

NBER WORKING PAPER SERIES

HOW RETRAINABLE ARE AI-EXPOSED WORKERS?

Benjamin G. Hyman
Benjamin Lahey
Karen Ni
Laura Pilossoph

Working Paper 34174
<http://www.nber.org/papers/w34174>

NATIONAL BUREAU OF ECONOMIC RESEARCH
1050 Massachusetts Avenue
Cambridge, MA 02138
August 2025, Revised May 2026

We thank Joe Altonji, David Card, Anders Humlum, Lisa Kahn, Fabian Lange, Michael Lee, Kyle Myers, Steve Raphael, Daniel Rock, Wilbert van der Klauuw, Till von Wachter, and seminar participants at the annual ASSA/AEA and SOLE conferences, Federal Reserve Banks of San Francisco and New York, MIT FutureTech, and UC Berkeley IRLE for helpful comments on an earlier version of the paper. We also thank Jonathan Lee for excellent research assistance. This research received support through Schmidt Sciences. Ni gratefully acknowledges support from the Institution of Education Sciences, US Department of Education, through grant R305B150012 to Harvard University. Any opinions, findings, and conclusions or recommendations expressed in this material are those of the authors and do not necessarily reflect the views of the Federal Reserve Bank of New York, the Federal Reserve System, or the National Bureau of Economic Research.

At least one co-author has disclosed additional relationships of potential relevance for this research. Further information is available online at <http://www.nber.org/papers/w34174>

NBER working papers are circulated for discussion and comment purposes. They have not been peer-reviewed or been subject to the review by the NBER Board of Directors that accompanies official NBER publications.

© 2025 by Benjamin G. Hyman, Benjamin Lahey, Karen Ni, and Laura Pilossoph. All rights reserved. Short sections of text, not to exceed two paragraphs, may be quoted without explicit permission provided that full credit, including © notice, is given to the source.

How Retractable are AI-Exposed Workers?

Benjamin G. Hyman, Benjamin Lahey, Karen Ni, and Laura Pilossoph

NBER Working Paper No. 34174

August 2025, Revised May 2026

JEL No. E0, E2, J6

ABSTRACT

As artificial intelligence (AI) capabilities advance, will workers best adapt by reskilling into AI-complementary work or by sorting into occupations less exposed to AI? To answer this question, we assemble a new dataset of 1.9 million occupational training spells funded by the U.S. Workforce Innovation and Opportunity Act from 2012–2024. We link pre- and post-training occupations to task-level AI exposure measures and estimate returns to training by comparing trainees to matched workers who sought workforce services but received only job search assistance. While trainees from low AI-exposure occupations earned high quarterly returns throughout the sample period, returns for workers from high-exposure occupations rose sharply—from about \$900 quarterly before 2020 to \$2,900 by 2022–2024. We attribute these gains primarily to transitions into less AI-exposed occupations and, to a lesser extent, to the expansion of training programs that build AI-complementary skills. To quantify when training into higher AI exposure work pays off, we construct a new AI Retraining Index (AIR) and find that a large share of occupations are “AI-retrainable,” pointing to broad potential for adaptation.

Benjamin G. Hyman
University of California, Los Angeles
benhym@gmail.com

Benjamin Lahey
New York University
bpl9631@nyu.edu

Karen Ni
Harvard University
Harvard Kennedy School
kni@g.harvard.edu

Laura Pilossoph
Duke University
and NBER
pilossoph@gmail.com

1 Introduction

The debate over whether advances in artificial intelligence (AI) will ultimately complement or substitute labor has drawn significant attention (Autor (2024); Deming et al. (2025); Humpole et al. (2025)). However, there is a dearth of research examining the role that existing job training programs might play in helping AI-exposed workers adapt to the evolving labor market. This gap is particularly striking in light of survey evidence showing that AI-adopting firms are retraining their workforces more rapidly than they are adjusting headcounts in response to AI technologies (Abel et al., 2025).¹

We study how workers adapt to AI-related pressure using detailed data from the Workforce Innovation and Opportunity Act (WIOA), formerly the Workforce Investment Act (WIA), the US’s flagship workforce development program. WIOA is a useful setting because it captures realized rather than hypothetical occupational transitions around reskilling, allowing us to evaluate the extent to which it is well suited to reskill workers in occupations facing AI-related pressure. Within this context, we assemble a novel dataset of over 1.9 million individual WIOA/WIA training spells from 2012 to 2024, spanning the transition into the modern AI era, and link them to administrative earnings records before and after training to estimate program returns. The earlier years in our data allow us to compare the training experiences of workers during the more recent generative AI boom to workers in the same occupations before frontier large language models (LLMs) were a salient force in the economy, providing a benchmark for how retraining outcomes have changed as AI has emerged. Crucially, the data include detailed occupation codes for workers’ jobs before and after training. These occupation codes allow us to track two leading measures of AI exposure throughout each worker’s job training spell: the *potential* AI exposure of tasks capturing what LLMs could do if the technology were adopted widely and without frictions (Eloundou et al. (2024)), as well as *realized* occupational AI exposure from observed task usage with Anthropic’s LLM, Claude (Handa et al. (2025)).

We use these measures to estimate heterogeneous earnings returns to training by initial AI exposure, and to test whether transitions in AI skill space mediate the earnings returns to training. In particular, we use the earnings returns to training to develop a test for whether transitions to more AI-exposed occupations reflect “AI upskilling.” Positive returns alongside movements to more exposed occupations are consistent with adjustment through building AI-complementary skills, while low or negative returns may indicate ineffective training or re-exposure to AI pressures. Alternatively, trainees may shift into less exposed occupations

¹In a large-scale survey of Danish workers most exposed to ChatGPT, Humlum and Vestergaard (2024) find that 43% of workers reported the need for more training as the largest barrier to adoption.

to avoid the shock altogether.

We find that workers initially displaced from low AI exposure occupations capture much larger returns to training than workers displaced from high AI exposure occupations. Across the sample period (2012-2024), low AI-exposure workers consistently garnered about \$2,200 in quarterly returns to training on average. However, we uncover a dramatic improvement in the returns to training for workers in high AI-exposure occupations over time: these workers' returns were relatively modest before 2020 (at about \$900 a quarter) yet grew to rival the returns captured by low-exposure workers during the AI boom.

In an ideal experiment, these returns would be identified by random assignment. Absent random assignment, we follow the approaches of [Rothstein et al. \(2022\)](#) and [Andersson et al. \(2024\)](#), matching each WIOA/WIA trainee to a similarly AI-exposed nearest-neighbor control worker who also sought out workforce services, but only received job search assistance. The design's credibility is strengthened by two features of the training programs. First, individuals are quasi-randomly assigned to WIOA caseworkers who gatekeep entry into training based on subjective interpretations of eligibility rules. While caseworkers are unobservable in the data, institutional details confirm that part of assignment into training is quasi-random. Second, because both training and control group units experience a similar decline in earnings prior to receiving services, the strategy also addresses the potential "Ashenfelter dip" concern that a decline in earnings prior to training could upwardly bias estimates of post-training earnings gains ([Andersson et al. \(2024\)](#)). The nearest neighbor approach is also well suited to separately examine outcomes for workers displaced from occupations with varying initial AI exposures. Though not without limitations, this selection-on-observables approach shines early light on the understudied intersection of AI and job training.

We document four main empirical findings. First, AI exposure among job training participants is strongly U-shaped, with clustering at both tails of the spectrum and a meaningful share coming from highly exposed occupations. This large degree of AI exposure is not unique to training participants, as their mean AI exposure closely tracks the AI exposure of respondents in the CPS Unemployed and Displaced Worker samples. Second, the average earnings returns to training were substantial from 2012-2024 — approximately \$900 - \$2,200 per quarter (\$USD 2010) for workers from high- and low-exposure occupations, respectively. Third, returns differ systematically by the direction of skill adjustment: workers from low-exposure occupations experience the largest gains when they "AI upskill" into more exposed work, while workers from high-exposure occupations benefit from transitioning toward less AI-exposed work. Finally, we decompose our training returns into occupational sorting effects and training returns holding sorting

constant, and find that occupational sorting became a stronger determinant of the earnings returns to training in the AI era (2022-2024). In particular, we show that workers in occupations thought to be highly AI-exposed today had limited returns to training before 2020, but much larger returns from 2022-2024.

We argue that this recent improvement for high-exposure workers is driven in part by increased sorting into low-exposure occupations, such as healthcare, and to a lesser extent by more successful upskilling to higher-exposure occupations. We show that the latter trend aligns with explicit efforts by WIOA governing bodies and training providers to equip trainees with AI-resilient skills. This period also unfolded alongside historically tight labor markets, raising the possibility that training credentials carry greater signaling value when firms are hiring more aggressively and reaching deeper into the skill distribution.

Turning our focus to workers who seem particularly adjustable to AI, we use the matched sample to construct an AI Retraining Index (AIR) which ranks occupations by the share of workers who successfully retrain into higher earnings occupations with more task-level exposure to AI. Using the index, we estimate that about 40%-45% of occupations are “AI retrainable” as measured by its workers receiving higher pay for moving to more AI-exposed occupations—a large magnitude given a relatively low-income sample of displaced workers. The index also enables us to decompose, for any given occupation, the extent to which successful AI retraining is driven by earnings gains while holding skills constant (e.g., gains from occupational licensing barriers) versus skill transitions from AI upskilling.

Our findings contribute to both the literature on the labor market effects of AI and the active labor market program evaluation literature. While personnel studies have found complementarities between AI adoption and worker productivity within individual firms, these studies generally do not attempt to isolate the effects of AI-upskilling incumbent workers from the outcomes of AI adoption (e.g., [Brynjolfsson et al. \(2025b\)](#) on ChatBots aiding less-experienced workers). They are also silent on the question of whether workers would be better off avoiding AI-exposed occupations altogether. Recent evidence also highlights the role of individual beliefs and barriers to reskilling: [Delfino et al. \(2026\)](#) show that jobseekers’ perceptions of monetary returns and occupational fit shape their demand for retraining into high-demand jobs. To our knowledge, our paper is the first to directly estimate the effects of AI retraining on worker earnings, and to do so at a nationally-representative scale.²

²In related work, [Manning and Aguirre \(2026\)](#) construct an occupation-level “adaptive capacity” index to characterize how well positioned workers are to manage a job transition if AI exposure ultimately results in displacement. The goal of their analysis is descriptive, and it provides valuable evidence on the association between occupational AI exposure and adaptive capacity, highlighting occupations with potential vulnerability. However, it does not estimate the causal effects of AI-induced displacement or retraining, nor

The implication that the labor demand environment—both evolving AI skill requirements and tight labor markets—is an important driver of positive returns also contributes to the literature on active labor market policy evaluations, especially studies focused on heterogeneity over the business cycle. In a meta-analysis of more than 200 evaluations, [Card et al. \(2018\)](#) provide suggestive evidence that active labor market policies are more effective in weak labor markets, consistent with the view that employers can be more selective in slack labor markets. We find evidence in favor of the opposing viewpoint that job training is less effective during recessions. Our results better align with the alternative explanation that training provides a particularly valuable signal when labor markets are tighter and firms have to recruit lower down in the skill distribution.

A secondary contribution of our work is that we document that modern AI exposure measures (e.g. [Eloundou et al. \(2024\)](#), [Handa et al. \(2025\)](#)) have a natural mapping to task content measures from the preceding computerization and automation literature ([Acemoglu and Autor \(2011\)](#)). Specifically, AI exposure measures align most strongly with routine cognitive task content, but also increasingly with non-routine cognitive tasks. This mapping is helpful because it underscores that estimates of the returns to job training by time-invariant AI exposure measures will reflect cognitive skill retraining in earlier years (when computerization and automation were rising) and increasingly non-routine cognitive skill retraining in later years (more aligned with generative AI). This broader alignment across both routine and non-routine cognitive tasks also suggests that AI exposure may not map neatly onto middle-skill hollowing documented in the prior job polarization literature ([Autor and Dorn \(2013\)](#)), but instead span a wider portion of the cognitive wage distribution. As the next phase in the long-term reallocation of cognitive work, understanding earlier evidence on cognitive skill retraining is informative for adaptation in the as yet uncertain AI era.

The paper proceeds as follows. [Section 2](#) provides background information and descriptive statistics on the WIOA/WIA training programs and defines AI exposure measures. [Section 3](#) details our nearest-neighbor matching identification strategy. [Section 4](#) presents our main results. [Section 5](#) outlines our AI retrainability (AIR) index which allows us to decompose the sources of successful AI retraining. [Section 6](#) concludes and discusses areas for future research.

does it speak to whether workers respond to AI exposure by reskilling within AI-intensive work versus moving away from it. In contrast, we directly estimate the earnings consequences of AI-related retraining transitions using a matched quasi-experimental design. Contemporaneously, [Jacobs and Canedy \(2026\)](#) produce a retrainability index that uses WIOA industry transitions away from high routine task content to quantify automation risk. Like [Manning and Aguirre \(2026\)](#), their index is also descriptive and does not use causal estimates on the returns to training as inputs as we do here.

2 Data and Measurement

Determining whether individuals can be retrained for AI-related work is challenging because analysts typically lack administrative data on how AI is used by workers within firms.³ We make progress by developing a dataset that can longitudinally track transitions of job training participants from and to occupations that vary in their susceptibility to AI; occupations are then categorized as AI-exposed using measures we describe below. We construct this new large-scale dataset of training participants by merging two publicly available data sources.

2.1 WIOA / WIA Participant Data

The US Department of Labor’s (USDOL) Participant Individual Record Layout (PIRL) is a national performance tracking dataset that contains one observation for each unique participation spell in any WIA, or its successor WIOA, program from 2005q1 to 2024q4.⁴ WIOA/WIA is a federally funded active labor market program in which qualified workers receive fully subsidized occupational job training (*Title I - “Workforce Development Activities”*) or job search assistance services (*Title III - Wagner-Peyser Act Employment Services*) among other services.⁵ We restrict our analysis to participants who exited WIOA/WIA between 2012q1 and 2024q4, as occupation codes are available throughout all stages of the job training spell during this period.

Our main focus is on Title I occupational training programs which span over 1.88 million training spells from 2012 through 2024. In some analyses we split the sample to compare training effects before the pandemic (2012-2019) and since the beginning of the generative AI boom (2022-2024). Overall, 58% of Title I participants are enrolled in the *Title I - Adult* program which prioritizes current public service recipients (e.g. Supplemental Nutrition Assistance Program (SNAP) recipients), lower-income individuals, and disadvantaged individuals with employment barriers. 34% of Title I participants are enrolled in the *Title I - Dislocated Worker* program which serves individuals who have been laid off, received a notice of termination, or are eligible for USDOL rapid response services directed to workers impacted by mass layoffs or plant closures. Unlike Adult

³While recent surveys provide firm AI usage rates (Bonney et al. (2024)) and some estimates of worker-level adoption across the income distribution (Hartley et al. (2024)) and by occupation (Humlum and Vestergaard (2024)), we are aware of only one study that has linked AI use to worker administrative data (Pizzinelli et al. (2023)) though it does not include individual-level job training data.

⁴WIA data adheres to the PIRL infrastructure, while its successor WIOA uses the updated [Workforce Investment Act Standardized Record Data](#) (WIASRD) files. Files are currently updated through 2024, however not enough time has yet elapsed to observe outcomes for all 2024 participants.

⁵Popular examples of training services are nursing degree programs and courses in data analysis with Microsoft Office and Python. See Andersson et al. (2024) for an involved description of available services.

program trainees who can be part-time employed while training, Dislocated Worker program trainees are generally enrolled full-time in training. Of the remaining Title I participants, less than 1% are in the program for young adults aged 14-24, while 8% are co-enrolled across the above programs.

To establish a control group of similarly AI-exposed workers, expanding on [Rothstein et al. \(2022\)](#) and [Andersson et al. \(2024\)](#), we use matched participants in Title III programs who receive job search assistance services through the Wagner-Peyser Act, but do not enroll in occupational retraining programs.⁶ Title III programs span over 31 million participants from 2012 to 2024 providing a rich set of characteristics from which to match Title I trainees. While both Title I and Title III workers can receive job search assistance services upon intake at American Job Centers, individuals must be evaluated by caseworkers to enroll in Title I training services (29 U.S.C. §3174(c)(3)(A)(i)).⁷ Because individuals are quasi-randomly assigned to caseworkers who may vary in how they interpret ambiguity in training eligibility rules (see [Appendix D](#)), the resulting variation in training assignment reduces concerns regarding selection into training programs based on unobservables (see [Appendix E](#)). To account for any lingering selection bias from endogenous sorting into Title I and Title III programs, we use the matched sample strategy described in detail in [Section 3](#).

We observe individual quarterly earnings data from state unemployment insurance administrative records (ES-202 forms) for three quarters before each participant enters a Title I or Title III program, and for three quarters after exiting the program. We deflate all earnings records to 2010 dollars using the consumer price index. Importantly, for a sizable subset of individuals we observe the quarterly 6-digit Standard Occupation Code (SOC) associated with each worker’s primary job if employed.⁸ One limitation of the PIRL dataset is that the requirement to report occupation codes is left to the discretion of local workforce development boards. After restricting to participants with complete occupation codes (available for 22% of training participation spells), earnings records, and covariate information, and dropping observations in the top 1% of pre-WIA/WIOA earnings, our final sample contains 119,463 total training spells, of which 113,935 spells (95.6%) are successfully matched to a nearest neighbor control group unit. In [Section 2.3](#), we discuss balance tests that confirm that individuals with missing occupation codes are qualitatively similar to those with observed occupation codes, consistent with arbitrarily distributed

⁶Wagner-Peyser Act services include basic and core services—job search assistance, referrals, labor market information and resume-building workshops—but can also span intensive career services like counseling.

⁷Because the case worker is unobserved, we cannot use an examiner design to estimate effects as in [Humlum et al. \(2023\)](#).

⁸All 6-digit SOC codes in our data are derived from 2010 detailed O*NET (Occupational Information Network) occupational codes. SOC codes represent the first 6 digits of the 8-digit O*NET codes.

reporting requirements for occupation codes.

Finally, for each individual, the PIRL also provides a rich set of demographic variables including sex, race, ethnicity, highest level of educational attainment, disability status, and social benefit distinctions including low-income status, Temporary Assistance for Needy Families (TANF) recipient, SNAP recipient, or other public assistance recipient, which are helpful for both description and matching.

2.2 Occupation-Level AI Exposure Measures

Firm- and individual-level AI use are rarely observed at scale in administrative data. This measurement constraint is well recognized in the digital economics literature, and a standard response is to proxy for AI impact with “AI exposure” measures using occupation-level indices for whether an occupation’s work tasks can be, or have been, accomplished by AI. To this end, we merge two leading occupational measures of AI exposure to WIOA/WIA participants’ pre-training occupations to study AI-related transitions.

Eloundou et al. (2024) – potential AI exposure. The first measure, from [Eloundou et al. \(2024\)](#), builds on the seminal “suitability for machine learning” (SML) framework of [Brynjolfsson et al. \(2018\)](#). In the original paper, human annotators evaluated 2,069 O*NET work activities (“tasks”) across 964 six-digit SOC occupations, rating each task–occupation pair on its suitability for machine learning.⁹ [Eloundou et al. \(2024\)](#) extend this approach to large language models (LLMs), asking whether generative AI systems can substantially reduce the time required to perform specific tasks. Annotators—and ChatGPT itself—assess whether (a) an LLM alone could reduce task completion time by at least 50%, and (b) additional software built on top of the LLM could achieve a 50% time reduction. We use the widely cited $AI_{Eloundou}^{\beta}$ option among their measures, which positively correlates with occupational earnings.¹⁰ A key advancement in [Eloundou et al. \(2024\)](#) is coverage: rather than the 2,069 activities used in the original SML work, they evaluate all 18,156 granular O*NET tasks, many of which are occupation-specific. As in [Brynjolfsson et al. \(2018\)](#), task-level exposure is aggregated to the occupation level using O*NET task-importance weights, which are derived from surveys of incumbent workers and occupational experts who report how critical each task is to performance in that occupation.

⁹Seven independent annotators scored each task–occupation pair on 23 Likert-scale questions capturing the likelihood that a machine could perform the task in that occupational context. A score of 5 indicates high suitability; 1 indicates low suitability.

¹⁰We use [Eloundou et al. \(2024\)](#)’s β weighting scheme, which defines AI exposure $AI_{Eloundou}^{\beta} \in [0, 1] = s_a + 0.5s_b$ where $s_a, s_b \in [0, 1]$ denote the shares of an occupation’s tasks for which conditions (a) and (b) hold. This measure captures the share of an occupation’s tasks that are LLM-exposed, and is validated by its positive correlation with average occupational earnings.

Among the unique occupations spanned by the WIOA/WIA dataset, the mean and median AI exposure scores (prior to program participation) are 0.338 and 0.343 with a standard deviation of 0.221. The distribution of AI exposure scores across all 6-digit occupations in the [Eloundou et al. \(2024\)](#) dataset (not conditioning on training participation) is very similar, having a mean of 0.326, a median of 0.329, and a standard deviation of 0.220. Consistent with [Brynjolfsson et al. \(2025a\)](#), customer service representatives—one of the most highly AI-exposed occupations—are also heavily represented in job training programs, while the lowest exposure quintile is dominated by manual occupations that are comparatively insulated from AI. [Table 1](#) illustrates this pattern by listing the three most common pre-training occupations within each AI exposure quintile among WIOA/WIA participants over the full sample from 2012–2024.

Table 1: Top Occupations Prior to Training by AI Exposure

Rank 6-Digit SOC Occupation			No. Trainees
<i>AI Exposure Quintile 5 (Highest Exposure)</i>			
1	434051	Customer service representatives	4,027
2	439061	Office clerks, general	1,390
3	434171	Receptionists and information clerks	850
<i>AI Exposure Quintile 4</i>			
1	435081	Stockers and order fillers	1,325
2	119199	Managers, all other	1,253
3	111021	General and operations managers	1,085
<i>AI Exposure Quintile 3</i>			
1	533032	Heavy and tractor-trailer truck drivers	2,078
2	292061	Licensed practical and licensed vocational nurses	1,535
3	412031	Retail salespersons	1,520
<i>AI Exposure Quintile 2</i>			
1	311014	Nursing assistants	5,393
2	412011	Cashiers	2,123
3	311011	Home health aides and personal care aides	1,234
<i>AI Exposure Quintile 1 (Lowest Exposure)</i>			
1	519199	Production workers, all other	2,635
2	537062	Laborers and freight, stock, and material movers, hand	2,093
3	519198	Helpers – production workers	1,598

Notes: This table shows the top three most prominent occupations in different AI exposure quintiles prior to entering WIOA/WIA training participation. Counts of number of training participants are shown in the final column. Table uses AI exposure measures from [Eloundou et al. \(2024\)](#). See text for details.

