

Perceived Exposure: A Pilot Report

Appendix (companion document)

Hang Yu
September 2026

This companion document holds the question wording, sample and cleaning details, the full belief-behavior matrix, the exposure-index results, the robustness checks, the survey comparison table, and the data notes for the exhibits of *Perceived Exposure: A Pilot Report*. The main report is written for a general reader; this document holds the technical material and the nuances the main text leaves out. Public reports are hyperlinked in the main text and academic sources are cited in full in its footnotes; here they are cited by author and year. Statistical conventions: OLS with heteroskedasticity-robust (HC1) standard errors unless stated; exposure indices standardized within sample; “controls” are AI-use frequency, the two-item AI-literacy score, education, and gender, and the belief-behavior matrix adds a sample indicator and 20 coarse occupation-group fixed effects. “Odds” and “chance” in the main text both refer to the stated percent probability.

A. Question wording

Both questionnaires opened with the same definition: “AI” refers to generative AI (DeepSeek, Doubao, Ernie Bot, Tongyi Qianwen, Kimi, ChatGPT, and similar tools), as well as robots and automation software.” The two instruments are otherwise identical except for one substitution, applied throughout: students (Sample A) are asked about their **target** job, workers (Sample B) about their **current or most recent** job. The table below gives exact fielded wording for the items this report uses.

Code	Chinese wording	English translation	Format / notes
C0	您[目标/目前]工作中，是否已经有任务由AI或自动化工具完成？	In your [target / current] job, are any tasks already performed by AI or automation tools?	Single choice: none at all / a few tasks / a sizable share / most tasks / don't know yet
C1	在您[目标/目前]工作的10项主要任务中，您认为未来5年AI能够独立完成几项？	Of the 10 main tasks in your [target / current] job, how many do you think AI will be able to complete on its own within the next 5 years?	0-10 scale, endpoints “0 tasks” / “10 tasks”
C2 (1/5/10 yr)	未来1年/5年/10年，AI	Within the next 1 / 5 / 10 years, what is the	Open-ended percent, 0-100, asked

Code	Chinese wording	English translation	Format / notes
	使您[目标/目前]岗位被取代、不再需要人来做的可能性有多大？	percent chance that AI displaces your [target / current] position, so it is no longer done by a person?	separately at each horizon
Self / similar others / society (E1, E2, E2b)	E1: 全社会广大劳动者；E2: 您自己[目标/目前]工作；E2b: 和您相似的大多数人会打几分	“Workers across society as a whole” (E1); “your own [target / current] work” (E2); “most people similar to you” self-rating on the same scale (E2b)	1-10 scale; 1 = no impact at all, 10 = vast majority of occupations fundamentally transformed or replaced
C3	未来5年，AI最可能如何改变您的[目标/目前]工作？	Over the next 5 years, how is AI most likely to change your [target / current] job?	Single choice: no change / AI-assisted efficiency gain, content unchanged / content and skills change, need to learn new things / new tasks or responsibilities emerge / most work replaced by AI / changes impossible to imagine today
J1 (worry)	想到AI对您未来工作和收入的影响，您的担忧程度更接近？	When you think about AI's effect on your future work and income, which best describes your level of worry?	Single choice, 1-5: not worried at all ... very worried, has affected my mood or decisions
Job-search direction	对AI的考虑是	Have	Students: less-

Code	Chinese wording	English translation	Format / notes
	否改变了您的求职方向 / 找工作换工作的方向？	considerations about AI changed the direction of your job search?	replaceable / AI-related / no change / hard to say. Workers: less-replaceable / AI-related / no change / not considering a change for now
G3 (1/3/6 mo, students)	您预计毕业后，在下列时间内找到一份您可以接受的工作的可能性有多大？请填写 0-100 的百分数。1 个月内 / 3 个月内 / 6 个月内	After graduation, what is the percent chance that you find a job you would accept within 1 month / 3 months / 6 months? (0 = impossible, 100 = almost certain)	Open-ended percent, 0-100, asked separately at each horizon; students only
I6 (free / ¥300 / ¥800)	若现在提供一门 20 小时的 AI 技能课程，你会报名吗？免费 / 收费 ¥300 / 收费 ¥800	If a 20-hour AI-skills course were offered right now, would you enroll if it were free / cost ¥300 / cost ¥800?	Three separate yes/no items, same course
I7 (WTP)	你最多愿意为一门能切实提升你 AI 技能的课程付多少钱？	What is the most you would be willing to pay for a course that genuinely improved your AI skills?	Open-ended amount, yuan
I10 (target AI mix)	你接下来最想投的 10 个岗位里，有几个是	Of the next 10 positions you would most like to apply to, how many are	0-10 scale

Code	Chinese wording	English translation	Format / notes
	AI 相关/前沿岗位 (其余为力求 AI-安全) ?	AI-related / frontier positions, with the rest AI-safe?	
I11 (pay cut to enter frontier)	为从一个 AI 暴露高的岗位，换到一个各方面都相同但 AI-安全的岗位，你愿意每月少拿多少钱？+ 方向	How much monthly pay would you give up to move from a high-AI-exposure position to an otherwise identical AI-safe one? + direction	Open amount (0 allowed) plus single choice: give up pay for the AI-safe job / give up pay instead to enter the AI-frontier job
G8 (premium to accept exposed job)	假设两份工作：A=AI 替代风险低；B=AI 暴露度高 (更可能是前沿/高成长岗位) 。下面哪一句最符合您？	Suppose Job A has low AI-replacement risk and Job B has high AI exposure (more likely to be a frontier, high-growth AI-related job). Which statement fits you best?	7-point ordinal from “choose A no matter how much more B pays” to “choose B no matter how much more A pays,” with 20% and 50% premium thresholds in between
Liquidity	如果家里突然需要一笔 ¥5000 的支出，您能否在一周内拿出？	If your household suddenly needed to cover an expense of ¥5,000, could you come up with the money within a week?	Single choice: yes, fairly easily / yes, but a strain / with difficulty / no
lit_B1 / lit_B2	B1: 关于 AI 是怎样得出回答的，哪项最准确？(正确：它根据大量文	B1: which statement about how AI systems such as ChatGPT and DeepSeek produce	Each single choice among 4 options, scored correct/incorrect; option order randomized

Code	Chinese wording	English translation	Format / notes
	本中的规律，推算出最可能的表述) B2: AI 给出一个看似可信但错误的电话号码，哪种理解最准确？(正确：AI 有时会自信地给出并不存在的信息)	answers is most accurate? (correct: it infers the most likely wording from patterns in text) B2: an AI returns a plausible but wrong phone number; which interpretation is most accurate? (correct: AI sometimes confidently provides information that does not exist)	
H_net	未来 10 年 AI 对整体就业岗位数量的影响？	Over the next 10 years, how will AI affect the total number of jobs in the economy?	Single choice, 6 options: creates far more jobs / creates slightly more / roughly equal / destroys slightly more / destroys far more / hard to say. "Net destroys" combines the two destroy options
income_effect	未来 5 年，AI 对您未来收入的影响？	Over the next 5 years, how will AI affect your future income?	Single choice, 5 options: decrease >20% / decrease ~10% / unchanged within 5% / increase ~10% /

Code	Chinese wording	English translation	Format / notes
			increase >20%
wellbeing	未来 5 年 , AI 对您整体生活幸福感的影 响 ?	Over the next 5 years, how will AI affect your overall life satisfaction?	Single choice, 5 options: substantially lower ... substantially higher
Open-ended	请用一句话描述 : 您觉得 AI 对您求职/未来工作 (目前工作) 最大的影响是什么 ?	In one sentence, what do you think is the biggest effect of AI on your job search and future work [current job]?	Open-ended text

B. Sample, cleaning, and representativeness

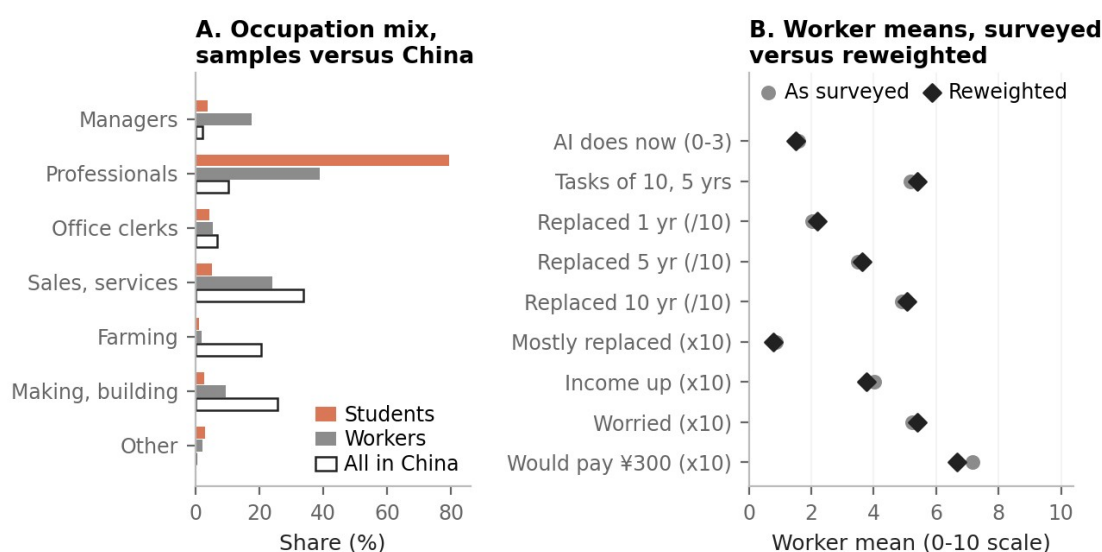
The survey ran on Credamo, a Chinese online panel, in July 2026, with soft quotas on city tier, occupation, and education. Students had to be full-time and within about a year of graduation, or graduated and job-hunting. Workers had to be employed or between jobs. The worker screen also asked for age in four bands (16 to 24, 25 to 29, 30 to 35, over 35); the intent was to admit only those 35 or under, but the platform admitted everyone, and 142 of 500 raw workers (138 of 485 retained, 28.5%) chose “over 35” on that first question. The later demographic item offers no bracket above 30 to 35, so 137 of the 138 appear there as 30 to 35 and one as 16 to 24; the “over 35” figure comes from the screen itself, and the true upper tail of the age distribution is unobservable. The cleaning rule does not remove them at a different rate from other workers (2.8% against 3.1%). I keep them, and Appendix F reports the main worker results with and without them. Older workers put their five-year replacement probability higher (39.2% against 33.1%, a difference of 6.1 points with an interval from 1.5 to 10.7) and are more often at 50% or above (40% against 27%); on the other headline items they do not differ. Because the demographic item has no bracket above 35, the main text calls the sample “workers” rather than “young workers.” The episode is also a signal about self-reports on a paid online panel: if more than a quarter of respondents answer an eligibility question in a way that should have excluded them, the other answers deserve the same skepticism.

Fielding gave a raw N of 500 per sample, in three student waves (raw sizes: a soft launch of 100 and a continuation of 256, both on July 2, and a top-up of 144 fielded July 2 to 13) and two worker waves (raw sizes 100 and 400, both on July 2). The student “waves” are therefore two batches a few hours apart plus a slow top-up, and the composition shift below is a recruitment shift rather than a calendar trend. Composition shifted across waves: the share of students without a bachelor’s degree rose from 8% to 16% to 76%, and among workers the manual and service share rose from 6% to 30% and the over-35 share from 11% to 33%. Appendix F reports what this does to the headline numbers.

