

College Graduates in the Labor Market: Geographic Mobility and Sorting into Firms and Occupations*

Chen Liu
NUS

Sibo Liu
HKBU

Yifan Zhang
CUHK

October 1, 2025

Abstract

This paper combines three online labor market platforms—LinkedIn user profiles, Burning Glass job postings, and Glassdoor wages—to construct a novel dataset on college graduates from U.S. universities that contains detailed information on their employers, occupations, and locations, as well as the tasks required in their jobs. We document new facts about the mobility of college graduates and show how mobility patterns differ across universities of varying rankings and locations. Estimating a model in which college graduates choose firms, occupations, and locations, we find evidence of positive sorting of higher-ranked universities matching to cognitively intensive jobs and high-amenity cities. Distance significantly reduces the probability of job matching for most universities, but has little effect for graduates from the most elite institutions. Our analysis also reveals a sizable geographic premium in labor market outcomes for universities located in major cities.

Keywords: the geographic premium of colleges

JEL Codes: I23, J61, R23

*Chen Liu: Department of Economics, National University of Singapore, ecsluoc@nus.edu.sg; Sibol Liu: Department of Accountancy, Economics and Finance, Hong Kong Baptist University, sibolliu@hkbu.edu.hk; Yifan Zhang: Department of Economics, The Chinese University of Hong Kong, yifan.zhang@cuhk.edu.hk. We thank seminar participants at NUS Brownbag for helpful comments. Generous research support was received from the National University of Singapore and the Chinese University of Hong Kong. All remaining errors are our own.

1 Introduction

In a recent graduate employability survey, recruiters from top companies were asked to rank the universities that best prepare students for the workplace. Unsurprisingly, the commonly cited elite institutions appear at the top of the list. While it is well known that graduates from elite universities often receive stronger training and skill development sought by top employers (Dale and Krueger, 2002) and enjoy a substantial wage premium (Chetty, Friedman, Saez, Turner and Yagan, 2020), many of the most prestigious U.S. universities (e.g., Yale, Cornell, University of Michigan) are located in small towns. In contrast, large firms and high-paying jobs are disproportionately concentrated in major metropolitan areas, but some of which host few top universities.

The geographic separation between the top universities and the major employers is based on historical, economic, and institutional factors. Many of the most prestigious universities were founded in the 18th and 19th centuries, often in small towns in the Northeast. The Morrill Land-Grant Acts of 1862 granted federal land to states for the creation of public universities, intentionally situated outside major cities to better serve rural communities. Many of these universities later became flagship public universities in their respective states. However, economic activities and high-paying jobs shifted geographically due to technological innovation, westward territorial expansion, and population growth. Once concentrated in the Northeast during the colonial era, economic hubs eventually expanded into the Midwest during the industrial age, while the South and West rose dramatically in the late 20th century.

This spatial mismatch between talents and employers suggests that graduates must be highly mobile to access elite firms and secure high-wage opportunities. In contrast, attending universities in large metropolitan areas may benefit from proximity to top employers. While both university ranking and location likely play important roles in shaping graduates' job search outcomes, addressing these questions requires rich microdata.

In this paper, we bring together three online labor market platforms—LinkedIn user profiles, Burning Glass Technologies (BGT), and Glassdoor—to construct a rich individual-level dataset. We study the geographic mobility and job search of fresh college graduates and analyze how these outcomes differ across universities of varying rankings and locations. Our LinkedIn data is a 2018 snapshot of user profiles that contains rich information on education and employment histories. The BGT data contain job posting information with the tasks required by each employer and detailed job ti-

tles. Our Glassdoor data provides entry-level wages by employer and job title. We standardize the information—including employer names, locations, and occupations—and construct a crosswalk to link these datasets.

We obtain a sample of more than 240 thousand fresh graduates who received their bachelor’s degrees between 2016 and 2018, from 266 U.S. universities ranked in the World University Rankings (WUR).¹ We show that our sample captures nearly 10% of all bachelor’s degree recipients from these institutions and that it is highly representative when compared with benchmark datasets such as the U.S. Census and the Integrated Postsecondary Education Data System (IPEDS). Compared to the matched employer–employee datasets commonly used in the literature, our dataset offers richer information: on the labor supply side, we observe the college each individual attended and their field of study; on the demand side, we observe the tasks required by employer and occupation; and on the spatial dimension, we observe the geographic location of the college and the current location of work.

Using this newly assembled dataset, we begin by documenting four new facts about the job search behavior of fresh college graduates. First, although more than half of recent graduates left their university cities, most moved within the same state or to neighboring states.² Second, graduates from higher-ranked universities are more mobile, with those from the Top 20 institutions being the most mobile group.³

Third, mobility patterns depend on the job opportunities available in university-hosting cities. High-wage cities retain a larger share of their locally educated graduates than low-wage cities. The magnitude of this difference is substantial: comparing Boston (one of the high-paid U.S. cities) with Pittsburgh (PA), the former retains 34.2 percentage points more of its graduates than the latter.

Fourth, among graduates from the same university, movers on average obtain jobs with higher (Glassdoor) posted wages, greater (BGT) cognitive and social task requirements, and locate in higher-amenity cities compared to stayers. The compensating differential is larger for universities in low-wage cities and smaller for those in high-wage cities. These findings highlight the importance of worker mobility and the geographic location of universities in shaping labor market outcomes.

Motivated by these facts, we estimate a model of college graduates sorting into firms and occupations to evaluate the relative importance of several factors in determining

¹See Section 2.1 for sample construction in detail.

²Throughout the paper, we use the terms city and commuting zone interchangeably.

³Graduates from top-ranked universities are more likely to move farther away and less likely to stay in the same city or state, even after controlling for local wages and the share of in-state students.

the job match. Throughout, we refer to a job as the combination of a firm and an occupation.⁴ Our model produces a gravity-type equation of worker sorting into jobs, where we regress log employment shares on university and firm–occupation fixed effects. These fixed effects account for the availability of alternative employment options for each group, as well as firm–occupation wage efficiency units and local prices, while the residual variation is explained by the observables of interest.

Our empirical analysis incorporates three sets of observed factors. In theory, the most emphasized determinant of worker-to-firm sorting is task complementarity (Eeckhout and Kircher, 2011). In our setting, it refers to the supermodular relationship between skills acquired at universities and the tasks demanded by firms. Positive sorting occurs if graduates from higher-ranked universities have a comparative advantage in working at more productive firms. Our rich dataset allows us to empirically measure university–job-specific productivity. Building on the literature that emphasizes the importance of education for skill acquisition (Hanushek and Kimko, 2000, Hanson and Liu, 2021) and the varying relevance of college-imparted skills across jobs (Acemoglu and Autor, 2011, Atalay et al., 2018), we model and estimate labor productivity using interactions between university rankings and four commonly used task requirements (cognitive, social, routine, and manual) obtained from BGT.

Second, we incorporate city amenities to examine how graduates from universities of different rankings sort based on amenities. Third, we account for moving costs to capture the potential heterogeneous effects of geography on job choice across universities, estimating interaction terms between university rankings and geographic distance. To account for potential nonlinear effects of distance and unobserved region-specific factors, we include indicators for whether a university and a job are located in the same city or the same state.

Our estimates reveal several findings. First, we find strong evidence of positive sorting into cognitive tasks. Specifically, as cognitive task requirements become one standard deviation higher (the difference of working as a computer scientist between Google and Sanmina corporation, an American electronics manufacturing services), graduates from Top 20 universities have a 0.199 higher in log points or about 19.9% more likely to match to the more cognitive-intensive job, relative to the benchmark group (outside Top 1000), holding other factors the same. Second, we find evidence of positive sorting on amenities: graduates from Top 20 universities, relative to the bench-

⁴In the paper, we refer to firms as the combination of a company name and the location of its establishment. Thus, establishments of the same company in different locations are treated as distinct firms in our analysis.

mark, are 6.8% more likely to choose a job in Seattle (99th percentile in amenities) than in Detroit (25th percentile in amenities). While prior studies have primarily examined spatial sorting and amenities between college- and non-college-educated workers (Diamond, 2016, Diamond and Gaubert, 2022), our findings complement this literature by highlighting positive sorting across graduates from different universities.

Third, we find that geographic distance reduces the likelihood of a job match for most universities, but has no effect on the job choices of Top 20 universities. The distance effect for lower-ranked universities is sizable: for example, comparing graduates from Boston College and UC Davis—both ranked outside the Top 20 but within the Top 200 globally—the former are 18.2% more likely to obtain a job in New York City than the latter. In contrast, no such difference exists between Harvard and Stanford graduates in terms of matching to the same job in New York. The results hold when we additionally control for the interaction between these geographic variables and variation in the origin composition of student bodies, measured by the fraction of in-state enrollment.

The Dictionary of Occupational Titles (DOT) and the Occupational Information Network (O*NET) both provide occupation-specific task content, have been widely used to study worker sorting across occupations (Acemoglu and Autor, 2011, Yamaguchi, 2012, Deming, 2017). Our task variables are taken from Burning Glass Technologies (BGT), which were first used by Hershbein and Kahn (2018) to study labor demand adjustments during the Great Recession, and have been used to study the effect of technological change on earnings dynamics (Deming and Noray, 2020), the consequences of job loss (Braxton and Taska, 2023), and in explaining regional disparities in earnings (Atalay, Sotelo and Tannenbaum, 2024). Our dataset enables us to estimate task-based sorting at a much more granular level, capturing heterogeneity across both firms and occupations.

To provide insight into the relative importance of sorting into firms versus sorting into occupations, our model also yields an alternative gravity-type specification that relates the share of graduates from each university employed in a given U.S. firm—conditional on the same occupation—to the factors emphasized above. This specification introduces additional controls for university–occupation fixed effects, further netting out the role of alternative employment options for each university–occupation pair and exploits residual variation across firms within the same occupation. Comparing these estimates with those from the baseline, the evidence suggests that sorting along both two dimensions—firms and occupations—is quantitatively important in

shaping the overall positive sorting on cognitive tasks.