Anthropic Economic Index (2025) – realized AI exposure. A limitation of the Eloundou et al. (2024) exposure measure is that it may not reflect how AI is being used in practice since it is a measure of what tasks LLMs could *potentially* complete. To address this, we also draw on *realized* usage data from Anthropic’s Economic Index (Handa et al. (2025)). This index (which we refer to as AI_{Claude}) estimates task-level LLM usage using a random sample of one million Claude conversations in late 2024 and early 2025. Its key advantage is that it reflects observed behavior rather than hypothetical task feasibility.¹¹ To map task-level Claude usage to occupations, we follow Anthropic’s preferred allocation procedure, dividing total usage on a task by the number of occupations that perform that task. In practice, most tasks are occupation-specific, so this adjustment affects only a minority of cases. To ensure comparability with Eloundou et al. (2024), we further weight tasks by their O*NET task-importance scores when aggregating to the occupation level.¹²

AI Exposure and Cognitive Skills

Occupational measures of AI exposure capture the extent to which tasks performed in an occupation could potentially be carried out by modern AI systems (Eloundou et al. (2024)), or are already observed being performed in practice (Handa et al. (2025)). Because the WIOA job training data span 2012–2024—a period that largely predates the recent generative AI surge—a natural question is what these exposure measures capture prior to the current AI wave. A novel secondary contribution of our paper is to show that two prominent measures used in the AI literature ($AI_{Eloundou}^{\beta}$ and AI_{Claude}) are strongly related to foundational cognitive task measures emphasized in the skill-biased technological change (SBTC) task framework (Acemoglu and Autor (2011)).

In Appendix Figure A.1, binscatter plots illustrate this relationship. Both AI exposure measures load most strongly on routine cognitive task content, while also increasing with non-routine cognitive activities. This mapping implies that estimates of returns to training by time-invariant AI exposure partly capture earlier episodes of cognitive skill retraining—when computerization and automation were rising—and increasingly reflect retraining in non-routine cognitive tasks that are more closely associated with modern AI capabilities. For this reason, at times we split the sample into the pre-pandemic period and the period marked by the beginning of the generative AI boom in 2022. More broadly, the fact that AI exposure aligns with both routine and non-routine cognitive work suggests that it may not map cleanly to the middle-skill

¹¹A limitation is that Anthropic’s user base may be disproportionately concentrated in certain industries, such as software development, which may not be fully representative of the national workforce.

¹²Analogous to Table 1, Appendix Table A.10 reports the three most prominent occupations among training participants within each quintile of Claude-based AI exposure.

hollowing effects emphasized in the job polarization literature ([Autor and Dorn \(2013\)](#)), but instead spans a wider portion of the cognitive wage distribution. Viewed through this lens, AI can be interpreted as the next phase in a longer process of reallocating cognitive work, making historical evidence on cognitive skill retraining informative for understanding worker adaptation in the emerging AI era.

2.3 Descriptive Statistics and Sample Restrictions

We split the sample of training participants into a “high” and “low” AI exposure group based on whether their pre-enrollment occupation was above or below the 55th percentile of the 6-digit occupation AI exposure score distribution. As we discuss in [Section 3](#), this cutoff is determined empirically based on the stability of estimates across potential cutoffs, however all estimates are robust to shifting the cutoff in either direction.¹³ [Appendix Table A.1](#) presents means and standard deviations of key worker characteristics at the time of training enrollment for these two groups.

Within the 113,935 pooled training spells in the final matched sample, the distribution is skewed toward workers with low pre-training AI exposure, comprising approximately 62% of spells, compared with 38% for their high-exposure counterparts. The average worker pooled across both samples made around \$40,000 in annualized 2010 earnings prior to training participation. In [Appendix Figure A.2](#), we also show that these workers had pre-separation AI exposure distributions that are almost identical to the distributions among the Current Population Survey (CPS) of unemployed workers when comparing occupations prior to separation in repeat cross sections over time. This suggests that WIOA/WIA training participants are representative of the national population of unemployed workers.¹⁴ [Appendix Table A.1](#) also reveals that the high AI exposure group has a larger female share (consistent with [Bollinger and Troske \(2025\)](#) who document greater success among females in coding-focused training programs), is more educated, and has slightly higher earnings prior to participation. The high AI exposure group is also more skewed toward participants in the dislocated worker program versus the adult (disadvantaged) worker program.

Sample Restrictions. In our summary statistics table ([Appendix Table A.1](#)) and throughout the analysis, we restrict to workers who have: (1) positive earnings in all six

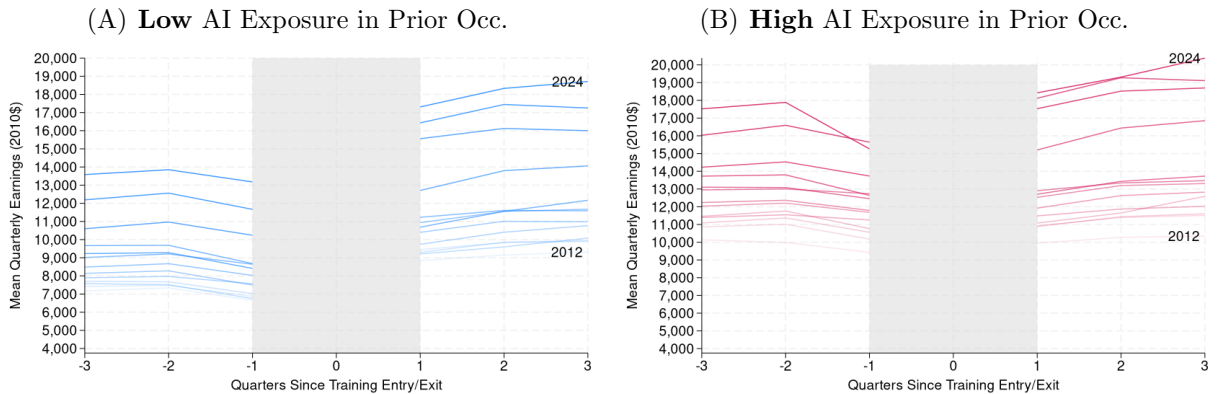
¹³We assign each training spell a single prior- and post-training occupation based on the SOC code associated with the plurality of earnings in the three quarters before or after training. For occupation codes that are self-reported at more aggregate levels, we use the average of all 6-digit codes within that occupation group.

¹⁴Occupations spanned by our full set of WIOA/WIA training participants account for around 93% of the CPS labor force on a monthly basis.

quarters of observation, (2) non-missing occupation codes in all six quarters of observation, and (3) non-missing covariates used for matching. The first restriction intentionally limits the sample to reemployed workers, allowing us to measure earnings responses to changes in the AI content of occupations insulated from any confounding extensive margin reemployment effects. Relaxing this restriction results in near-identical earnings returns to WIOA/WIA (see [Appendix Figure A.5](#)). The second restriction ensures we can measure flows in AI skill space, however, results in dropping a significant share of training participants. This is because the reporting of occupation codes varies by state and may be missing “due to language in the WIOA statute and a variety of practical/quality reasons” or if the “occupation code is not a part of the UI wage record in most states” (correspondence with USDOL Employment and Training Administration, November 2024). In [Appendix Table A.2](#), we show that compositional differences between workers in states with and without recorded occupation codes are mostly negligible. Ultimately, our matching estimator will compare training participants with nearest-neighbor matched workers within occupation bins while matching on state location.

Earnings Selection Patterns. Before turning to our matching design, [Figure 1](#) presents earnings patterns around training participation spells by calendar year of program exit subject to the above sampling restrictions. Panels (A) and (B) show real quarterly earnings for the three quarters before training program entry and three quarters after training program exit, separately by low and high AI exposure in the worker’s occupation prior to participation based on the [Eloundou et al. \(2024\)](#) occupational measure.

Figure 1: Earnings Patterns by AI Exposure Prior to Job Training



Notes: Training participants are classified as “high” AI exposure if their 6-digit SOC exposure score is above the 55th percentile of the [Eloundou et al. \(2024\)](#) 6-digit AI exposure score distribution. Heavier lines correspond to more recent years. The shaded gray region emphasizes that quarterly earnings are observed in the quarters prior to entry, and after exit, but not during training. All earnings are in real 2010 dollars.

The primary selection concern in the active labor market literature is the “Ashenfelter

dip”: earnings declines prior to training that can mechanically upward bias comparisons with post-training earnings (Ashenfelter, 1978). The event-study patterns show a modest dip in earnings immediately before training, followed by weakly higher earnings in most years and substantially larger post-training differences in more recent cohorts (shown with heavier lines), when AI technologies were most prevalent.¹⁵ Participants from highly AI-exposed occupations begin with higher baseline earnings but display the same overall time pattern—larger post-training earnings differences in later years. The data do indicate slightly lower post-training earnings among highly exposed groups, but this difference is small relative to the dominant time pattern.

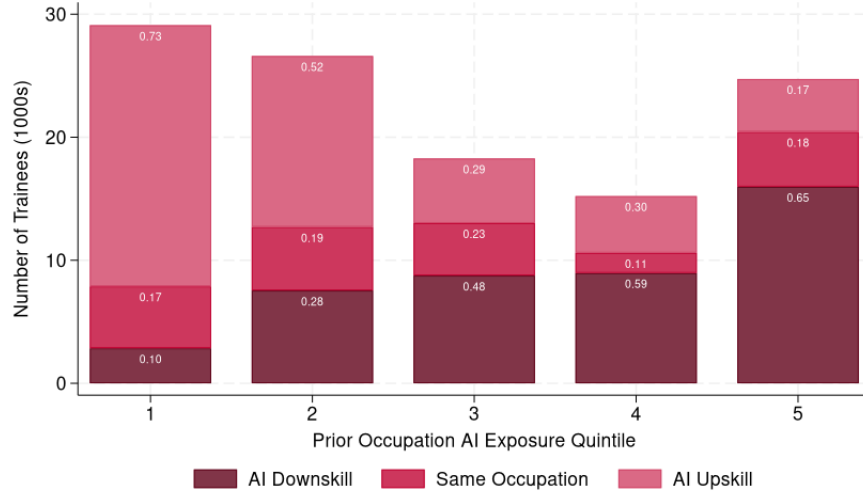
Our main strategy to handle the Ashenfelter dip is to compare training participants with nearest-neighbor matched program entrants that only receive job search assistance, potentially for idiosyncratic reasons related to caseworker assignments (see Appendix D). Indeed, when we plot the earnings patterns for matched workers who receive only job search assistance, they exhibit an almost identical pre-enrollment earnings dip (see Appendix Figure A.8). This indicates that the dip reflects shared job instability upon any service entry rather than differential selection into training. These patterns suggest that Ashenfelter-type selection is unlikely to be a significant driver of the observed post-training earnings returns we find using the nearest-neighbor matching strategy described in Section 3 below.

Job Flows Across AI Skill Space. Before jointly analyzing the earnings returns to training coupled with transitions in AI skills across occupations, Figure 2 first describes the flows of job trainees in AI skill space. On the x-axis, we split trainees in our sample into five quintiles of AI exposure prior to training participation based on the Eloundou et al. (2024) AI exposure of occupations. We use the term “AI skill space” synonymously with AI exposure.¹⁶

¹⁵The stronger patterns in later years may also reflect tighter labor markets, in which training credentials carry greater signaling value when firms must reach further down the skill distribution.

¹⁶This interpretation is supported by the fact that workers’ post-training occupational exposure is closely aligned with their targeted occupation’s exposure, implying that movements in AI exposure reflect intentional reskilling directions rather than unexpected reexposure to AI. Targeted occupations may differ from realized destination occupations, but in practice they mostly overlap. Within the main analysis sample of the paper, 43% of trainees have the same targeted training and destination occupations at the detailed occupation (6-digit) level of SOC specificity. 64% of trainees have the same targeted training and destination occupations at the major group (2-digit) level of SOC specificity. We also show in Appendix Figure A.3 that when workers transition into occupations different from their target, the resulting differences in AI exposure remain small.

Figure 2: Transitions in AI Skill Space Before and After Training



Notes: AI exposure is based on Eloundou et al. (2024) occupation-level AI exposure scores. The flows characterized in the figure correspond to the main analysis sample of the paper, representing transitions from pre-WIOA (prior to job loss) to post-WIOA (realized) occupations.

Of note, there are large shares of both high- and low-exposure workers among job training participants.¹⁷ A substantial share of workers are displaced from top quintile occupations in AI exposure prior to training enrollment. This underscores that in the context of workers interfacing with public job training programs, WIOA/WIA trainees are highly exposed to AI forces despite their lower income and education levels. Consistent with this finding, prominent occupations in the top exposure quintile include trainees who were formerly customer service representatives, cashiers, and office clerks (as shown in Table 1).

Further decomposition of the bars into the AI skills associated with destination occupations reveals that initially low AI exposed workers disproportionately upskill to higher AI content occupations, while high AI exposed workers move down in AI exposure.¹⁸ Although part of this pattern reflects mechanical floors and ceilings of the measure, it also mirrors the structure of task demand: AI appears to reshape occupational sorting by steering workers toward mid-level cognitive work. Because AI exposure is strongly correlated with earnings, these flows imply movement toward the middle of the skill and wage distribution, consistent with AI having a leveling effect on inequality as discussed in Autor (2024). Overall, these flows point to substantial mobility across occupations after training. In subsequent sections, we disentangle the extent to which this

¹⁷Hartley et al. (2024) find a similar bimodal pattern among users of generative AI tools.

¹⁸In Appendix Figure A.4, we show that while the share of workers flowing to higher AI occupations has remained stable over time, those making the transition have received growing earnings returns over time.

reflects successful reskilling or alternative explanations.

3 Empirical Methodology

While positive earnings differentials for training participants and a high degree of AI reskilling are consistent with successful AI retraining, WIOA/WIA workers may be positively selected, either by virtue of the Ashenfelter dip which depresses earnings prior to training participation, or due to selection into occupational transitions. To address these concerns, we follow Rothstein et al. (2022) and Andersson et al. (2024) who use a matching design to compare WIOA/WIA trainees with Wagner-Peyser job search assistance recipients that are observationally similar, but do not enroll in occupational retraining programs. The key assumption is that, in the absence of training, both earnings and occupational trajectories for trainees would have evolved similarly to those of the matched job-search-only group.

The credibility of this design is strengthened by the fact that, although all participants can receive job search assistance, entry into training is mediated by caseworkers, introducing a quasi-random element to the assignment into training (see details in Appendix D and Appendix E). To control for any lingering selection concerns, our control group of nearest-neighbor matched workers who only receive job search assistance further restricts the donor pool of control workers to begin in occupations with similar AI exposure and for each nearest-neighbor pair to exit their programs in the same calendar year—the latter of which controls for the labor demand environment.¹⁹

Because our goal is to attain heterogeneous estimates across AI exposure subgroups, we opt for a simple nearest-neighbor matching strategy that matches each training participant in our sample to a unique Wagner-Peyser participant, and then pools estimates across nearest-neighbor pairs within AI exposure subgroups.²⁰ We use the semiparametric matching estimator described in Abadie and Imbens (2002), where we desire a treatment effect for each treated trainee n , $\theta_n = Y_n(1) - Y_n(0)$, but only one potential outcome is observable for each unit. We implement the algorithm of Abadie et al. (2004), using a Mahalanobis distance statistic to normalize covariates on different scales and generate a single similarity score between candidate matches. This score is then used as an inverse weight in regressions to give greater influence to closer distances.²¹

¹⁹All results are robust to using the year of program entry instead of the year of program exit.

²⁰We show in a potential outcomes framework that heterogeneous comparisons across groups help difference out some of the remaining selection bias from the matching procedure in any given group (see Appendix F).

²¹The estimator also has the useful property that violators of a geometric version of the “covariate overlap” assumption can be dropped (i.e. observations outside a specified caliper excluded), resulting in treated units

We match each trainee to a Wagner-Peyser nearest-neighbor (with replacement) by choosing pairs that minimize pairwise distance calculated from a set of “fuzzy” match covariates with the requirement that a matched pair belong to the same “exact” pre-program occupation group, quintile of AI exposure, employment status at program entry, and exit WIA/WIOA in the same calendar year. Approximately 4.6% of treated observations are dropped at the exact-matching stage due to a lack of suitable matches, a small fraction consistent with the large donor pool of comparable control observations from the Wagner-Peyser program. Conditional on this restriction, matching proceeds by assigning each treated observation a nearest neighbor control based on fuzzy-matched covariates. The covariates include:

- *Exact Match Covariates:* calendar year of program exit (controls for demand environment), AI exposure quintile in prior occupation (ensures nearest-neighbor pairs are comparable *within* similar AI exposure bins), SOC major occupation group (first two digits of 6-digit code) in prior occupation, and employment status at entry.
- *Fuzzy Match Covariates:* sex, age at program enrollment, race and ethnicity dummies, adult vs. dislocated program indicator, limited-English flag, low-income flag, Temporary Assistance for Needy Families (TANF) recipient, veteran status, Supplemental Security Income and Social Security Disability Insurance (SSI/SSDI) recipient status, highest education attained at time of program enrollment, and state centroid longitude/latitude as proxy for geographic distance of program location.

The matching routine achieves pre-program balance in the main outcome of interest (earnings) even without matching on the dependent variable in pre-program periods. This is likely due to the institutional advantage in our setting in which all nearest neighbor pairs are already self-selected into seeking workforce services.

Standardized distances among our matches are tightly distributed, with the interquartile range lying between between -0.75 and 0.42 standard deviations. Imposing a caliper (i.e., a maximum allowable distance that would further restrict estimation to regions of common support) thus has little effect on the results. To maximize the sample for heterogeneity analysis, we include all remaining matched pairs regardless of their distances. In heterogeneous subsamples, we add the more stringent matching requirement that nearest-neighbor pairs come from the same pre-program earnings quintile to allow us to credibly examine many subgroups (details below). Matching on earnings histories has been shown in the literature to improve matching performance ([Lechner et al. \(2011\)](#)), and

only being used if the covariates share a sufficiently common covariate support. See [Appendix G](#) for details.

in our case, allow us to credibly estimate separate effects across the distribution of prior AI exposure and for different time periods.²²

After matching each training participation spell to a nearest-neighbor job search assistance spell, we estimate earnings regressions by pooling together pre-training quarters $-3, -2, -1$ prior to program entry (placebo estimates) and post-training quarters $1, 2, 3$ after program exit (treatment estimates). We track outcomes for each workforce program participant n in matched pair p , with pre-entry occupation i and post-exit occupation j . We use occupation i to split the sample into “high” and “low” AI exposure subgroups prior to program entry based on a data-driven approach to the AI exposure distribution (discussed in [Section 4](#)).

Let $Y_{ni\tau}$ denote average quarterly earnings by prior AI exposure $i \in \{high, low\}$ in period $\tau \in \{pre, post\}$. In our main specification, we estimate

$$Y_{n(p)i\tau} = \alpha_{i\tau} + \beta_{i\tau} D_{n(p)} + \varepsilon_{n(p)i\tau} \quad (1)$$

using weights ω_p for each matched pair equal to the inverse Mahalanobis distance between each treated individual ($D_{n(p)} = 1$) and their matched control ($D_{n(p)} = 0$) based on their vector of fuzzy covariates X_n within each exact-matched covariate group.²³

To validate this design, [Appendix Table A.3](#) and [Appendix Table A.4](#) present balance test results from estimating equation (1) on a variety of outcomes prior to program entry by low and high prior AI exposure matched samples respectively. In all regressions, standard errors are clustered at the control unit level to account for the possibility that the same matched control unit is selected by multiple treated units. While there are some imbalances (as would be expected by chance when estimating 44 t-tests without adjusting for multiple hypothesis testing), differences in mean values for treated and control groups are, for the most part, economically trivial or quantitatively small. The exception to this is *total quarters enrolled*, which is mechanically different because trainees remain enrolled in the program for up to two more quarters than their Wagner-Peyser nearest-neighbors. This longer training duration rules out the possibility that positive earnings returns from training programs are

²²A rich literature in non-experimental matching estimators weighs the costs and benefits of different matching estimators. We show in [Appendix Table A.13](#) that our training returns estimates are unchanged when we use two alternate matching procedure discussed in [Andersson et al. \(2024\)](#): propensity-score matching and nearest-neighbor matching on the propensity score.

²³The Mahalanobis weight is calculated as

$$\omega_p = \sqrt{(X_{D_{n(p)}=1} - X_{D_{n(p)}=0})' \Sigma_X^{-1} (X_{D_{n(p)}=1} - X_{D_{n(p)}=0})}$$

and accounts for covariates measured on different scales via the inverse covariance matrix Σ_X^{-1} . We estimate effects without imposing a caliper (maximum distance restriction), however large distances receive the lowest weight in estimation. See [Appendix G](#) for full details and robustness to alternative matching estimators.

driven by quicker reemployment, as longer program durations would lead us to understate rather than overstate effects. The table also reports pre-participation control group means, highlighting that high AI exposure trainees begin from a higher earnings baseline, which helps contextualize differences in post-training gains.

Decomposition Analysis. We also estimate a related statistical decomposition of the main earnings effect to shed light on how the skills acquired from the training program, as proxied by the destination occupation j 's AI exposure, matter for earnings returns. The basic idea (expressed diagrammatically in [Appendix B](#)) is that there are two ways for earnings to increase through training: one is through changing occupational AI-exposure, and one is through experiencing earnings growth while maintaining the same occupational AI-exposure as one's nearest neighbor. The former is captured by an additional "training choice" (see [Appendix B](#))—the AI content of the destination occupation—which mediates the impact of training on outcomes.

Formally, we decompose the estimated overall earnings returns to training within each exposure group i in time τ by writing

$$\hat{\beta}_{i\tau} = \underbrace{\sum_{g \in G} \hat{\beta}_{gi\tau} \hat{\sigma}_{gi\tau}}_{\text{returns holding reskilling fixed}} + \underbrace{\sum_{g \in G} \hat{\beta}_{gi\tau} \hat{\gamma}_{gi\tau}}_{\text{returns from training-induced reskilling}} \quad (2)$$

where $g \in G \equiv \{\text{Upskill, Downskill, Same}\}$ represents whether the trainee upskilled, downskilled, or remained at the same level of occupational AI exposure post-training (regardless of how their nearest neighbor may have reskilled).²⁴ $\hat{\beta}_{gi\tau}$ denotes average earnings returns among trainees in group g , measured as a pre/post difference. $\hat{\sigma}_{gi\tau}$ denotes the share of matched neighbors that fall into reskilling group g without training, such that $\sum_{g \in G} \hat{\beta}_{gi\tau} \hat{\sigma}_{gi\tau}$ can be interpreted as counterfactual earnings if trainees had the same reskilling distribution as their matched controls. $\hat{\gamma}_{gi\tau}$ captures the causal effect of training on reskilling, defined as the change in the probability of being in group g relative to the control group.²⁵

For example, if trainees disproportionately upskill relative to their matched controls, then $\hat{\gamma}_{g\tau}$ for $g = \text{upskill}$ will be positive, shifting weight toward the returns for the upskilling

²⁴The decomposition follows from the law of iterated expectations and is analogous to Kitagawa-Oaxaca-Blinder decompositions, separating within-group returns from composition (reallocation) effects.