Cleaning. A rule fixed before the analysis reported here, but after a first data-quality review, dropped a respondent who tripped any of five hard flags (speeding, a failed attention check, identical answers across the belief items, a shared user ID across the two surveys, or a duplicated open-text answer) or three or more of

seven major flags (a failed mid-survey attention check or numeracy check, a five-year probability more than ten points above a longer horizon, a lower bound above the upper bound, a budget allocation more than ¥500 off its total, enrollment rising with price, or a retest answer four or more tasks away from the original). Applied literally, the rule drops 5 students and 15 workers. Three user IDs appear in both survey files. Two pairs were dropped on both sides. For the third, the student record is kept by a hand exemption and the worker record dropped: the two records are internally consistent descriptions of the same job seeker taking both surveys a minute apart. The analysis sample is therefore 496 students and 485 workers. An earlier, broader candidate rule with ten major flags, including an age contradiction between the screen and the demographic item, and an additional “two major plus three minor” clause would have dropped 5 students and 20 workers; every respondent the final rule drops is on that list, so the two rules disagree only about the six the final rule keeps. Same-session test-retest correlations are 0.75 (students) and 0.74 (workers) for the ten-task item and 0.80 and 0.87, Spearman, for training willingness to pay. These are answers given minutes apart in the same sitting; they show consistency, not stability over time, and they are computed after a cleaning rule that itself uses the retest gap as one of its flags, which raises them somewhat.

Composition against the 2020 census benchmark



Appendix Figure A1: The panel over-represents professionals and under-represents farmers and production workers; reweighting to the census moves the worker means by a few points. Panel A: occupation mix of the two samples against all employed people in the 2020 census, in six broad groups. Panel B: nine worker means as surveyed and after post-stratification to the national occupation mix, on a common 0 to 10 scale.

Occupation group	Students (target job), %	Workers (current job), %	National employment, %	Worker weight
Managers	3.8	17.5	2.2	0.12
Professionals and technicians	79.4	39.0	10.4	0.26
Clerical	4.4	5.6	6.9	1.21
Commerce, services and transport	5.2	24.1	33.9	1.38
Agriculture	1.2	2.1	20.5	9.75
Production, construction and repair	2.8	9.5	25.8	2.67

Occupation group	Students (target job), %	Workers (current job), %	National employment, %	Worker weight
Other / not classified	3.0	2.3	0.2	0
Total of printed cells	99.8	100.1	99.9	

Columns do not sum to exactly 100 because of rounding. National shares are computed from the published counts in the 2020 population census long-form tabulation (国家统计局, 《中国人口普查年鉴—2020》, 表 4-6 各地区分性别、职业中类的就业人口, <https://www.stats.gov.cn/sj/pcsj/rkpc/7rp/zk/html/B0406.xls>).

The census uses the 2015 occupational classification, in which road-transport drivers and couriers (中类 4-02) belong to the services major group while operators of trains, ships, aircraft and cranes (中类 6-30) stay in production; the sample's Transport group consists of drivers and delivery riders, so it is mapped to services. The panel over-represents professionals and managers and under-represents agriculture and production at both ends. Which way these biases push the results is not obvious. Over-sampling professionals compresses the range of true exposure, which by itself flattens any belief-exposure relationship; over-sampling daily AI users (two-thirds of each sample) selects on the one variable that predicts worker beliefs most strongly. Students cannot be reweighted, since a target occupation is not an employment outcome.

Respondent characteristics	Students	Workers
N	496	485
Female (%)	69.8	52.8
Bachelor's degree or above (%)	68.1	61.9
Uses generative AI daily or more (%)	68.1	65.8
Answered "over 35" at the age screen (%)	0.0	28.5
Rural hukou (%)	55.8	38.1
Could not easily raise ¥5,000 within a week: "fairly difficult" or "unable" (%)	28.4	3.7
Of which "unable" (%)	9.3	0.4

Reweighting the worker sample

Worker post-stratification weights are computed within the six classified groups above: the national shares are renormalized to sum to one across those six groups, each is divided by the group's share of the 474 classified workers, the result is clipped at 10, the 11 workers with an unclassified occupation receive zero weight, and the positive weights are normalized to mean one. The largest raw weight is 9.75, for agriculture, so the clip never binds. The effective sample size (Kish) is 144.7 of 474 classified workers, and the 10 agriculture respondents carry 20.6% of the total weight.

Worker outcome	Unweighted	Rewighted	Difference
How much AI already does (0-3)	1.574	1.492	-0.082
Tasks of 10 AI could do in 5 years	5.165	5.394	+0.229
Chance replaced within 1 year (%)	20.347	22.084	+1.737
Chance replaced within 5 years (%)	34.865	36.402	+1.537
Chance replaced within 10 years (%)	48.954	50.550	+1.596

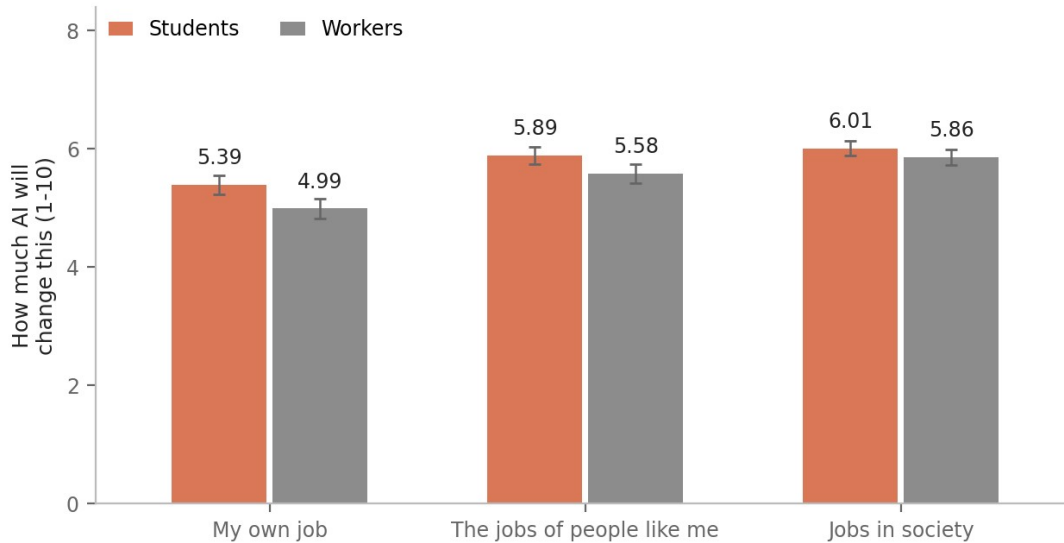
Worker outcome	Unweighted	Reweighted	Difference
Expects to be mostly replaced (share)	0.087	0.079	-0.008
Expects income to rise (share)	0.402	0.378	-0.025
At least somewhat worried (share)	0.524	0.539	+0.015
Would pay ¥300 for a 20-hour AI course (share)	0.718	0.668	-0.049

The reweighted column carries no standard errors; with a design effect of about 3.3 they would be nearly twice the unweighted ones. The reweighted column moves by at most a few points, but that is not evidence that the unweighted numbers are close to national ones. With an effective sample of 145 and ten online-panel farmers standing in for a fifth of national employment, the exercise has little power to detect a larger true gap. The reweighted column is indicative only and should not be read as a national estimate.

C. Secondary patterns

The shape of the five-year probability. Answers heap at multiples of 5 and 10 (95.5% and 73.7% of five-year answers), and the median respondent's own interval around the five-year guess is 20 points wide, so the report does not read differences of less than 5 points. Of respondents at 50% or above, 38% answered exactly 50, which on an open scale reads as a toss-up or "I don't know" as much as a forecast; only 10% of all respondents put the chance at 70% or higher. The probability item and the categorical outlook item (C3) rank people the same way: those who expect to be "mostly replaced" put their five-year chance near 55 to 60% and those who expect AI to assist near 25%. Students and workers are statistically indistinguishable on the five-year probability (35.9 against 34.9), the ten-year probability, and the expected task share, but not on the one-year probability, where students are higher by about 3.5 points (23.8 against 20.3, $p = 0.008$). The probabilities were not incentivized and no calibration question was asked.

Self versus similar others versus society. On the 1-10 exposure scale, both samples rate their own job as least affected and society as most affected: students rate their own job 5.39, people like them 5.89, and society as a whole 6.01; workers rate their own job 4.99, people like them 5.58, and society 5.86. Every step is significant in a paired test (the smallest, students' society-minus-similar-others gap of 0.12, has $p = 0.048$), and all six paired tests survive a Benjamini-Hochberg cut at $q = 0.05$ as one family. The first step, self below others, matches Anthropic's Economic Index Survey, where respondents rate the chance of job loss higher for peers and colleagues than for themselves (*Cadences*, Figure 3.5); that survey has no "society" tier.



Appendix Figure A2: Respondents rate the threat to their own job below the threat to people like them, and both below the threat to workers in general. Mean rating on a 1 to 10 scale of how much AI will change one's own job, the jobs of people like oneself, and jobs in society as a whole, for students (orange) and workers (grey), with 95% confidence intervals.

Horizon increments by exposure quartile. Splitting each sample into quartiles of Eloundou-style LLM exposure and requiring all three horizons answered, the table gives the mean increment in the stated replacement probability from one to five years and from five to ten. Increments run from 10 to 14 points among students and 12 to 17 among workers. Only students' five-to-ten-year increment rises monotonically from the least to the most exposed quartile; the other three series do not, and the worker one-to-five-year increment is largest in the most exposed quartile and second largest in the third.