We show that our estimated results are not contaminated by potential confounding factors in four ways. First, we augment our model with college selectivity criteria, which proxy for students' average ability prior to entering college, rather than academic standards or the quality of education or institution reputation that are captured in college rankings. Second, we consider the university's field emphasis in education, the share of graduates in STEM majors. Third, we use different sources of university rankings or more disaggregated groups. Including these factors, our messages remain the same as the baseline analysis. Fourth, we show that our results are not contaminated by university-to-firm network effects. Under common assumptions where networks are a function of historical university-to-job matching probabilities, we show that accounting for network effects does not alter the model specification but instead requires only a re-interpretation of the regression coefficients. Empirically, we also control for a network measure capturing university alumni exposure to specific firms and examine which sorting channel—task, amenity, or geography—the network primarily operates through. Our findings indicate that it operates mainly through the geographic channel, with little effect on sorting by task or amenity.

While students attending universities in large metropolitan areas may benefit from easier access to high-paying jobs, causally identifying the wage premium due to a university's geographic location is challenging.⁵ Importantly, our estimated model informs how the assignment of talent to jobs is shaped by various factors, providing a platform to quantify how geography shapes the labor market outcomes of fresh graduates from each university.

Through a counterfactual exercise that shuts down geographic factors, we estimate the job-matching probabilities and wages that graduates from each university would have experienced. Geography plays an important role: for universities located in the Bay Area, geography increases the likelihood of matching to top 5% paid U.S. jobs by about 2 percentage points (ppts). By contrast, for universities in Midwestern and Southwestern cities, geography reduces the same likelihood by more than 3 ppts. These difference amounts to 5 ppts in access to top 5% paid jobs.

We also find that universities in the Bay Area enjoy a geographic premium in annual salaries, which remains sizable even after adjusting for regional price differences. Notably, this premium is nontrivial when benchmarked against the wage premium associated with university rankings. Relative to Midwestern and Southwestern cities

⁵Identification might require random assignment of students to universities and that universities across cities be identical in training and all other respects.

that host major universities—such as Lafayette (LA), Ann Arbor (MI), Bloomington (IN), Lansing (MI), and Columbus (OH)—the Bay Area premium, in both nominal and real terms, exceeds half of the average premium from attending a Top 200 globally ranked university and amounts to between one-fifth and one-third of the premium from attending a Top 20 university.⁶

Our paper relates to the literature on labor market outcomes across various universities. Since information on the universities attended is rarely publicly available in large samples, [Chetty, Friedman, Saez, Turner and Yagan \(2020\)](#) link multiple administrative datasets to examine how attending elite universities shapes income segregation and inter-generational mobility. [Chetty, Deming and Friedman \(2023\)](#) show that attending an Ivy-Plus college instead of the average flagship public university triples their chances of working in a prestigious firm. [Zimmerman \(2019\)](#) shows that attending elite universities substantially increases mobility into top-paid jobs and raises income in Chile. We are not aware of any prior study that has examined the geographic mobility of recent college graduates across a representative set of U.S. universities. Complementing this literature, we estimate the assignment function of talent to jobs and provide a model-based quantification that highlights a sizable geographic premium associated with university location.

The implications of labor mobility on aggregate productivity have received increasing attention in developing countries ([Lagakos and Waugh, 2013](#), [Tombe and Zhu, 2019](#), [Bryan and Morten, 2019](#)). [Pellegrina and Sotelo \(2021\)](#) study how the migration of farmers to western Brazil shaped regional comparative advantage in agriculture. In the United States, geographic mobility has been studied in response to the Great Recession ([Cadena and Kovak, 2016](#)) and import competition ([Bound and Holzer, 2000](#), [Greenland et al., 2019](#), [Autor et al., 2025](#)), and in determining the aggregate productivity ([Albert and Monras, 2022](#)). We complement the literature to analyze patterns of geographic mobility across universities of different rankings and locations.

Finally, the literature on worker–firm sorting often relies on the Abowd–Kramarz–Margolis (AKM) framework ([Abowd et al., 1999](#), [Card et al., 2013](#), [Lopes de Melo, 2018](#), [Song et al., 2019](#), [Bonhomme et al., 2023](#)). By combining multiple datasets, our sample provides richer information on both the supply and demand sides, allowing us to study worker–firm sorting in task space and across geographic dimensions within a unified framework. Our findings provide empirical validation for the theoretical framework of worker–firm matching ([Eeckhout, 2018](#)).

⁶These premiums are relative to institutions ranked outside the Top 1000.

The paper is organized as follows. Section 2 briefs the data and presents motivating facts. Section 3 describes the model. Section 4 discusses the estimation results, and Section 5 tests potential confounders. Section 6 estimates the geographic premium. Section 7 concludes.

2 Data and Facts

This section describes our data sources and presents new facts on the job-matching and spatial sorting of college graduates.

2.1 Data Sources

Our data come from multiple sources, which are briefly described below, with further details provided in Appendix A.

LinkedIn. Our first data source is purchased from Revelio Labs, which processes LinkedIn profiles containing detailed résumés for over 52 million users, captured in a 2018 snapshot.⁷ The dataset provides rich self-reported information on the firms and job titles individuals have held, the institutions they attended, the degree awarded, fields of study, and the start and end dates of each job and degree. Some users report multiple universities for their undergraduate studies, which might involve exchange programs or institutions where they completed minors. We define the primary university as the institution where the individual spent the longest time.

Because we focus on job matching among fresh college graduates, we restrict the sample to individuals who earned a bachelor’s degree (as their highest degree) between 2016 and 2018, graduated from a U.S. institution, and are currently employed by a U.S. firm.⁸ As fresh graduates first enter the job market, they are likely to maintain accurate and up-to-date information on their LinkedIn profiles. Our objective is to measure individuals’ first “primary” job immediately after graduation. We use firm and occupation information from their 2018 job record. For those reporting multiple jobs in their profile, we select the position in which they had the longest tenure. We exclude individuals currently working as interns or enrolled in master’s or Ph.D. programs.

Burning Glass Technology. Our second data source is the universe of job postings,

⁷For our research purpose, we restrict the full sample to individuals with ongoing job experience in 2018.

⁸We include 2016 graduates to increase the sample size. Since most LinkedIn users report the same job in both 2017 and 2018, we focus on employment outcomes in 2018.

which measures task content by occupation and employer.⁹ We use BGT job posting data from 2018. Following [Spitz-Oener \(2006\)](#) and [Atalay et al. \(2020\)](#), we apply text analysis to construct four commonly used task measures from job advertisements: cognitive, social, routine, and manual. For each task, we compute percentile rankings across all postings and then average them across postings that share the same firm name and occupation code.¹⁰ All task measures are standardized to range between 0 and 1. In Section 4, we further standardize each BGT task variable by its standard deviation to ease interpretation.

Glassdoor. Firm-occupation-occupation-specific wage information is obtained from Glassdoor.¹¹ Glassdoor is an online platform where users voluntarily and anonymously report wages and review employers. Workers are incentivized to contribute through a “give-to-get” policy, whereby access to information provided by others requires contributing one’s own first ([Martellini et al., 2024](#)).

We obtain a snapshot of Glassdoor data collected between September and October 2024, which includes detailed wage information by firm, occupation, location, and years of experience. To integrate this dataset with our other sources, we match job titles to Standard Occupational Classification (OCCSOC) occupation codes, firm names to those in Burning Glass and LinkedIn, and job locations to commuting zones (hereafter, CZ). We use the entry-level wages (0-3 years of experience), and at each firm, CZ, and two-digit SOC occupation levels.

In our Glassdoor data, 16% of entry-level wages are reported as exact values, while 84% are reported as intervals. Importantly, since these intervals are generally narrow among entry-level wages, we choose the midpoint as the wage for each job.¹² Finally, we deflate wages to 2018 dollars using the Consumer Price Index (CPI) from the Bureau of Labor Statistics.

Amenity Index. Our city amenity measure uses a single index taken from [Diamond](#)

⁹Notably, the BGT data are based on online job postings and are available beginning in 2007. For earlier periods, see [Atalay, Phongthientham, Sotelo and Tannenbaum \(2020\)](#), who measure job tasks using newspaper postings dating back to 1960.

¹⁰Because many postings list multiple possible locations with identical task descriptions, our BGT measures do not capture variation across locations within the same firm–occupation combination, in order to maintain consistency.

¹¹In a recent study, [Martellini et al. \(2024\)](#) use an individual-level sample of Glassdoor data to estimate college quality and assess its role in explaining cross-country variation in entrepreneurship and innovation. By contrast, our data are accessed directly from the public Glassdoor platform and are less granular, containing wage information at the firm–occupation–experience level.

¹²For jobs with salaries reported in intervals, we calculate the interval range relative to the midpoint. On average, this ratio is 11%, with a maximum of 20%, indicating that reported entry-level wages fall within a relatively narrow range.

(2016). The measure is based on a collection of rich amenity variables from six different categories: the retail environment, transportation infrastructure, crime, environmental quality, school quality, and local skill demand. The single amenity index is calculated as the first component of the principal component analysis. We convert it into a percentile ranking, normalized to range from 0 to 1.

University Rankings. We group universities into groups based on the World University Rankings (WUR), which is widely recognized as a proxy-based ranking and has been recently used in Martellini et al. (2024). The ranking is based on factors such as academic reputation, employer reputation, faculty-student ratio, and citations per faculty, with more weight given to the latter two factors. We also use rankings from US News as an alternative measure.

2.2 Data Processing

We extensively process data from the three online labor market platforms, as detailed in Appendix A.2. In brief, our work involves four main tasks. First, we standardize employer names across the three datasets. Second, for LinkedIn and BGT, the Standard Occupational Classification (OCCSOC) codes are directly available. We use a large language model (ChatGPT-4o) to map job titles in Glassdoor to the OCCSOC code. In our analysis, we aggregate the detailed OCCSOC occupations into 22 broad categories based on their first two-digit codes. Third, for LinkedIn data, we standardize self-reported university names to align with institutional names in the World University Rankings and U.S. News. Finally, we process geographic information from LinkedIn profiles to construct commuting zone codes, identifying both the location of universities and their place of work.