²⁵The identifying assumption underlying the decomposition is more demanding because the counterfactual occupation choice absent training is unobserved and must be proxied by the matched control unit's realized occupation. This requires assuming that, absent training, trainees would have sorted into occupations similar to those of their matched Wagner-Peyser counterparts. Because occupational sorting is a post-treatment outcome that is not balanced in the matching stage, differences in sorting may also reflect residual selection on unobservables related to transition propensities. Accordingly, the decomposition should be interpreted as descriptive of margins of adjustment rather than a fully causal attribution of sorting itself.

group. The ratio $\sum_{g \in G} \hat{\beta}_{gi\tau} \hat{\sigma}_{gi\tau} / \hat{\beta}_{i\tau}$ therefore measures the share of total earnings gains attributable to returns that would arise even in the absence of training-induced occupational movement. If mechanisms such as credentialing (e.g., occupational licensing) or signaling raise earnings independently of changes in AI exposure, this share will be large. Conversely, if training primarily operates by shifting trainees into different occupations, the contribution of this component will be smaller. Finally, if observed occupational movements primarily reflect reexposure rather than true reskilling, then group-specific estimates will conflate these effects, and we should correspondingly observe attenuated returns within reskilling groups.

4 Main Results

4.1 Returns to Training by AI Exposure

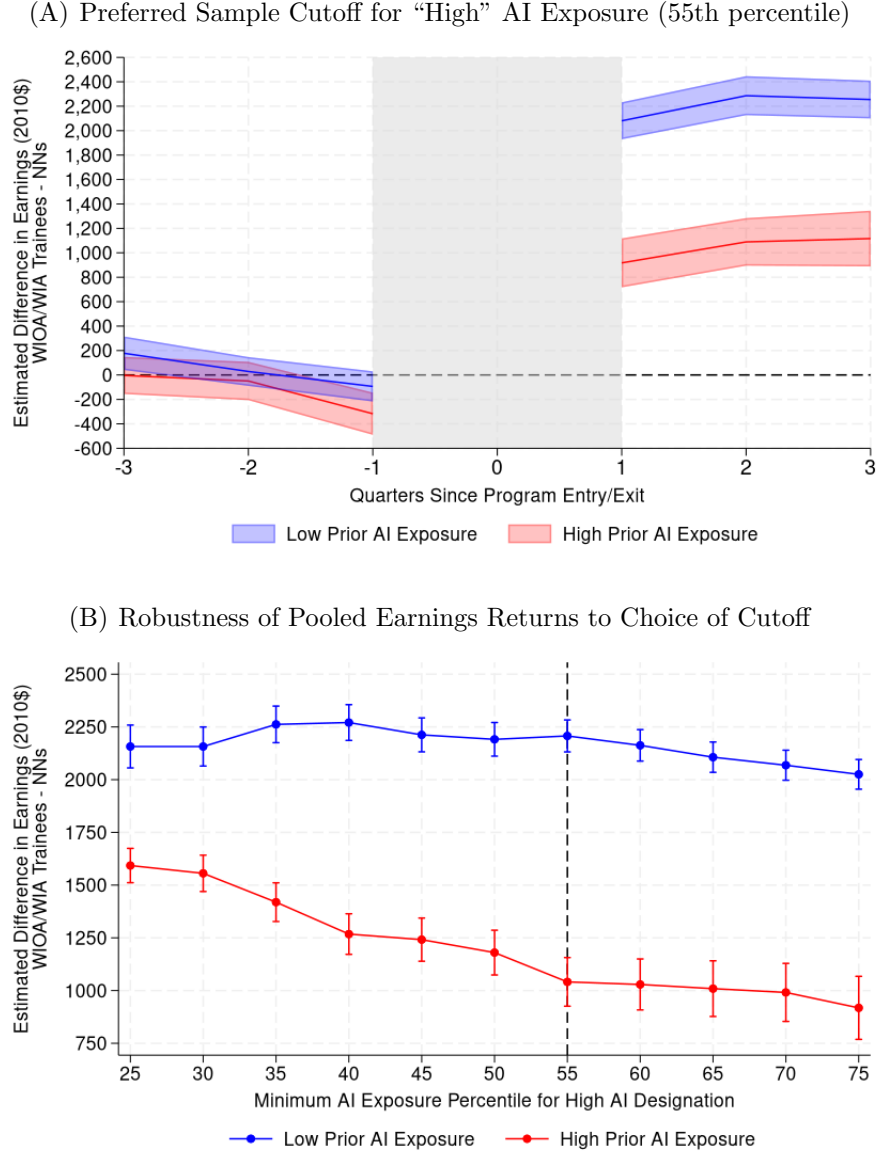
In [Figure 3](#), we graphically present the main empirical findings of the paper for our full sample. Each series shows estimates of quarterly differences in mean earnings between WIOA/WIA training participants and their nearest-neighbor matched Wagner-Peyser job search assistance recipients separately by each AI exposure subgroup. Panel (A) shows estimates of $\hat{\beta}_{i\tau}$ from equation (1) for high and low AI exposure groups prior to program enrollment, separately estimated in each quarter. Earnings differences are weighted by inverse match distance, and standard errors are clustered at the control unit level to account for error correlation when control units match to multiple treated units.

Panel A shows that prior earnings differences are well balanced in both sub-groups prior to program participation. This mitigates concerns about selection on unobserved heterogeneity, and is consistent with caseworkers adding an element of quasi-random gatekeeping into training. Estimated earnings returns for high initial AI exposure training participants are about half the returns for low AI exposure participants in the three quarters after program exit. However, overall training returns in both groups remain quite high in quantitative terms. In the three quarters after training exit, high AI exposure workers have quarterly earnings that are \$850-1,231 higher relative to their matched control group pairs, or on average \$4,164 annually on a base of roughly \$48,000 prior to participation. Among high AI exposure workers, this translates to those who train having about 9% higher short-run earnings relative to similar workers who only receive job search services.²⁶ Panel B further shows that the patterns of both large effects and stronger effects for low initial AI-exposure workers is robust across cutoffs for defining high exposure. We

²⁶See [Appendix Table A.1](#) for baseline real wages, and [Appendix Table A.5](#) for pooled regression estimates for high and low AI exposure groups before and after participation.

choose the 55th percentile as an expositional heuristic for separating high and low exposure participants. In subsequent heterogeneity analyses, we show effects across all exposure bins.

Figure 3: NN-Matched Earnings Returns to Job Training by Prior AI Exposure, 2012-2024



Notes: AI exposure is based on [Eloundou et al. \(2024\)](#) occupation-level AI exposure scores. Each series in Panel A represents the earnings return relative to nearest-neighbor matched twins weighted by inverse distance in matches. Panel B shows mean estimates from pooling together quarterly earnings across the three quarters after program exit. The 55th percentile is chosen as the preferred cutoff for high and low based on estimate stability beyond this cutoff (see text for details).

4.2 The Role of Reskilling Before and During the AI Boom

Training participants may pursue general skill training that enables transitions away from AI-exposed occupations, or occupation-specific AI training that deepens human capital in ways that complement AI technology and make the worker more AI-resilient. In this section we consider how AI skill movements mediate the earnings returns from training, and how that mediation has evolved over time as AI emerged.

Panels A and B of [Figure 4](#) group workers into 12 quintile-spaced bins of AI exposure before training. The average AI exposures after training for each pre-training exposure bin are plotted in circles using the left y-axis, and the mean earnings returns to training are plotted in triangles using the right y-axis. The figure also plots the pre-program earnings differences between trainees and their non-training nearest-neighbors in squares to underscore that we continue to see balance between treatment and control units across the AI exposure distribution.²⁷ The 45-degree line in each panel allows us to distinguish which workers AI-upskill on average – meaning their post-training occupations are more AI-exposed than their pre-training occupations – from those who AI-downskill on average. The figure therefore links earnings gains not only to where workers start in the AI exposure distribution, but also to the direction of their movement in AI skill space following training.

Panel A contains workers who began training between 2012-2019 and had exited before 2020. Panel B contains workers who began training between 2022-2024 and exited by 2024q4 (the end of our sample). We define the latter time period to capture the beginning of the generative AI boom, set off by the watershed releases of GitHub Copilot in 2021 and later OpenAI’s ChatGPT 3.5 in late 2022. We omit 2020 and 2021 altogether to avoid pandemic-era dynamics and, consistent with earlier heterogeneity analysis, continue to match on pre-program earnings quintiles for precision. The resulting early and late trainee subsamples represent about 58% and 10% of the original estimation sample, respectively. As before, low AI exposure bins below the 55th percentile cutoff are colored blue and high exposure bins are colored red.

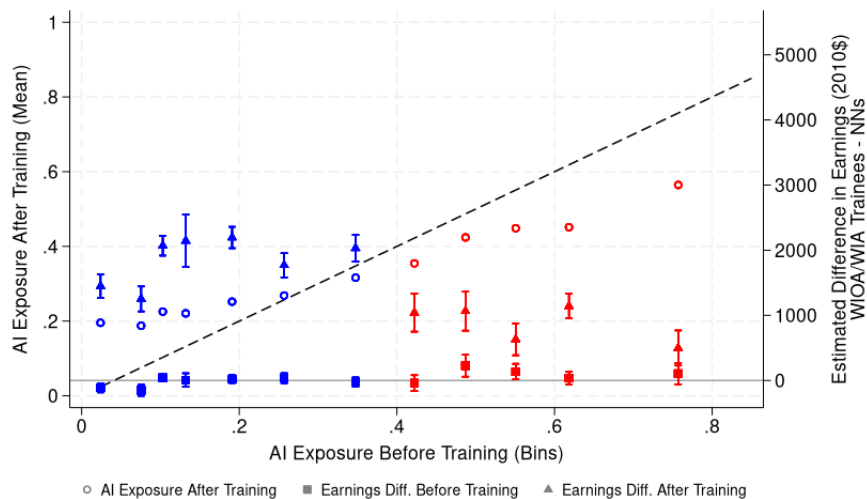
In the early period in Panel A we find strong asymmetries in returns by initial AI exposure. Workers originating in low AI-exposure occupations capture the largest earnings gains and tend to move above the 45-degree line to more AI-exposed work, consistent with positive returns to AI-related upskilling from a low baseline. Workers leaving highly AI-exposed occupations earn lower returns and tend to transition toward less AI-exposed jobs.

²⁷As noted in the methodology section, heterogeneity analyses include exact-matching on vintiles of prior earnings. In [Appendix Figure A.7](#), we show that balance remains strong without matching on earnings. Additionally, in [Appendix Figure A.6](#) we split [Figure 3](#) (not matching on earnings as in the previous section) into the early and late period.

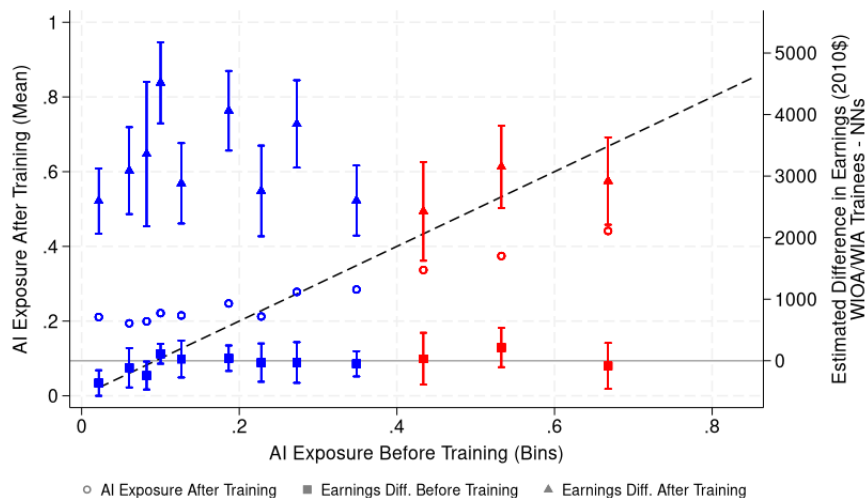
Returns for some high-exposure workers, in particular the most exposed bin, are quite small while at least 95% of workers in each low-exposure bin have quarterly returns above \$1,000 and some earn more than \$2,000 per quarter more than their non-training neighbors after exit. There appears to be a negative relationship between training returns and AI exposure before training.

Figure 4: NN-Matched Earnings Returns to Job Training by Prior AI Exposure

(A) Pre-Pandemic Years: 2012-2019



(B) AI Boom Years: 2022-2024



Notes: Figure overlays pooled estimates for the returns to training across the AI exposure spectrum (triangles), placebo estimates for pooled estimates prior to entry into training/job search services for the control group (squares), and average movements in AI skill before and after training (circles). Circles above the 45 degree line represent AI upskilling, while below the line represent AI downskilling. Red and blue correspond to our prior preferred high and low exposure distinction. A horizontal gridline is drawn at zero earnings returns. AI exposure uses the measure from [Eloundou et al. \(2024\)](#).

Panel B reveals a stark improvement in earnings returns for both groups – but especially high-exposure workers – during the AI boom period of 2022-2024. The bulk of the low-exposure workers earn \$2,500 - \$4,000 more per quarter than their nearest-neighbors while almost 95% of the high-exposure workers now have returns over \$2,000. The relatively larger shift down in the red circles relative to the 45-degree line during the later period demonstrates that high-exposure workers were more likely to AI-downskill during the AI boom. There is also no longer a clear relationship between returns and prior exposure for most of the exposure distribution. There are also fewer trainees coming from low-exposure occupations, as shown by the clustering of the 12 quintile bins in the late period lower down on the AI exposure distribution.

One natural first interpretation is that the improvement in training returns in recent years could be attributed to exceptionally-tight labor markets across the AI exposure distribution. However, since returns to training are estimated within pairs of trainees and nearest-neighbors who exit into the same labor markets, the improvement must be driven by something specific to training. It remains possible that training has higher signaling value when markets are tight as firms are forced to hire further down the productivity ladder, though we find it unlikely that signaling alone would produce the demonstrated catch-up in earnings returns for high AI exposure workers.²⁸ For the remainder of this section, we consider evidence for non-signaling explanations by decomposing earnings returns into training and sorting effects, and interpret the sorting component in light of the growing prevalence of training programs that build AI-complementary skills.

4.3 Statistical Decomposition

The previous section showed that skill sorting became more prominent during the AI boom. To better understand the contribution of occupational sorting to trainees’ earnings returns, we begin by showing the causal effects of training on workers’ likelihoods of AI-downskilling (moving into occupations less exposed to AI) or AI-upskilling (moving into higher-exposure occupations), estimated relative to the occupation transitions made by the non-training nearest-neighbors. These are the $\hat{\gamma}_{gi\tau}$ terms from equation (2), whose estimates are shown in Panels A and B of Figure 5.

Panel A shows that from 2012-2019, trainees in low exposure occupations were 16pp more likely to upskill (purple bars) after WIOA participation, and this effect rose to 22pp

²⁸In complementary work, Hyman et al. (2025) use a mismatch model following the work of Şahin et al. (2014) to quantify the degree to which the recent effects emanate from a strong labor market, or from improved alignment between the demand and supply for specific skills (especially AI skills) achieved via job training programs.

by 2022-2024. The training effect on downskilling (green bars) also rose modestly for low-exposure workers. Panel B demonstrates that during the AI boom, the share of high-exposure trainees who AI-downskilled rose to over 20 percentage points relative to nearest-neighbor matched pairs. Meanwhile the effect on the share upskilling remained roughly the same. Together, Panels A and B illustrate that during the AI boom, training induced more low-exposure workers to both downskill and upskill (leaving fewer trainees in the occupations they started), but induced far more high-exposure workers to attempt downskilling transitions. This is suggestive of an effort to avoid occupations expected to be at high risk of AI-related displacement.

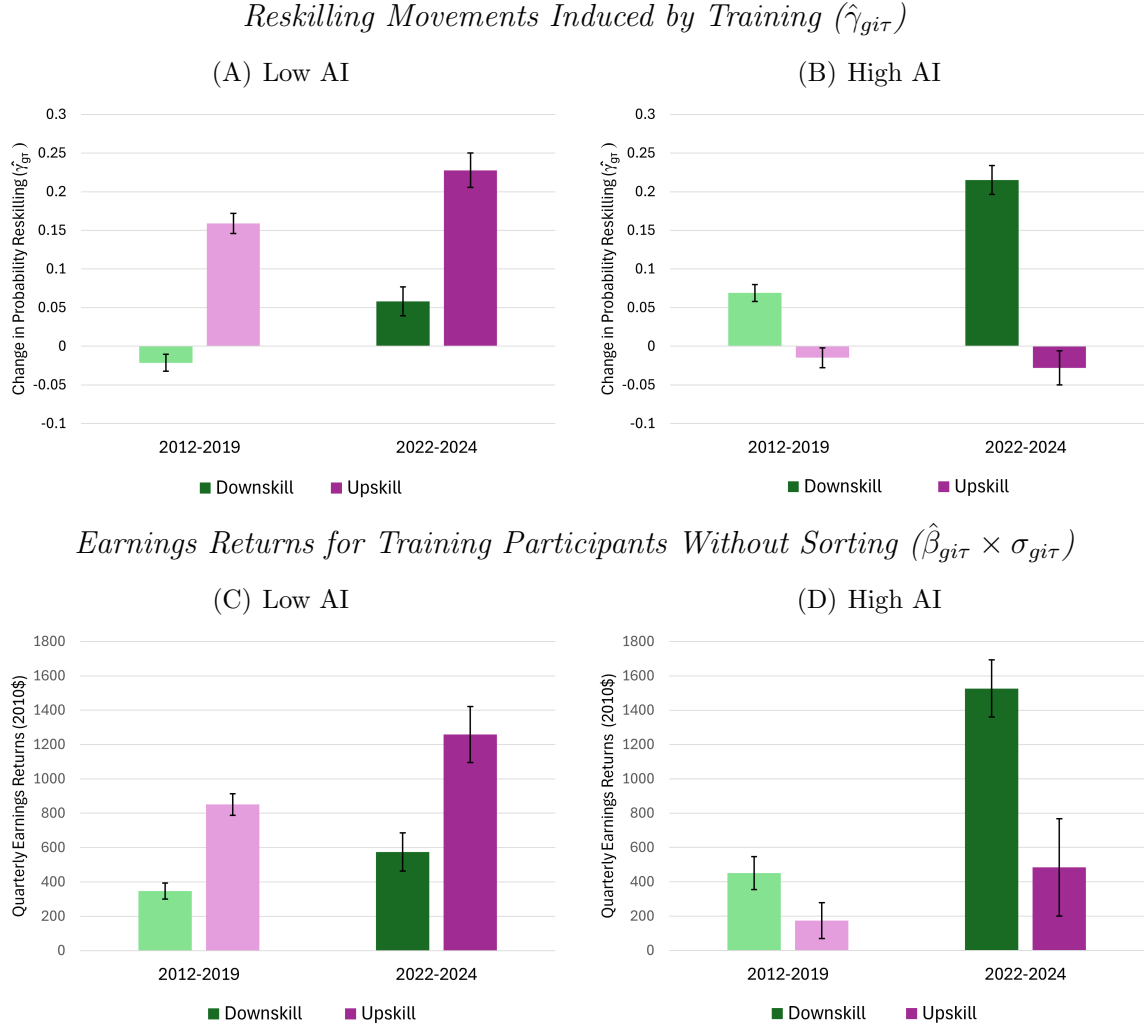
Panels C and D are statistical objects, not causal estimates, but provide important context for the sorting effects underlying the results in Panels A and B. The bars in these bottom two panels show the earnings returns for trainees when holding sorting components constant ($\hat{\beta}_{gi\tau} \times \sigma_{gi\tau}$ from equation (2)). Panel C shows that earnings returns for low-exposure occupations exhibited small increases in the late period for both downskilling and upskilling workers. Panel D clarifies that the recent rise in downskilling among high AI exposure workers from Panel B coincided with about a \$1,075 higher relative gain in the returns to training for those workers choosing to downskill.

Overall, [Figure 5](#) suggests that the improvement in earnings returns for highly-exposed workers during the AI boom was associated with increased movements of those workers into less-exposed occupations and a large, concurrent improvement in the earnings returns to downskilling. In other words, workers who might have tried to upskill before the AI boom were less likely to do so after 2022, thus earning more, and workers who would have downskilled anyway still captured higher returns than they would have before the pandemic. We consider the mechanisms behind these changes in the next section.

As a final quantification exercise, in [Appendix Tables A.6](#) and [A.7](#) we show that summed over exposure groups, about 3.2% of earnings returns were driven by sorting prior to the AI boom, while during the AI boom, the sorting component rose to 9.3%.²⁹ While still modest relative to the majority of returns to training attributable to non-sorting mechanisms, this represents a near-tripling of the sorting contribution, consistent with evolving AI-driven changes in occupational labor demand as AI technology has begun to diffuse throughout the economy.

²⁹The reported sorting component shares are computed by aggregating the sorting effect and total columns across high- and low-AI exposure workers across both [Appendix Tables A.6](#) and [A.7](#), separately for the pre-AI period (2012–2019) and the AI boom (2022–2024). Specifically, we sum the sorting effects across both groups and divide by the corresponding sums of total returns. This yields shares of \$86.05 / \$2,721.28 $\approx$ 3.2% prior to the AI boom and \$568.21 / \$6,137.20 $\approx$ 9.3% during the AI boom.

Figure 5: Decomposition of Earnings Returns by AI Reskilling



Notes: Figures show the reskilling movements and earnings returns induced by WIOA job training for High and Low AI exposed workers, split by time period of WIOA participation. Trainees exiting WIOA during the height of the COVID-19 pandemic (2020-2021) are excluded. Panels in the top row detail trainees' propensities to upskill or downskill in AI space, while panels in the bottom row detail trainees' earnings returns while holding sorting constant. 95% confidence interval bands are plotted with each panel, where standard errors in Panels C and D are calculated using the delta method.

4.4 Discussion: Potential Mechanisms

We have shown that rates of AI downskilling among high-exposure trainees and AI upskilling among low-exposure trainees rose during the AI boom, with greater increases among downskillers. In this section, we examine the extent to which these patterns reflect changes in labor demand versus adjustments in labor supply, as reflected in the occupations that WIOA programs prepare trainees for. First, we consider shifts in occupational demand across AI exposure groups.

Labor Demand. Should AI reduce the demand for particular occupations, one might observe that the equilibrium occupations workers transition to may shift over time. In the WIOA data, transitions into traditionally low AI-exposure occupations—particularly in healthcare and transportation—are associated with substantially higher returns during the AI boom (see [Appendix Table A.8](#)). These are occupations that WIOA programs have long specialized in training for, suggesting that existing program infrastructure is well aligned with labor market opportunities in service sectors that have historically absorbed displaced workers. Consistent with this, many high-exposure workers who downskill transition into these roles, where returns have risen markedly over time. This pattern points to a combination of shifting labor demand and program-driven steering toward occupations where WIOA has a comparative advantage.