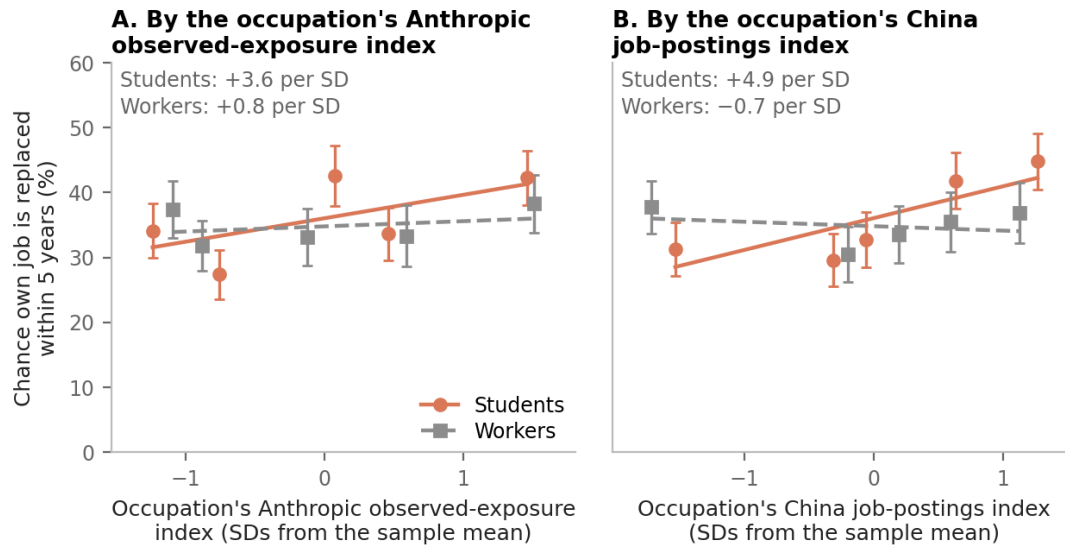
sample	quarti le	n	mean index	1→5		5→10		mean C2_1y r	mean C2_5y r	mean C2_10y r
				yr incre ment	95% CI	yr incre ment	95% CI			
worker	Q1	115	0.374	14.63	[12.22, 17.05]	14.90	[12.89, 16.90]	22.1	36.7	51.6
worker	Q2	114	0.638	11.96	[9.97, 13.96]	13.52	[11.49, 15.54]	16.1	28.1	41.6
worker	Q3	114	0.726	14.88	[11.95, 17.80]	15.26	[12.85, 17.68]	22.9	37.8	53.0
worker	Q4	115	0.791	17.35	[14.65, 20.05]	13.45	[11.38, 15.52]	18.5	35.9	49.3
student	Q1	118	0.510	10.26	[8.31, 12.22]	10.41	[8.33, 12.48]	21.3	31.6	42.0
student	Q2	117	0.636	11.32	[9.46, 13.19]	11.98	[10.25, 13.72]	18.7	30.0	42.0
student	Q3	117	0.737	13.74	[11.68, 15.81]	12.74	[10.91, 14.56]	23.2	36.9	49.6

sample	quarti le	n	mean index	1→5	5→10		mean C2_1y r	mean C2_5y r	mean C2_1 0yr	
				yr incre ment	95% CI	yr incre ment				95% CI
student	Q4	118	0.816	13.32	[11.04, 15.60]	13.14	[11.31, 14.97]	31.7	45.0	58.2



Appendix Figure A3: The expected increment of AI progress from one to five to ten years is similar across quartiles of occupational exposure. Mean stated percent chance one's own job is replaced at one, five, and ten years, for each within-sample quarter of the Eloundou-style exposure of the coded occupation (Q1 least exposed, Q4 most), students in panel A and workers in panel B.

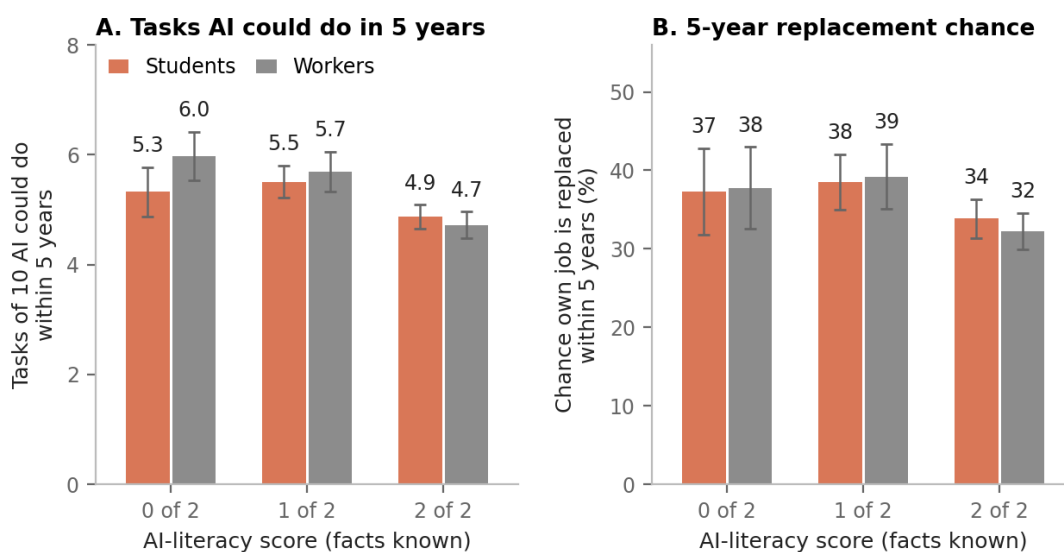
The Anthropropic and Chinese posting indices. Exhibit 5 in the main text shows the expected five-year replacement probability against the Eloundou-style index. The same picture against Anthropropic's observed-exposure index and against the index built from Chinese online job postings (Zhang, Yu, and coauthors, 2025) is below. Both are flat for workers (individual-level $r = +0.04$ and -0.03). For students both rise ($r = +0.17$ and $+0.23$), but not smoothly: across fifths of the Anthropropic index the student means run 34, 27, 43, 34, and 42, a sawtooth whose swings exceed their confidence intervals, whereas the Chinese posting index gives a nearly monotone 31, 30, 33, 42, 45.



Appendix Figure A4: Students' replacement expectations rise with both indices; workers' are flat on both.

Mean stated percent chance that AI replaces one's own job within five years, by fifth of Massenkoff and McCrory's observed-exposure index (panel A) and of the index built from Chinese online job postings (panel B), plotted at each fifth's mean score with the individual-level fit line and slope, for students (orange) and workers (grey), with 95% confidence intervals.

AI literacy. The two-item test asks how language models produce answers (correct: they infer the most likely wording from patterns in text) and how to read a confidently stated wrong answer (correct: models sometimes state information that does not exist). The relationship is not a dose response. Pooled, the expected number of automatable tasks is 5.7 with no correct answer, 5.6 with one, and 4.8 with two; the five-year replacement probability is 37.5%, 38.8%, and 33.0%. In a pooled regression with controls for AI use, education, sample, gender, and 20 occupation-group fixed effects, answering both items correctly goes with 0.90 fewer tasks (interval 0.54 to 1.26) than answering neither, and one correct answer with 0.05 fewer (interval -0.44 to +0.34). On the five-year probability, both correct goes with a 6.4-point lower probability among workers in the same specification (interval 0.3 to 12.4) and 2.4 points lower among students (interval -3.3 to +8.1). Replacing the occupation fixed effects with an exposure index weakens the worker estimate (a linear slope of -2.4 per correct answer with the Eloundou minor-group index, $p = 0.08$, against -4.1 with fixed effects). Given the heaping in the probability item, the task-count result is the one to rely on, and it is a worker result. Split by sample, answering both items correctly goes with 1.28 fewer tasks among workers (interval 0.74 to 1.82) and 0.31 fewer among students (interval -0.19 to +0.81, $p = 0.23$). The two-item test has no competence control of its own; adding the numeracy check and time spent on the survey to the pooled specification leaves the coefficient at 0.86 fewer tasks (interval 0.51 to 1.22), so the result is not inattention. Literacy also goes with a slightly lower answer to how much AI already does in the job (0.09 fewer steps on the 0-3 scale per correct answer, $p = 0.004$), so it is not only about forecasts. Education (main text, Section 1) runs about 4 points lower per step of the education scale in both samples, but on different scales: workers span five levels, while 85% of students sit at either 大专 or 本科, so the student figure is essentially the gap between those two.



Appendix Figure A5: People who answered both questions about how language models work expect AI to take fewer of their tasks; one correct answer looks no different from none. Mean expected number of the job's ten main tasks AI could do within five years (panel A) and mean five-year replacement chance (panel B), by the number of two conceptual questions answered correctly, with 95% confidence intervals. In panel B the intervals overlap substantially.

D. Beliefs and behaviors, the full matrix

All pooled specifications include a sample dummy, occupation-group fixed effects (the 20 coarse groups built from the respondent's closed-list occupation), and controls for AI-use frequency, the two-item AI-literacy score, education, and gender, with HC1 standard errors in parentheses. The last two rows of each table are asked of only one sample (the six-month job-finding chance of students, the switching expectation of workers), so those rows carry no sample dummy. Both belief and behavior are standardized to mean 0 and SD 1 within the estimation sample, so a coefficient is the SD change in the behavior per SD of the belief. Table D1 estimates each belief on its own (one regression per cell); Table D2 enters all seven beliefs together in a single regression per behavior. Stars: * $p < 0.10$, ** $p < 0.05$, *** $p < 0.01$. The two rows marked as expectations have another stated expectation, not a stated action, on the left-hand side; they are the two highest- R^2 rows in D2 (0.35 and 0.21), which is what one expects when two subjective items answered minutes apart are regressed on each other. Every variable in these tables is self-reported in one sitting, so common-method variance and reverse causality (someone who has already moved toward AI work may describe the job as more AI-touched) are live in every cell. Two structural points. The two direction outcomes, moved toward and moved away, are mutually exclusive answers to one item (correlated at -0.49 by construction), so a belief that predicts one will tend to predict the other with the opposite sign; the direction margin among movers is estimated separately below. And the two income indicators, expects a rise and expects a fall, come from one five-point item, so their meaning changes between D1 and D2: in D1 "rise" is rise versus flat or fall, in D2, with "fall" also entered, it is rise versus flat. Occupation fixed effects are post-treatment for the direction outcomes among students, whose closed-list occupation is the target they now name. The cleaning rule keeps 7 students and 3 workers whose enrollment answers rise with price; monotonicity of the ladder is therefore not enforced.

D1. One belief at a time (pooled), coefficient and (HC1 SE)

	C0			C3				
Behavior (standardized)	already does	C1	C2 5-yr risk	replace d	income down	income up	worry	n
Moved job search TOWARD AI-related work	+0.161 (0.034)	+0.111 (0.030)	+0.079 (0.033)	+0.009 (0.034)	-0.063 (0.033)	+0.181 (0.034)	-0.028 (0.033)	981
Wants AI in the job's day-to- day (I10, 0-10)	+0.273 (0.032)	+0.278 (0.033)	+0.218 (0.033)	+0.082 (0.031)	+0.044 (0.031)	+0.163 (0.033)	+0.045 (0.034)	981
Would take a pay cut for an AI-frontier job	+0.078 (0.036)	+0.046 (0.031)	+0.006 (0.032)	-0.043 (0.031)	-0.051 (0.033)	+0.136 (0.035)	-0.135 (0.034)	981
Would enroll in the 20-hr course at ¥300	+0.115 (0.030)	+0.095 (0.030)	+0.052 (0.032)	+0.000 (0.031)	+0.022 (0.031)	+0.158 (0.030)	-0.004 (0.031)	981
Willingness to pay for the course (¥)	+0.157 (0.031)	+0.085 (0.030)	+0.050 (0.031)	-0.013 (0.029)	-0.042 (0.030)	+0.099 (0.032)	-0.072 (0.032)	967
Hours per week learning AI	+0.177 (0.037)	+0.097 (0.032)	+0.098 (0.033)	+0.016 (0.030)	-0.009 (0.032)	+0.096 (0.035)	-0.043 (0.034)	970
Moved job search AWAY from AI-exposed work	-0.059 (0.033)	+0.027 (0.033)	+0.112 (0.034)	+0.058 (0.034)	+0.212 (0.034)	-0.137 (0.032)	+0.245 (0.031)	981
Premium demanded to take the AI-exposed job (G8)	-0.035 (0.036)	+0.021 (0.035)	+0.005 (0.036)	+0.056 (0.031)	+0.077 (0.033)	-0.090 (0.033)	+0.123 (0.036)	981
Planned job applications (z, within sample)	+0.000 (0.035)	+0.005 (0.033)	+0.072 (0.034)	+0.048 (0.035)	+0.099 (0.035)	-0.006 (0.033)	+0.106 (0.035)	975
Own chance of finding a job within 6 months	-0.134 (0.047)	-0.119 (0.049)	-0.168 (0.051)	-0.214 (0.050)	-0.212 (0.046)	+0.255 (0.039)	-0.165 (0.048)	492
Expects switching employer to leave me worse off	+0.176 (0.047)	+0.345 (0.048)	+0.408 (0.045)	+0.217 (0.047)	+0.373 (0.042)	-0.189 (0.045)	+0.277 (0.043)	485

Of the 77 tests in D1, 47 are significant at $p < 0.05$ (3.9 expected by chance), and 46 survive a Benjamini-Hochberg false-discovery-rate cut at $q = 0.05$. The two expectation rows contribute 14 of the 47; on the 63 stated-behavior cells alone, 33 are significant and 30 survive the cut. Neither family was pre-specified. Table D1 changed by at most 0.03 SD in any cell relative to the version that omitted the 46 respondents lost to the coding error described in Appendix F, and ten cells moved across a significance boundary, all near it.