Putting all together, we restrict the sample to LinkedIn users who: (1) received their bachelor’s degree (as their highest degree) between 2016 and 2018 from a U.S. institution ranked in the Top 2000 in WUR; (2) were working in the United States in 2018; (3) have employer names and occupational titles that are clearly identified and can be matched to BGT data; and (4) have identifiable employer geographic locations. To improve estimation precision, we further limit the sample to U.S. universities with at least 100 LinkedIn users.

We also restrict to 266 U.S. universities that are ranked in the WUR.¹³ We group universities into four tiers using the WUR: the Top 20 globally ranked (Top 20), those ranked 21–200 globally (Top 21–200), those ranked 201–1000 globally (Top 201–1000),

¹³WUR ranks the top 2,000 universities globally. Among them, 348 are U.S. universities.

and all others (including those ranked outside the Top 1000). In robustness checks, we also use more detailed grouping.

The sample covers 266 universities, 25492 distinct firms, and a total of 244,632 LinkedIn users.¹⁴ A firm appears in our sample if we observe at least one LinkedIn user employed there.

According to IPEDS, the 266 universities in our study awarded 2.52 million bachelor’s degrees (including international students) between 2016 and 2018, implying that our sample covers approximately 10% of this student population. Because some bachelor’s degree recipients do not enter the labor force (e.g., pursue further study, are unemployed, or leave the U.S.), our sample captures an even larger share of graduates who enter the U.S. labor force.

2.3 Sample Validation

Since the BGT data have been extensively used and validated in prior studies ([Hershbein and Kahn, 2018](#), [Atalay et al., 2024](#)), we focus on validating our LinkedIn and Glassdoor samples. Below, we briefly describe the six validation exercises we performed, with full details provided in Appendix B.

We validate the LinkedIn data in three ways. First, we assess spatial representativeness by comparing each CZ’s share of national college-graduate employment between LinkedIn and the American Community Survey (ACS). Second, we assess occupational representativeness by comparing employment shares across two-digit SOC occupations between LinkedIn and the ACS. In both cases, correlations exceed 0.9, with OLS regression slopes close to one and R^2 values above 0.9. Third, we evaluate the representativeness of graduating class sizes across U.S. universities by comparing the national share of graduates by university between LinkedIn and IPEDS. These exercises show that LinkedIn data are broadly representative of the U.S. college-educated workforce and graduating class sizes.

As a recent study documents that online job postings contain little wage information ([Batra, Michaud and Mongey, 2023](#)), we conduct extensive validation to show that Glassdoor wage data provide meaningful information for college graduates. Specifically, we compare average annual wages from Glassdoor and the ACS across occupations, commuting zones (CZs), and CZ–occupation pairs. In all cases, we find strong correlations, indicating that Glassdoor wages capture meaningful variation across re-

¹⁴Throughout the paper, we define a firm as the combination of a company name and the location of its establishment.

gions and occupations.

2.4 Empirical Facts

Top U.S. universities and high-paying jobs are not clustered in the same locations. On the one hand, many of the most prestigious universities are located outside major metropolitan areas. On the other hand, many metropolitan cities that host large firms and high-paying jobs have few, if any, elite universities. To compare the geographic distribution of top U.S. college graduates (supply) with that of high-paying jobs (demand), Figure 1 plots the relationship between each city’s share of top-paid U.S. jobs (y-axis) and its share of locally educated bachelor’s degree recipients from top US universities (x-axis), each expressed as a share of the national total. We use data from the 2018 American Community Survey (Ruggles et al., 2010) to define top-paid jobs as those that belong to the top 5% of the income distribution among all college-educated wage earners. To measure the supply, we use the World University Rankings (WUR) to define top U.S. universities as those ranked among the Top 20 globally or belonging to the Ivy League. The number of bachelor’s degree recipients from each top university is obtained from the IPEDS. The dashed line in the figure represents the 45-degree line.

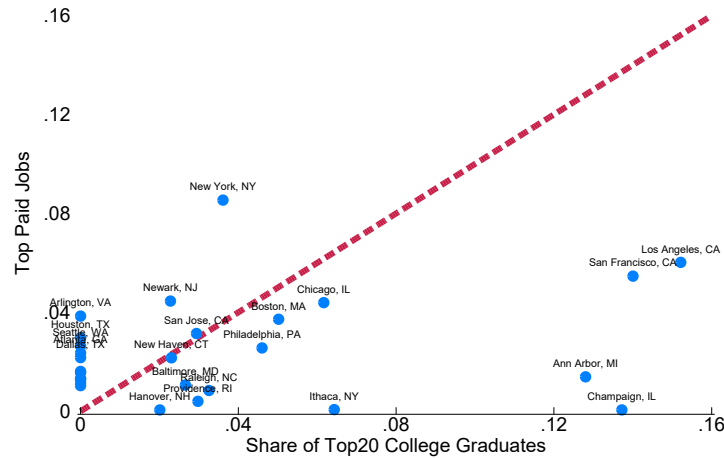


Figure 1: Spatial Distribution of Top College Graduates and Top Paid Jobs.

Notes: The y-axis city’s share of top jobs in the United States. We define top jobs as those that belong to the top 5% of the income distribution among all college-educated wage earners. The x-axis shows each city’s college graduates as a share of the national total, among universities ranked in the global Top 20 or belonging to the Ivy League.

Figure 1 shows that most cities are clustered either near the y-axis or the x-axis. For clarity, we restrict the plot to cities that either account for more than 1% of the

top U.S. jobs or host at least one top-20 university. On the one hand, cities such as Arlington (Virginia), Seattle, Dallas, Houston, and Atlanta lie near the y-axis. These cities host a sizable share of high-paying jobs, but produce few top college graduates. Similarly, New York City and Newark are located close to the y-axis, reflecting that they host a relatively large share of top jobs compared with the share of top graduates produced locally.

On the other hand, several cities near the x-axis—such as Ann Arbor (MI), Champaign (IL), and Ithaca (NY)—are home to prominent universities but offer relatively fewer high-paying job opportunities. In addition, many cities fall between the 45-degree line and the x-axis, representing locations that produce a disproportionately large share of top college graduates while hosting a smaller share of top jobs. Similar patterns persist when top-paid jobs are defined using the 75th or 90th percentile of the income distribution (see Appendix Figure C.1).

This geographic separation between the “best jobs” and the “best talents” suggests that college graduates must move in order to access elite firms and better job opportunities. We now present four new facts on the job search for fresh college graduates.

Fact 1. Fresh graduates are highly mobile across cities, with the majority moving within the same state or to neighboring states.

Column (1) of Panel A in Table 1 reports the fraction of graduates who remain in the same CZ as their college. Overall, more than half of graduates leave their college city for employment, with 47.7% staying and working in the same CZ.

Despite the low retention rate within the city, column (2) shows that 70% of graduates are employed within the same state. On average, graduates travel 320 miles to secure a job, a distance similar to that between Boston and Philadelphia, or from Phoenix to San Diego. Compared to the vast distances spanned across the U.S. continent (nearly 3,000 miles from northeast to northwest, and over 3,300 miles from southeast to northwest), these figures indicate that, on average, graduates tend to work relatively close to their place of study.

Fact 2. Graduates from higher-ranked universities are more mobile, with the most notable group being graduates from the Top 20 universities.

Column (1) shows that within-CZ retention is 56.1% for universities outside the Top 1000 but drops sharply to 37.6% for Top 20. Column (2) shows that within-state retention is as high as 80% for universities outside the Top 1000, falling to 70% for Top 201–1000 universities, 65.4% for Top 21–200, and only 55.1% for Top 20. Column (3)

shows that graduates from Top 20 universities travel, on average, 605 miles for jobs, compared to 380 miles for Top 21–200 universities, 279 miles for Top 201–1000, and 243 miles for universities outside the Top 1000.¹⁵

Table 1: Retention Rates by College Rankings

Panel A: Summary Statistics				
College	(1) Retention Rate Commuting Zones	(2) Retention Rate State	(3) Distance Traveled (Miles)	(4) Weekly Wages of Hosting City
All	0.477	0.703	320	1547
Top20	0.376	0.551	605	1733
Top21-200	0.431	0.654	380	1520
Top201-1000	0.475	0.704	279	1489
Top1001-2000	0.561	0.800	243	1623

Panel B: Mobility on Rankings and Wages				
	(1) Retention Rate	(2) Retention Rate	(3) Distance	(4) Distance
Top20	-0.124** (0.057)	-0.111*** (0.039)	263.652*** (37.609)	260.575*** (48.866)
Top21-200	-0.075* (0.039)	-0.052 (0.032)	94.000*** (25.853)	51.963 (40.353)
Top201-1000	-0.069** (0.030)	-0.027 (0.026)	27.434 (19.858)	9.886 (31.920)
Log wage	1.423*** (0.106)		-414.403*** (70.108)	
In-state Enrollment	0.424*** (0.055)	0.425*** (0.042)	-537.340*** (36.581)	-520.504*** (52.034)
Observations	258	153	258	153
R^2	0.57	0.85	0.72	0.78

Notes: In Panel B, the dependent variable is within-CZ retention rate in Columns (1) and (2), and average distance traveled in Columns (3) and (4). Columns (1) and (3) control for the log of average wages in the university-hosting city and for state fixed effects, whereas Columns (2) and (4) include CZ fixed effects.

This pattern persists when we examine the distance traveled by movers.¹⁶ Figure 2 shows that, for graduates of lower-ranked universities, more than half relocate within 300 miles, and the vast majority move within 600 miles. In contrast, graduates of the top 20 universities are far more geographically mobile: 25.7% relocate more than 1,800 miles for a job.