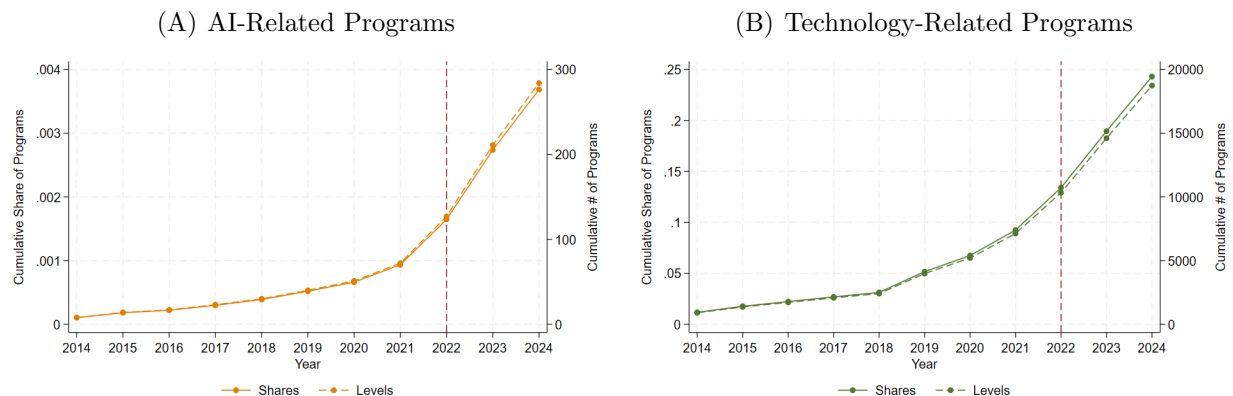
Labor Supply. While we do not observe the specific programs in which training participants enroll, we can use unlinked training provider information to characterize the evolution of AI content in workforce development programs using textual program descriptions from the WIOA eligible training provider lists (ETPLs).³⁰ This can shed light on how prevalent technological training became over this period, as well as the evolution of technological training content before and during the AI era. These WIOA ETPLs contain detailed curricular and aggregate WIOA performance information on approved training programs that entered between 2014 and 2024. We use an NLP to identify common AI- and general technology-related keywords within a subsample of WIOA ETPL program descriptions, and apply those lists of keywords to tag training programs as “AI” related, “Technology” related, or “Other” across the full dataset of WIOA training programs (see [Appendix C](#) for details).

In [Figure 6](#), we show the growth in AI-specific and general Technology related training programs in terms of their cumulative levels and shares. Panel B shows that the share of general technology trainings grew to account for nearly 25% of all WIOA-eligible training programs by end of 2024, representing a significant share of the overall job training landscape. Contrastingly, AI-specific training content remains a very small share of the overall job training programs, but have begun to grow at a rapid rate since the 2022 GPT boom. At the end of 2021, there were only about 70 total WIOA programs that contained AI-content, yet just 3 years later, the total number of AI programs nearly quadrupled. These growth trends are not just evident in the number of AI and technology training programs in the post-2022 period, but are also reflected in the number of WIOA training *participants* who enrolled in such programs (which we observe on an aggregate program basis, but not by individuals). [Appendix Figure A.9](#) shows similar trend breaks at 2022 in AI-related programs when we

³⁰The WIOA ETPL data is provided by the U.S. Department of Labor on [TrainingProviderResults.gov](https://www.dhs.gov/training-provider-results).

weight these figures by the number of WIOA trainee participants.

Figure 6: Cumulative Growth in U.S. AI- and Technology-Related Job Training Programs

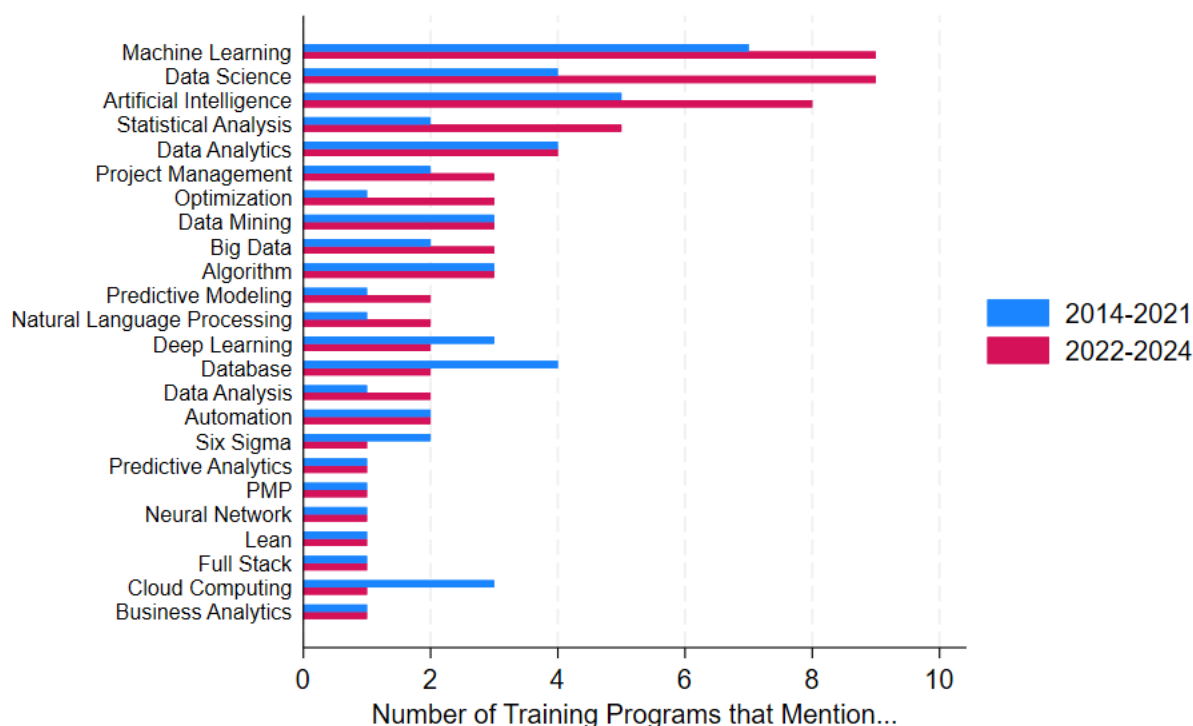


Notes: Panel A shows AI-related programs (e.g. keywords: *Artificial Intelligence*, *Natural Language Processing*, *Machine Learning*), and Panel B shows non-AI general Technology-related programs (e.g. keywords: *Programming*, *Web Development*, *Python*) drawn from WIOA ETPL lists over time. Shares are calculated as the cumulative number of AI or Technology programs introduced over the cumulative number of all WIOA ETPL programs within a given year. Dashed vertical line at the end of 2022 is meant to represent the release of ChatGPT 3.5 in November 2022, viewed as a watershed innovation in the literature. See [Appendix C](#) for details.

Beyond the general growth of AI-specific training programs in the post-2022 era, we also seek to understand the evolution of AI content included within such training programs. To do this, we analyze the textual descriptions of program curricula and targeted competencies within the WIOA ETPL database to draw insights on the trends in AI training content over time. [Figure 7](#) shows a tabulation of the number of WIOA training programs containing different AI-related keywords across the pre-2022 period (2014Q1-2021Q4) and the post-2022 period (2022Q1-2024Q4)³¹. Across the board, there are more mentions of core AI terms such as “machine learning” and “artificial intelligence” within the post-2022 period. These results highlight a growing supply in AI content by training providers as well as a growing interest in AI-training take up by WIOA participants, sharply coinciding with the advent of GPT-3.5’s public release in 2022. However, as evident from the lower overall counts, AI-specific training programs are still not widely available which may be creating some upskilling frictions, particularly for high-exposure trainees.

³¹See [Appendix Table A.12](#) for a complete list of keywords used for tagging AI training content.

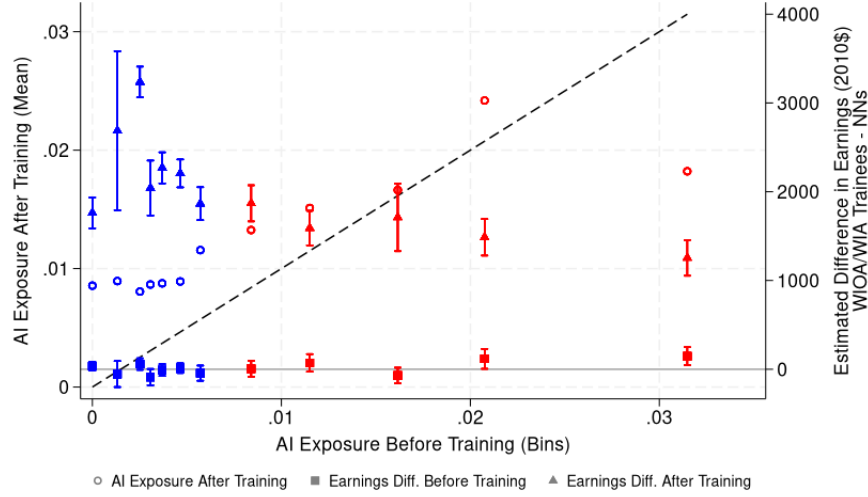
Figure 7: Common Keywords in AI Program Descriptions (Pre- vs. Post-2022)



One final piece of evidence can be attained by estimating the earnings returns based on occupations which have been shown to use LLMs intensively in practice (rather than potentially, as captured by our baseline measure). Here, we use the Claude-based measure from Anthropic and reproduce our main estimates of the heterogeneous returns to training relative to different skill transitions. [Figure 8](#) shows that the results based on exposure measures from Claude usage data are remarkably similar to what was found using the potential AI exposure measure. This helps confirm that those who exit occupations that actually used AI both upskilled and downskilled while retaining large returns from training.

Taken together, these patterns suggest that WIOA is not merely expanding AI skills at the margin, but reallocating workers toward segments of the labor market where returns remain high despite rising AI exposure. In this sense, the program acts as a buffer against AI-driven disruption—facilitating both exit from vulnerable occupations and targeted entry into resilient, service-oriented roles where demand continues to grow.

Figure 8: Earnings Returns to Job Training vs. Skill Transitions: Claude-Based Exposure



Notes: AI exposure is based on the Anthropic Economic Index Claude-based occupation-level AI exposure scores (Handa et al., 2025). Plot overlays pooled estimates for the returns to training across the AI exposure spectrum (triangles), placebo estimates for pooled estimates prior to entry into training/job search services for the control group (squares), and average movements in AI skill before and after training (circles). Circles above the 45 degree line represent AI upskilling, while below the line represent AI downskilling. Red and blue correspond to our prior preferred high and low exposure distinction. The Claude use distribution skews left; one estimate to the right of the support is suppressed for exposition.

5 AI Retrainingability (AIR) Index

We have established that the earnings returns to training are positive across the AI spectrum, but also that the payoffs to training depend on transitions in AI skill space. In this section, we shift focus to workers who appear most resilient to AI: those who jointly move up in AI skill space after job training, and receive higher earnings for doing so.

5.1 What share of workers and occupations are AI retrainable?

Using our nearest-neighbor matched estimates, we start by calculating the total number of training participants for whom both earnings returns and the AI content of their next occupation are higher than those of their matched pair counterparts. In Table A.9, we show the share of training spells for which both AI intensity and earnings are weakly greater, separately for the training group relative to matched nearest neighbors. We find that about 40% of workers are AI retrainable, with this estimate declining slightly to 39% when using the LLM measure developed by Handa et al. (2025). We also perform the same calculation by asking for any given 6-digit occupation, whether on average workers displaced from that occupation move to higher earnings and AI exposure scores, and find 40%-45% of occupations

doing so. These AI retrainable occupations cover 39% of the CPS labor force, though this is likely a lower bound as some occupations in the CPS do not appear in WIOA/WIA. The same table shows that AI retrainability is high even among occupations that are initially highly exposed to AI.

5.2 AI Retraining Index

To further unpack the mechanisms underlying the estimated large degree of adaptability to AI shocks through retraining, we develop an “AI Retraining Index” (AIR) that ranks occupations based on the share of job training participants who receive positive earnings returns from moving to more AI-exposed work, relative to their matched control group pairs. This allows us to decompose whether training success emanates from joint movements towards higher earnings and AI skill, versus one dimensional movements favoring one component over the other. Higher earnings returns without significant AI upskilling would be consistent with alternative stories unrelated to AI, such as overcoming occupational licensing barriers through job training. Such occupations may be considered generally retrainable, but not “AI retrainable”. As noted previously, AI upskilling may not always be indicative of successful AI skill adjustment if workers are just reexposing themselves to AI-related shocks without also gaining positive earnings returns. Thus, our index utilizes information from both dimensions to assess the degree of successful AI retrainability at the occupation level.

We use information from each nearest-neighbor matched pair p . Training participants (and their matched pairs) begin in origin occupation i prior to training and can flow to any destination occupation j after training.³² The index is as follows

$$AIR_i = \sum_j \frac{1}{n_{ij}} \Delta AI_{ij}^\sigma * \Delta \log(\overline{Y}_{ij})^\sigma \quad (3)$$

where n_{ij} is the number of trainees flowing from occupation i to j , capturing probability of reemployment in occupations of varying AI intensities. The relative change in AI skill term $\Delta AI_{ij}^\sigma = (AI_j - AI_i)_{p=1}^\sigma - (AI_j - AI_i)_{p=0}^\sigma$ is the difference in standardized Eloundou et al. (2024) AI exposure scores associated with each trainee’s occupational transition from i to j relative to matched control group units. $\Delta \log(\overline{Y}_{ij})^\sigma = \frac{1}{p_{ij}} \sum_p \Delta \log(Y_{p=1}^{i \rightarrow j})^\sigma - \Delta \log(Y_{p=0}^i)^\sigma$ is the difference in standardized mean returns to training for each treated unit in the nearest neighbor pair p relative to each control group unit $p=0$ that started in same occupation i

³²In this descriptive exercise, we further strip out any participants who were dual enrolled in WIOA job training and Wagner-Peyser search assistance to hone in on the mechanisms driving the high returns to job training in particular.

but is free to move to any subsequent sector j . Ordinarily, we would be concerned about the potential of multiplying negative wage returns by negative AI exposure changes, however within our sample of nearest-neighbor matched trainees, all 2-digit occupation-level mean earnings returns are positive, allowing us to express both terms in standardized Z-score form. Having both components set on the same standardized scale has the benefit of weighting WIOA-induced AI exposure movements and earnings returns equally when determining a ranking of AI retrainability.³³

The functional form in Equation 3 is chosen such that the index has the largest value when both the training-induced change in the AI content of work (i.e. AI exposure) and earnings returns to training are highly positive. Given earnings returns are on average positive across occupations, the index only turns negative when the change in the AI content of work is negative. This allows us to hone in on whether the highest ranked occupations in AI retrainability are driven by much higher earnings gains or instead large leaps in AI skill after job training. By differencing out effects with respect to nearest neighbor pairs, this functional form also minimizes any mechanical bias that may arise from occupations with more or less scope to move in AI skill space due to starting at initially low or high exposure values.

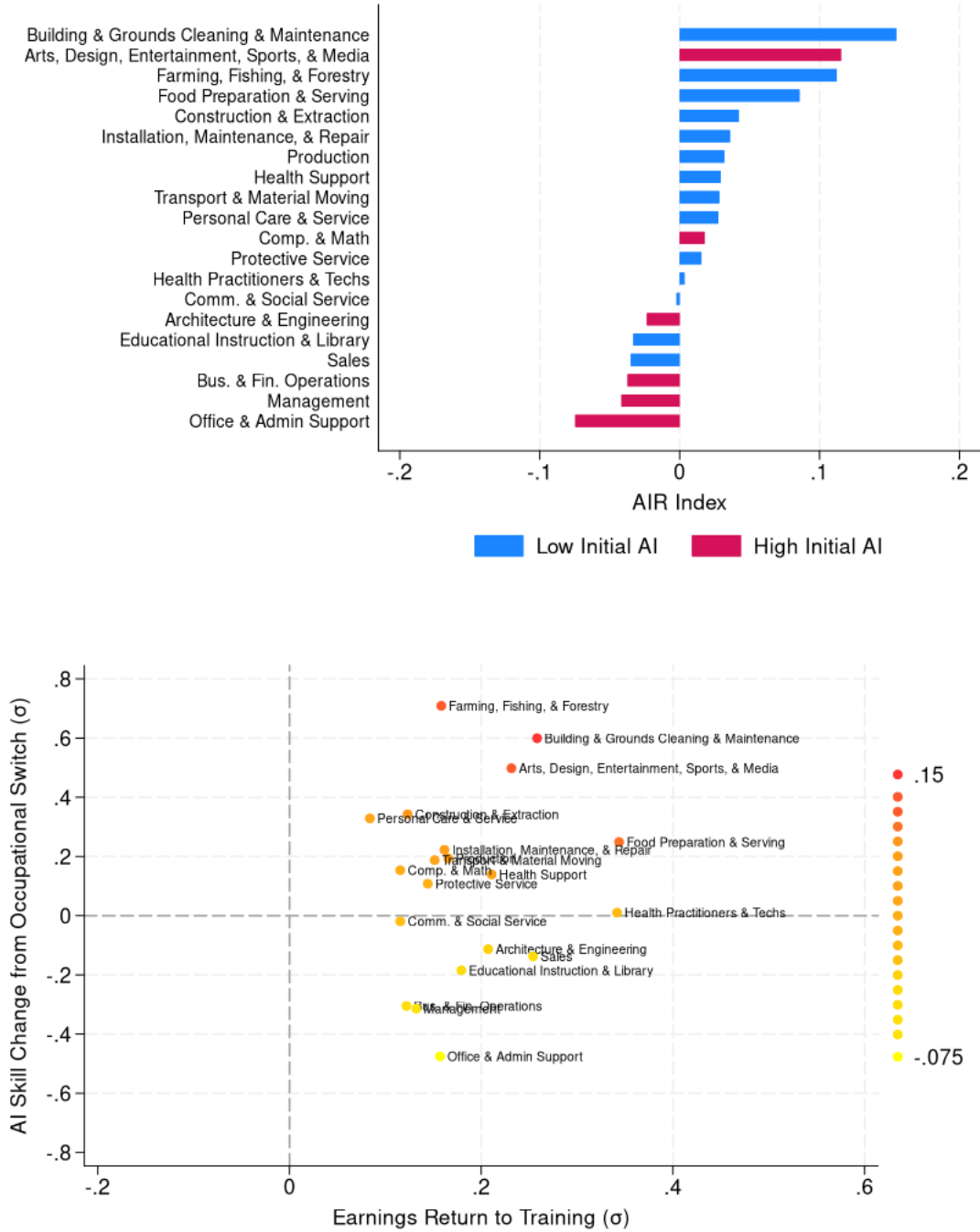
To see this clearly, Figure 9 presents the overall 2-digit rankings of occupations by their AI retrainability. We focus on trainees who entered WIOA training programs between 2022q1 and 2024q4 to capture occupation-level AI retrainability for workers who would have been most exposed to developments during the AI boom. Panel (A) shows that a large majority of occupations have positive AIR values. Panel (B) decomposes the contribution of AI skill and earnings returns components to the index, where index values are expressed as a heatmap with top ranked occupations in darker reds. The bottom panel clearly shows that the rankings are more highly influenced by variation in AI skill changes (larger vertical movements along the y-axis) rather than variation in earnings returns (horizontal movements along the x-axis). Positive AIR values are associated with training-induced upskilling, while negative values are associated with downskilling.

³³We also consider an alternative functional form of AIR_i which utilizes a MinMax transformation on both the AI and earnings terms, allowing us to circumvent the problem of multiplying negative wage returns by negative AI exposure changes by converting the most negative earnings returns into low positive values. Specifically, where $\mathcal{M} \in [0, 1]$, the MinMax transformation is

$$\mathcal{M}(\Delta \log(\bar{Y}_{ij})) = \frac{\Delta \log(\bar{Y}_{ij}) - \min(\Delta \log(\bar{Y}_{ij}))}{\max(\Delta \log(\bar{Y}_{ij})) - \min(\Delta \log(\bar{Y}_{ij}))} \quad (4)$$

This alternative form yields similar rankings overall (see Appendix Figure A.10).

Figure 9: AI Retraining Index during the AI Boom (2022-2024)



Notes: Panel (A) shows 2-digit rankings of AI retrainability index (AIR) using [Eloundou et al. \(2024\)](#) AI exposure measure, where higher values reflect joint movements up in earnings AI skill space, inclusive of all trainees who entered WIOA training programs between 2022q1 and 2024q4. Red and blue bars indicate high and low initial AI exposure prior to training. Panel (B) shows a decomposition with a heat map for higher index values, in which trainees leaving occupations above the horizontal dashed line are moving up in AI skill space.

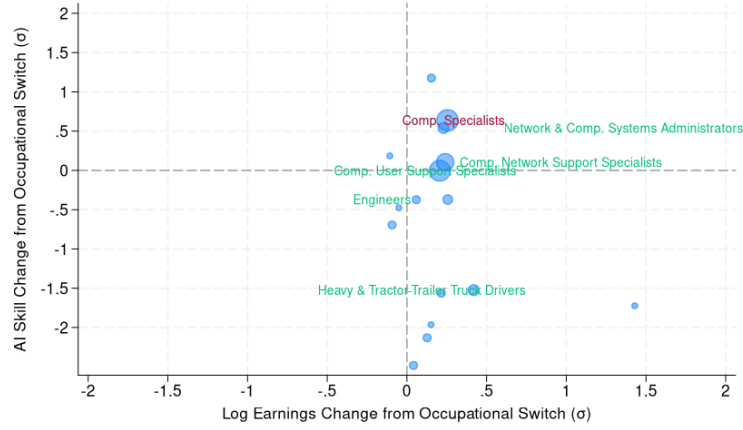
All but two of the positive AIR occupations originate from low prior AI exposure occupations. This is consistent with our earlier findings that low AI exposure workers are more likely to upskill through WIOA training, which allows them to capture higher earnings returns. Categories such as *Building & Grounds Cleaning & Maintenance*, *Farming, Fishing & Forestry*, and *Construction & Extraction* are highly manual at baseline, and workers in such categories appear to benefit significantly from WIOA training, allowing them to transition into more cognitively-intensive roles such as administrators, sales professionals, and supervising managers. Among the same subset of positive AIR rankings, only two occupational categories originate from highly AI exposed fields: *Arts, Design, Entertainment, Sports, & Media* and *Computer & Mathematics*. Unlike many other highly exposed AI occupations (e.g. *Business and Financial Operations* and *Office and Admin Support*), these occupations stand out as trainees in these categories appear to be successfully AI upskilling despite already starting from a high AI position.

In a final exercise, [Figure 10](#) shows decompositions for three examples of disaggregated 6-digit occupations that illustrate the mechanisms underlying the 2-digit rankings. These include: *Computer User Support Specialists* within the 2-digit *Computer & Math* occupation, *Medical Assistants* within the 2-digit *Health Support* occupation, and *Customer Service Representatives* within the 2-digit *Office and Administrative Support* occupation. To assist with the interpretation of these figures, we plot the nearest-neighbor matched earnings returns on the x-axis (interpreted as average returns to WIOA training), and the standardized difference in AI exposure for only trainees on the y-axis (interpreted as the magnitude of up or downskilling).