D2. All seven beliefs entered together, pooled, coefficient and (HC1 SE)

Behavior (standardized)	C0		C2 5-yr risk	C3 replaced	income		worry	n	R2
	already y does	C1			down	income e up			
Moved job search TOWARD AI-related work	+0.132 (0.038)	+0.032 (0.039)	+0.064 (0.040)	-0.010 (0.037)	+0.002 (0.040)	+0.176 (0.037)	-0.035 (0.035)	966	0.119
Wants AI in the job's day-to-day (I10, 0-10)	+0.178 (0.035)	+0.145 (0.039)	+0.112 (0.040)	+0.007 (0.031)	+0.081 (0.036)	+0.204 (0.035)	-0.032 (0.033)	966	0.223
Would take a pay cut for an AI-frontier job	+0.069 (0.039)	+0.028 (0.041)	+0.030 (0.042)	-0.042 (0.033)	+0.050 (0.041)	+0.131 (0.037)	-0.147 (0.037)	966	0.080
Would enroll in the 20-	+0.091	+0.042	+0.008	-0.028	+0.123	+0.207	-0.029	966	0.214

Behavior (standardized)	C0 already y does	C1	C2 5-yr risk	C3 replac ed	incom e down	incom e up	worry	n	R2
hr course at ¥300	(0.033)	(0.038)	(0.038)	(0.033)	(0.038)	(0.034)	(0.033)		
Willingness to pay for the course (¥)	+0.145 (0.034)	+0.027 (0.038)	+0.042 (0.041)	-0.025 (0.031)	+0.015 (0.041)	+0.088 (0.039)	-0.092 (0.034)	952	0.160
Hours per week learning AI	+0.154 (0.039)	+0.003 (0.038)	+0.086 (0.039)	-0.011 (0.032)	+0.043 (0.038)	+0.107 (0.038)	-0.086 (0.037)	955	0.121
Moved job search AWAY from AI-exposed work	-0.074 (0.035)	-0.020 (0.040)	+0.034 (0.040)	-0.000 (0.034)	+0.127 (0.042)	-0.034 (0.036)	+0.198 (0.034)	966	0.112
Premium demanded to take the AI-exposed job (G8)	-0.048 (0.041)	+0.036 (0.046)	-0.060 (0.046)	+0.042 (0.035)	+0.011 (0.043)	-0.064 (0.037)	+0.117 (0.041)	966	0.065
Planned job applications (z, within sample)	-0.006 (0.038)	-0.051 (0.041)	+0.055 (0.040)	+0.023 (0.037)	+0.072 (0.041)	+0.050 (0.035)	+0.079 (0.037)	960	0.067
Own chance of finding a job within 6 months	-0.081 (0.053)	+0.006 (0.060)	-0.072 (0.061)	-0.169 (0.050)	-0.030 (0.055)	+0.226 (0.045)	-0.060 (0.048)	484	0.209
Expects switching employer to leave me worse off	+0.097 (0.048)	+0.138 (0.057)	+0.198 (0.053)	+0.029 (0.044)	+0.213 (0.051)	-0.039 (0.048)	+0.083 (0.046)	482	0.348

Of the 77 coefficients in D2, 30 are significant at $p < 0.05$ and 23 survive a Benjamini-Hochberg cut at $q = 0.05$; on the 63 stated-behavior cells, 24 and 18. The matrix does not support a simple two-cluster story of “approach” versus “avoidance” beliefs. Direct experience of AI in one’s own job (C0) mirrors when all seven enter together, positive on approach and negative on avoidance. The expected direction of one’s own income mirrors when examined on its own (D1) but loses its avoidance half in D2. That is not because the beliefs overlap: the largest pairwise correlation among the seven is 0.52 (C1 with C2) and their variance inflation factors run from 1.1 to 1.6. It is a change of reference category. Adding only the “expects a fall” indicator to the D1 regression of moving away on “expects a rise” takes the rise coefficient from -0.134 ($p < 0.001$) to -0.039 ($p = 0.28$); adding C0 alone leaves it at -0.132 and adding worry alone at -0.094 . The avoidance content of “expects a rise” in D1 was mostly “does not expect a fall.” A single signed income expectation (-2 to $+2$, from the same five-point item) entered with the other five beliefs predicts both directions: $+0.164$ SD on moving toward and -0.150 SD on moving away (both $p < 0.001$). Worry goes with more avoidance in both tables and with less of the money-denominated approach behaviors (the frontier pay cut and willingness to pay) in both; it is unrelated to enrollment and to the wish for AI in the job. The five-year replacement probability goes with activity of every kind rather than acting as a pure threat variable: in D2 its two direction coefficients ($+0.064$ and $+0.034$) are both indistinguishable from zero.

In natural units (one belief at a time, the D1 specification). With the D1 controls, each step up the four-point “how much AI already does” scale goes with 8 points more enrollment at ¥300 (interval 4 to 12), ¥110 more willingness to pay (68 to 152), 10 points more likelihood of having moved toward AI work (6 to 15), and 4 points less of having moved away (-8.5 to $+0.4$). Each point on the five-point worry scale goes with 15 points more likelihood of having moved away (11 to 19), ¥46 less willingness to pay (6 to 86), and no detectable change in enrollment at ¥300 (-4 to $+4$). Expecting a pay rise goes with 17 points more likelihood of having moved toward AI work (11 to 24) and 14 points less of having moved away (8 to 21); expecting a cut goes with 22 points more of having moved away (15 to 29). Ten points of five-year replacement probability go with 1.6 points more likelihood of having moved toward (0.3 to 3.0) and 2.5 more of having moved away (1.0 to 4.0).

Three of these twelve figures (the income-rise effect on moving away and both replacement-probability figures) do not survive the joint specification of D2.

The direction margin among movers. Because the two direction dummies are mechanically opposed, the clean estimand is the share moving toward AI work among the 650 respondents (332 students, 318 workers) who moved at all; 44% of them moved toward. The table gives the linear-probability coefficient of each belief, standardized, on that share, and beside it the coefficient on moving at all, with the D1 controls.

Belief	P(toward, among movers), points per SD	p	P(moved at all), points per SD	p
AI already does part of my job (C0)	+8.4	<0.001	+4.5	0.004
Tasks AI could do in 5 years (C1)	+3.7	0.056	+6.4	<0.001
Chance replaced in 5 years (C2)	-0.2	0.923	+9.0	<0.001
Expects most work replaced (C3)	-2.0	0.329	+3.2	0.013
Expects income to fall	-8.8	<0.001	+7.4	<0.001
Expects income to rise	+11.2	<0.001	+1.6	0.300
Worry	-8.5	<0.001	+10.6	<0.001

The participation margin is driven by the threat beliefs (replacement probability, worry, expected fall); the direction margin by experience, expected pay, and worry. The replacement probability, which predicts moving most strongly, predicts direction not at all.

Clustering. Re-estimating the six most-cited D1 cells with standard errors clustered by SOC minor group (86 clusters: the 71 coded minor groups plus 15 singleton clusters for respondents without a code; two of the six cells have 84) changes no significance category; the clustered standard errors are between 0.77 and 1.28 times the HC1 ones. Clustering on the 20 coarse groups would rest on too few clusters to be reliable, which is why HC1 is the reported specification.

D3. Gender

Means by gender within each sample, with a Welch t-test p-value on the difference. The willingness-to-pay row uses the winsorized amount and adds a Mann-Whitney rank test.

Variable	Student		Welch p	Rank p	Worker		Welch p	Rank p
	Student s: men (n)	women (n)			Worker s: men (n)	women (n)		
AI already does part of my	1.680	1.571	0.099*		1.590	1.561	0.669	

Variable	Student s: men (n)	Student s: women (n)	Welch p	Rank p	Worker s: men (n)	Worker s: women (n)	Welch p	Rank p
job (C0, 0-3)	(150)	(338)			(229)	(255)		
Tasks of 10 AI could do in 5 yrs (C1)	5.160 (150)	5.133 (346)	0.886		5.096 (229)	5.227 (256)	0.496	
Chance replaced within 5 yrs, % (C2_5yr)	36.4 (149)	35.6 (343)	0.729		35.0 (228)	34.7 (255)	0.877	
Worry about AI (1-5)	2.773 (150)	2.864 (346)	0.232		2.550 (229)	2.648 (256)	0.177	
Moved search toward AI work	0.313 (150)	0.237 (346)	0.086*		0.336 (229)	0.301 (256)	0.404	
Moved search away from AI-exposed work	0.400 (150)	0.413 (346)	0.782		0.336 (229)	0.340 (256)	0.934	
Wants AI in the job, target AI mix (I10, 0-10)	4.253 (150)	3.694 (346)	0.012* *		3.638 (229)	3.492 (256)	0.547	
Would take a pay cut for a frontier job	0.240 (150)	0.199 (346)	0.324		0.245 (229)	0.191 (256)	0.159	
Would enroll at ¥300	0.493 (150)	0.370 (346)	0.012* *		0.738 (229)	0.699 (256)	0.344	
Willingness to pay, winsorized (¥, I7_wtp_win)	398.9 (148)	274.1 (341)	<0.001 ***	<0.001	560.2 (227)	550.2 (251)	0.856	0.619
Chance of finding a job in 6 months, % (G3_6mo, students only)	73.0 (149)	67.2 (343)	0.0096 ***					

The gap (women minus men) with its 95% interval, within sample:

Variable	Students: gap	95% CI	p	Workers: gap	95% CI	p
AI already does part of my job (C0, 0-3)	-0.109	[-0.239, +0.021]	0.099	-0.029	[-0.161, +0.103]	0.669
Tasks of 10 AI could do in 5 yrs (C1)	-0.027	[-0.399, +0.345]	0.886	+0.130	[-0.246, +0.507]	0.496
Chance replaced within 5 yrs, % (C2_5yr)	-0.7	[-5.0, +3.5]	0.729	-0.3	[-4.2, +3.6]	0.877
Worry about AI (1-5)	+0.091	[-0.059, +0.240]	0.232	+0.098	[-0.044, +0.241]	0.177
Moved search toward AI work	-0.076	[-0.164, +0.011]	0.086	-0.035	[-0.119, +0.048]	0.404
Moved search away from AI-exposed work	+0.013	[-0.081, +0.108]	0.782	+0.004	[-0.081, +0.088]	0.934
Wants AI in the job, target AI mix (I10, 0-10)	-0.560	[-0.994, -0.125]	0.012	-0.145	[-0.620, +0.329]	0.547

Variable	Students:			Workers:		
	gap	95% CI	p	gap	95% CI	p
Would take a pay cut for a frontier job	-0.041	[-0.121, +0.040]	0.324	-0.053	[-0.127, +0.021]	0.159
Would enroll at ¥300	-0.123	[-0.219, -0.028]	0.012	-0.039	[-0.119, +0.042]	0.344
Willingness to pay, winsorized (¥, I7_wtp_win)	-124.8	[-193.2, -56.4]	<0.001	-10.0	[-118.2, +98.3]	0.856
Chance of finding a job in 6 months, % (G3_6mo, students only)	-5.7	[-10.1, -1.4]	0.010			

Whether the gap differs between the two samples: coefficient on female × worker in a pooled regression, without and with the standard controls and occupation fixed effects.