The last column of Panel A reports the average weekly wage in the university-hosting cities.¹⁷ Top 20 universities, on average, are situated in the highest-paying cities than other groups. Intuitively, local economic opportunities should influence retention rates and mobility, which we analyze next.

¹⁵These average distances traveled are estimated using both movers and stayers.

¹⁶Movers are defined as graduates who relocate to a different CZ for employment after graduation.

¹⁷We compute the average weekly wage using the ACS, restricting the sample to individuals who work full-time and hold a college degree.

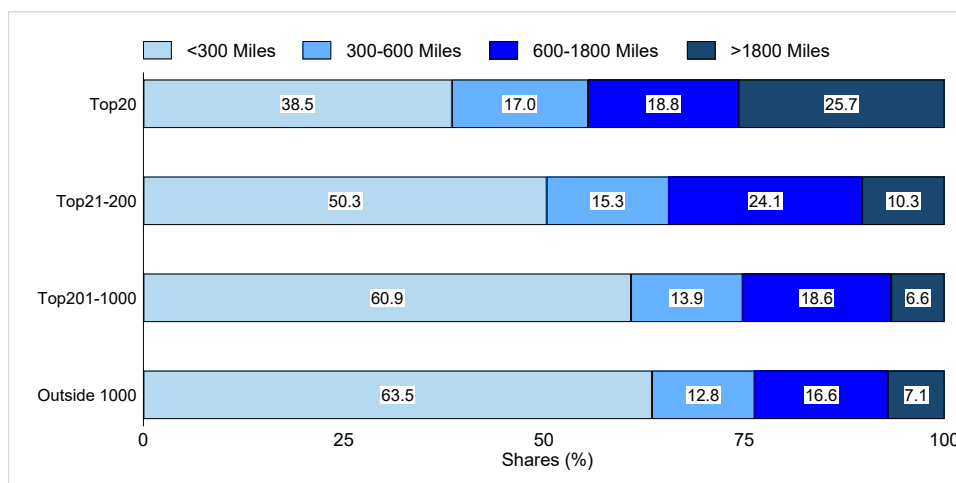


Figure 2: Distribution of Distance Traveled by University Group, Among Movers

Notes: The plot displays the share of graduates who match to jobs in each distance category, conditional on movers (graduates who relocate to a different CZ for employment after graduation). County-to-county distances are obtained from the NBER County Distance Database, and CZ-to-CZ distances are computed as a population-weighted average.

Fact 3. High-wage cities retain a larger share of their locally educated graduates than do low-wage cities.

To systematically summarize the mobility patterns, we regress CZ-level retention rates on three dummy variables—Top 20, Top 21–200, and Top 1000, as well as the log of average weekly wages of the university-hosting city. Universities outside the Top 1000 serve as the benchmark group.

Next, we use the IPEDS to compute the share of undergraduate students who are in-state enrollees. Controlling for this variable accounts for the origin composition of student bodies, thereby addressing potential home bias in location choice (Kennan and Walker, 2011). In addition, we include state fixed effects to absorb cross-state differences in employment opportunities for college graduates. Eight out of 266 universities that are the sole institution (in our sample) within their state are automatically excluded.

Column (1) reports a coefficient of -0.124 (s.e.= 0.057) on the Top 20 dummy, indicating that graduates from Top 20 universities are 0.124 percentage points (ppts) less likely to remain in the same city compared with graduates from universities outside the Top 1000. In addition, the in-state enrollment share also matters for retention rates: a 10 ppts increase in the share of in-state students is associated with a 4.24 ppts higher retention rate.

The estimated wage coefficient is positive, at 1.423 (s.e. = 0.106), suggesting that

higher wages increase the likelihood that cities retain locally educated college graduates. The magnitude of the wage effect on retention rate is large: comparing Boston (among the top 10 highest-paid US cities) with Pittsburgh (PA), the former has a 34.2 percentage points higher retention rate for their college graduates than the latter.¹⁸ In Column (2), we replace the log wage of the university-hosting CZ and state fixed effects with CZ fixed effects. The CZ fixed effects capture all unobserved, CZ-specific factors that influence employment outcomes.¹⁹ We find similar results under this specification.

In Columns (3) and (4), the outcome variable is replaced with the distance traveled between the college city and the job city. Compared to the benchmark group, the Top 20 university graduates travel an additional 264 miles on average to secure employment. Graduates from Top 21–200 universities travel more than 90 miles farther than the benchmark group. Column (4), again, reports similar estimates for rank dummies and log wages when including CZ fixed effects.

Fact 4. There is a mover premium: movers, compared with stayers, tend to relocate to access better employment opportunities and higher amenities. The mover premium is high for universities located in low-wage cities and low for universities located in high-wage cities.

Using individual-level data, we regress the cognitive, social tasks (BGT), and log wages (Glassdoor) on a binary variable of mover, which equals one if an individual’s CZ of work differs from the CZ of his/her university, and zero otherwise. For all regressions, we control for university fixed effects, and we compare movers and stayers within the same university.

Panel A of Table 2, Columns (1) and (2), shows that movers tend to obtain jobs with task requirements 3.0 percentiles higher in cognitive and 2.2 percentiles higher in social task intensity, compared to stayers. Because occupations or jobs that are cognitive and social intensive are often the ones that pay higher wages (Deming, 2017, Atalay et al., 2024), we consider better jobs as those that have high cognitive or social task requirements. Column (3) shows that movers also tend to land in jobs that pay significantly higher, about 13.8%, relative to stayers. Column (4) shows that movers settle in cities with local amenities 7.3 percentiles higher than stayers. The evidence points to the existence of a compensating differential: migration is indeed costly, and college graduates tend to move farther to secure better employment opportunities or to

¹⁸The value is computed as $1.423 \times (7.52 - 7.28)$, where 7.52 and 7.28 are the log weekly college wages for Boston and Pittsburgh (PA), respectively.

¹⁹Because we condition on CZ-level fixed effects, we compare universities of different rankings located within the same CZ, and universities that are the sole institution within their CZ are excluded.

access high-amenity locations.

Because task variables and wages vary across firms and occupations, we also report estimates that condition on university and occupation fixed effects. When comparing movers and stayers from the same university who enter the same occupation but different firms, the mover premium remains statistically significant (albeit attenuated) for social tasks and wages, while becoming imprecisely estimated for cognitive tasks; see Columns (5)–(7).

Table 2: Job Outcomes of Movers and Stayers

	(1)	(2)	(3)	(4)	(5)	(6)	(7)
	Cognitive	Social	Log Wages	Amenity	Cognitive	Social	Log Wages
Panel A: The Mover Premium							
Movers	0.030*** (0.010)	0.022*** (0.007)	0.138*** (0.021)	0.073*** (0.006)	0.008 (0.007)	0.011*** (0.003)	0.078*** (0.014)
Occupation FE					✓	✓	✓
Observations	244,632	244,632	244,552	243,195	244,632	244,632	244,552
R^2	0.07	0.03	0.13	0.42	0.29	0.20	0.56
Panel B: The Mover Premium by University Hosting Cities							
Movers	0.032*** (0.011)	0.024*** (0.008)	0.147*** (0.023)	0.086*** (0.009)	0.009 (0.007)	0.013*** (0.004)	0.085*** (0.015)
Movers $\times \ln W_g^{\text{Host}}$	-0.090** (0.036)	-0.097*** (0.018)	-0.417*** (0.064)	-0.558*** (0.019)	-0.051** (0.025)	-0.066*** (0.015)	-0.310*** (0.054)
Occupation FE					✓	✓	✓
Observations	244,632	244,632	244,552	243,195	244,632	244,632	244,552
R^2	0.07	0.03	0.14	0.43	0.29	0.20	0.56

Notes: All regressions are estimated using individual LinkedIn users and include controls for university fixed effects. Columns (1) to (4) include university fixed effects. Columns (5) to (7) include university and occupation fixed effects. The dependent variable in column (3) is log wages, measured using Glassdoor data for the firm–occupation in which a LinkedIn user is employed. $\ln W_g^{\text{Host}}$ is the demeaned log weekly wage in university-hosting cities obtained from the ACS. Standard errors in parentheses are clustered at the firm-occupation level.

Mover premium depends on job opportunities available in university-hosting cities. We augment the model by including an interaction between the mover dummy and $\ln W_g^{\text{Host}}$, the demeaned log weekly wage in university-hosting cities obtained from the ACS.²⁰ We find negative and precisely estimated interaction coefficients that are both sizable and economically meaningful across all columns. For example, a 0.1 log-point lower in the hosting city’s log wage increases the mover premium of the matched job by 0.9 percentiles in cognitive tasks (column 1), 0.97 percentiles in social tasks (column 2), and 4.17% in wages (column 3).

Motivated by these documented patterns, we next estimate a model of how gradu-

²⁰We demean the log wage to ease interpretation of the main effect: in Panel B, the coefficient on the mover dummy captures the mover premium for cities at the sample mean of the log wage.

ates sort into firms, cities, and occupations, enabling a comprehensive assessment of the relative importance of multiple factors in shaping labor market sorting.

3 The Model

We present a model where college graduates decide which firm (indexed by f) and occupation (indexed by o) to pursue, taking into account wages, search costs, the amenities of the city (c) where the job is located, and individual tastes. We group workers by their university (indexed by g). Again, a firm refers to the combination of an employer name and the location of its establishment. Since the firm index f and the university index g already nest information in the location information for jobs and schools, respectively, we omit the city subscript c when its exclusion does not create confusion.²¹

3.1 Preference

For individuals who attend university g , their utility of living and working in city c , firm f and occupation o is

$$U_{fo}^i = \left(\frac{C_{fo}^g}{1 - \kappa} \right)^{1-\kappa} \left(\frac{H_{fo}^g}{\kappa} \right)^\kappa a_c^g \tau_{fo}^g \varepsilon_{fo}^i. \quad (1)$$

In Equation (1), C_{fo}^g represents consumption of tradable goods, and H_{fo}^g represents consumption of non-tradable goods. κ denotes the expenditure share allocated to non-tradables. a_c^g captures the utility obtained from the local amenity in city c 's felt by group g . τ_{fo}^g is the costs incurred during job searching, specific to university g graduates to firm f and occupation o . To ease interpretation, we model search cost as an iceberg cost (or take-home utility) with a higher value indicating lower costs. ε_{fo}^i is idiosyncratic preferences for job $\{f, o\}$ which allows workers to have heterogeneous preferences over the work environments of different potential employers and occupations.