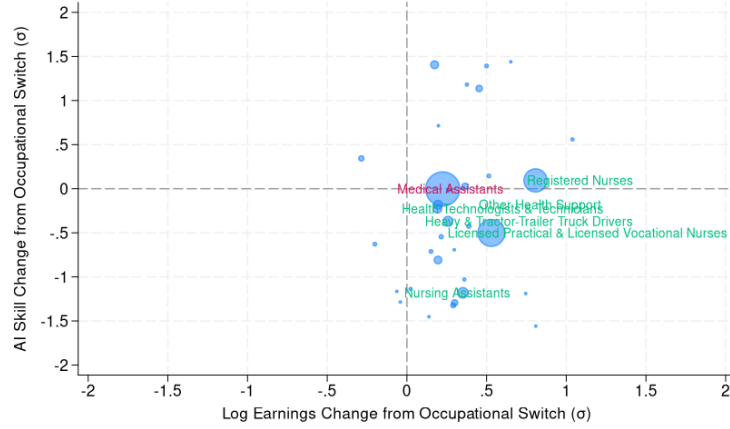
Panel (A) shows that *Computer User Support Specialists* (e.g., IT specialists), an initially highly AI exposed occupation, transition to a wide range of post-WIOA occupations across the AI skill distribution. While some trainees end up significantly downskilling into roles like *Heavy & Tractor-Trailer Truck Drivers*, there remains a significant share that manages to upskill into higher AI exposure roles within the broad *Computer & Mathematics* occupational category, including *Network & Computer Systems Administrators*. Such workers may be utilizing WIOA training to augment complementary technological skills that would allow them to remain in highly exposed technology occupations while gaining in earnings. The fact that we do not observe all workers in this category upskilling could be indicative of upskilling frictions arising from the still-nascent AI training content within WIOA’s current infrastructure. However, if the rising trends in AI- and Tech-based WIOA ETPLs (as documented in [Figure 6](#)) are to continue, it is possible that future years of WIOA programming may see a greater capacity for upskilling within technologically-focused occupations.

Figure 10: AIR Index: Decomposition by Illustrative 6-Digit SOC Examples

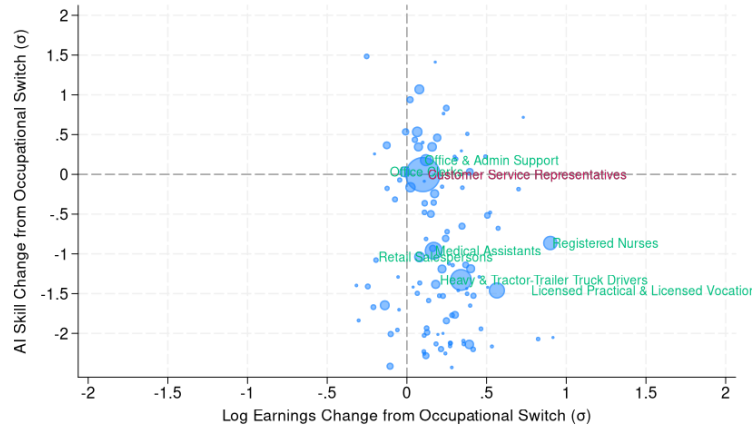
(A) Computer User Support Specialists (2-digit: Computer & Mathematics)



(B) Medical Assistants (2-digit: Healthcare Support)



(C) Customer Service Representatives (2-digit: Office & Admin Support)



Notes: Figures show flows in occupational AI exposure (y-axis) and earnings returns from training (x-axis) relative to nearest neighbor matched pairs, for three illustrative occupations (prior to program participation) along with destination occupations. Red text indicates most popular destination occupation for the origin 6-digit occupation.

While *Computer User Support Specialists* highlights the nuanced dynamics for high AI exposure workers who attempt to upskill, Panel (B) illustrates a different story that relates more to credentialing and possibly occupational licensing. *Medical Assistants* who enter WIOA training programs tend to experience relatively high earnings returns while holding AI skills constant by transitioning into higher level healthcare roles such as *Licensed Vocational Nurses* and *Registered Nurses*. This is consistent with a narrative of workers utilizing WIOA training resources to overcome credentialing barriers (also documented by [Jacobson and Davis \(2017\)](#)).

Finally in Panel (C), we show an example of a high AI-exposure occupation where the majority of trainees end up transitioning towards lower-exposed jobs. *Customer Service Representatives*, one of the largest occupational categories of WIOA trainees and one that has received much attention (see [Brynjolfsson et al. \(2025a\)](#)), are far more likely to downskill into roles in more physical professions such as *Truck Drivers and Movers*. This is consistent with an AI-avoidance narrative, where workers who are unable to successfully adapt in the face of AI pressure instead opt to move into more insulated professions. It is interesting to note the contrast between *Computer User Support Specialists* and *Customer Service Representatives*, as both are considered to be highly AI-exposed, yet show very different scopes for successful AI upskilling. Overall these decompositions illustrate that occupations can vary tremendously in their adaptability to AI, with some workers showing potential to deepen AI-complementary skills and others avoiding them.

6 Discussion

One of the most pressing policy questions today is the effects that AI technologies will have in the labor market, specifically their effects on workers. One way that workers can respond to these new technologies, thus far understudied, is to retrain themselves – to either reskill to use the technology, or reskill to avoid it. In the context of a virtually nonexistent AI re-skilling literature, our main contribution is to provide a national and large-scale lens into publicly funded AI reskilling efforts that goes beyond individual firm training studies.

How retrainable are AI-exposed workers? We find substantial but heterogeneous returns to training that vary systematically with initial AI exposure and have evolved over time. Workers from low AI-exposure occupations consistently experience the largest gains from training—on the order of roughly \$2,200 per quarter on average—primarily by upskilling into more AI-exposed, higher-paying work. In contrast, workers from high AI-exposure occupations initially earn smaller returns, but these returns grow markedly over time, rising from under \$1,000 per quarter prior to 2020 to levels that approach those

of low-exposure workers during the AI boom. These gains are closely tied to skill sorting: low-exposure workers tend to move up the AI exposure ladder, while high-exposure workers increasingly benefit from transitioning into less AI-exposed occupations. Taken together, the evidence suggests that successful retraining does not mean moving uniformly toward more AI-intensive work, but rather adjusting in different directions depending on where workers start.

Using a newly constructed AI retrainability index (AIR), we estimate that about 40-45% of occupations are “AI-retrainable”—workers leaving these occupations are compensated with higher earnings from moving up the AI skill ladder – a large share given our relatively low-income sample of displaced workers. Our evidence thus indicates that with appropriate investment in training infrastructure, workers can successfully adapt to AI technologies, lending support to more optimistic scenarios in which policy interventions enable adjustment to AI instead of displacement by it.

References

- Abadie, Alberto and Guido Imbens**, “Simple and bias-corrected matching estimators for average treatment effects,” 2002.
- , **David Drukker, Jane Leber Herr, and Guido W Imbens**, “Implementing matching estimators for average treatment effects in Stata,” *The stata journal*, 2004, 4 (3), 290–311.
- Abel, Jaison R, Richard Deitz, Natalia Emanuel, Benjamin Hyman, and Nick Montalbano**, “Are Businesses Scaling Back Hiring Due to AI?,” *Liberty Street Economics*, 2025.
- Acemoglu, Daron and David Autor**, “Skills, tasks and technologies: Implications for employment and earnings,” in “Handbook of labor economics,” Vol. 4, Elsevier, 2011, pp. 1043–1171.
- Andersson, Fredrik, Harry J Holzer, Julia I Lane, David Rosenblum, and Jeffrey Smith**, “Does Federally Funded Job Training Work?: Nonexperimental Estimates of WIA Training Impacts Using Longitudinal Data on Workers and Firms,” *Journal of Human Resources*, 2024, 59 (4), 1244–1283.
- Ashenfelter, Orley**, “Estimating the Effect of Training Programs on Earnings,” in Orley Ashenfelter and James Blum, eds., *Evaluating the Labor Market Effects of Social Programs*, Princeton University Press, 1978, pp. 95–124.
- Autor, David**, “Applying AI to rebuild middle class jobs,” Technical Report, National Bureau of Economic Research 2024.
- Autor, David H. and David Dorn**, “The Growth of Low-Skill Service Jobs and the Polarization of the U.S. Labor Market,” *American Economic Review*, 2013, 103 (5), 1553–1597.
- Bollinger, Christopher R. and Kenneth R. Troske**, “Evaluation of a New Job Training Program: Code Louisville,” *Journal of Policy Analysis and Management*, 2025, —.
- Bonney, Kathryn, Cory Breaux, Cathy Buffington, Emin Dinlersoz, Lucia S Foster, Nathan Goldschlag, John C Haltiwanger, Zachary Kroff, and Keith Savage**, “Tracking firm use of AI in real time: A snapshot from the Business Trends and Outlook Survey,” Technical Report, National Bureau of Economic Research 2024.
- Brynjolfsson, Erik, Bharat Chandar, and Ruyu Chen**, “Canaries in the Coal Mine? Six Facts about the Recent Employment Effects of Artificial Intelligence,” Working Paper, Stanford Digital Economy Lab 2025. Accessed 1 Feb 2026.
- , **Danielle Li, and Lindsey Raymond**, “Generative AI at work,” *The Quarterly Journal of Economics*, 2025, p. qjae044.
- , **Tom Mitchell, and Daniel Rock**, “What can machines learn and what does it mean for occupations and the economy?,” in “AEA papers and proceedings,” Vol. 108 American Economic Association 2014 Broadway, Suite 305, Nashville, TN 37203 2018, pp. 43–47.
- Card, David, Jochen Kluve, and Andrea Weber**, “What works? A meta analysis of recent active labor market program evaluations,” *Journal of the European Economic Association*, 2018, 16 (3), 894–931.
- Crump, Richard K., V. Joseph Hotz, Guido W. Imbens, and Oscar A. Mitnik**, “Dealing with Limited Overlap in Estimation of Average Treatment Effects,” *Biometrika*, 2009, 96 (1), 187–199.
- Delfino, Alexia, Andrea Garnero, Sergio Infrerra, Marco Leonardi, and Raffaella Sadun**, “Unwilling to Reskill? Experimental Evidence from Real-World Jobseekers,” Technical Report w34633, National Bureau of Economic Research 2026.
- Deming, David J, Christopher Ong, and Lawrence H Summers**, “Technological disruption

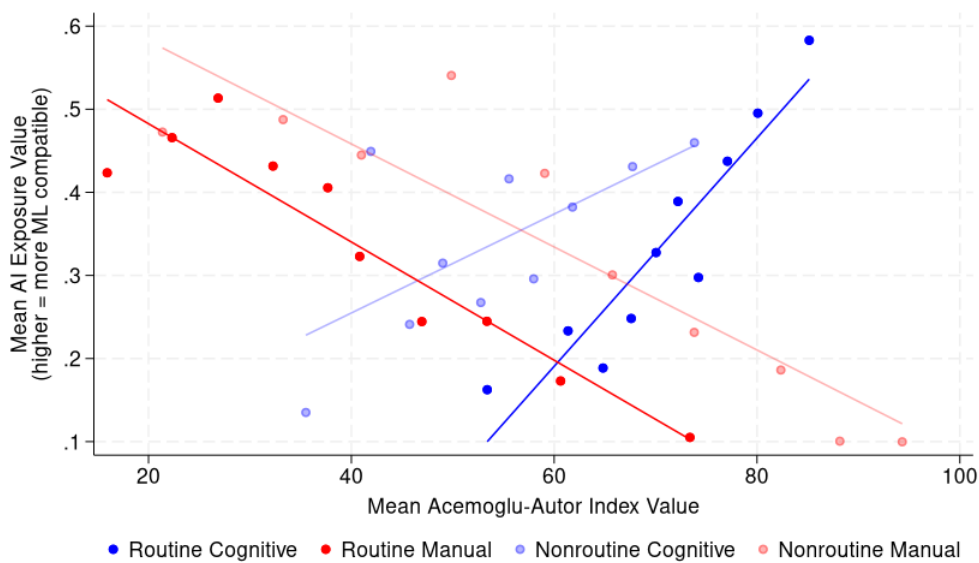
- in the labor market,” Technical Report, National Bureau of Economic Research 2025.
- Eloundou, Tyna, Sam Manning, Pamela Mishkin, and Daniel Rock**, “GPTs are GPTs: Labor market impact potential of LLMs,” *Science*, 2024, *384* (6702), 1306–1308.
- Hampole, Menaka, Dimitris Papanikolaou, Lawrence Schmidt, and Bryan Seegmiller**, “Artificial Intelligence and the Labor Market,” *Available at SSRN*, 2025.
- Handa, Kunal, Alex Tamkin, Miles McCain, Saffron Huang, Esin Durmus, Sarah Heck, Jared Mueller, Jerry Hong, Stuart Ritchie, Tim Belonax et al.**, “Which economic tasks are performed with ai? evidence from millions of claude conversations,” *arXiv preprint arXiv:2503.04761*, 2025.
- Hartley, Jonathan, Filip Jolevski, Vitor Melo, and Brendan Moore**, “The Labor Market Effects of Generative Artificial Intelligence,” *Available at SSRN*, 2024.
- Heckman, James J., Hidehiko Ichimura, Jeffrey A. Smith, and Petra E. Todd**, “Characterizing Selection Bias Using Experimental Data,” *Econometrica*, 1998, *66* (5), 1017–1098.
- Humlum, Anders and Emilie Vestergaard**, “The adoption of ChatGPT,” Technical Report, IZA Discussion Papers 2024.
- , **Jakob Munch, and Mette Rasmussen**, “What works for the unemployed? Evidence from quasi-random caseworker assignments,” 2023.
- Hyman, Ben, Ben Lahey, Karen Ni, and Laura Pilossoph**, “Job Training Mismatch,” 2025.
- Jacobs, Julian and Jordan Canedy**, “Do US Retraining Programs Protect Against Technological Automation,” Technical Report 2026.
- Jacobson, Louis and Jonathan Davis**, “The relative returns to workforce investment act–supported training in florida by field, gender, and education and ways to improve trainees’ choices,” *Journal of Labor Economics*, 2017, *35* (S1), S337–S375.
- Lechner, Michael, Ruth Miquel, and Conny Wunsch**, “Long-Run Effects of Public Sector Sponsored Training in West Germany,” *Journal of the European Economic Association*, 2011, *9* (4), 742–784.
- Manning, Sam and Juan Aguirre**, “Measuring U.S. Workers’ Capacity to Adapt to AI-Driven Job Displacement,” NBER Working Paper XXXX, National Bureau of Economic Research, Cambridge, MA 2026.
- Pizzinelli, Carlo, Augustus J Panton, Ms Marina Mendes Tavares, Mauro Cazzaniga, and Longji Li**, *Labor market exposure to AI: Cross-country differences and distributional implications*, International Monetary Fund, 2023.
- Rothstein, Jesse, Robert Santillano, Till Von Wachter, Wahid Khan, and Mary Yang**, “CAAL-Skills: Study of Workforce Training Programs in California,” 2022.
- Şahin, Aysegül, Joseph Song, Giorgio Topa, and Giovanni L Violante**, “Mismatch unemployment,” *American Economic Review*, 2014, *104* (11), 3529–3564.

Appendices

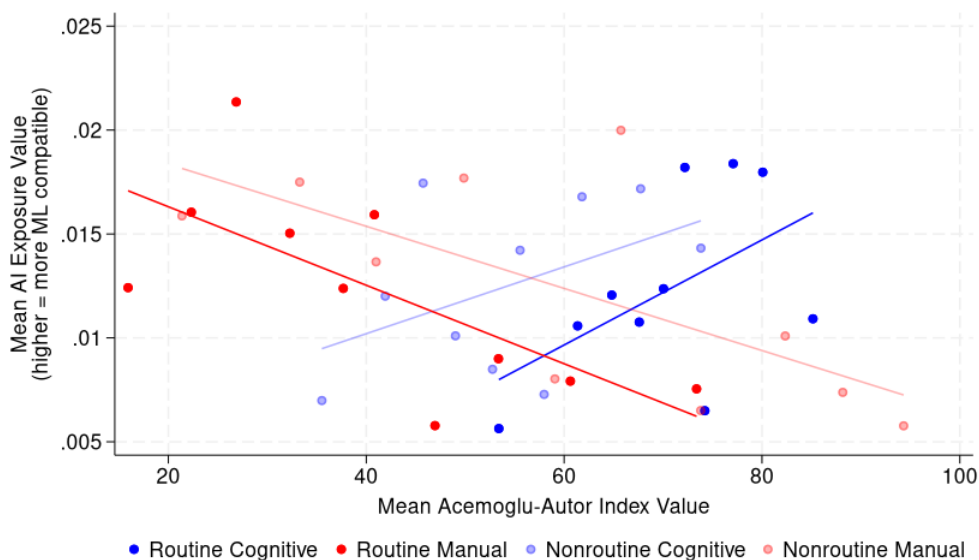
Appendix A. Additional Results

Figure A.1: AI Exposure Correlations with Acemoglu and Autor (2011) Task Measures

(A) Eloundou AI Exposure v. AA Index

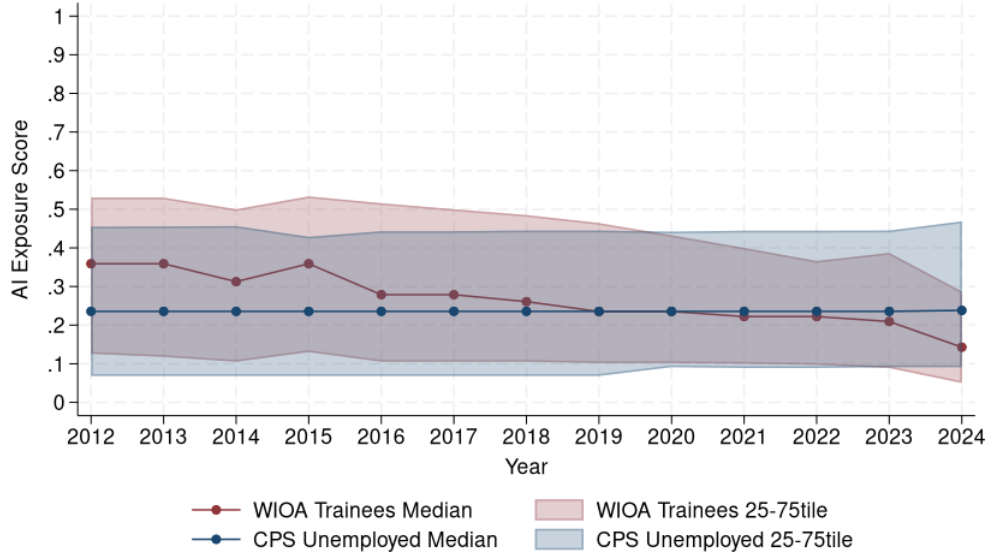


(B) Claude AI Exposure vs. AA Index



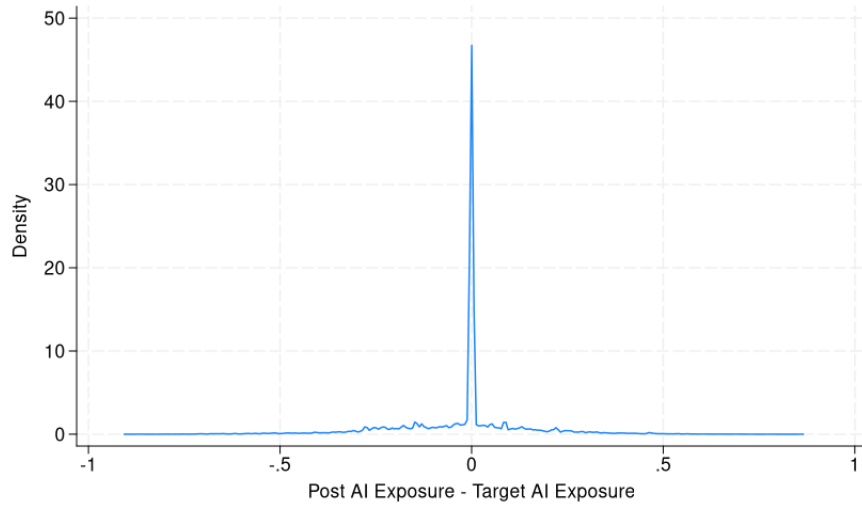
Notes: The x-axis in both panels measures the mean [Acemoglu and Autor \(2011\)](#) (AA) task share for each 6-digit occupation. The y-axis in panel (A) is our main AI exposure measure—[Eloundou et al. \(2024\)](#) at the 6-digit occupation level. The y-axis for panel (B) represents the secondary AI exposure measure based on Claude task usage from Anthropic’s Economic Index ([Handa et al., 2025](#)). Correlations between AA task shares and AI exposure measures are expressed as binscatter plots.

Figure A.2: WIOA and CPS-Unemployed AI Exposure Distributions Over Time



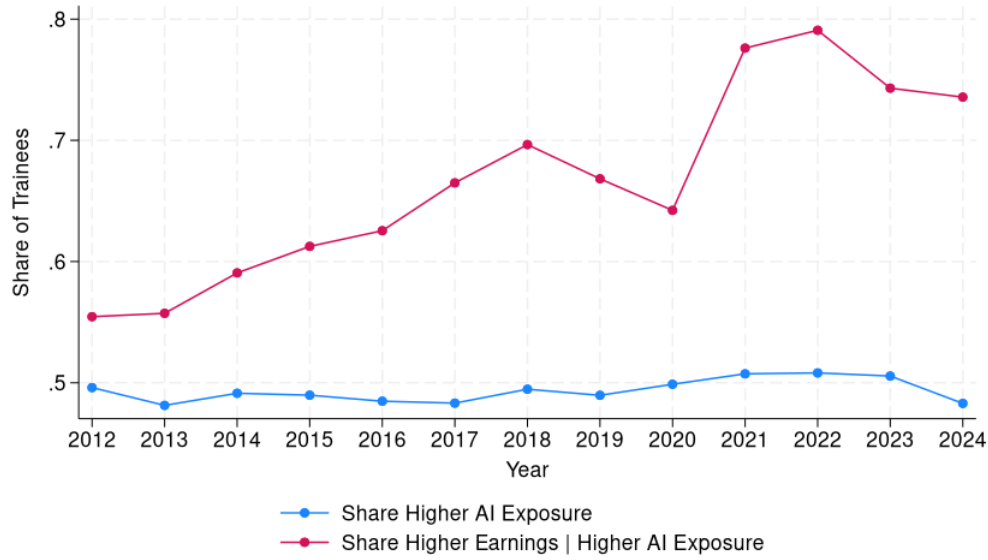
Notes: Figure plots the inter-quartile range of occupation AI exposure scores (Eloundou et al. (2024)) for WIOA/WIA trainees' occupations prior to training, and CPS unemployed workers' occupations prior to their current unemployment spell.