Variable	Unconditional			With controls		
	gap	95% CI	p	gap	95% CI	p
AI already does part of my job (C0, 0-3)	+0.080	[-0.104, +0.264]	0.393	+0.044	[-0.136, +0.224]	0.632
Tasks of 10 AI could do in 5 yrs (C1)	+0.158	[-0.369, +0.685]	0.558	+0.015	[-0.508, +0.539]	0.954
Chance replaced within 5 yrs, % (C2_5yr)	+0.4	[-5.3, +6.2]	0.879	+0.2	[-5.3, +5.6]	0.952
Worry about AI (1-5)	+0.007	[-0.198, +0.213]	0.944	+0.024	[-0.179, +0.227]	0.816
Moved search toward AI work	+0.041	[-0.079, +0.161]	0.506	+0.013	[-0.106, +0.132]	0.827
Moved search away from AI-exposed work	-0.010	[-0.136, +0.117]	0.881	+0.003	[-0.125, +0.130]	0.966
Wants AI in the job, target AI mix (I10, 0-10)	+0.414	[-0.226, +1.055]	0.205	+0.195	[-0.441, +0.830]	0.549
Would take a pay cut for a frontier job	-0.013	[-0.122, +0.097]	0.822	-0.025	[-0.136, +0.087]	0.664
Would enroll at ¥300	+0.085	[-0.040, +0.209]	0.182	+0.053	[-0.070, +0.176]	0.399
Willingness to pay, winsorized (¥, I7_wtp_win)	+114.8	[-12.8, +242.4]	0.078	+80.1	[-46.7, +206.9]	0.216
Chance of finding a job in 6 months, % (G3_6mo, students only)	— student-only outcome, no interaction possible					

Within the student sample, women shift their search toward AI work less (at the 10% level), want less AI in the target job, enroll at ¥300 less, state a lower willingness to pay, and expect a lower chance of finding a job

within six months, while their beliefs about the target job are statistically indistinguishable from men's (the one marginal case is how much AI already does, $p = 0.10$). Holding the seven beliefs constant, the student gaps shrink by between a tenth and two-fifths: willingness to pay $-\text{¥}91$ (interval -161 to -22 , a 27% shrinkage), job-finding chance -5.1 points (-9.1 to -1.1 , 12%), target AI mix -0.3 jobs (-0.7 to $+0.1$, 39%), enrollment at $\text{¥}300$ -7.5 points (-16.8 to $+1.8$, 39%). Adding AI use, literacy, and education changes little; adding the 20 occupation-group fixed effects shrinks all four further ($-\text{¥}68$, -4.7 points, -0.3 jobs, -6.4 points, with only the job-finding gap still at $p < 0.05$), so part of the gap runs through which occupations women target. Wave indicators leave the four gaps unchanged ($-\text{¥}90$, -5.1 , -0.4 , -7.0). Within the worker sample the same gaps are smaller and none is distinguishable from zero. The interaction test shows that the pilot cannot distinguish the student gaps from the worker gaps: not one of the ten female $\times$ worker coefficients has $p < 0.05$, unconditionally or with controls. Two limits on the "indistinguishable beliefs, different choices" reading. The minimum detectable gap at 80% power is the same in every row, 0.27 SD among students and 0.25 among workers, because the group sizes are identical across rows; three of the four significant student behavior gaps (enrollment, target AI mix, job-finding chance) are 0.25 SD, at that floor, and the willingness-to-pay gap is 0.38 SD. And the difference between a belief gap and a behavior gap is never tested directly. The gender pattern is therefore a within-sample finding about gaps of a quarter of a standard deviation or more.

E. Objective exposure benchmarks and occupation coding

Eight objective exposure measures from five families are used, merged onto respondents through an occupation crosswalk. Eloundou-style LLM exposure ("theory") rates occupational tasks the way Eloundou and coauthors rate them, aggregated here to the US ONET/SOC 2018 *minor-group level*; the index is defined for 60 of the 71 minor groups observed in the sample. China posting-based LLM exposure, built from Zhilian job postings in a paper of which I am a coauthor (Zhang, Yu, Li, Hu, Mo, and Li, 2025), applies the same logic at SOC minor (72 groups) and six-digit (247 codes) level, weighted by posting counts; the version used here is the paper's own index, not the later extended release documented on the NSD website (于航、张丹丹、张润博、李泓亨 2025, covering the 60 largest SOC minor groups and 100 largest detailed occupations). Respondents whose six-digit code is not among the 247 (25 codes) receive the mean of their minor group. The Felten-Raj-Seamans AI occupational exposure index (2021) links ten AI application areas from the Electronic Frontier Foundation's AI Progress Measurement project to 52 ONET abilities through a crowd-sourced relatedness survey; their language-modeling variant (2023) restricts the applications to language modeling alone. Neither uses patents. Both are at SOC 2010, crosswalked to SOC 2018 via the BLS crosswalk. Anthropic's observed exposure index (Massenkoff and McCrory, 2026) is the time-weighted share of an occupation's tasks that see meaningful use in Claude, combining Claude.ai conversations filtered to work-related use with first-party API traffic, which is upweighted; it counts only tasks that the Eloundou rating deems doable by a language model, counts augmenting use at half the weight of automating use, and covers 756 SOC 2018 six-digit occupations. Its authors state that it is not a pure percentage of task coverage and does not measure the intensive margin of use. It is the authors' measure, built from Economic Index data, and the main text attributes it to them. Webb's (2020) patent-text-based AI, software, and robot exposure scores run ONET-SOC 2010 through the ONET-SOC 2019 and SOC 2018 six-digit crosswalks, converted to percentile ranks (0-100).

Three occupation groupings appear in this report and should not be confused. The 20 coarse groups used as fixed effects in Appendix D and to build the group-level benchmarks come from the respondent's closed-list occupation choice. The SOC minor (71 groups) and six-digit (202 codes) occupations used for every index in Section 5, Exhibit 5, and Appendix F come from the free-text job title coded by a language model as described below. The seven census rows in Appendix B are the 20 coarse groups collapsed to the census scheme.

Respondents wrote their job title as free text: 591 distinct Chinese titles among 981 respondents. Each unique title, closed-list category, and sample was coded once to a SOC 2018 six-digit occupation by an LLM (GPT-5.4), given the 756-line SOC list plus worked examples for ambiguous titles, flagging confidence as high, medium, or low. Of the 981 respondents, 966 (98.5%) received a valid code, spanning 202 six-digit occupations and 71 minor groups; among coded respondents, confidence was high for 48%, medium 34%, low 18% (at the level of unique titles, 98.0% coded, with 39%, 40%, and 21%). The coded major group agreed with the respondent's closed-list choice 72.7% of the time among respondents whose closed-list answer named a group (71.0% including the 26 whose answer was "Other"). Benchmark coverage among the 966 coded respondents ranged from 97% (theory) to 99% (Anthropic, Webb). Cell sizes are thin at six digits (median 1.5; 73% of respondents in a cell of 5+) but adequate at the minor-group level (median 4; 92% in a cell of 5+), which is why the report compares minor or major groups rather than six-digit detail.

Second-coder check. To gauge the reliability of the coding, a stratified random subsample of 100 titles (50 flagged high confidence, 50 medium or low) was re-coded by a second, independent language model (Claude), given the same title, closed-list category, and sample as the first coder but blind to the first coder's code, confidence flag, and notes. The two coders agree on 80% of six-digit codes, 90% of minor groups, and 94% of major groups. Agreement at six digits is 92% among high-confidence assignments and 68% among medium and low ones, so the first coder's confidence flag is informative. Ten of the 100 titles are uninformative on their own (公务员, 体制内人员, 服装, and the like); nine of the ten were flagged low confidence. Disagreements concentrate in two classes: under-specified public-sector titles pushed to residual "all other" clerical codes, and titles where the two coders followed different but defensible conventions (数据分析师 as statistician or data scientist, 初中教师 as secondary or middle-school teacher). Agreement does not differ between students' target occupations and workers' current ones (80% and 80%). Because the second coder is also a language model, this is an inter-coder reliability check, not validation against a human standard, and it cannot detect error common to both models. A human-coded subsample remains to be done before the larger wave.

F. Robustness

A coding error in the AI-use control, and its correction. An earlier map from the AI-use frequency item to a 0-5 scale used two option labels that were never fielded, so the 46 respondents (13 students, 33 workers) who reported using AI a few times a month or less, or about once a week, were missing on that control and dropped from every regression that included it. The map is corrected; every specification in Appendices C, D, and F runs on the full sample. Point estimates in Appendix D moved by at most 0.03 SD. The worker exposure slope with controls moved more: with the 33 low-use workers absent, the five-year probability rose 2.4 points per standard deviation of the Eloundou index ($p = 0.04$); with them present it rises 1.8 (interval -0.5 to $+4.0$, $p = 0.13$).

Beliefs on each index, workers. Slope of the stated five-year replacement probability (points), the expected task share (0-1), and worry (1-5) on each index, per standard deviation of the index within the sample, first on its own and then with controls for AI use, literacy, education, and gender. HC1 standard errors; 95% intervals in brackets.

index	5-yr prob., per SD	95% CI	p	task share, per SD	95% CI	p	worry, per SD	95% CI	p
Eloundou (minor group)	-0.113	[-2.101, +1.874]	0.911	+0.003	[-0.017, +0.023]	0.779	-0.053	[-0.132, +0.027]	0.197

index	5-yr prob., per SD	95% CI	p	task share, per SD	95% CI	p	worry, per SD	95% CI	p
China postings	-0.685	[-2.660, +1.290]	0.497	+0.001	[-0.018, +0.020]	0.898	-0.032	[-0.114, +0.050]	0.450
Felten AIOE	-1.109	[-3.101, +0.883]	0.275	+0.004	[-0.016, +0.024]	0.691	-0.032	[-0.112, +0.048]	0.434
Felten language-model	-0.936	[-2.940, +1.068]	0.360	+0.003	[-0.017, +0.022]	0.778	-0.032	[-0.112, +0.048]	0.429
Anthropic observed	+0.798	[-1.215, +2.810]	0.437	+0.008	[-0.012, +0.027]	0.444	+0.003	[-0.073, +0.079]	0.940
Webb AI	-1.090	[-3.024, +0.844]	0.269	+0.007	[-0.011, +0.026]	0.439	-0.026	[-0.101, +0.049]	0.501
Webb software	+0.420	[-1.611, +2.450]	0.685	+0.009	[-0.010, +0.027]	0.363	-0.049	[-0.126, +0.027]	0.208
Webb robot	+0.601	[-1.308, +2.510]	0.537	+0.006	[-0.013, +0.025]	0.544	-0.030	[-0.106, +0.045]	0.433