Individuals face the following budget constraints $P_c C_{fo}^g + R_c H_{fo}^g \leq W_{fo}^g$, where P_c is the price for tradables and R_c is the rent for non-tradables. W_{fo}^g is the wage that g -workers would earn from job $\{f, o\}$. Utility optimization implies that the demand is

$$C_{fo}^g = (1 - \kappa) \frac{W_{fo}^g}{P_c}, \quad H_{fo}^g = \kappa \frac{W_{fo}^g}{R_c}. \quad (2)$$

²¹In our model, students from the same university differ only in idiosyncratic tastes, ε_{fo}^i . We analyze sorting across universities of different rankings, while remaining silent on heterogeneity within each university.

We can then obtain the indirect utility as

$$V_{fo}^i = \frac{W_{fo}^g}{p_c} \tau_{fo}^g a_c^g \varepsilon_{fo}^i, \quad (3)$$

where $p_c = P_c^{1-\kappa} R_c^\kappa$ is the price index in city c .

3.2 Labor Allocation

For tractability, we assume ε_{fo}^i is drawn from i.i.d. *Fréchet* distribution with shape parameter θ and scale parameter 1. In equilibrium, the fraction of g -graduates who choose firm f and o can be expressed as

$$\Pi_{fo}^g = \frac{\left(W_{fo}^g a_c^g \tau_{fo}^g / p_c \right)^\theta}{\sum_{f'o'} \left(W_{f'o'}^g a_{c'}^g \tau_{f'o'}^g / p_{c'} \right)^\theta}. \quad (4)$$

By applying the Law of Conditional Probability, we can also derive the fraction of g -graduates who choose firm f , conditional on choosing occupation o as

$$\Pi_{f|o}^g = \frac{\left(W_{fo}^g a_c^g \tau_{fo}^g / p_c \right)^\theta}{\sum_{f'} \left(W_{f'o}^g a_{c'}^g \tau_{f'o}^g / p_{c'} \right)^\theta}. \quad (5)$$

3.3 Estimating Equations

We assume that the wage has two components as follows:

$$W_{fo}^g = \omega_{fo} \times T_{fo}^g, \quad (6)$$

where ω_{fo} is wage per efficiency unit paid at firm f and occupation o , and T_{fo}^g is labor productivity (or efficiency units) for g -group if working at f and o . We also consider search cost τ_{fo}^g consist of two parts as follows

$$\tau_{fo}^g = \tau_c^{g,\text{Geo}} \times \tau_{fo}^{g,\text{UBV}}. \quad (7)$$

Here, $\tau_c^{g,\text{Geo}}$ is the component related to moving costs, which we will measure as a function of geographic variables. $\tau_{fo}^{g,\text{UBV}}$ is the unobserved component, highlighted using superscript UBV.

Sorting into Jobs. We can use Equations (4), (6), and (7) to derive a log-linear estima-

tion equation for the determination of the log share of g -group who choose firm f and occupation o ,

$$\ln \Pi_{fo}^g = \lambda_g + \lambda_{fo} + \theta \ln T_{fo}^g + \theta \ln a_c^g + \theta \ln \tau_c^{g, \text{Geo}} + \theta \ln \tau_{fo}^{g, \text{UBV}}. \quad (8)$$

The term $\lambda_g = -\ln \sum_{f'o'} (W_{f'o'}^g a_{f'}^g \tau_{f'o'}^g / p_{c'})^\theta$ is group fixed effects that capture overall employment opportunities for g -workers. This term will also absorb any reputation effects that are specific to university g and common across firms and occupations.

The term $\lambda_{fo} = \theta \ln \omega_{fo} - \theta \ln p_c$ is the firm-occupation-specific fixed effects. It absorbs the wage per efficiency unit. Since index f nests city c , it also absorbs the local price index, p_c . $\theta \ln \tau_{fo}^{g, \text{UBV}}$ will be treated as the structural residual in the estimation.

Sorting across firms conditional on occupation. Equation (5) can be used to derive a log-linear estimating equation for the determination of the log share of g -group who choose firm f , conditional on occupation o ,

$$\ln \Pi_{f|o}^g = \lambda_{go} + \lambda_{fo} + \theta \ln T_{fo}^g + \theta \ln a_c^g + \theta \ln \tau_c^{g, \text{Geo}} + \theta \ln \tau_{fo}^{g, \text{UBV}}. \quad (9)$$

The only difference from Equation (8) is $\lambda_{go} = -\theta \ln \sum_{f'} (W_{f'o'}^g a_{f'}^g \tau_{f'o'}^g / p_{c'})$, which is a group-occupation fixed effect absorbing the average employment opportunities conditional on a specific group and an occupation. Note that this term will also absorb any reputation effects that are specific to university g and occupation o , but common across firms.

3.4 Parameterization

To estimate the model, we impose parametric assumptions on $\ln T_{fo}^g$, $\ln a_c^g$, and $\ln \tau_c^g$.

Labor Productivity (Task Complementarity). Note that T_{fo}^g is specific to each university, firm, and occupation, and is high-dimensional. Our unique dataset enables us to reduce this dimensionality. Specifically, we assume university-job-specific labor productivity as follows:

$$\ln T_{fo}^g = \sum_{j=1}^J \sum_{k=1}^K \beta_{jk}^{\text{BGT}} X_j^g Y_{foK}^{\text{BGT}}, \quad (10)$$

where index j refers to the four university tier groups. $X^g = \{X_1^g, \dots, X_4^g\}$ is therefore a set of binary variables representing these groups. $Y_{fo}^{\text{BGT}} = \{Y_{fo1}^{\text{BGT}}, \dots, Y_{foK}^{\text{BGT}}\}$ represents the BGT task requirements that are specific to firm f and occupation o . The index k refers to the type of tasks, which include cognitive, social, routine, and manual tasks.

β_{jk}^{BGT} governs the sign and strength of sorting. Taking cognitive tasks as an example and suppose g_1 is the group of better-ranked universities, and g_4 is the group of lowest-ranked universities. $\beta_{1k}^{\text{BGT}} > \beta_{2k}^{\text{BGT}} > \beta_{3k}^{\text{BGT}} > \beta_{4k}^{\text{BGT}} > 0$ implies that graduates from better-ranked universities are more likely to work in cognitive-intensive jobs

$$\frac{\partial \ln \Pi_{fo}^{g_1}}{\partial Y_{fo,\text{cog}}^{\text{BGT}}} > \frac{\partial \ln \Pi_{fo}^{g_2}}{\partial Y_{fo,\text{cog}}^{\text{BGT}}} > \frac{\partial \ln \Pi_{fo}^{g_3}}{\partial Y_{fo,\text{cog}}^{\text{BGT}}} > \frac{\partial \ln \Pi_{fo}^{g_4}}{\partial Y_{fo,\text{cog}}^{\text{BGT}}}. \quad (11)$$

The above inequalities extend the positive sorting defined in Costinot and Vogel (2010) to the probabilistic version that satisfies the Monotone Likelihood Ratio Property as studied in Costinot and Vogel (2015).

In Equation (10), $\beta_{j,\text{cog}}^{\text{BGT}} \times Y_{fo,\text{cog}}^{\text{BGT}}$ can be considered as the marginal productivity for tier- j graduates of performing cognitive task in firm f and occupation o . When $\beta_{1k}^{\text{BGT}} > \beta_{2k}^{\text{BGT}} > \beta_{3k}^{\text{BGT}} > \beta_{4k}^{\text{BGT}} > 0$, it then implies that labor productivity increases as cognitive task requirements are higher, and increases more for graduates from better-ranked universities. The specification of Equation (10) is in line with micro-foundation task complementarity (or supermodular property), giving better-ranked universities the “productivity premium” to work in better firms, and is the source generating positive assortive matching (Becker, 1973, Eeckhout, 2018).²³

Amenity. We measure the group-specific utility derived from local amenities as

$$\ln a_c^g = \sum_{j=1}^J \beta_j^{\text{Amen}} X_j^g Y_c^{\text{Amen}}, \quad (12)$$

where Y_c^{Amen} is the measure of local amenity. β_j^{Amen} captures the strength of sorting for tier- j universities on amenities, with a larger positive value corresponding to stronger sorting. Similarly, $\beta_1^{\text{Amen}} > \beta_2^{\text{Amen}} > \beta_3^{\text{Amen}} > \beta_4^{\text{Amen}}$ implies that graduates from better-ranked universities sort in cities with better amenities.

Mobility costs. We measure group-city-specific mobility costs using the interaction between X^g and geographic variables, Y_{fk}^{Geo} . Specifically, we assume

$$\ln \tau_c^{g,\text{Geo}} = \sum_{j=1}^J \sum_{k=1}^K \beta_{jk}^{\text{Geo}} X_j^g Y_{fk}^{\text{Geo}}, \quad (13)$$

²²The results hold because $\frac{\partial \ln \Pi_{fo}^{g_j}}{\partial Y_{fo,\text{cog}}^{\text{BGT}}} = \theta \beta_{j,\text{cog}}^{\text{BGT}}$ where θ is constant across all groups.

²³Since we include interactive terms for four types of tasks, our estimation will reveal task complementarity for each of the four dimensions.

where Y_{fk}^{Geo} include three variables, the logarithm of the geographic distance between university and firm, a dummy variable CZ_{fg} that equals one if university g and the firm f are in the same city; and a dummy $State_{fg}$ that equals one if the university and firm are in the same state. These dummies can pick up any non-linear features between costs and log distance. These interactions pick up potential differential impacts of geography on job choice across groups.