Figure A.3: WIOA Target vs. Destination Difference in AI Exposure



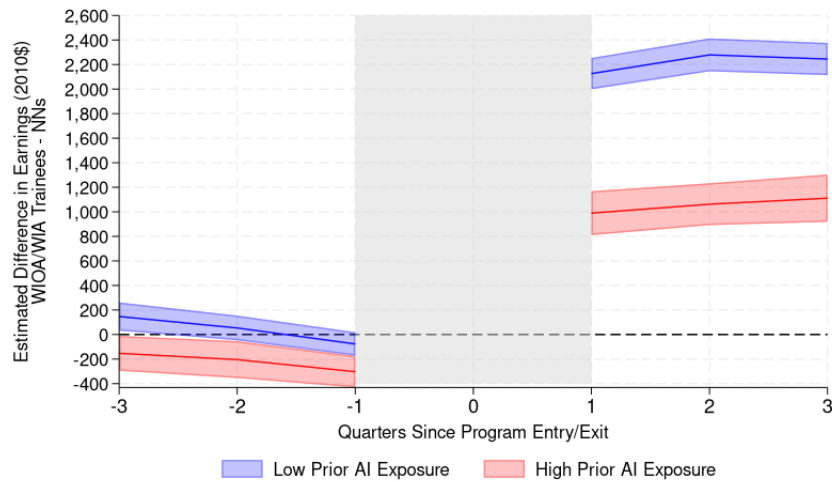
Notes: Figure plots the kernel densities of the difference in occupation AI exposure scores (Eloundou et al. (2024)) for the destination occupation and target training occupation of WIOA job training participants, where differences are calculated at the trainee level as $AI_{post} - AI_{target}$. We utilize an Epanechnikov kernel function and limit the sample to WIOA trainees within our main analysis sample who have non-missing target and post-WIOA SOC occupation information.

Figure A.4: Share of Trainees With Positive AI Skill and Earnings Changes over Time



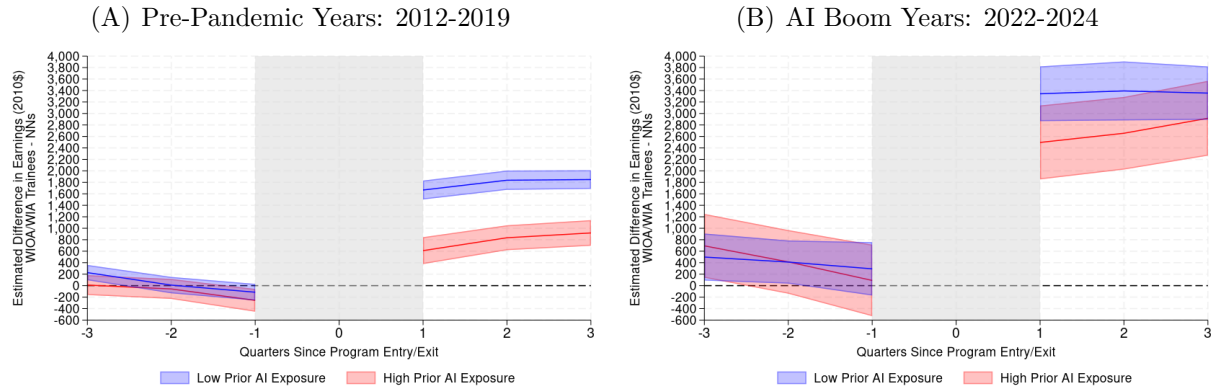
Notes: Figure plots share of WIOA/WIA training participants moving to higher AI exposure occupations (Eloundou et al. (2024)) after training by calendar year of participation, and jointly receiving higher earnings from moving to higher AI exposure values.

Figure A.5: Nearest-Neighbor Matched Returns to WIOA/WIA: Pooled Across Years: Robustness to Including Zero Earnings



Notes: Figure is equivalent to main results Figure 3, but includes observations in which earnings are recorded as zero. When including zero earnings observations, we additionally exact-match on employment in the quarter prior to program entry to ensure comparisons are within workers with similar employment histories.

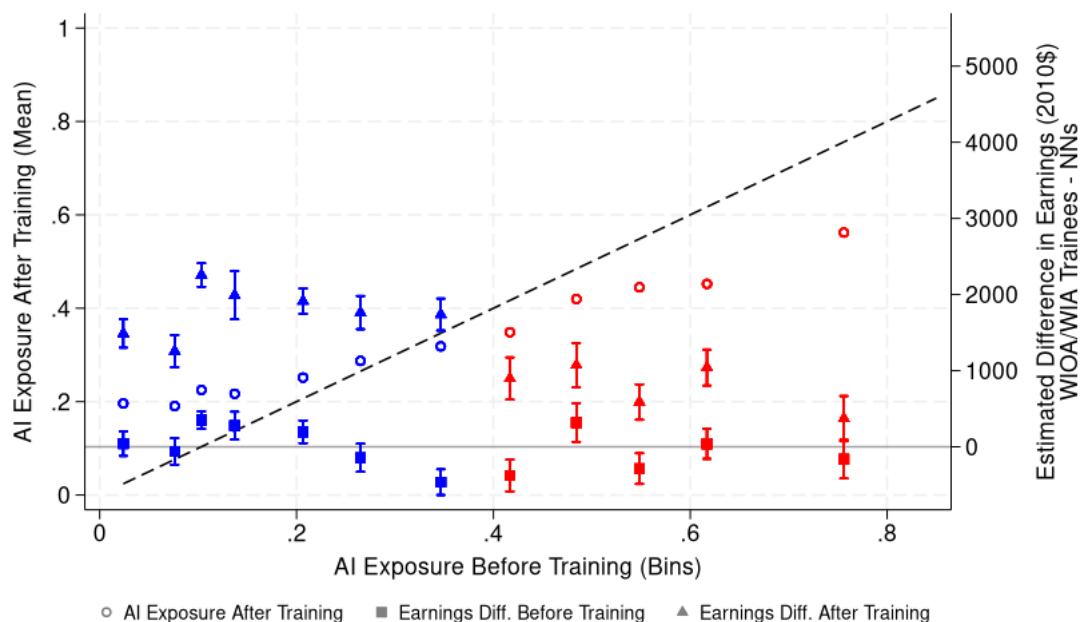
Figure A.6: NN-Matched Earnings Returns to Job Training by Prior AI Exposure, 2012-2019 and 2022-2024



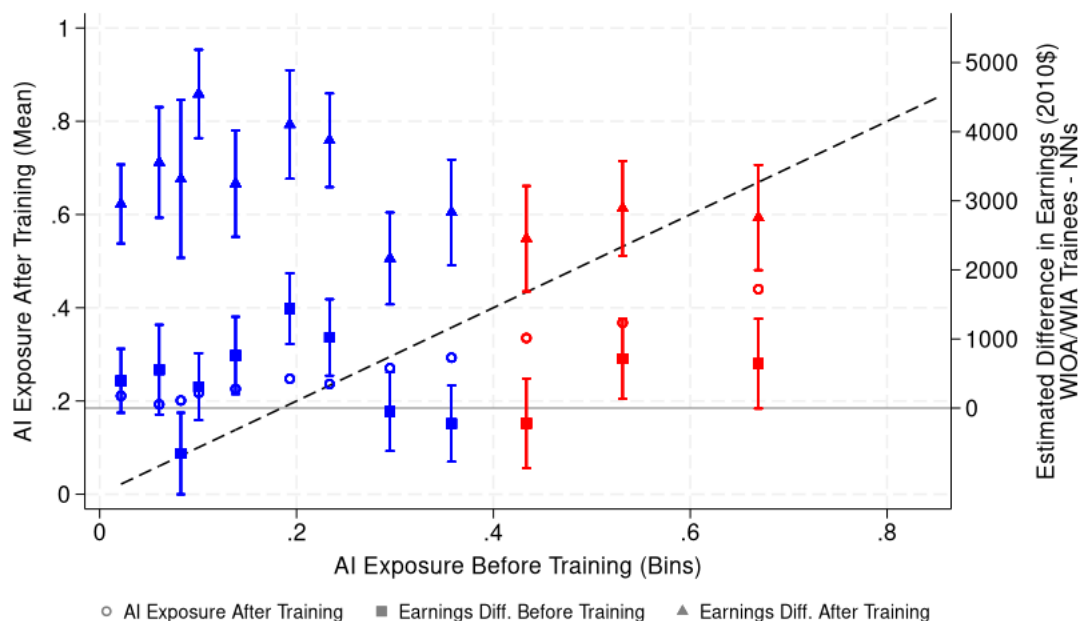
Notes: AI exposure is based on [Eloundou et al. \(2024\)](#) occupation-level AI exposure scores. Each series represents the earnings return relative to nearest-neighbor matched twins weighted by inverse distance in matches. The 55th percentile is chosen as the preferred cutoff for high and low based on estimate stability beyond this cutoff (see text for details).

Figure A.7: NN-Matched Earnings Returns to Job Training by Prior AI Exposure, Without Matching on Earnings

(A) Pre-Pandemic Years: 2012-2019

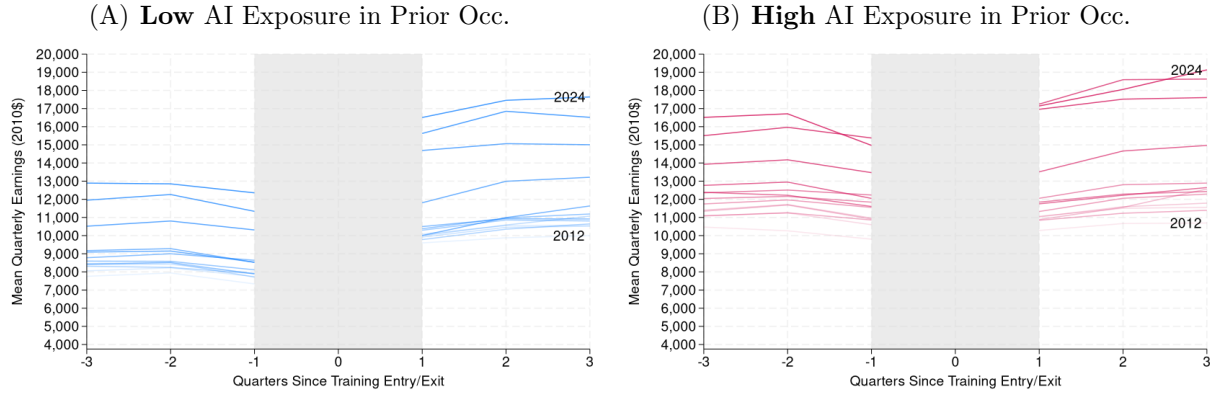


(B) AI Boom Years: 2022-2024



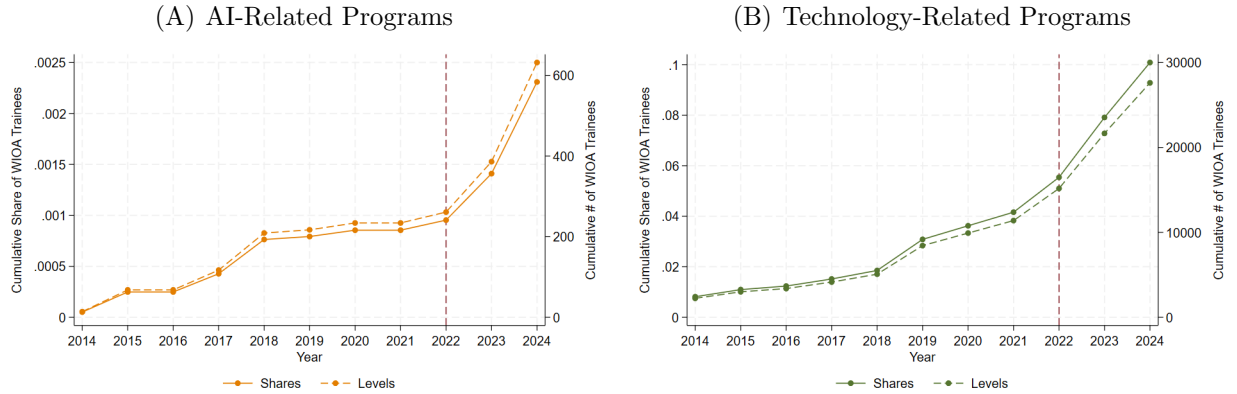
Notes: Figure overlays pooled estimates for the returns to training across the AI exposure spectrum (triangles), placebo estimates for pooled estimates prior to entry into training/job search services for the control group (squares), and average movements in AI skill before and after training (circles). Circles above the 45 degree line represent AI upskilling, while below the line represent AI downskilling. Red and blue correspond to our prior preferred high and low exposure distinction. A horizontal gridline is drawn at zero earnings returns. AI exposure uses the measure from [Eloundou et al. \(2024\)](#).

Figure A.8: Earnings Patterns by AI Exposure Prior to Job Training, NN-Matched Control Group



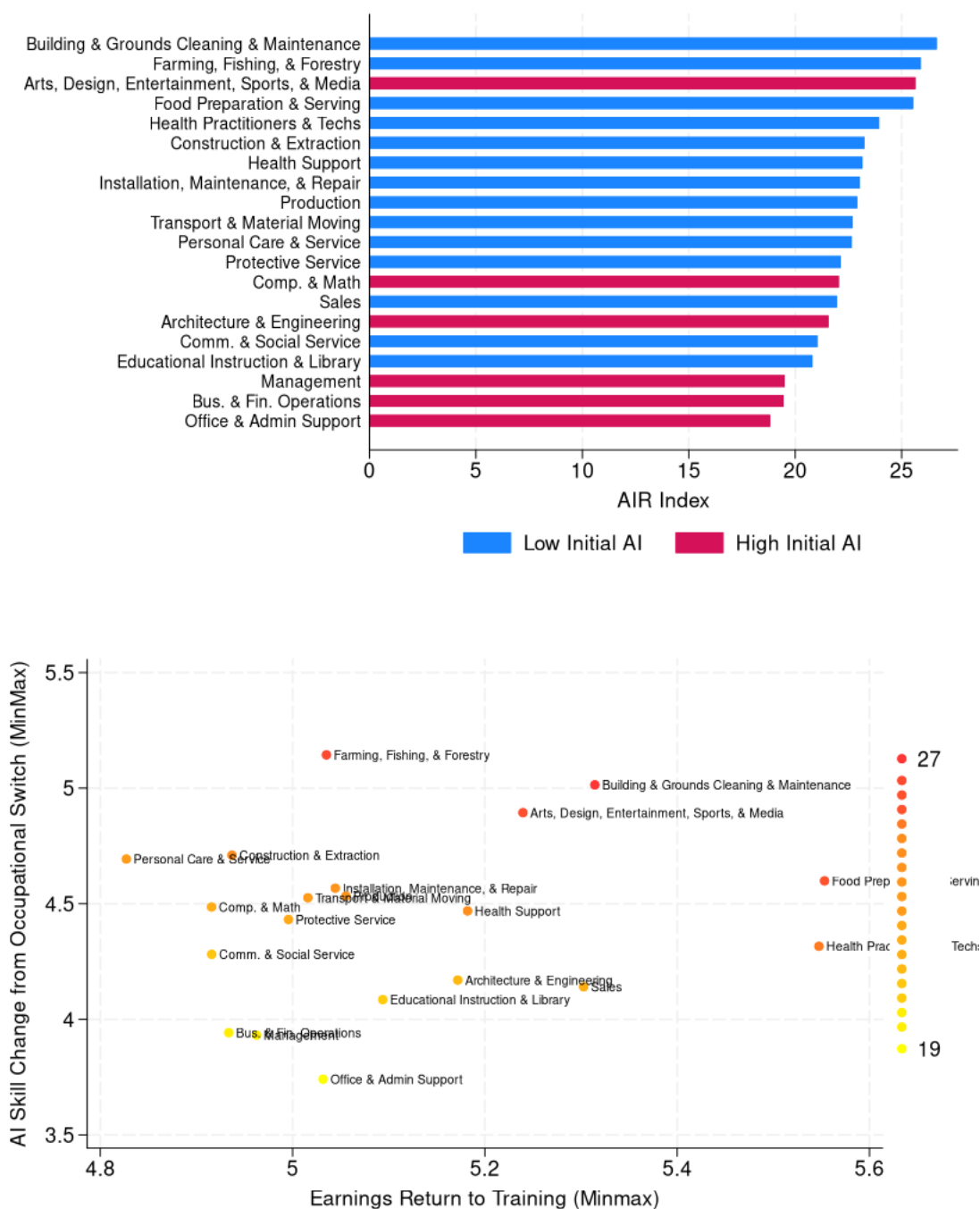
Notes: Plots are computed for Wagner-Peyser matched nearest neighbors. Participants are classified as “high” AI exposure if their 6-digit SOC exposure score is above the 55th percentile of the [Eloundou et al. \(2024\)](#) 6-digit AI exposure score distribution, and “low” if below the 55th percentile. Heavier lines correspond to more recent years. The shaded gray region emphasizes that quarterly earnings are observed in the quarters prior to entry, and after exit, but not during program receipt. All earnings are in real 2010 dollars.

Figure A.9: Cumulative Growth in U.S. AI- and Technology-Related Job Training Participants



Notes: Panel A shows AI-related programs (e.g. keywords: *Artificial Intelligence, Natural Language Processing, Machine Learning*), and Panel B shows non-AI general Technology-related programs (e.g. keywords: *Programming, Web Development, Python*) drawn from WIOA ETPL lists over time, and weighted by the number of WIOA job training participants enrolled. Shares are calculated as the cumulative number of AI or Technology programs (trainee-weighted) introduced over the cumulative number of all WIOA ETPL programs (trainee-weighted) within a given year. See [Appendix C](#) for details.

Figure A.10: AI Retraining Index during the AI Boom (2022-2024)- MinMax Version



Notes: Panel (A) shows 2-digit rankings of AI retrainability index (AIR) using [Eloundou et al. \(2024\)](#) AI exposure measure, where higher values reflect joint movements up in earnings AI skill space, inclusive of all trainees who entered WIOA training programs between 2022q1 and 2024q4. The functional form of the index utilizes a MinMax transformation for both the earnings returns and AI skill changes (see [Equation 4](#)). Red and blue bars indicate high and low initial AI exposure prior to training. Panel (B) shows a decomposition with a heat map for higher index values, in which trainees leaving occupations above the horizontal dashed line are moving up in AI skill space.

Table A.1: Summary Statistics by AI Exposure Prior to Training Participation

	High Prior AI Exposure		Low Prior AI Exposure		High AI - Low AI	
	Mean/SD (1)	# Training Spells (2)	Mean/SD (3)	# Training Spells (4)	Δ /SE (5)	% Diff (6)
A. Demographic						
Female	0.60 [0.49]	43,472	0.45 [0.50]	70,463	0.15 (0.0030)	25.2
Age	41.7 [11.1]	43,472	36.8 [10.4]	70,463	4.93 (0.065)	11.8
Asian	0.043 [0.20]	43,472	0.026 [0.16]	70,463	0.017 (0.0011)	40.3
Black	0.23 [0.42]	43,472	0.30 [0.46]	70,463	-0.072 (0.0027)	-31.0
White	0.61 [0.49]	43,472	0.56 [0.50]	70,463	0.051 (0.0030)	8.36
Hispanic/Latino	0.16 [0.36]	43,472	0.16 [0.37]	70,463	-0.0039 (0.0022)	-2.49
Disability Status	0.030 [0.17]	43,472	0.026 [0.16]	70,463	0.0038 (0.0010)	12.6
No HS Diploma/GED	0.026 [0.16]	43,472	0.058 [0.23]	70,463	-0.032 (0.0013)	-125.1
HS Diploma/GED	0.37 [0.48]	43,472	0.55 [0.50]	70,463	-0.18 (0.0030)	-49.8
Some College	0.19 [0.40]	43,472	0.19 [0.39]	70,463	0.0046 (0.0024)	2.38
College Degree Plus	0.41 [0.49]	43,472	0.20 [0.40]	70,463	0.21 (0.0027)	51.1
B. Social Benefits						
Disadvantaged Adult Status	0.40 [0.49]	43,472	0.63 [0.48]	70,463	-0.24 (0.0030)	-60.4
Dislocated Worker Status	0.61 [0.49]	43,472	0.36 [0.48]	70,463	0.25 (0.0029)	41.2
Low-income Status	0.38 [0.49]	43,472	0.49 [0.50]	70,463	-0.11 (0.0030)	-28.2
TANF Recipient	0.0082 [0.090]	43,472	0.011 [0.10]	70,463	-0.0026 (0.00060)	-32.1
SNAP Recipient	0.18 [0.38]	30,306	0.25 [0.44]	54,994	-0.074 (0.0030)	-40.7
Other Public Assistance Recipient	0.22 [0.41]	32,654	0.29 [0.45]	60,803	-0.070 (0.0030)	-32.3
C. Pre-Separation Employment						
Employed at Participation	0.22 [0.41]	43,472	0.41 [0.49]	70,463	-0.19 (0.0028)	-88.1
Real Earnings 1Q Before	11605.1 [9707.6]	43,472	8433.5 [6692.1]	70,463	3171.7 (48.7)	27.3
Real Earnings 2Q Before	12270.6 [8819.5]	43,472	9133.4 [6403.5]	70,463	3137.2 (45.2)	25.6
Real Earnings 3Q Before	12132.9 [8628.1]	43,472	8979.0 [6301.5]	70,463	3153.8 (44.4)	26.0
Quarters in WIOA/WIA	5.17 [3.45]	43,472	4.55 [3.23]	70,463	0.62 (0.020)	12.0
Observations	43,472		70,463			

Notes: Table presents means, standard deviations, and sample counts for main sample of job training participants by High and Low AI exposure prior to job training participation (see text for details). The pooled sample of trainees contains 113,935 participation spells. T-test standard errors in column (5) are adjusted for unequal variances (heteroskedasticity) across groups. Participants are weighted by the inverse distance to their nearest-neighbor from Wagner-Peyser used for results estimation.

Table A.2: Attrition Balance for WIOA/WIA Trainees with Missing Occupation Codes

	Nonmissing Occ Codes		Missing Occ Codes		Nonmiss - Miss	
	Mean/SD (1)	# Training Spells (2)	Mean/SD (3)	# Training Spells (4)	Δ /SE (5)	% Diff (6)
A. Demographic						
Female	0.51 [0.50]	113,935	0.51 [0.50]	499,779	-0.00050 (0.0016)	-0.099
Age	38.6 [11.0]	113,935	37.9 [11.0]	418,907	0.71 (0.037)	1.84
Asian	0.041 [0.20]	90,723	0.035 [0.18]	428,960	0.0059 (0.00068)	14.4
Black	0.33 [0.47]	94,595	0.32 [0.47]	433,178	0.012 (0.0017)	3.73
White	0.67 [0.47]	98,637	0.63 [0.48]	447,855	0.035 (0.0017)	5.27
Hispanic/Latino	0.17 [0.37]	108,709	0.16 [0.37]	451,859	0.0043 (0.0012)	2.57
Disability Status	0.029 [0.17]	111,190	0.028 [0.17]	471,451	0.00052 (0.00055)	1.81
No HS Diploma/GED	0.045 [0.21]	113,935	0.059 [0.24]	499,779	-0.014 (0.00076)	-30.0
HS Diploma/GED	0.48 [0.50]	113,935	0.47 [0.50]	499,779	0.010 (0.0016)	2.09
Some College	0.19 [0.39]	113,935	0.19 [0.39]	499,779	0.0038 (0.0013)	2.01
College Degree Plus	0.28 [0.45]	113,935	0.28 [0.45]	499,779	0.0028 (0.0015)	0.98
B. Social Benefits						
Disadvantaged Adult Status	0.54 [0.50]	113,935	0.65 [0.48]	499,779	-0.11 (0.0016)	-20.4
Dislocated Worker Status	0.46 [0.50]	113,935	0.37 [0.48]	499,779	0.086 (0.0016)	18.9
Low-income Status	0.45 [0.50]	113,935	0.44 [0.50]	468,234	0.015 (0.0016)	3.35
TANF Recipient	0.011 [0.11]	98,429	0.014 [0.12]	421,787	-0.0021 (0.00040)	-18.9
SNAP Recipient	0.23 [0.42]	85,300	0.22 [0.41]	267,580	0.0087 (0.0016)	3.83
Other Public Assistance Recipient	0.26 [0.44]	93,457	0.22 [0.42]	399,977	0.040 (0.0015)	15.0
C. Pre-Separation Employment						
Employed at Participation	0.34 [0.47]	113,935	0.41 [0.49]	499,779	-0.070 (0.0016)	-20.7
Real Earnings 1Q Before	9643.6 [8125.7]	113,935	9874.9 [13027.1]	499,779	-231.2 (40.3)	-2.40
Real Earnings 2Q Before	10330.4 [7573.6]	113,935	10352.2 [12133.6]	499,779	-21.8 (37.5)	-0.21
Real Earnings 3Q Before	10182.4 [7437.0]	113,935	10129.9 [11143.0]	499,779	52.5 (34.6)	0.52
Quarters in WIOA/WIA	4.79 [3.33]	113,935	4.36 [3.14]	499,779	0.43 (0.010)	8.91
Observations	113,935		499,779			

Notes: Columns (1) and (2) show summary statistics for the main analysis sample of job trainees (N=113,935 training spells). Columns (3) and (4) present summary statistics for training spells with partially missing occupation data either prior to training entry or after training exit. T-test standard errors in column (5) are adjusted for unequal variances (heteroskedasticity) across groups.