With controls:

index	5-yr prob., per SD	95% CI	p	task share, per SD	95% CI	p	worry, per SD	95% CI	p
Eloundou (minor group)	+1.759	[-0.487, +4.004]	0.125	+0.007	[-0.015, +0.029]	0.521	-0.009	[-0.102, +0.083]	0.841
China postings	+1.097	[-1.057, +3.251]	0.318	+0.007	[-0.013, +0.027]	0.502	+0.011	[-0.078, +0.100]	0.806
Felten AIOE	+1.048	[-1.229, +3.325]	0.367	+0.011	[-0.011, +0.033]	0.334	+0.032	[-0.060, +0.125]	0.496
Felten language-model	+1.136	[-1.114, +3.386]	0.322	+0.009	[-0.012, +0.030]	0.411	+0.024	[-0.066, +0.114]	0.598
Anthropic observed	+1.652	[-0.447, +3.752]	0.123	+0.010	[-0.010, +0.030]	0.332	+0.019	[-0.057, +0.095]	0.624
Webb AI	-0.575	[-2.570, +1.419]	0.572	+0.002	[-0.017, +0.021]	0.846	+0.007	[-0.069, +0.082]	0.864
Webb software	+0.178	[-1.830, +2.186]	0.862	+0.003	[-0.015, +0.021]	0.729	-0.043	[-0.119, +0.033]	0.267
Webb robot	-0.844	[-2.854, +1.166]	0.410	+0.001	[-0.019, +0.020]	0.958	-0.072	[-0.152, +0.008]	0.078

Beliefs on each index, students.

index	5-yr prob., per SD	95% CI	p	task share, per SD	95% CI	p	worry, per SD	95% CI	p
Eloundou (minor group)	+5.318	[+3.450	<0.001	+0.037	[+0.020	<0.001	+0.057	[-0.010,	0.097

index	5-yr prob., per SD	95% CI	p	task share, per SD	95% CI	p	worry, per SD	95% CI	p
		, +7.186]			, +0.054]			+0.125]	
China postings	+4.909	[+2.897	<0.001	+0.034	[+0.017	<0.001	+0.046	[-0.022, +0.113]	0.184
		, +6.920]			, +0.050]				
Felten AIOE	+1.187	[-0.763, +3.136]	0.233	+0.006	[-0.011, +0.024]	0.484	-0.024	[-0.091, +0.044]	0.489
Felten language-model	+0.014	[-1.870, +1.899]	0.988	+0.001	[-0.016, +0.019]	0.887	-0.042	[-0.107, +0.023]	0.202
Anthropic observed	+3.605	[+1.646	<0.001	+0.022	[+0.006	0.007	+0.053	[-0.013, +0.118]	0.115
		, +5.564]			, +0.038]				
Webb AI	+3.508	[+1.439	<0.001	+0.017	[+0.000	0.046	+0.010	[-0.062, +0.082]	0.781
		, +5.577]			, +0.033]				
Webb software	+4.169	[+2.186	<0.001	+0.014	[-0.003, +0.031]	0.108	+0.044	[-0.023, +0.111]	0.198
		, +6.153]							
Webb robot	+2.311	[+0.330	0.022	+0.008	[-0.009, +0.025]	0.360	+0.032	[-0.032, +0.097]	0.324
		, +4.292]							

With controls:

index	5-yr prob., per SD	95% CI	p	task share, per SD	95% CI	p	worry, per SD	95% CI	p
Eloundou (minor group)	+5.400	[+3.534	<0.001	+0.038	[+0.021	<0.001	+0.066	[-0.003, +0.134]	0.060
		, +7.266]			, +0.055]				
China postings	+4.950	[+2.929	<0.001	+0.033	[+0.016	<0.001	+0.051	[-0.018, +0.119]	0.145
		, +6.971]			, +0.050]				
Felten AIOE	+2.044	[+0.015	0.048	+0.014	[-0.004, +0.032]	0.124	-0.004	[-0.074, +0.065]	0.903
		, +4.074]							
Felten language-model	+0.932	[-1.042, +2.906]	0.355	+0.010	[-0.008, +0.028]	0.288	-0.024	[-0.093, +0.044]	0.487
Anthropic observed	+3.767	[+1.787	<0.001	+0.023	[+0.007	0.005	+0.059	[-0.007, +0.125]	0.080
		, +5.748]			, +0.038]				
Webb AI	+3.381	[+1.290	0.002	+0.015	[-0.001, +0.032]	0.067	+0.010	[-0.062, +0.083]	0.780
		, 							

index	5-yr prob., per SD	95% CI	p	task share, per SD	95% CI	p	worry, per SD	95% CI	p
		+5.471]							
Webb software	+4.010	[+2.007 , +6.014]	<0.001	+0.011	[-0.005, +0.028]	0.183	+0.045	[-0.024, +0.114]	0.199
Webb robot	+1.948	[-0.044, +3.940]	0.055	+0.004	[-0.013, +0.021]	0.634	+0.027	[-0.040, +0.094]	0.426

Worker upper bounds against student slopes (five-year probability, per SD of the index). The last three columns give the student $\times$ index interaction from a pooled regression in which the index is standardized over the pooled sample, so the interaction is not the arithmetic difference of the two within-sample slopes beside it; the table after the next gives the interaction under the within-sample convention, where it is.

Bivariate:

index	worker b/SD	worker 95% CI upper	student b/SD	worker upper below student point?	student $\times$ index	95% CI	p
Eloundou (minor group)	-0.11	+1.87	+5.32	yes	+6.45	[+3.60, +9.29]	<0.001
China postings	-0.68	+1.29	+4.91	yes	+6.33	[+3.40, +9.26]	<0.001
Felten AIOE	-1.11	+0.88	+1.19	yes	+2.57	[-0.59, +5.74]	0.111
Felten language-model	-0.94	+1.07	+0.01	no	+0.83	[-2.21, +3.87]	0.593
Anthropic observed	+0.80	+2.81	+3.61	yes	+3.33	[+0.48, +6.19]	0.022
Webb AI	-1.09	+0.84	+3.51	yes	+4.79	[+1.90, +7.67]	0.001
Webb software	+0.42	+2.45	+4.17	yes	+3.91	[+1.07, +6.75]	0.007
Webb robot	+0.60	+2.51	+2.31	no	+2.27	[-0.68, +5.22]	0.132

With controls:

index	worker b/SD	worker 95% CI upper	student b/SD	worker upper below student point?	student $\times$ index	95% CI	p
Eloundou (minor group)	+1.76	+4.00	+5.40	yes	+4.92	[+2.03,	<0.001

index	worker b/SD	worker 95% CI upper	student b/SD	worker upper below student point?	student × index	95% CI	p
						+7.81]	
China postings	+1.10	+3.25	+4.95	yes	+4.85	[+1.93, +7.78]	0.001
Felten AIOE	+1.05	+3.33	+2.04	no	+1.70	[-1.50, +4.89]	0.298
Felten language-model	+1.14	+3.39	+0.93	no	+0.13	[-2.93, +3.18]	0.936
Anthropic observed	+1.65	+3.75	+3.77	yes	+2.69	[-0.17, +5.55]	0.066
Webb AI	-0.58	+1.42	+3.38	yes	+4.24	[+1.34, +7.13]	0.004
Webb software	+0.18	+2.19	+4.01	yes	+4.00	[+1.21, +6.79]	0.005
Webb robot	-0.84	+1.17	+1.95	yes	+3.04	[+0.05, +6.02]	0.046

The interaction under the within-sample convention, and clustered standard errors. The table repeats the student × index interaction with the index standardized within each sample, so that it equals the difference between the student and worker slopes printed above, and adds p-values with standard errors clustered by SOC minor group (CR1; 39 to 58 clusters depending on the index and sample, respondents without a code as singletons). Clustering changes no conclusion: every worker slope stays null, the student Eloundou, China, Anthropic, Webb AI, and Webb software slopes keep $p \leq 0.003$ in both specifications, the student Webb robot slope stays at $p = 0.025$ bivariate and 0.048 with controls, and the only change of category is the student Felten AIOE slope with controls ($p = 0.048$ to 0.18). The count of interactions significant at 5% is five of eight under the pooled convention (six with clustering, with controls) and four of eight under the within-sample convention (five with clustering); the main text says “about half.”

index	within- sample interactio n, bivariate	HC1 p	clustered p	with controls	HC1 p	clustered p
Eloundou (minor group)	+5.43	<0.001	<0.001	+3.66	0.011	0.008
China postings	+5.59	<0.001	<0.001	+3.91	0.007	0.015
Felten AIOE	+2.30	0.107	0.112	+0.79	0.586	0.624
Felten language-model	+0.95	0.498	0.537	-0.31	0.826	0.854
Anthropic observed	+2.81	0.050	0.025	+2.06	0.151	0.114
Webb AI	+4.60	0.001	0.002	+4.03	0.006	0.006
Webb software	+3.75	0.010	0.002	+3.86	0.007	<0.001
Webb robot	+1.71	0.223	0.176	+2.70	0.057	0.042

Common support. The two samples occupy different parts of the occupation space: 86% of students' coded target occupations fall in SOC majors 11 to 29 (management, business, computing, engineering, science, education, health), against 60% of workers'. The table restricts both samples to the 38 minor groups both occupy, to majors 11 to 29, and to the other majors. Slopes are per standard deviation of the index in the full sample of the same group, so they are comparable to the headline slopes. Inside the shared range, workers' beliefs do rise with the Eloundou index, at about the students' rate once controls enter; for the Anthropic and Chinese posting indices the restricted worker slopes are positive but do not reach conventional significance. Outside the professional majors neither sample's beliefs track any index, on small cells. The student gradient is itself concentrated: dropping computing (SOC 15) and business and finance (SOC 13) students takes the student Eloundou slope from +5.3 to +2.6 ($p = 0.045$, $n = 323$) and the Anthropic slope from +3.6 to +2.0 ($p = 0.14$).

index	restriction	students b/SD (p, n)	workers b/SD (p, n)
Eloundou	full sample, bivariate	+5.33 (<0.001, 470)	-0.11 (0.911, 460)
Eloundou	shared minor groups, bivariate	+5.31 (<0.001, 461)	+3.12 (0.057, 391)
Eloundou	SOC majors 11-29, bivariate	+5.99 (<0.001, 408)	+6.18 (0.011, 277)
Eloundou	other majors, bivariate	+0.00 (0.999, 62)	-0.02 (0.990, 183)
Eloundou	full sample, with controls	+5.42 (<0.001, 470)	+1.77 (0.125, 460)
Eloundou	shared minor groups, with controls	+5.39 (<0.001, 461)	+5.56 (0.003, 391)
Eloundou	SOC majors 11-29, with controls	+5.98 (<0.001, 408)	+7.58 (0.004, 277)
Eloundou	other majors, with controls	+1.40 (0.715, 62)	+0.73 (0.638, 183)
Anthropic	full sample, bivariate	+3.61 (<0.001, 483)	+0.80 (0.437, 465)
Anthropic	shared minor groups, bivariate	+3.56 (<0.001, 465)	+1.75 (0.126, 393)
Anthropic	SOC majors 11-29, bivariate	+4.72 (<0.001, 421)	+2.29 (0.136, 284)
Anthropic	other majors, bivariate	-0.04 (0.987, 62)	-0.52 (0.705, 181)
Anthropic	full sample, with controls	+3.77 (<0.001, 483)	+1.65 (0.123, 465)
Anthropic	shared minor groups, with controls	+3.80 (<0.001, 465)	+2.13 (0.070, 393)
Anthropic	SOC majors 11-29, with controls	+4.87 (<0.001, 421)	+3.16 (0.050, 284)
Anthropic	other majors, with controls	+0.65 (0.799, 62)	-0.06 (0.968, 181)
China postings	full sample, bivariate	+4.92 (<0.001, 479)	-0.69 (0.497, 460)
China postings	shared minor groups, bivariate	+4.93 (<0.001, 463)	+0.74 (0.587, 390)
China postings	SOC majors 11-29, bivariate	+6.54 (<0.001, 420)	+3.09 (0.153, 284)
China postings	other majors, bivariate	-0.57 (0.813, 59)	-0.50 (0.749, 176)
China postings	full sample, with controls	+4.96 (<0.001, 479)	+1.10 (0.318, 460)
China postings	shared minor groups, with controls	+5.01 (<0.001, 463)	+2.50 (0.098, 390)
China postings	SOC majors 11-29, with controls	+6.47 (<0.001, 420)	+4.67 (0.040, 284)
China postings	other majors, with controls	+0.06 (0.982, 59)	+0.15 (0.929, 176)

