*, 74(1), 1–13.
- Cui, Z. K., et al. (2025). The effects of generative AI on high-skilled work: evidence from three field experiments with software developers. Working paper.
- Danaher, J. (2019). *Automation and Utopia*. Harvard University Press.
- David, P. A. (1990). The dynamo and the computer. *American Economic Review*, 80(2), 355–361.
- Davis, S. (2026). AI is simultaneously aiding and replacing workers, wage data suggest. Federal Reserve Bank of Dallas, 24 February.
- Davis, S. J., & von Wachter, T. (2011). Recessions and the costs of job loss. *Brookings Papers on Economic Activity*.
- Di Tella, R., & Rodrik, D. (2020). Labour market shocks and the demand for trade protection: evidence from online surveys. *The Economic Journal*, 130(628), 1008–1030.
- Emanuel, N., Harrington, E., & Pallais, A. (2023). The power of proximity to coworkers. Working paper / Federal Reserve Bank of New York.
- Engels, F. (1845/1892). *The Condition of the Working Class in England*; and the preface to the English edition, 11 January 1892.
- Fast, J., & Gould, E. (2026). Class of 2026: young college graduates face a weaker labor market. Economic Policy Institute, 7 May.
- Federal Reserve Board (2026). *Distributional Financial Accounts*, Q1 2026.

- Flaaen, A., & Pierce, J. (2019). Disentangling the effects of the 2018–2019 tariffs on a globally connected US manufacturing sector. Federal Reserve Board FEDS 2019-086.
- Fraser, N. (2017). *Why two Karls are better than one*. Working Paper 1/2017, Jena.
- Gallardo-Albarrán, D., & de Jong, H. (2021). Optimism or pessimism? *European Review of Economic History*, 25(1), 1–19.
- Gauthier, D. (1986). *Morals by Agreement*. Oxford University Press.
- George, H. (1879). *Progress and Poverty*. [Books VI–VIII.]
- Gheaus, A., & Herzog, L. (2016). The goods of work (other than money!). *Journal of Social Philosophy*, 47(1), 70–89.
- Guerreiro, J., Rebelo, S., & Teles, P. (2022). Should robots be taxed? *Review of Economic Studies*, 89(1), 279–311.
- Hampole, M., Papanikolaou, D., Schmidt, L. D. W., & Seegmiller, B. (2025). *Artificial intelligence and the labor market*. NBER Working Paper 33509.
- Heinrich, M. (2013). The "Fragment on Machines": a Marxian misconception. In Bellofiore, Starosta & Thomas (eds.), *In Marx's Laboratory*, 197–211. Brill.
- Hobbes, T. (1651). *Leviathan*. [Chapters XIII–XV, XVII.]
- Hosseini Maasoum, S. M., & Lichtinger, G. (2025). Generative AI as seniority-biased technological change: evidence from US résumé and job posting data. Working paper.
- Hume, D. (1751). *An Enquiry Concerning the Principles of Morals*, §III. And *A Treatise of Human Nature*, III.ii.2.
- Jones, C. I. (2026). *A.I. and our economic future*. NBER Working Paper 34779.
- Jonker, J. D. (2025). Is data labor? *Business Ethics Quarterly*, 35(2), 153–186.
- Jordà, Ò., Knoll, K., Kuvshinov, D., Schularick, M., & Taylor, A. M. (2019). The rate of return on everything. *Quarterly Journal of Economics*, 134(3).
- Karabarbounis, L., & Neiman, B. (2014). The global decline of the labor share. *Quarterly Journal of Economics*, 129(1), 61–103.
- Karger, E., et al. (2026). *Forecasting the economic effects of AI*. NBER Working Paper 35046.
- Keynes, J. M. (1930). Economic possibilities for our grandchildren. In *Essays in Persuasion*, Macmillan, 1931.
- Korinek, A. (2024). *Economic policy challenges for the age of AI*. NBER Working Paper 32980.
- Korinek, A., & Juelfs, M. (2022). *Preparing for the (non-existent?) future of work*. NBER Working Paper 30172.
- Korinek, A., & Suh, D. (2024). *Scenarios for the transition to AGI*. NBER Working Paper 32255.
- Kuziemko, I., Norton, M. I., Saez, E., & Stantcheva, S. (2015). How elastic are preferences for redistribution? *American Economic Review*, 105(4).