Table A.3: Balance Table: Trainees vs. Matched WP Neighbors (Low Initial AI)

	Trainees		Matched Wagner Peyser		Trainees - Wagner Peyser	
	Mean/SD (1)	# Spells (2)	Mean/SD (3)	# Spells (4)	Δ /SE (5)	% Diff (6)
A. Demographic						
Female	0.45 [0.50]	70,463	0.47 [0.50]	70,463	-0.022 (0.0027)	-4.84
Age	36.8 [10.4]	70,463	38.1 [10.7]	70,463	-1.30 (0.056)	-3.54
Asian	0.026 [0.16]	70,463	0.020 [0.14]	70,463	0.0059 (0.00080)	22.7
Black	0.30 [0.46]	70,463	0.36 [0.48]	70,463	-0.056 (0.0025)	-18.5
White	0.56 [0.50]	70,463	0.53 [0.50]	70,463	0.026 (0.0027)	4.70
Hispanic/Latino	0.16 [0.37]	70,463	0.11 [0.32]	70,463	0.047 (0.0018)	29.2
Disability Status	0.026 [0.16]	70,463	0.020 [0.14]	70,463	0.0060 (0.00081)	22.7
No HS Diploma/GED	0.058 [0.23]	70,463	0.048 [0.21]	70,463	0.0096 (0.0012)	16.7
HS Diploma/GED	0.55 [0.50]	70,463	0.59 [0.49]	70,463	-0.044 (0.0026)	-7.93
Some College	0.19 [0.39]	70,463	0.18 [0.38]	70,463	0.0087 (0.0021)	4.61
College Degree Plus	0.20 [0.40]	70,463	0.18 [0.38]	70,463	0.026 (0.0021)	12.9
B. Social Benefits						
Disadvantaged Adult Status	0.63 [0.48]	70,463	0.53 [0.50]	70,463	0.11 (0.0026)	16.7
Dislocated Worker Status	0.36 [0.48]	70,463	0.32 [0.47]	70,463	0.037 (0.0025)	10.3
Low-income Status	0.49 [0.50]	70,463	0.46 [0.50]	70,463	0.029 (0.0027)	5.92
TANF Recipient	0.011 [0.10]	70,463	0.0072 [0.084]	70,463	0.0036 (0.00050)	33.7
SNAP Recipient	0.25 [0.44]	54,994	0.20 [0.40]	55,512	0.059 (0.0025)	23.1
Other Public Assistance Recipient	0.29 [0.45]	60,803	0.28 [0.45]	51,772	0.0092 (0.0027)	3.19
C. Pre-Separation Employment						
Employed at Participation	0.41 [0.49]	70,463	0.41 [0.49]	70,463	0 (0.0026)	0
Real Earnings 1Q Before	8433.5 [6692.1]	70,463	8532.1 [6598.6]	70,463	-98.6 (35.4)	-1.17
Real Earnings 2Q Before	9133.4 [6403.5]	70,463	9113.6 [6635.8]	70,463	19.8 (34.7)	0.22
Real Earnings 3Q Before	8979.0 [6301.5]	70,463	8806.9 [6615.1]	70,463	172.1 (34.4)	1.92
Quarters in WIOA/WP	4.55 [3.23]	70,463	2.83 [2.45]	70,463	1.72 (0.015)	37.8
Observations	70,463		23,962			

Notes: Columns (1) through (4) show summary statistics for low initial AI exposure trainees (using [Eloundou et al. \(2024\)](#)) and nearest neighbor matches, while columns (5) and (6) estimate t-tests and standard errors adjusted for unequal variances (heteroskedasticity) across groups. Participants are weighted by the inverse distance to their nearest-neighbor from Wagner-Peyser. **A-10**

Table A.4: Balance Table: Trainees vs. Matched WP Neighbors (High Initial AI)

	Trainees		Matched Wagner Peyser		Trainees - Wagner Peyser	
	Mean/SD (1)	# Spells (2)	Mean/SD (3)	# Spells (4)	Δ /SE (5)	% Diff (6)
A. Demographic						
Female	0.60 [0.49]	43,472	0.63 [0.48]	43,472	-0.029 (0.0033)	-4.82
Age	41.7 [11.1]	43,472	43.1 [10.9]	43,472	-1.42 (0.075)	-3.39
Asian	0.043 [0.20]	43,472	0.035 [0.18]	43,472	0.0080 (0.0013)	18.4
Black	0.23 [0.42]	43,472	0.23 [0.42]	43,472	-0.0035 (0.0029)	-1.53
White	0.61 [0.49]	43,472	0.63 [0.48]	43,472	-0.026 (0.0033)	-4.31
Hispanic/Latino	0.16 [0.36]	43,472	0.12 [0.33]	43,472	0.032 (0.0024)	20.4
Disability Status	0.030 [0.17]	43,472	0.025 [0.16]	43,472	0.0052 (0.0011)	17.2
No HS Diploma/GED	0.026 [0.16]	43,472	0.021 [0.14]	43,472	0.0046 (0.0010)	17.8
HS Diploma/GED	0.37 [0.48]	43,472	0.37 [0.48]	43,472	0.0013 (0.0033)	0.35
Some College	0.19 [0.40]	43,472	0.18 [0.39]	43,472	0.010 (0.0027)	5.34
College Degree Plus	0.41 [0.49]	43,472	0.43 [0.49]	43,472	-0.016 (0.0033)	-3.78
B. Social Benefits						
Disadvantaged Adult Status	0.40 [0.49]	43,472	0.33 [0.47]	43,472	0.060 (0.0033)	15.2
Dislocated Worker Status	0.61 [0.49]	43,472	0.60 [0.49]	43,472	0.0090 (0.0033)	1.48
Low-income Status	0.38 [0.49]	43,472	0.35 [0.48]	43,472	0.039 (0.0033)	10.1
TANF Recipient	0.0082 [0.090]	43,472	0.0058 [0.076]	43,472	0.0024 (0.00056)	29.5
SNAP Recipient	0.18 [0.38]	30,306	0.15 [0.35]	30,547	0.034 (0.0030)	18.8
Other Public Assistance Recipient	0.22 [0.41]	32,654	0.20 [0.40]	29,821	0.019 (0.0033)	8.91
C. Pre-Separation Employment						
Employed at Participation	0.22 [0.41]	43,472	0.22 [0.41]	43,472	0 (0.0028)	0
Real Earnings 1Q Before	11605.1 [9707.6]	43,472	11944.3 [9294.4]	43,472	-339.2 (64.5)	-2.92
Real Earnings 2Q Before	12270.6 [8819.5]	43,472	12347.0 [8585.9]	43,472	-76.5 (59.0)	-0.62
Real Earnings 3Q Before	12132.9 [8628.1]	43,472	12159.0 [8441.9]	43,472	-26.1 (57.9)	-0.21
Quarters in WIOA/WP	5.17 [3.45]	43,472	3.21 [2.59]	43,472	1.96 (0.021)	38.0
Observations	43,472		18,195			

Notes: Columns (1) through (4) show summary statistics for high initial AI exposure trainees (using [Eloundou et al. \(2024\)](#)) and nearest neighbor matches, while columns (5) and (6) estimate t-tests and standard errors adjusted for unequal variances (heteroskedasticity) across groups. Participants are weighted by the inverse distance to their nearest-neighbor from Wagner-Peyser used for results estimation.

Table A.5: Regression Estimates for Quarterly Earnings Returns to WIOA/WIA Training

	Prior AI Exposure (Eloundou)		Prior AI Exposure (Claude)	
	High	Low	High	Low
A. Pooled Post:	1,041 (97.29)	2,207 (71.65)	1,480 (84.84)	2,113 (81.91)
B. Pooled Pre (Placebo):	-123 (71.74)	38 (54.41)	2 (66.94)	35 (57.58)
N	43,472	70,463	50,966	62,257
Mean Trainee Quarterly Earnings:	12,002	8,849	11,086	9,227

Notes: Each cell reports a separate mean quarterly earnings regression estimate. All estimates are weighted by the inverse of the Mahalanobis distance for each nearest-neighbor matched pair. Mean baseline earnings reflect average quarterly earnings pooled across three quarters prior to WIOA/WIA participation. Prior AI exposure is the [Eloundou et al. \(2024\)](#) exposure of the occupation worked prior to WIOA/WIA participation. Standard errors are clustered by nearest-neighbor and reported in parentheses.

Table A.6: Decomposition of Earnings Returns to Training: High AI Workers

	$\hat{\beta}_{g\tau}$	$\hat{\sigma}_{g\tau}$	$\hat{\gamma}_{g\tau}$	Earnings Effect $\hat{\beta}_{g\tau} \times \hat{\sigma}_{g\tau}$	Sorting Effect $\hat{\beta}_{g\tau} \times \hat{\gamma}_{g\tau}$	Total
<i>Panel A: 2012-2019</i>						
Upskill ($\Delta AI > 0$)	691.10 (210.27)	0.25 (0.005)	-0.01 (0.005)	174.67 (53.31)	-10.25 (4.92)	164.42
Same ($\Delta AI = 0$)	1,220.36 (328.34)	0.22 (0.005)	-0.05 (0.006)	267.86 (72.16)	-65.98 (17.77)	201.89
Downskill ($\Delta AI < 0$)	854.35 (91.25)	0.53 (0.006)	0.07 (0.006)	450.90 (48.68)	58.86 (8.26)	509.76
$\Sigma_g(\cdot)$ (RHS)				893.43	-17.37	876.06
$\hat{\beta}_{i\tau}$ (LHS)						877.10
Percent of Total				102.00%	-2.00%	
<i>Panel B: 2022-2024</i>						
Upskill ($\Delta AI > 0$)	2,334.54 (683.78)	0.21 (0.013)	-0.03 (0.015)	483.75 (144.56)	-65.37 (38.13)	418.38
Same ($\Delta AI = 0$)	2,461.77 (729.17)	0.29 (0.016)	-0.19 (0.016)	723.20 (222.78)	-460.92 (147.09)	262.28
Downskill ($\Delta AI < 0$)	3,058.62 (316.73)	0.50 (0.017)	0.22 (0.018)	1,526.29 (166.69)	658.32 (86.83)	2,184.60
$\Sigma_g(\cdot)$ (RHS)				2,733.24	132.03	2,865.27
$\hat{\beta}_{i\tau}$ (LHS)						2,867.74
Percent of Total				95.39%	4.61%	

Notes: This table shows the estimated components of workers' earnings returns to training, decomposed following Equation 2, for High AI workers who entered WIOA between 2012q1 and 2024q4. Each cell represents the estimated component for workers who fall into each of the three AI reskilling groups g . Standard errors, clustered at the control-id level, are presented in parentheses under each estimate for $\hat{\beta}_{g\tau}$, $\hat{\sigma}_{g\tau}$ and $\hat{\gamma}_{g\tau}$. Standard errors for $\hat{\beta}_{g\tau} \times \hat{\sigma}_{g\tau}$ and $\hat{\beta}_{g\tau} \times \hat{\gamma}_{g\tau}$ are estimated using the delta method. The overall Earnings Effect is the product of $\hat{\beta}_{g\tau} \times \hat{\sigma}_{g\tau}$, while the Sorting Effect is the product of $\hat{\beta}_{g\tau} \times \hat{\gamma}_{g\tau}$. The table also shows that the aggregate total of the decomposed effects are very close to the overall earnings returns to training $\beta_{i\tau}$ obtained from Equation 1.

Table A.7: Decomposition of Earnings Returns to Training: Low AI Workers

	$\hat{\beta}_{g\tau}$	$\hat{\sigma}_{g\tau}$	$\hat{\gamma}_{g\tau}$	Earnings Effect $\hat{\beta}_{g\tau} \times \hat{\sigma}_{g\tau}$	Sorting Effect $\hat{\beta}_{g\tau} \times \hat{\gamma}_{g\tau}$	Total
<i>Panel A: 2012-2019</i>						
Upskill ($\Delta AI > 0$)	2,182.20 (79.32)	0.39 (0.006)	0.16 (0.007)	850.83 (32.06)	346.70 (20.17)	1,197.53
Same ($\Delta AI = 0$)	1,563.60 (129.91)	0.35 (0.006)	-0.14 (0.006)	544.01 (45.95)	-214.98 (20.18)	329.04
Downskill ($\Delta AI < 0$)	1,323.34 (87.44)	0.26 (0.005)	-0.02 (0.006)	346.95 (23.98)	-28.30 (7.69)	318.65
$\Sigma_g(\cdot)$ (RHS)				1,741.80	103.42	1,845.22
$\hat{\beta}_{i\tau}$ (LHS)						1845.26
Percent of Total				94.40%	5.60%	
<i>Panel B: 2022-2024</i>						
Upskill ($\Delta AI > 0$)	3,818.26 (228.88)	0.33 (0.011)	0.23 (0.011)	1,258.18 (83.47)	869.47 (69.27)	2,127.66
Same ($\Delta AI = 0$)	2,112.80 (294.87)	0.47 (0.012)	-0.29 (0.012)	1,003.01 (143.61)	-604.00 (89.25)	399.01
Downskill ($\Delta AI < 0$)	2,935.07 (261.22)	0.20 (0.009)	0.06 (0.009)	574.55 (57.01)	170.71 (31.73)	745.26
$\Sigma_g(\cdot)$ (RHS)				2,835.75	436.18	3,271.93
$\hat{\beta}_{i\tau}$ (LHS)						3,272.23
Percent of Total				86.67%	13.33%	

Notes: This table shows the estimated components of workers' earnings returns to training, decomposed following Equation 2, for Low AI workers who entered WIOA between 2012q1 and 2024q4. Each cell represents the estimated component for workers who fall into each of the three AI reskilling groups g . Standard errors, clustered at the control-id level, are presented in parentheses under each estimate for $\hat{\beta}_{g\tau}$, $\hat{\sigma}_{g\tau}$ and $\hat{\gamma}_{g\tau}$. Standard errors for $\hat{\beta}_{g\tau} \times \hat{\sigma}_{g\tau}$ and $\hat{\beta}_{g\tau} \times \hat{\gamma}_{g\tau}$ are estimated using the delta method. The overall Earnings Effect is the product of $\hat{\beta}_{g\tau} \times \hat{\sigma}_{g\tau}$, while the Sorting Effect is the product of $\hat{\beta}_{g\tau} \times \hat{\gamma}_{g\tau}$. The table also shows that the aggregate total of the decomposed effects are very close to the overall earnings returns to training $\beta_{i\tau}$ obtained from Equation 1.

Table A.8: Popular Occupation Transitions

	Origin	Destination	Avg. Return (2010\$)
<i>Panel A: 2012–2019</i>			
Low Exposure Downskill	Personal Care Aides	Nursing Assistants	941
	Light Truck Drivers	Heavy Truck Drivers	1,789
	Driver/Sales Workers	Heavy Truck Drivers	3,888
Low Exposure Upskill	Nursing Assistants	LPN/LVN Nurses	4,699
	Nursing Assistants	Registered Nurses	8,022
	LPN/LVN Nurses	Registered Nurses	6,151
High Exposure Downskill	Stock Clerks & Order Fillers	Heavy Truck Drivers	3,247
	Customer Service Reps.	Nursing Assistants	264
	Customer Service Reps.	Heavy Truck Drivers	2,182
High Exposure Upskill	Customer Service Reps.	Computer User Support	467
	Computer & IS Managers	Computer Occ., All Other	–2,691
	Accountants & Auditors	Bookkeeping Clerks	–99
<i>Panel B: 2022–2024</i>			
Low Exposure Downskill	Light Truck Drivers	Heavy Truck Drivers	3,902
	Driver/Sales Workers	Heavy Truck Drivers	2,783
	Personal Care Aides	Nursing Assistants	2,340
Low Exposure Upskill	Laborers & Freight Movers	Heavy Truck Drivers	4,766
	Nursing Assistants	LPN/LVN Nurses	10,431
	Construction Laborers	Heavy Truck Drivers	4,314
High Exposure Downskill	Customer Service Reps.	Heavy Truck Drivers	4,552
	Shipping & Receiving Clerks	Heavy Truck Drivers	4,585
	General & Operations Mgrs.	Heavy Truck Drivers	–308
High Exposure Upskill	Customer Service Reps.	Office & Admin. Support	–558
	Computer Occ., All Other	Computer Occupations	4,854
	Office Clerks, General	Office & Admin. Support	4,427

Notes: This table contains the three most popular occupation transitions for each transition type (low exposure downskill, low exposure upskill, high exposure downskill, high exposure upskill) in the early and late periods. The far right column shows the real average quarterly earnings return by period for each transition, where an individual's return is the difference between the trainee's post-program earnings and their nearest-neighbor's post-program earnings.

Table A.9: Headline Results for Share of Participants and Occupations “AI-Retrainable”

		Trainee - Matched	
		Full Sample (1)	High Initial AI Exposure (2)
AI Retraining Statistic			
Sh. Spells	$\Delta \overline{Eloundou} \geq 0 \ \& \ \Delta Y \geq 0$	39.8%	27.1%
Sh. Spells	$\Delta \overline{Claude} \geq 0 \ \& \ \Delta Y \geq 0$	38.7%	29.7%
Sh. Occups	$\Delta \overline{Eloundou} \geq 0 \ \& \ \Delta \overline{Y} \geq 0$	44.0%	26.3%
Sh. Occups	$\Delta \overline{Claude} \geq 0 \ \& \ \Delta \overline{Y} \geq 0$	41.1%	27.1%

Notes: Table shows counts of shares of workers and occupations for whom both the earnings returns from job training and the AI content of their subsequent occupation are higher after participation. Columns (1) and (2) show our preferred measures that calculate the number of workers for whom both outcomes are higher relative to the outcomes for matched nearest neighbor pairs. A 6-digit occupation is labeled AI-retrainable if the mean worker in that occupations AI-retrainable. See text for further details.

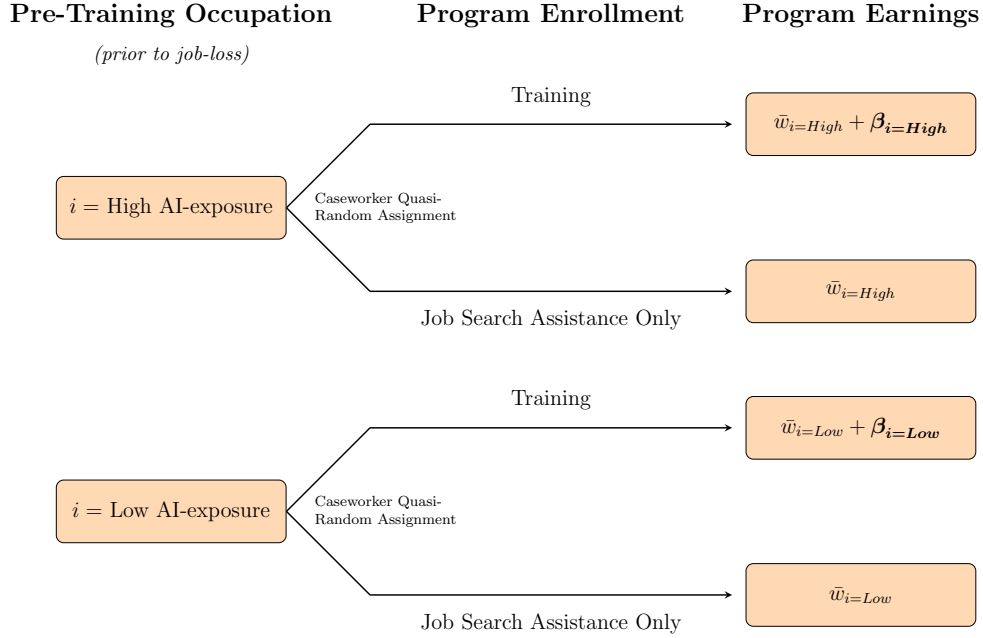
Table A.10: Top Occupations Prior to Training by AI Exposure (Claude-Use Based)

Rank	6-Digit SOC	Occupation	No. Trainees
<i>AI Exposure Quintile 5 (Highest Exposure)</i>			
1	412011	Cashiers	2,123
2	412031	Retail salespersons	1,520
3	439061	Office clerks, general	1,390
<i>AI Exposure Quintile 4</i>			
1	434051	Customer service representatives	4,027
2	533032	Heavy and tractor-trailer truck drivers	2,078
3	119199	Managers, all other	1,253
<i>AI Exposure Quintile 3</i>			
1	435081	Stockers and order fillers	1,325
2	353031	Waiters and waitresses	1,077
3	499071	Maintenance and repair workers, general	864
<i>AI Exposure Quintile 2</i>			
1	311014	Nursing assistants	5,393
2	537062	Laborers and freight, stock, and material movers, hand	2,093
3	519198	Helpers – production workers	1,598
<i>AI Exposure Quintile 1 (Lowest Exposure)</i>			
1	519199	Production workers, all other	2,635
2	292061	Licensed practical and licensed vocational nurses	1,535
3	533033	Light truck drivers	1,436

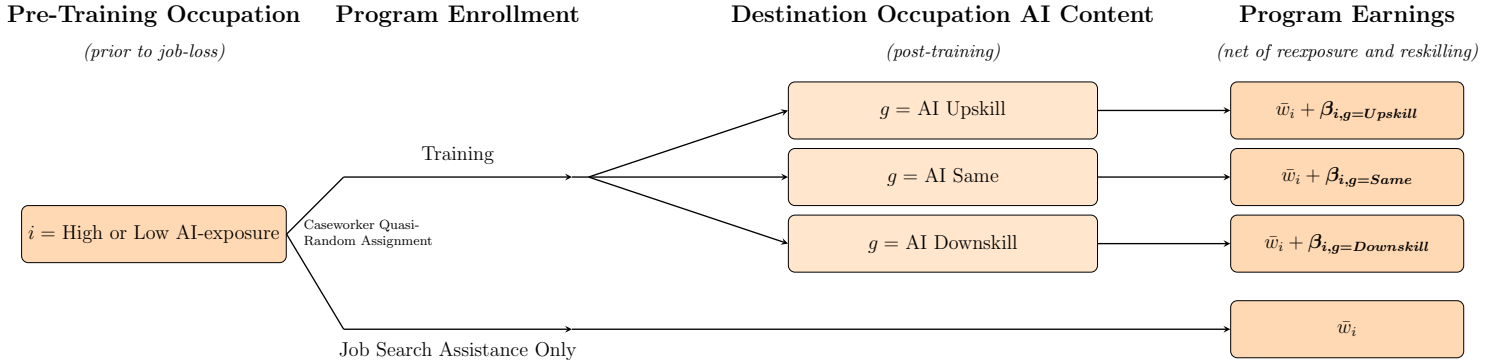
Notes: This table shows the top three most prominent occupations in different AI exposure quintiles prior to entering WIOA/WIA training participation. Counts of number of training participants are shown in the final column. Table uses the AI exposure measure from Anthropic’s Economic Index ([Handa et al., 2025](#)). See text for details.