Is “how much AI already does” an occupation attribute? Regressing the C0 answer (0-3) on occupation dummies shows that for workers it is not: the 58 SOC minor groups leave an adjusted R^2 of zero (F-test $p = 0.49$), the 20 coarse groups likewise ($p = 0.54$), and no index predicts it (the closest, Eloundou, has $p = 0.07$). General and at-work AI-use frequency alone explain 13% of it (the correlation with at-work use is 0.36). Among students, by contrast, occupation does structure the answer (adjusted R^2 0.06, $p = 0.005$ for the minor groups)

and six of the eight indices predict it (all but the two Felten indices). This is the basis for the main text's reading that workers' experience variable is a report about their own AI use, while students' is a report about the occupation they name.

sample	regressors	n	R ²	adjusted	F-test p
				R ²	
workers	SOC minor dummies	472	0.123	0.000	0.487
workers	20 coarse groups	484	0.037	0.000	0.537
workers	general AI-use frequency only	484	0.066	0.064	<0.001
workers	general + at-work AI use	484	0.132	0.128	<0.001
workers	AI use + minor dummies	472	0.211	0.096	<0.001
students	SOC minor dummies	485	0.156	0.061	0.005
students	20 coarse groups	488	0.086	0.051	<0.001
students	general AI-use frequency only	488	0.044	0.042	<0.001

The C0 slope against the index slope among students. Entering the standardized C0 answer and the standardized index together in one student regression of the five-year probability with controls, the C0 coefficient (+5.9 to +6.8 across indices) exceeds the index coefficient for every index, and the difference is significant at 5% for six of the eight; the exceptions are the Eloundou index (index coefficient +4.7, difference p = 0.42) and the Chinese posting index (+4.6, p = 0.26), the two that come closest to the experience variable.

Students' target occupation is endogenous to the belief. Two-thirds of students say AI has already changed the kind of job they are looking for, so the occupation whose index is attached to them was chosen partly on the belief being predicted. Splitting students by whether they moved (with controls), the index slope is nearly twice as large among movers as among non-movers, and steepest among those who moved away from exposed work. Conditioning on how much AI the student wants in the next job (I10) changes the full-sample slopes little (+4.8 Eloundou, +3.7 Anthropic). The non-mover slopes are positive, so endogenous choice does not explain the whole student gradient, but it explains enough that the main text names it.

index	all students	movers	non-movers	moved toward AI	moved away
Eloundou	+5.42 (<0.001, 470)	+5.27 (<0.001, 318)	+2.78 (0.141, 152)	+2.81 (0.092, 123)	+7.16 (<0.001, 195)
Anthropic	+3.77 (<0.001, 483)	+3.78 (0.001, 329)	+2.41 (0.252, 154)	+0.18 (0.900, 129)	+5.66 (<0.001, 200)
China postings	+4.96 (<0.001, 479)	+5.00 (<0.001, 325)	+2.18 (0.301, 154)	+2.77 (0.113, 127)	+6.44 (<0.001, 198)

Measurement error in the coding, and waves. On high-confidence occupation codes only, the Eloundou slope with controls is +6.6 among students (n = 248) and +2.3 among workers (p = 0.18, n = 204); on the Anthropic index, +6.4 and +3.1 (p = 0.06). Coding error therefore does not explain the worker null. Adding wave indicators leaves the student slopes unchanged: +5.4 (Eloundou), +3.9 (Anthropic), +5.0 (China postings).

Robot exposure. Workers' five-year probability on Webb's robot percentile: +0.60 per SD (interval -1.31 to +2.51, p = 0.54, n = 465); on Webb's software percentile: +0.42 (-1.61 to +2.45). Students: +2.31 (0.33 to 4.29, p = 0.02) and +4.17 (2.19 to 6.15). Across respondents' occupations the robot percentile correlates at -0.73 with the Felten language-model index and at -0.38 with the Eloundou minor-group index; across the benchmark occupations, -0.76 (756 six-digit occupations) and -0.62 (the 603 with an Eloundou minor-group

value). Entered jointly with the Eloundou index and controls, the robot index has a worker coefficient of -0.21 ($p = 0.86$) and the Eloundou index $+1.86$ ($p = 0.17$).

Occupation fixed effects. Fixed effects for the LLM-coded SOC minor groups explain 20.3% of the raw variance in workers' five-year probability (58 groups, $n = 471$) and 19.7% among students (50 groups, $n = 490$), but raw R^2 with that many dummies is inflated: adjusted, the figures are 9.3% and 10.8%, and randomly permuting the occupation labels gives a raw R^2 of 12% among workers and 10% among students by chance alone. The fixed effects are nonetheless jointly significant in both samples (F-test $p < 0.001$). The 20 coarse groups explain 7.0% raw (3.2% adjusted, $p = 0.016$) among workers and 13.4% (10.1%) among students, so at the coarse level occupation structures student beliefs more than worker beliefs. No single index explains more than 0.3% among workers, against 5.9% for the Eloundou index among students. The main text's statement that occupation matters to workers rests on the F-test, not on the raw R^2 .

Occupation cells among workers. Drivers (SOC 53-3, $n = 25$), whose occupations sit near the bottom of every language-model index and near the top of the robot index, put their five-year probability at 40.0%; accountants (13-2011, $n = 31$) at 37.5% and software developers (15-125x, $n = 22$) at 41.9%. Bundling drivers, waiters, cooks, and all production occupations (51-xxxx; $n = 68$) against accountants and developers ($n = 53$) gives 39.6% against 39.4%, a difference of 0.2 points (interval -7.5 to $+7.9$). The bottom fifth of the Felten language-model index puts the probability at 38.0% and the top fifth at 36.6%, a bottom-minus-top difference of $+1.4$ points (interval -5.1 to $+7.9$). All of these cells are small and all differences are inside the 5-point noise band.

Worker exposure slope under restrictions (five-year probability on the Eloundou minor-group index, with controls). Full sample: $+1.76$ (-0.49 to $+4.00$, $p = 0.13$, $n = 460$). Workers 35 or under: $+1.77$ (-0.73 to $+4.27$, $p = 0.17$, $n = 334$). High-confidence occupation codes only: $+2.42$ (-1.10 to $+5.95$, $p = 0.18$, $n = 204$). The restricted samples give larger point estimates with wider intervals, which is what smaller samples do; they do not confirm the null so much as fail to overturn it. On the Anthropic index the corresponding figures are $+1.65$ ($p = 0.12$), $+1.95$ ($p = 0.09$), and $+3.11$ ($p = 0.06$, $n = 205$). Without controls all six are between -0.1 and $+2.1$ with $p > 0.12$.

Headline worker claims, all workers, workers 35 or under, and workers over 35. The test compares the over-35 and 35-or-under groups.

Worker outcome	All workers (n)	35 or under (n)	Over 35 (n)	Over 35 minus 35 or under, 95% interval	p
5-year replacement chance, mean %	34.9 (483)	33.1 (345)	39.2 (138)	$+6.1$ [$+1.5$, $+10.7$]	0.009
Share at 50% or higher, %	30.8	27.2	39.9	$+12.6$ [$+3.1$, $+22.1$]	0.010
How much AI already does (0-3)	1.57 (484)	1.56 (347)	1.61 (137)	$+0.04$ [-0.10 , $+0.19$]	0.560
Expects income to rise, %	40.2 (485)	41.8 (347)	36.2 (138)	-5.6 [-15.2 , $+4.1$]	0.257
At least somewhat worried, %	52.4	51.0	55.8	$+4.8$ [-5.1 , $+14.7$]	0.341
Would enroll at ¥300, %	71.8	71.8	71.7	0.0 [-9.0 ,	0.997

Worker outcome	All workers (n)	35 or under (n)	Over 35 (n)	Over 35 minus 35 or under, 95% interval +8.9]	p
Willingness to pay, median ¥	474 (478)	500 (345)	400 (133)		0.222

Excluding the midpoint answer. The midpoint (5 of 10 tasks) was chosen by 13% of students and 11% of workers; dropping it leaves the expected-task-share slope on the Eloundou index essentially unchanged for workers (+0.008 per SD, $p = 0.54$, $n = 408$, against +0.007 with it).

Waves. Wave sizes here are after cleaning. Student enrollment at ¥300 by wave: 47.5% (wave 1, $n = 99$), 47.0% (wave 2, $n = 253$), 25.0% (wave 3, $n = 144$). In a regression with wave dummies the wave-3 gap is -22.5 points; with education added, -18.0 (interval -32 to -4); with the full control set and occupation fixed effects, -16.0 (-30 to -2). The five-year probability is 34.4, 34.0, and 40.1 across the three student waves, and winsorized willingness to pay ¥339, ¥340, and ¥243 (raw means ¥431, ¥400, ¥247). Among workers, enrollment at ¥300 is 82.8% in wave 1 ($n = 93$) and 69.1% in wave 2 ($n = 392$), and the five-year probability 31.2 and 35.7. Replacing the sample dummy with five sample-by-wave dummies moves the four headline gradients by less than a point: the liquidity gap on enrollment goes from -19.8 to -19.5 points, the experience gradient from +8.1 to +8.2 per category, the worry gradient on moving away from +15.2 to +15.2 per point, and the expected-raise gradient on moving toward from +17.4 to +17.6 points. Among students alone the corresponding pairs are -16.2 and -17.0, +3.4 and +3.3, +16.9 and +16.8, +21.1 and +20.6.

Power for the asserted nulls. Enrollment at ¥300 in the top against the bottom third of the five-year replacement probability: students 38.4% against 38.4% (difference 0.0, interval -10.5 to +10.5); workers 73.9% against 70.8% (+3.1, -6.7 to +12.9). The 95% interval half-width of that comparison is about 10 points (1.96 standard errors) and the minimum detectable difference at 80% power is about 15 points among students and 13 among workers (2.80 standard errors); the main text's "rules out a gap larger than about 10 points" and "would have detected about 15" are the same comparison read two ways. As a continuous slope with the Section 4 controls (education, income band, liquidity, experience, worry, expected income), enrollment moves by +0.2 points per 10 points of probability among students (interval -2.0 to +2.3) and -0.2 among workers (-2.3 to +1.8). Students' expected starting salary in the top against the bottom third: -¥303 raw (interval -902 to +296) and -¥395 winsorized (-857 to +67); the design could detect a gap of about 16% of the mean.