Similarly, β_{jk}^{Geo} captures the strength of sorting for tier- j universities on geographic factors. For example, parameter values of $\beta_{1,\text{dis}}^{\text{Geo}} > \beta_{j,\text{dis}}^{\text{Geo}}, j = 2, 3, 4$ would capture the fact that graduates from Top 20 universities are more likely to match to jobs that are further away than other groups.

4 The Estimation Results

We group universities into four tiers as in Section 2, covering 18 U.S. universities are ranked among the top 20 globally or in the Ivy League (Top 20), 44 universities are ranked between 21 and 200 globally (Top 21-200), which are typically the best universities within each state, 115 U.S. universities are ranked between Top 201 and 1000 (Top 201-1000), and 89 universities from Top 1000 and 2000 (outside Top 1000).²⁴ Our estimation includes three dummy variables for Top 20, Top 21-200, and Top 201-1000, with universities outside Top 1000 considered as the benchmark group.

In estimating Equations (8) and (9), our sample is formed by a combination of 266 universities, 25492 firms, and 22 occupations that have non-zero matches.

4.1 Task Complementarity

We begin by estimating Equation (8) as our baseline model. We first estimate an OLS regression for the interactive coefficients between our four BGT task variables (cognitive, social, routine, and manual) with rank dummies (Top 20, Top 21-200, and Top 201-1000). To ease interpretation, we further standardize each BGT task variable by its standard deviation.²⁵

Following Equation (8), we control for the group fixed effects, λ_g , and the firm-occupation fixed effects, λ_{fo} . Since the inclusion of firm-occupation fixed effects absorbs the main effects of BGT task variables, only the interaction effects are included in the

²⁴Top 20 universities include Harvard, MIT, Stanford, Columbia, Princeton, UC Berkeley, Penn, Chicago, Yale, Cornell, Northwestern, UCLA, Michigan, Johns Hopkins, UIUC, Duke, Dartmouth, and Brown.

²⁵The variables in Table 2 are expressed as percentiles, ranging from 0 to 1. In this section, the BGT task variable is further standardized to have a standard deviation of one.

regression. Column (1) in Table 3 reports the coefficients from the OLS regression of $\ln \Pi_{fo}^g$ on the interactive coefficients. We find positive and precisely estimated interactive coefficients for cognitive tasks, suggesting positive sorting of university ranking into cognitive tasks. Specifically, the cognitive-Top 20 coefficient equals 0.199 (s.e. = 0.043), and the cognitive-Top 21-200 coefficient remains positive, although it decreases to 0.073 (s.e. = 0.020), and further decreases to 0.029 (s.e. = 0.015) for cognitive-Top 1000.

The estimates indicate that, holding other factors constant, a standard deviation increase in cognitive task requirements raises the likelihood of working as a computer scientist at Google rather than at Sanmina Corporation (an American electronics manufacturing services provider), graduates from Top 20 universities have a 0.199 log points or about 19.9% more likely to match to the cognitive-intensive job (25.6% of a standard deviation of $\ln \Pi_{fo}^g$), relative to the benchmark group.²⁶ The result shows a 0.073 higher in log points for graduates from Top 21-200, and a 0.029 higher than the benchmark group.

Additionally, we observe small, negative, but mostly imprecisely estimated coefficients for the interactions with routine and manual tasks, indicating little systematic sorting into these tasks across university rankings among college graduates. The interaction terms between social tasks and Top 20 is 0.065 and statistically significant at 90% confidence.

4.2 Local Amenities

We augment our model to estimate the effects of the city’s amenities in determining job matching. To ease interpretation, we further standardize the variable to have a unit standard deviation. The amenity index is city-specific and does not vary across firms that are situated in the same city. We include interaction terms between the amenity and rank dummies. Again, the main effect of amenities is again absorbed by the firm–occupation fixed effect.

Column (2) in Table 3 reports a sizable interactive coefficient equal to 0.091 (s.e. = 0.036) for amenity-Top 20. For example, consider Detroit, MI, which sits at the 25th percentile in terms of amenities among all U.S. cities, and Seattle, which ranks at the 99th percentile in amenities. Holding job characteristics constant, graduates from Top 20 universities are 6.8% more likely to choose a job in Seattle than in Detroit,

²⁶The outcome variable, $\ln \Pi_{fo}^g$, has a standard deviation of 0.778. The 25.6% is computed as the ratio between 0.199 and 0.778.

relative to graduates from universities outside Top 1000.²⁷ These results also indicate positive sorting on amenities among graduates from Top 20 universities. Relative to the benchmark group, we find no pattern in differential sorting on amenity for universities in Top 21-200, and a small negative coefficient for universities in Top 201-1000.

Our estimated results for positive sorting into cognitive tasks remain robust with the inclusion of these amenity interactions, although the estimates are slightly smaller in magnitude.

4.3 Geographic Factors

Next, we estimate the extent to which geographic proximity determines job matching and sorting. We augment our model by incorporating the interaction of ranking dummies with three geographic variables: the logarithm of the geographic distance between a university and the firm location, a dummy variable CZ_{fg} that equals one if university g and firm f are in the same city; a dummy $State_{fg}$ that equals one if university g and firm f are in the same state.

Column (3) of Table 3 reports the coefficient estimates of OLS regression with geography variables included. The coefficients for log distance and its interaction terms are identified from out-of-state movers. We see a negative and statistically significant coefficient for -0.067 (s.e. = 0.011) for log distance, indicating that a log point increase in geographic distance reduces job matching probability by a 0.067 log point for movers graduate from outside the Top 1000 universities. The interaction terms between log distance and tier dummies pick up any differential effects of how geographic distance affects job matching relative to the baseline group. Importantly, we find a sizable and positive coefficient for the interaction of log distance and Top 20, equaling to 0.061 (s.e. = 0.024). The interaction coefficients with Top21-200 and Top 1000 are small and statistically insignificant.

These interaction coefficients show that, among movers, distance has no effect on job matching for Top 20 graduates but has notable effects for graduates from all other tiers. To interpret this using a concrete example, it means that graduates from Harvard and Stanford (both Top 20) are equally likely to secure a job in New York City, despite Boston's geographic proximity to New York. In contrast, distance matters significantly for graduates outside the Top 20 universities. For example, those who studied in Boston (e.g., Boston College) are 18.2% more likely to obtain a job in New York City than those who studied in the Bay Area (e.g., UC Davis).²⁸

²⁷We estimate the value as $0.091 \times (0.99 - 0.25) = 6.8\%$, using our preferred estimates in Column (4)

²⁸We estimate the value as $-0.067 \times (\ln(2557) - \ln(171)) = -0.182$, where the distances between the

Unsurprisingly, we find positive, sizable, and statistically significant coefficients for CZ_{fg} and $State_{fg}$: these coefficients pick up the fact that college graduates are more than proportional to remain in university city or state.²⁹ The estimates indicate that for universities outside top 1000, their graduates have a 44.8% higher share or are more likely to find a job in the same city of their alma mater, relative to other cities, holding other factors the same; and have a 8.6% higher share or are more likely to find a job within the state relative to other states.

Column (3) also reports positive and precisely estimated interaction effects for $CZ \times \text{Top 21–200}$, $\text{State} \times \text{Top 20}$, and $\text{State} \times \text{Top 21–200}$. These positive interactive coefficients may capture two mechanisms. The first is related to unobserved economic factors: how graduates from each tier of universities capitalize on job opportunities within the CZ or state of their university’s location, relative to the benchmark group. The second is preferences: students from institutions of different rankings may exhibit varying tendencies to remain near their university’s city or state.

4.4 Composition of the Student Body

Likewise, the home bias studied in migration literature, students may also prefer to remain close to their hometown after graduation. This suggests that universities with a higher fraction of in-state students should exhibit higher in-state retention rates. To potentially isolate the preference factor arising from the difference in the composition of their student bodies, we augment six additional interaction terms between three geographic variables (log distance, CZ dummy, and state dummy) and two institutional characteristics: a dummy for public universities and the share of in-state student enrollment. These shares of in-state enrollment vary widely across institutions, ranging from less than 10% at elite private institutions such as Columbia, MIT, and Yale, to more than 90% at public universities such as Texas A&M University and the University of California, Riverside.

Column (4) of Table 3 reports a precisely estimated coefficient of 0.234 (s.e. = 0.234) for the CZ–public interaction and 0.292 (s.e. = 0.117) for the interaction of state dummy and in-state enrollment share. These estimates validate the home-bias hypothesis, indicating a stronger local attachment for universities with a high share of in-state students.

two city pairs are 2,557 and 171 miles.

²⁹Note that although only 47.7% of college graduates stay in their university city after graduation, the share is much higher than the city’s employment share in the national total.

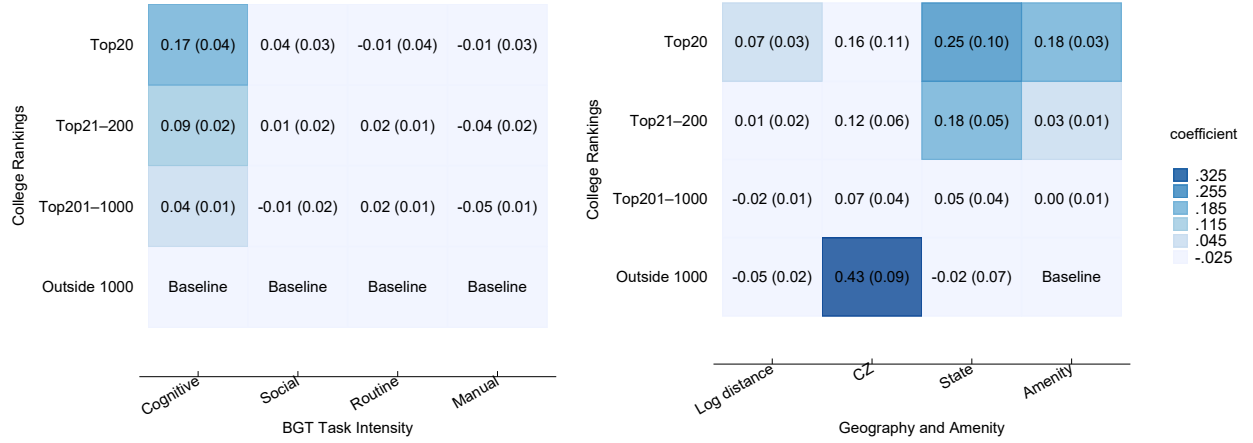


Figure 3: Summary for All Estimated Interactive Coefficients of Equation (8)

Notes: The coefficients are based on column (4) of Table 3. Darker blue colors indicate positive and statistically significant estimates; small or statistically insignificant estimates are plotted in light blue. Reported values are coefficients with standard errors in parentheses.