Appendix B. Diagram Representations of Research Design

Main Design. Heterogeneous Effects by Prior Exposure



Decomposition. Heterogeneous Effects by Skill Transition (Choice of Training)



Notes: Within the main design framework, returns to job training are estimated as a simple comparison of the WIOA program returns between trainees and their nearest neighbors who only receive job search assistance, producing $\beta_{i=High}$ and $\beta_{i=Low}$ for the two pre-training subgroups. The identifying assumption underlying the decomposition is more demanding because the counterfactual occupation choice absent training is unobserved and must be proxied by the matched control unit's realized occupation. The identifying assumption in the decomposition analysis is therefore that, had trainees not exercised their training-induced choice, they would have sorted into occupations similar to those of their matched Wagner–Peyser counterparts.

Appendix C. WIOA ETPL Training Content Classification

To better understand how workforce development training programs evolve in their AI and technology content over time, we utilize the U.S. Department of Labor’s comprehensive database of WIOA Eligible Training Provider Lists (ETPL), which contain detailed curricular and performance information on approved WIOA training programs (as reported by individual states on Form ETA-9171). The database covers training programs approved to be included in the WIOA ETPLs between January 2014 and October 2024, totaling 77,085 training programs across that time frame.

We utilize textual program descriptions reported by training providers to characterize each training program within this data as having (1) AI-specific content or (2) general technology content. To do this, we first fed a random sample of ~ 200 training program descriptions from the WIOA ETPLs to an NLP tool which helped us identify common keywords and phrases related to AI or general technology training content. We then went through the comprehensive list of programs and tagged the relevant keywords within their program descriptions, creating an AI-specific tag and general Technology tag.³⁴ The keywords used for this tagging exercise are reported in [Appendix Table A.12](#)

To guard against erroneous tagging of non-AI and non-tech program descriptions, we matched keywords using case-insensitive regular expressions with word boundary detection, and accounted for exclusion patterns of common keywords that may appear in non-technical programs. For instance, the term “coding” which appears in some electronic health record administration programs that train on medical billing codes would not be tagged as having general technology content unless if it also contains a strong technology component (e.g. database management, software usage, systems administration). Of the 77,085 training programs, we flagged about 0.4% as having AI-specific content, 24% as having general technology content without AI, and 75% as having neither.

³⁴These tags are not mutually exclusive.

We validate the results of our AI and technology classifications by linking the SOC occupation(s) associated with each training program to measures of AI exposure based on Eloundou et al. (2024) and Handa et al. (2025) ($AI_{Eloundou}^{\beta}$ and AI_{Claude} respectively) as well as task-based automation indices from the Acemoglu and Autor (2011) task framework. We convert each measure to percentiles in order to compare the relative intensities across indices.

We show that, in line with expectations, AI-specific training programs are associated with target occupations that have higher AI-exposure, and are also associated with high degrees of routine and non-routine cognitive tasks. General technology training programs (excluding those that also contain AI content) are also associated with above median AI-exposure in their target occupations, although with lower magnitudes compared to those from AI-specific trainings. In contrast, other non-tech, non-AI trainings have are associated with target occupations with lower levels of AI-exposure, hovering around the median levels of the overall distribution.

Table A.11: Occupational Task Content of WIOA ETPL Programs

	Mean Percentile of Associated Occupations		
	AI Trainings	Tech Trainings	Other Trainings
A. AI Exposure Measures			
$AI_{Eloundou}^{\beta}$	83.9	75.6	50.6
AI_{Claude}	71.0	63.3	51.2
B. Automation Task Framework			
Routine Cognitive	60.4	68.2	55.2
Nonroutine Cognitive	69.6	57.4	56.4
Routine Manual	39.7	50.1	48.8
Nonroutine Manual	37.9	48.7	50.1

Notes: AI-exposure and automation task indices are reported as percentiles in terms of the overall distribution across all available 6-digit SOC codes. Higher percentiles pertain to higher levels of AI exposure or higher compatibility with each of the automation framework indices. $AI_{Eloundou}^{\beta}$ refers to SOC-level potential AI-exposure measures as described in Eloundou et al. (2024), while $AI_{Anthropic}$ refers to SOC-level realized AI exposure measures based on Anthropic’s Economic Index Handa et al. (2025). Automation Task Framework Indices are based on SOC-level aggregations of the task-level measures described in Acemoglu and Autor (2011).

Table A.12: WIOA ETPL Program Tagging Keywords

Category	Keywords and Patterns
Panel A: Artificial Intelligence Keywords	
Core AI Terms	artificial intelligence, ai, machine learning, deep learning, neural network, neural net
AI Techniques	natural language processing, nlp, computer vision, image recognition, predictive analytics, predictive modeling, reinforcement learning
AI Applications	chatgpt, gpt, llm, large language model, generative ai, ai automation
Robotics	autonomous robot, intelligent robot, robot ai, autonomous system, unmanned system
Panel B: General Technology Keywords	
Programming & Software	programming, coding, software development, software engineering, web development, mobile development, app development, full stack, front-end, back-end
Programming Languages	python, java, javascript, c++, c#, ruby, php, sql, r, html5, css, react, angular, vue
Computer Science & IT	computer science, computer systems, information technology, it support, help desk, systems administration, systems analyst, technical support, it infrastructure
Cybersecurity & Networks	cybersecurity, information security, network security, security+, cisco, computer network, network administration, tcp/ip, firewall, cryptography, penetration test
Data & Analytics	data science, data analytics, data analysis, big data, data mining, data visualization, business intelligence, data engineer
Cloud & DevOps	cloud computing, aws, azure, google cloud, gcp, devops, docker, kubernetes
Databases	database administrator, database management, mongodb, postgresql, mysql, nosql, sql server
Design & Engineering	cad, cam, cad-cam, autocad, catia, solidworks, computer-aided design, engineering software
Medical Technology	ehr, electronic health record, medical coding software, healthcare software, health information
Industrial Automation	plc programming, programmable logic controller, cnc programming, computer control, industrial automation, scada
IT Certifications	comptia, a+ certification, network+, microsoft certified, cisco certified, ccna
Operating Systems	linux, unix, windows server

Appendix D. Quasi-Random Assignment to WIOA Caseworkers

Caseworkers are the gatekeepers of training under WIOA. Frontline WIOA staff—typically titled “case managers,” “caseworkers,” “career planners,” or “employment specialists”—serve as the primary gatekeepers for training services under WIOA. When a worker enters an American Job Center (AJC), they first receive an initial assessment as required under 20 CFR § 678.430(a)(2). Training may be approved only if these staff determine, after an interview, evaluation, or assessment and career planning, that the participant meets the federal statutory criteria in 20 CFR § 680.210(a):

1. career services alone are insufficient for the worker to obtain self-sufficient employment;
2. training is needed to achieve the employment goal;
3. the worker has the ability to participate successfully in training.

These criteria are evaluated subjectively by each frontline staff member. Through Individual Employment Plans (IEPs) under 20 CFR § 680.170, check-ins at regular intervals provide multiple opportunities for frontline staff to recommend, reassess, or newly assign the participant to training as circumstances evolve. Otherwise, workers instead receive Wagner–Peyser basic career services, such as job search assistance, résumé preparation, and labor-exchange services.

Caseworkers differ substantially in how often they assign training. Local workforce development areas (LWDAs) layer additional requirements on to federal standards, such as in-demand occupation lists, local self-sufficiency wage thresholds, cost or duration caps on Individual Training Accounts (ITAs) that can be applied toward training, prerequisite or suitability requirements, and supervisory or committee approval—giving caseworkers considerable discretion over which workers qualify for training. This discretion produces large observable differences across administrative units: in the “WIA Gold Standard Evaluation” (Fortson et al., 2017), a randomized trial conducted across 28 randomly selected Local Workforce Investment Areas in 19 states, the share of eligible customers receiving staff-approved training ranged from 3 percent to 86 percent. Comparable patterns arise in the Trade Adjustment Assistance (TAA) program, where Hyman (2018) shows that variation in petition adjudication meaningfully affects which displaced workers ultimately receive TAA-funded training. Together, these findings demonstrate that frontline staff exercise substantial influence in determining access to publicly funded job training.

Assignment to caseworkers is effectively random due to operational constraints. The particular caseworker a jobseeker meets is typically determined by operational factors rather than participant characteristics, creating quasi-random variation in who conducts the assessment and training determination. AJCs generally assign walk-ins and scheduled appointments to the next available staff member based on real-time availability, rotating intake rosters, daily staffing matrices, or automated scheduling systems. These procedures are documented explicitly in local operating manuals. The San Bernardino County AJCC Desk Manual states that “every Workforce Development Specialist will be assigned new

cases on a rotation basis” and that customers are served by the “next WDS available” (Section 4, pp. 2–5). The Chicago–Cook Workforce Partnership Case Assignment Re-Assignment Instructions direct providers to assign customers using a rotation designed to “balance caseloads across staff” and to ensure assignment to “the first available career planner” (pp. 1–2). The San Francisco OEWD Workforce Services Protocols describe an intake flow in which jobseekers scheduling online or entering by walk-in are automatically routed to “the first open appointment slot” without regard to staff identity (Workforce Development Plan, 2017–2022, pp. 23–29). Because customers rarely request specific staff and AJCs prioritize rapid service delivery, these operational practices create a plausibly exogenous assignment mechanism that determines which caseworker applies the training-eligibility criteria specified under 20 CFR § 680.210(a).

Appendix D References

- Chicago–Cook Workforce Partnership**, *Case Assignment & Re-Assignment Instructions* Chicago–Cook Workforce Partnership. Internal Operating Guidance, pp. 1–2.
- Fortson, Kenneth, David H. Greenberg, Annalisa Mastri et al.**, “Providing Public Workforce Services to Job Seekers: 30-Month Impact Findings from the WIA Gold Standard Evaluation,” Report, Mathematica Policy Research 2017.
- Hyman, Benjamin**, “Can Displaced Labor Be Re-employed? Evidence from Quasi-Random Assignment to Trade Adjustment Assistance,” 2018. Working paper.
- San Bernardino County Workforce Development Department**, *AJCC Desk Manual* San Bernardino County Workforce Development Department. Section 4, pp. 2–5.
- San Francisco Office of Economic and Workforce Development**, *Workforce Development Plan, 2017–2022* San Francisco Office of Economic and Workforce Development. Workforce Services Protocols, pp. 23–29.
- U.S. Department of Labor**, “Title 20, Code of Federal Regulations,” Code of Federal Regulations. Sections 678.430, 680.170, and 680.210.

Appendix E. Caseworker Variation and Selection on Unobservables

In this appendix, we show that quasi-random assignment to unobserved caseworkers with heterogeneous assignment leniencies attenuates selection on unobservables, thereby strengthening the credibility of matching on observables.

Access to training is mediated by caseworker discretion in interpreting admission rules. Let individuals n be (as-good-as-randomly) assigned to caseworkers $c(n)$ within offices and time periods. Caseworkers differ in *leniency* toward allowing individuals into training programs. We model the decision to assign worker n to training (D_n) using a latent threshold rule:

$$D_n = \mathbf{1}(\sigma_n + U_n + Z_{c(n)} > \xi_{c(n)}) \quad (\text{A.1})$$

where:

- σ_n : determinants of training assignment observed by both the caseworker and the econometrician (used in matching),
- U_n : determinants of training assignment and potential outcomes ($Y(0)$ and $Y(1)$) unobserved by the econometrician, but potentially observed by the caseworker
- $Z_{c(n)}$: caseworker-specific leniency component (unobserved to the econometrician),
- $\xi_{c(n)}$: caseworker-specific assignment threshold reflecting heterogeneity in interpreting training approval rules (see [Appendix D](#)).

We assume that, conditional on observed characteristics (used in matching),

$$Z_{c(n)} \perp (U_n, Y_n(0), Y_n(1)) \mid \sigma_n \quad (\text{A.2})$$

reflecting quasi-random assignment of workers to caseworkers. Under decision rule (A.1), treatment depends on both U_n and the independent component $Z_{c(n)}$. Conditional on σ_n ,

$$D_n = 1 \Leftrightarrow U_n > \xi_{c(n)} - \sigma_n - Z_{c(n)}.$$

Because $Z_{c(n)}$ is independent of U_n , assignment contains an exogenous component. Intuitively, if $Z_{c(n)}$ had zero variance, assignment would be a deterministic function of U_n and σ_n , producing strong selection on U_n . As the variance of $Z_{c(n)}$ increases, the role of U_n in determining D_n weakens. Formally, parameterize the dispersion of caseworker leniency with a scale parameter ($Z_{c(n)} = \lambda \tilde{Z}_{c(n)}$) where $\tilde{Z}$ has mean zero and unit variance, so that $\text{Var}(Z) = \lambda^2$. Then

$$\frac{\partial}{\partial \lambda} \underbrace{\left(\mathbb{E}[U_n \mid D_n = 1, \sigma_n] - \mathbb{E}[U_n \mid D_n = 0, \sigma_n] \right)}_{\text{Selection on unobservables}} < 0 \quad (\text{A.3})$$

so greater quasi-random caseworker variation dampens selection on unobservables, easing the identification burden placed on the matching variables σ_n .

Appendix F. Differencing Out Common Matching Bias

Many of the main findings in the paper are statements about heterogeneous effects. For example, we compare estimates for high AI-exposed training participants versus their nearest-neighbor (NN) control unit pairs, against low AI-exposed participants and their control unit pairs. While Figure 3 showed no pretrends in earnings for all subgroups relative to their nearest-neighbor control pairs, here we formalize how comparisons across groups allow us to further difference out any lingering imbalances, similar in spirit to a triple difference design.

Potential Outcomes Setup and Notation

Let n index individuals. Let $H_n \in \{0, 1\}$ denote AI exposure group membership for individual n (e.g., $H=1 \equiv$ “High”, $H=0 \equiv$ “Low”). Let $D_n \in \{0, 1\}$ denote treatment (non-random). Potential outcomes are $Y_n(d, h)$ for $d \in \{0, 1\}$, $h \in \{0, 1\}$. The observed outcome is

$$Y_n = D_n Y_n(1, H_n) + (1 - D_n) Y_n(0, H_n) \quad (\text{A.4})$$

We want group-specific causal ATE parameters:

$$\tau_h \equiv \mathbb{E}[Y_n(1, h) - Y_n(0, h) \mid H_n = h], \quad h \in \{0, 1\} \quad (\text{A.5})$$

We estimate within-group “one-way” differences (treated – NN control within h):

$$\delta_h \equiv \mathbb{E}[Y_n(1, h) \mid D_n = 1] - \mathbb{E}[Y_n(0, h) \mid D_n = 0] \quad (\text{A.6})$$

Lingering Selection Bias within Each Group

Add and subtract $\mathbb{E}[Y_n(0, h) \mid D_n = 1]$ to get the standard decomposition:

$$\delta_h = \underbrace{\tau_h}_{\text{causal ATE}} + \underbrace{\left(\mathbb{E}[Y_n(1, h) \mid D_n = 1] - \mathbb{E}[Y_n(0, h) \mid D_n = 0] \right)}_{\equiv B_h \text{ (lingering selection bias within } h)} \quad (\text{A.7})$$

Each within-group estimate equals the true effect τ_h plus a group-specific selection bias B_h that remains even after nearest-neighbor matching.

Differencing Out Common Bias from Nearest-Neighbor Estimation

Consider the difference across the two groups of the within-group estimates:

$$\Delta \equiv \delta_1 - \delta_0 = \underbrace{(\tau_1 - \tau_0)}_{\text{heterogeneous effect of interest}} + \underbrace{(B_1 - B_0)}_{\text{residual/differential bias}} \quad (\text{A.8})$$

The simple insight is that any component of the selection bias that is *common across groups* cancels when the focus is on differences in estimates. That is, if

$$B_h = \bar{B} + \varepsilon_h \quad (\text{A.9})$$

with component $\bar{B}$ the matching bias shared by both groups h , and ε_h denotes group-specific selection effects across groups, then the bias becomes

$$\Delta - (\tau_1 - \tau_0) = (B_1 - B_0) = \varepsilon_1 - \varepsilon_0 \quad (\text{A.10})$$

which in principle could be much smaller than the *level* bias $B_h = \bar{B} + \varepsilon_h$ that contaminates δ_h . Thus, the heterogeneous comparisons for Δ may be less biased than either one-way estimate δ_h . Formally, if the following assumption holds:

Assumption (Common selection on unobservables across H).

$$\mathbb{E}[Y_n(0, h) \mid D_n = 1, H_n = h] - \mathbb{E}[Y_n(0, h) \mid D_n = 0, H_n = h] = \bar{B} \quad \text{for } h \in \{0, 1\} \quad (\text{A.11})$$

then Δ may yield a less biased estimate for $\tau_1 - \tau_0$, and comparisons between groups yield improvements in bias over one-way differences.

Appendix G. Matching Estimator Details

We estimate treatment effects using the semiparametric nearest-neighbor matching estimator developed by [Abadie and Imbens \(2002\)](#). For each treated trainee n , the estimand is the individual treatment effect $\theta_n = Y_n(1) - Y_n(0)$, but only one potential outcome is observed. The estimator imputes the missing counterfactual using outcomes from observationally similar control units.

Similarity between treated and control units is defined over pre-treatment covariates X . For each treated unit, we identify the closest control unit in covariate space using the Mahalanobis distance, which rescales variables and accounts for their covariance structure:

$$\|X_{D_{n(p)}=1} - X_{D_{n(p)}=0}\|_{\text{Mahalanobis}} = \sqrt{(X_{D_{n(p)}=1} - X_{D_{n(p)}=0})' \Sigma_X^{-1} (X_{D_{n(p)}=1} - X_{D_{n(p)}=0})}. \quad (\text{A.12})$$

Following [Abadie et al. \(2004\)](#), we enforce overlap directly in the exact-matched covariate space by pairing each treated observation with its nearest neighbor control. In practice, approximately 4.6% of treated observations are dropped at the exact-matching stage due to a lack of suitable matches, a small fraction consistent with the large donor pool of comparable control observations from the Wagner–Peyser program.

Conditional on this restriction, matching proceeds by assigning to each treated observation its nearest neighbor control. To maximize the sample for heterogeneity analysis, we include all remaining matched pairs regardless of their distances, rather than imposing a caliper (i.e., a maximum allowable distance) that would further restrict estimation to regions of common support. The interquartile range of standardized distances lies between -0.75 and 0.42 standard deviations, suggesting that distances are relatively tightly distributed in our setting; imposing a caliper thus has little effect on the results. This is consistent with earlier balance table tests showing that covariate imbalances are, for the most part, confined to a narrow band around zero.

Unlike propensity score approaches that trim observations based on estimated treatment probabilities (e.g., [Crump et al. \(2009\)](#)), overlap here is ensured through the

matching procedure itself—via exact matching followed by nearest-neighbor assignment—rather than through ex-post trimming based on propensity scores.

Alternative Matching and Weighting Approaches

While our baseline follows the Abadie–Imbens nearest-neighbor framework, alternative approaches highlight several tradeoffs. First, as the dimension of X increases, distances can grow nonlinearly, potentially weakening match quality. Moreover, caliper-based trimming in covariate space does not necessarily align with trimming based on lack of overlap in treatment probabilities, which is the target of propensity-score methods.

At the same time, match quality can be improved by incorporating richer pre-treatment earnings histories. Prior earnings trajectories capture persistent unobserved heterogeneity and have been shown to absorb a substantial share of selection bias in workforce program evaluations (Heckman et al., 1998; Lechner et al., 2011). Including lagged earnings in X therefore provides a theory-guided expansion of the conditioning set.

Propensity-score methods offer an alternative approach by matching or weighting observations based on estimated treatment probabilities $P(D = 1 | X)$, thereby enforcing overlap in probability space rather than directly in covariate space. However, these methods require specification of a first-stage model and can be sensitive when estimated propensities approach zero or one.

To assess robustness, we follow Andersson et al. (2024) and implement both (i) a standard propensity-score approach using inverse probability weights after dropping overlap violations using a default caliper (ii) a hybrid approach that performs nearest-neighbor matching on the estimated propensity score. The latter leverages the dimension reduction of the propensity score while retaining the nearest-neighbor structure. As is shown in Appendix Table A.13, across specifications, the choice of estimator does not meaningfully affect the magnitude or qualitative interpretation of the estimated treatment effects. We therefore adopt the Mahalanobis nearest-neighbor estimator as our preferred baseline.

Table A.13: Robustness of Training Return Estimates to Alternate Matching Procedures

Matching Procedure:				
Training Propensity-Score	Prior AI Exposure (Eloundou)		Prior AI Exposure (Claude)	
	High	Low	High	Low
A. Pooled Post:	1,393 (88.93)	2,430 (66.22)	1,701 (82.96)	2,294 (64.39)
B. Pooled Pre (Placebo):	34 (80.62)	-17 (54.01)	137 (73.86)	-172 (57.77)
N	44,525	71,394	52,564	63,355
Mean Trainee Baseline Wages:	11,974	8,846	11,054	9,213
Matching Procedure:				
Nearest-Neighbor on P-Score	Prior AI Exposure (Eloundou)		Prior AI Exposure (Claude)	
	High	Low	High	Low
A. Pooled Post:	1,191 (90.62)	2,228 (82.41)	1,574 (81.02)	2,049 (79.05)
B. Pooled Pre (Placebo):	90 (74.91)	145 (61.22)	218 (66.74)	44 (57.99)
N	43,472	70,463	51,461	62,256
Mean Trainee Baseline Wages:	12,002	8,849	11,064	9,210

Notes: Each cell reports a separate mean quarterly earnings regression estimate. Estimates under propensity score matching are weighted by the inverse of each trainee-neighbor pair's propensity score distance and estimates under nearest-neighbor matching on propensity score are weighted by the inverse Mahalanobis distance. Mean baseline earnings reflect average quarterly earnings pooled across three quarters prior to WIOA/WIA participation. Prior AI exposure is measured prior to program enrollment—see text for further details. Standard errors are clustered by nearest-neighbor and reported in parentheses.