The liquidity gap among students under successive controls. The liquidity variable combines "fairly difficult" and "unable" to raise ¥5,000 within a week; 95 of the 141 constrained students chose the former. Because most students report no stable income, the income-band control in the main text controls for little, so the specification adds rural hukou (56% of students), city tier, and province. The gap halves and remains.

specification	gap, point s	95% interval	p	n
raw	-22.2	[-31.1, -13.4]	<0.001	496
+ education, income band	-18.2	[-27.7, -8.8]	<0.001	492
+ beliefs (experience, replacement probability, worry, expected income)	-14.3	[-23.9, -4.8]	0.003	481

specification	gap, point s	95% interval	p	n
+ rural hukou, city tier	-13.0	[-22.7, -3.2]	0.009	481
+ province	-12.6	[-22.8, -2.3]	0.016	481
+ gender, AI use, literacy, occupation fixed effects, wave	-11.2	[-21.6, -0.9]	0.033	481

Treated as a four-step ordinal (easy, a strain, fairly difficult, unable), each step goes with 7.8 points less enrollment at ¥300 with the full control set ($p = 0.005$). Rural hukou on its own, with education and income band, goes with 7.9 points less ($p = 0.08$); 34% of rural students are constrained against 22% of urban ones. Among workers only 18 are constrained (2 “unable”), of whom 5 enroll at ¥300; the raw gap is 28% against 73%, a difference of -46 points with an interval from -67 to -25, and the cell cannot carry a controlled comparison.

Experience, worry, and risk as predictors of training demand. With education, income band, and liquidity held constant, workers who say AI already does a sizable part or most of their job are 16.7 points more likely to enroll at ¥300 (interval 9.1 to 24.4; raw 82% against 62%) and state a willingness to pay ¥305 higher (205 to 405; raw ¥734 against ¥387). Among students the same contrasts are +2.9 points (-5.7 to +11.5) and +¥27 (-30 to +84). Worry, with the Section 4 controls (education, income band, liquidity, experience, replacement probability, expected income): among students, the worried are 11.6 points less likely to enroll at ¥300 (-21.2 to -2.1) and state ¥102 less (-167 to -36); among workers, enrollment is 7.0 points higher among the worried (-1.0 to +15.0) and rises 5.7 points per step of the five-point worry scale (0.6 to 10.9), while willingness to pay is ¥50 lower (-160 to +60). Worry therefore lowers demand only among students, and the main text says so. Replacement probability, with the same controls, predicts enrollment in neither sample (see the power paragraph above).

The enrollment ladder. Between ¥300 and ¥800 the arc elasticity of the enrollment share is -1.5 among students and -1.2 among workers; the free-to-¥300 segment has no elasticity, since the price starts at zero, and the survey offers no comparison good, which is why the main text says demand “falls steeply with price” rather than “is elastic.” The ladder is internally consistent with the open willingness-to-pay item: 97% of students and 95% of workers who would enroll at ¥300 state a maximum of ¥300 or more, and 93% and 91% of those who would not enroll at ¥300 state a maximum below ¥300; the median stated maximum is ¥1,000 among those who would enroll at ¥800, ¥500 (workers) or ¥500 (students) among those who would enroll at ¥300 but not ¥800, and ¥100 among those who would enroll only if free. Seven students and three workers say yes at a higher price after no at a lower one; the cleaning rule flags this but drops a respondent only on three flags.

The frontier pay cut. The share who would take a pay cut to enter frontier work (21% in each sample) counts the direction chosen on the I11 item without requiring a positive amount; requiring an amount above zero gives 13% of students and 15% of workers, with median amounts of ¥500 and ¥1,200 a month. Cross-tabulated against the G8 premium item, those choosing the frontier direction are far more likely to sit at the “prefer B” end of G8 (24% of students and 34% of workers at -1, against 2% and 6% of those choosing the safe direction), but a sizable minority of them still demand a premium for the exposed job, so the two items measure related but distinct things.

Student-worker gaps with composition held constant. Students are 70% female, 28% liquidity-constrained, and more often without a bachelor's degree. Holding gender, education, liquidity, AI use, and literacy constant, and then occupation fixed effects, the gaps in worry and expected income survive; the five-year probability and task-count gaps stay at zero; the one-year probability gap of about 4 points is significant throughout.

outcome	raw student minus worker	+ composition controls	+ occupation fixed effects
At least somewhat worried (points)	+15.0 (p <0.001)	+14.8 (p <0.001)	+12.0 (p 0.001)
Expects income to rise (points)	-13.8 (p <0.001)	-11.2 (p <0.001)	-10.7 (p 0.003)
Expects income to fall (points)	+13.3 (p <0.001)	+12.9 (p <0.001)	+15.1 (p <0.001)
Five-year probability (points)	+1.0 (p 0.476)	+1.3 (p 0.394)	+1.8 (p 0.259)
One-year probability (points)	+3.5 (p 0.008)	+4.3 (p 0.002)	+4.1 (p 0.009)
Tasks of ten (count)	-0.0 (p 0.852)	-0.1 (p 0.555)	-0.1 (p 0.691)

Headline shares by wave. Raw shares by wave (sizes after cleaning). The third student wave, three-quarters sub-bachelor, is higher on the five-year probability and lower on enrollment; the direction shares and the worry and income shares move by a few points at most.

sample	wave	n	5-yr probab ility ≥ 50 (%)	mean 5-yr probabili ty	moved toward (%)	moved away (%)	worried (%)	income up (%)	income down (%)	enroll at ¥300 (%)
students	1	99	30	34.4	37	35	62	28	38	47
students	2	253	30	34.0	24	43	68	28	37	47
students	3	144	40	40.1	22	41	69	22	42	25
workers	1	93	24	31.2	40	31	54	38	25	83
workers	2	392	32	35.7	30	34	52	41	26	69

Denominators. Shares in the main text are computed among respondents who answered the item. Item non-response is at most 1.6% on any item used (the frontier pay-cut amount), under 1% on the belief items, and zero on the closed-choice items, so the two conventions differ by at most a few tenths of a point; the one-year probability share for workers, for instance, is 13.4% on either denominator.

G. Comparable surveys

None of the surveys below asks exactly the same question of the same population, so it is the pattern rather than any one gap that matters. The Ipsos item is paraphrased in the table; its exact wording is “And how likely, if at all, do you think it is that AI will replace your current job in the next 5 years?” The 2023 Ipsos press materials report a “likely to replace your job” figure among workers without stating a five-year horizon, and the questionnaire could not be located, so that row is a looser match. “Very or somewhat likely,” in Ipsos and Gallup, and “likely or very likely,” in Anthropic’s survey, are looser categories than a stated probability of 50 or higher, and the Anthropic population, Claude users concentrated in computing and management, differs most from this pilot’s.

Source and date	Population	Item	Result
This pilot, July 2026	981 workers and graduating students,	Percent chance AI replaces your job within	31% of workers and 33% of students at 50 or

Source and date	Population	Item	Result
	Chinese online panel	five years (0 to 100)	higher; means 35% and 36%
Ipsos AI Monitor, March to April 2026	About 1,000 online adults in mainland China, part of a 32-country survey of 23,532	Likelihood that AI will replace your current job in the next five years	52% very or somewhat likely in China (43% in 2024); 35% across 32 countries; 25% in the United States
Ipsos, May to June 2023	22,816 adults in 31 countries, China not included	Likelihood that AI will replace your job, among workers (horizon not stated in the press release)	36%
Gallup, February 2026	23,717 U.S. employees	Job eliminated within five years due to AI or automation	18% very or somewhat likely
Anthropic Economic Index Survey, 2026	About 9,700 Claude users, mostly in computing and management	Chance of losing your own job next year	10% likely or very likely; this pilot's one-year probability is at 50 or higher for 13% of workers and 19% of students
OECD AI surveys, January to February 2022	Finance and manufacturing workers in seven countries	Worried about losing your job to AI within two and ten years	19% and 14% very or extremely worried at two years; 22% and 19% at ten
Pew, October 2024	5,273 U.S. workers	Worried about future AI use in the workplace; AI will bring fewer opportunities	52% worried; 32% fewer opportunities

Ipsos describes its Chinese online sample as more urban, educated, and affluent than the general population, which is also true of this pilot's. I do not compare with exposure estimates such as the World Bank's figure that 4.5% of existing jobs in low- and middle-income countries are technically automatable by generative AI (*World Development Report 2026*, press release of August 4, 2026): that figure has no time horizon, measures technical feasibility rather than expected displacement, and rests on a 135-country background paper (Gmyrek, Viollaz, and Winkler, ILO Working Paper 166, 2026) that excludes China for lack of microdata.

H. Data notes for the exhibits

Every exhibit is one row of at most two panels, 6.5 by 3.4 inches, students in orange and workers in grey; in Exhibit 3 panel B only, colour marks direction (blue toward AI-related work, dark red away from it), not sample. Numbers are recomputed from the analysis files (496 students, 485 workers) and match the fact files behind this report. Confidence intervals on shares are 95% Wilson score intervals, which stay inside 0 to 100 and are not zero-width when an observed rate is 0% or 100%; intervals on means are normal (t-based).

- *Exhibit 1.* Panel A plots the mean stated percent chance the job is replaced at one, five, and ten years and the mean expected share of the job's ten tasks AI could do in five years. An imputed "today" task share (mapping the four-category C0 answer to 0, 20, 50, and 80%) is not plotted: it is 38% under that map and

26% or 46% under alternative maps, so it is a recode, not a measurement. Panel B gives the C3 distribution as shares of each sample.

- *Exhibit 2.* Workers are split by whether they use AI at work daily or more ($n = 310$) or less often. Fill (solid versus hatched), not colour, marks the rise/fall distinction.
- *Exhibit 3.* Panel A collapses the three residual answers (no change, not considering a job change, cannot say) into one bar because the two surveys did not offer identical residual options. Panel B pools the samples; the “how much AI already does” split treats the nine “don’t know yet” answers as missing, the same rule as Exhibit 5; the income splits are separate yes/no splits (expects a rise; expects a fall), since not rising and falling are not the same thing.
- *Exhibit 4.* Panel A’s thin lines split students by whether they could easily raise ¥5,000 within a week; the worker split (18 constrained) is not drawn. Panel B’s four splits are experience (C0 sizable or most versus little), replacement probability (top versus bottom third within sample), worry (at least somewhat versus less), and liquidity (students only).
- *Exhibit 5.* Panel A folds the thin “none” category into “a few tasks.” Panel B groups respondents into within-sample fifths of the Eloundou-style index of their LLM-coded occupation, plots each fifth at its mean standardized score, and overlays the bivariate individual-level fit line; the printed slope is per within-sample standard deviation. Bin positions are uneven because the index is discrete at the occupation level. Appendix Figure A4 gives the same picture for the Anthropic and Chinese posting indices.
- *Appendix figures.* A1 uses the six census groups of Appendix B and the reweighting of that section. A3 forms Eloundou exposure quartiles within sample and requires all three horizons answered. A5’s panel B is flat then down, with overlapping intervals, so the caption does not describe it as a clear literacy effect on the replacement probability.