- Leontief, W. (1983). National perspective: the definition of problems and opportunities. *Population and Development Review*, 9(3), 403–410.
- Locke, J. (1689). *Two Treatises of Government*. [Second Treatise ch. V, §§27–51; First Treatise §42. Section numbers are Laslett-standard.]
- Macpherson, C. B. (1962). *The Political Theory of Possessive Individualism*. Clarendon Press.
- Marx, K. (1857–58/1973). *Grundrisse*, trans. M. Nicolaus. Penguin. [Fragment on Machines, pp. 690–712.]
- Marx, K. (1867/1887). *Capital, Volume I*, trans. Moore & Aveling. [Chs. 15, 25.]
- METR (2025, 2026). *Measuring the impact of early-2025 AI on experienced open-source developer productivity*, arXiv:2507.09089; *Time Horizon 1.1*; and the February 2026 experiment redesign.
- METR (2026). *Time Horizon 1.1*, 29 January; and *Measuring the impact of early-2025 AI on experienced open-source developer productivity*, February 2026 follow-up.
- Mutz, D. (2021). *Winners and Losers: The Psychology of Foreign Trade*. Princeton University Press.
- New York State Department of Labor (2025–2026). WARN notice AI-attribution disclosure, implemented March 2025; first-year results as reported May 2026.
- Nozick, R. (1974). *Anarchy, State, and Utopia*. Basic Books. [Ch. 7.]
- Nussbaum, M. C. (2006). *Frontiers of Justice*. Harvard University Press.
- Okishio, N. (1961). Technical changes and the rate of profit. *Kobe University Economic Review*, 7, 85–99.
- OpenResearch (2024). *The employment effects of a guaranteed income*. NBER Working Paper 32719.
- Paine, T. (1797). *Agrarian Justice*.

-
- Pasquinelli, M. (2019). On the origins of Marx's general intellect. *Radical Philosophy*, 2.06, 43–56. And (2023) *The Eye of the Master*. Verso.
- Piketty, T. (2014). *Capital in the Twenty-First Century*. Harvard University Press. And (2015) *Journal of Economic Perspectives*, 29(1).
- Polanyi, K. (2001 [1944]). *The Great Transformation*, 2nd Beacon ed.
- Polanyi, M. (1966). *The Tacit Dimension*. Doubleday.
- Rawls, J. (1988). The priority of right and ideas of the good. *Philosophy & Public Affairs*, 17(4), 251–276. [The 'Malibu surfers' remark is at p. 257 n.7.] And (2001) *Justice as Fairness: A Restatement*. Harvard University Press.
- Rodrik, D. (2021). Why does globalization fuel populism? *Annual Review of Economics*, 13, 133–170.
- Rognlie, M. (2015). Deciphering the fall and rise in the net capital share. *Brookings Papers on Economic Activity*, Spring.
- Ross, M. L. (2001, 2015). Does oil hinder democracy? and subsequent revisions. *World Politics; Annual Review of Political Science*.
- Rousseau, J.-J. (1755). *Discourse on the Origin of Inequality*, Part Two.
- Saez, E., & Zucman, G. (2016, 2020); and Smith, M., Zidar, O., & Zwick, E. (2023). Competing estimates of US top wealth shares.
- Scheidel, W. (2017). *The Great Leveler*. Princeton University Press.
- Sreenivasan, G. (1995). *The Limits of Lockean Rights in Property*. Oxford University Press.
- Trammell, P., & Korinek, A. (2023, rev. 2026). *Economic growth under transformative AI*. NBER Working Paper 31815.
- Tully, J. (1980). *A Discourse on Property*. Cambridge University Press.
- US Bureau of Labor Statistics (2026). *Union members — 2025*; CPS Table 53; *Employment Situation*, June 2026.
- US Census Bureau (2026). *Firm use of artificial intelligence*. Center for Economic Studies Working Paper 26-25.
- Van Parijs, P. (1995). *Real Freedom for All*. Oxford University Press. [Ch. 4, 'Jobs as Assets'.]
- Waldron, J. (1979). Enough and as good left for others. *The Philosophical Quarterly*, 29(117), 319–328.
- World Inequality Database (2026). Wealth share series.
- Yotzov, I., et al. (2026). *Firm data on AI*. NBER Working Paper 34836.