Conditional on these interactions, we find statistically insignificant estimates for rank-CZ interactive coefficients. The coefficients for the state-Top 20 and state-Top 21-200 interaction terms remain positive and statistically significant. Importantly, these interactive coefficients do not contradict the pattern of within-state retention rates documented in Table 1. The variation we explore here already nets out sorting by observables (tasks and amenities, the origin composition of universities, and firm-occupation-specific wages). Thus, the positive state-Top 20 and state-Top 21-200 interaction coefficients likely capture unobserved economic benefits and amenities, which on average are greater in states that host higher-ranked universities (e.g., California, New York, Massachusetts), and to which top-university graduates are closely proximate.

Students enrolled in public and private universities differ systematically in family income background (Chetty et al., 2020), and such differences may shape labor market outcomes. To account for this possibility, we augment the model with four interaction terms between the BGT task measures (cognitive, social, routine, and manual) and a public-university dummy. Because our regression already controls for university fixed effects and ranking-related interactions, these public-task interaction terms capture differential sorting between public and private universities within the same ranking tier—likely reflecting differences in family background. If students from higher-income families are more likely to obtain better jobs, we would expect negative coefficients on

the cognitive–public or social–public interactions. Column (5), however, shows that none of these additional controls matter. Importantly, the coefficient estimates remain nearly identical with or without these controls (columns 4 vs. 5). In a recent study, [Chetty et al. \(2023\)](#) find that the factors giving children from high-income families an admissions advantage are uncorrelated with post-college outcomes. Our findings are in line with theirs.

To summarize the coefficient estimates with the full set of controls, Figure 3 displays the university ranking-related interaction coefficients based on the preferred specification reported in Column (4). Darker blue colors indicate positive and statistically significant estimates; small or statistically insignificant estimates are plotted in light blue. Reported values are coefficients with standard errors in parentheses.

The previously highlighted interaction coefficients change only modestly but remain statistically significant. For example, we estimate a cognitive–Top 20 coefficient of 0.17 (s.e. = 0.04), an amenity–Top 20 coefficient of 0.18 (s.e. = 0.03), a distance–Top 20 coefficient of 0.07 (s.e. = 0.03), and a distance coefficient of -0.05 (s.e. = 0.02).

4.5 Positive Sorting into Firms or Occupations?

Unlike the DOT or O*NET data, the BGT task requirements vary across both firms and occupations, the positive sorting in cognitive tasks that we estimate may reflect two distinct channels: (1) firm sorting—graduates from higher-ranked universities tend to work in firms with greater cognitive task requirements; and (2) occupational sorting—graduates from higher-ranked universities tend to enter more cognitively intensive occupations. This section examines the relative importance of these two channels in driving the observed positive sorting.

Previous literature has analyzed sorting into occupations using O*NET data ([Deming, 2017](#), [Burstein et al., 2019](#)) and sorting into firms using the AKM (1999) approach and employer–employee matched data ([Card et al., 2013](#), [Lopes de Melo, 2018](#), [Song et al., 2019](#), [Bonhomme et al., 2023](#)). Our approach estimates sorting into both firms and occupations within a unified framework, allowing us to shed light on the relative strength of firm versus occupation dimensions.

We then estimate Equation (9), in which the outcome variable is worker sorting into firms conditional on occupation, which requires controlling for university-occupation and firm-occupation fixed effects. Since the university-occupation fixed effects absorb the aggregate occupational employment opportunities specific to a university, the coefficients for the rank-task interaction terms now characterize the strength of sorting

across firms, holding the same occupation.

For comparison, Figure 4 displays the university ranking–related interaction coefficients. The complete sets of coefficients are reported in Table 4 with different specifications. Because university-occupation pairs that are only observed in one firm would not contribute to identification, the sample size differs slightly from the baseline. The estimates are notably smaller. As cognitive task requirements increase by one standard deviation, graduates from Top 20 universities are 9% more likely to work in a cognitively intensive firm relative to the benchmark group. Compared to Figure 3, the estimates suggest that sorting into firms accounts for about $0.09/0.17 = 52\%$, or more than half, of the overall sorting into jobs (firms and occupations combined). The interaction coefficients for Cognitive $\times$ Top 21–200 and Cognitive $\times$ Top 201–1000 also decline, falling to roughly half of the corresponding estimates in Table 3. Positive sorting into both firms and occupations is quantitatively important in shaping the overall positive sorting on cognitive tasks.³⁰

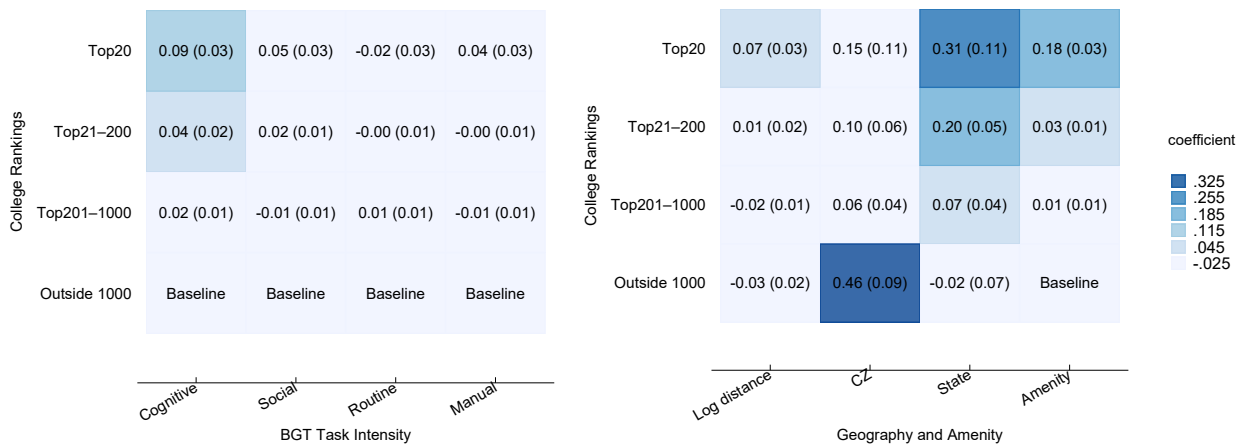


Figure 4: Summary for All Estimated Interactive Coefficients of Equation (9)

Notes: The coefficients are based on column (4) of Table 4. Darker blue colors indicate positive and statistically significant estimates; small or statistically insignificant estimates are plotted in light blue. Reported values are coefficients with standard errors in parentheses.

5 Sorting on Other Mechanisms

This section tests several potential confounding mechanisms that might influence our findings.

³⁰In addition, the interaction coefficients related to geography and amenities are similar between the estimates of the two figures. The results are expected, since spatial variation is captured in our definition of a firm rather than in occupation.

5.1 Admission Selectivity

The sorting we have found based on university ranking can be a byproduct of two primary factors: the effect of attending a specific university and the selection—students admitted into better universities on average might have higher abilities who might perform better regardless of the universities they attend. This section provides evidence that sheds light on the relative importance of the two factors.

To this end, we use Barron’s Profiles of American universities (2017 edition), which classifies universities into six tiers based on SAT/ACT scores and other admission requirements such as high school transcripts and class rank ([Barron’s Educational Series, 2017](#)) (see Table C.1). Unlike the WUR or U.S. News rankings that we use elsewhere, these six categories assess the selectivity of college admissions and serve as proxies for students’ ability prior to entering college, rather than providing ratings based on academic standards, educational quality, or institutional reputation.

Importantly, the selectivity criterion does not align with the WUR rankings. For example, UC Berkeley (ranked 8th in WUR), Tulane University (407th in WUR), and Davidson College (1866th in WUR) are all classified among the most selective universities. To increase precision in the estimation, we aggregate the admission criteria into four groups: most selective, highly and very selective, selective, and less and non-selective.

Column (2) of Table 5 augments the baseline model with interaction between the selectivity dummies and Y_{fok} , where Y_{fok} includes the three sets of variables (BGT tasks, amenities, and geographical variables). In comparison, Column (1) reports the baseline estimates. As shown in Column (2), there is no evidence that suggests labor market sorting based on the college selectivity; university ranking and related interaction terms are also included. Importantly, we see only modest changes in the ranking-related interactive coefficients, compared to the baseline.

We find little evidence of positive sorting into cognitive jobs based on differential selection of college admissions. Rather, they reflect systematic differences associated with universities’ training, reputation, and educational quality that shape the career outcomes of their graduates.

5.2 STEM vs. Non-STEM Majors

Graduates who majored in fields such as STEM and quantitative disciplines are known to be more likely to obtain high-paid jobs ([Deming and Noray, 2020](#)). To examine

whether our result is driven by variation in major composition across institutions, we augment the model with interaction terms between the share of graduates in STEM majors and Y_{fok} (BGT tasks, amenities, and geographical variables). The share of STEM majors for each university is calculated from our LinkedIn data, where STEM majors include one of the following broad fields: Science, Technology, Engineering, or Mathematics.

In Column (3) of Table 5, we find a sizable and positive coefficient for the cognitive–STEM interaction, equal to 0.224 (s.e. = 0.067), and a negative coefficient for the social–STEM interaction, -0.14 (s.e. = 0.066). These results suggest that universities with a stronger STEM orientation tend to place their graduates in more cognitively intensive but less socially intensive jobs. Interestingly, we also find positive and precisely estimated coefficients for the distance-STEM and amenity-STEM interactions, suggesting that STEM graduates are more geographically mobile than non-STEM graduates and are more likely to relocate to cities with better amenities.

Importantly, while the estimates underscore the role of major composition in shaping job outcomes, the positive sorting on cognitive-ranking related interactive coefficients are similar (some fall modestly) and remain statistically significant when major composition is considered.

5.3 University-to-Firm Network

The literature has documented that migration networks not only improve the likelihood of job matching but also enhance job outcomes through referrals (Munshi, 2003). Our dataset allows us to measure the university-to-firm network. This section provides evidence, from two aspects, that our estimated results hold even when university-to-firm network effects are considered.

A Theoretical Re-Interpretation. First, we show that by incorporating network effects into the model under a certain structure, one can derive an estimating equation isomorphic to our baseline Equation (8). Specifically, suppose that unobserved job search costs are related to the university-to-job (firm and occupation) matching pattern in the previous period. In particular, we assume that $\ln \Pi_{fo,t}^g$ follows an $AR(1)$ process driven by the unobserved search cost component, that is,

$$\ln \tau_{fo,t}^{g,UBV} = \rho \ln \Pi_{fo,t-1}^g + \varepsilon_{fo,t}, \quad \rho \in (0, 1), \quad (14)$$

where $\Pi_{fo,t-1}^g$ is the share of g-graduates who work in firm f and o in the past or that

of earlier graduates (alumni). The parameter ρ captures the extent to which the job matching of alumni influences current job search costs. The term $\varepsilon_{fo,t}$ represents white noise—factors unrelated to $\ln \Pi_{fo,t-1}^g$ but affecting job search costs $\ln \tau_{fo,t}^{g,UBV}$.

In the data, there is a strong persistence pattern of university-to-firm job matches (see Appendix Figure C.2). The pattern is similar to the well-documented literature on the immigrant enclave, where studies have documented that for US immigrants from a particular origin country, their location settlement (Card, 2001) and occupational choices (Hanson and Liu, 2016) tend to persist over time.

Assuming T_{fo}^g , a_c^g , and $\ln \tau_c^g$ are all time invariant, we can substitute Equation (14) into (8), recursively, to obtain

$$\ln \Pi_{fo,t}^g = \lambda_g + \lambda_{fo} + \tilde{\theta} \ln T_{fo}^g + \tilde{\theta} \ln a_c^g + \tilde{\theta} \ln \tau_c^{g,Geo} + \theta \sum_{m=0}^t \rho^m \varepsilon_{fo,t-m}, \quad (15)$$

where $\tilde{\theta} = \frac{\theta}{1-\rho}$ is a function of the job-matching elasticity to productivity, θ , and the path dependence of university-specific job matching, $1 - \rho$. Equation (15) is therefore isomorphic to our baseline equation, differing only in the interpretation of the reduced-form parameter and the structural residuals.

Equation (15) further implies that, as long as idiosyncratic shocks to job matching at any time t , $\varepsilon_{fo,t-m}$, are independent of the observed controls, the OLS estimate of the reduced-form parameter is unlikely to suffer from the omitted variable bias. The main difference lies in the interpretation of the coefficients.

Empirical Results. The second piece of evidence we provide is empirical. Using the LinkedIn dataset, we can control firm exposure to alumni from the university g as

$$\text{Network}_f^g = \frac{L_f^{g,\text{alumni}}}{L_f}, \quad (16)$$

capturing the share of firm f 's current employees that graduated from g before 2014 (we refer to as alumni), $L_f^{g,\text{alumni}}$, relative to L_f , the current employment size of firm f . Both L_f and $L_f^{g,\text{alumni}}$ are estimated from the LinkedIn data.

Since unobserved factors that lead to a job match of alumni might also promote current graduates, we do not claim any causal estimates for the impact of alumni networks on job matching. The empirical exercise below aims to shed light on which of the sorting channels we have estimated so far—task, amenity, or geography—the network primarily operates through.

Column (4) of Table 5 augments the model with the network measure and shows that the university-to-firm network is a strong predictor of employment outcomes, with a coefficient of 1.675 (s.e. = 0.048), indicating that a one-standard deviation increase in network exposure is associated with a 167.5% higher likelihood of graduates working at that firm.

Despite the sizable network coefficient, the cognitive- and amenity-related interaction coefficients remain statistically significant, though somewhat reduced in magnitude. For example, the coefficient on Cognitive $\times$ Top 20 declines from 0.147 to 0.119, while the coefficient on Cognitive $\times$ Top 21–200 decreases from 0.071 to 0.047.

With this network control, the largest changes occur in the interaction terms involving geographic variables. This pattern suggests that the university-to-firm network operates primarily through geographic channels.

5.4 Other Robustness Checks

Using Disaggregated Groupings. Our baseline estimation divides the universities into four groups. We show that the results remain similar when using more disaggregated groupings. Specifically, we split the Top 21–200 category into two subgroups: universities ranked 21–100 (labeled Top 21–100) and those ranked 101–200 (labeled Top 101–200). Likewise, we split the Top 201–1000 category into two subgroups: universities ranked 201–500 (labeled Top 201–500) and those ranked 501–1000 (labeled Top 501–1000). Table 6 reports the estimated coefficients, showing that most interaction terms are similar across the subgroups within each broader category. These findings suggest that using broader groupings, as in the baseline estimation, does not obscure substantial heterogeneity in the sorting behavior.

U.S. News Rankings. While our baseline estimates rely on the WUR rankings, we find similar results when using the U.S. News rankings. Unlike WUR, which covers the Top 2000 universities globally, U.S. News provides rankings only for the Top 1000. For U.S. universities that appear in both sources, the correlation between the two measures is 0.862. In the estimation, we classify U.S. universities into three groups: Top 20, Top 21–200, and Top 201–1000 (the benchmark group). As reported in Table 7, estimates based on U.S. News rankings are similar in magnitude and lead to the same conclusions.

6 Geography and Labor Market Outcomes

Our model estimates the assignment of talent to jobs as a function of university ranking, job characteristics, amenities, and geography. In this section, we use the estimated model to quantify how geography shapes the labor market outcomes of fresh graduates.

6.1 Counterfactual Job Matching Probability

We first use the estimated model to compute the model-based benchmark probabilities as

$$\Pi_{fo}^{g, \text{Benchmark}} = \frac{\exp(\lambda_{fo}) (T_{fo}^g a_c^g \tau_c^{g, \text{Geo}} \tau_{fo}^{g, \text{UBV}})^\theta}{\sum_{f'o'} \exp(\lambda_{f'o'}) (T_{f'o'}^g a_{c'}^g \tau_{c'}^{g, \text{Geo}} \tau_{f'o'}^{g, \text{UBV}})^\theta}, \quad (17)$$

where each component is estimated based on our preferred specification, reported in Column (4) of Table 3. Specifically, λ_{fo} is the estimated firm-occupation fixed effects, and $\exp(\lambda_{fo})$ measures $(\omega_{fo}/p_c)^\theta$. T_{fo}^g , a_c^g , and $\tau_c^{g, \text{Geo}}$ are estimated as

$$\begin{aligned} T_{fo}^g &= \exp \left(\sum_{j=1}^J \sum_{k=1}^K \hat{\beta}_{jk}^{\text{BGT}} X_j^g Y_{fok}^{\text{BGT}} \right), \quad a_c^g = \exp \left(\sum_{j=1}^J \hat{\beta}_j^{\text{Amen}} X_j^g Y_c^{\text{Amen}} \right), \\ \tau_c^{g, \text{Geo}} &= \exp \left(\sum_{j=1}^J \sum_{k=1}^K \hat{\beta}_{jk}^{\text{Geo}} X_j^g Y_{fk}^{\text{Geo}} + \sum_{j=1}^J \sum_{k=1}^K \hat{\gamma}_{jk}^{\text{Geo}} C_j^g Y_{fk}^{\text{Geo}} \right). \end{aligned} \quad (18)$$

$\hat{\beta}_{jk}^{\text{BGT}}$, $\hat{\beta}_j^{\text{Amen}}$, $\hat{\beta}_{jk}^{\text{Geo}}$, and $\hat{\gamma}_{jk}^{\text{Geo}}$ denote the estimated coefficients. C_j^g includes two variables of university characteristics: whether g is a public university, and the fraction of in-state student enrollment.

In computing T_{fo}^g , a_c^g , and $\tau_c^{g, \text{Geo}}$, we rely on coefficients that are statistically significant at the 95% confidence level, setting statistically insignificant parameters to zero. $\tau_{fo}^{g, \text{UBV}}$ is obtained as the exponent of the regression residuals. Since the estimated model has an R^2 of 0.75—implying that unobserved components account for 25% of the variation in outcomes—we thus decide to incorporate the regression residual (through $\tau_{fo}^{g, \text{UBV}}$) when computing the benchmark probability.

Note that, because we compute $\Pi_{fo}^{g, \text{Benchmark}}$ using statistically significant coefficients, the predicted probabilities do not match exactly the observed allocations but almost perfectly. Appendix Figure C.3 shows that the predicted allocations fit the observed ones very well: a simple OLS regression of observed shares on predicted shares yields a coefficient of 0.99 and an R^2 of 0.998.

We then compute the counterfactual probabilities when geographic forces are absent. The difference from Equation (17) is that we set $\hat{\beta}_j^{\text{Geo}} = \hat{\gamma}_j^{\text{Geo}} = 0$, which implies

$\tau_c^{g,\text{Geo}} = 1$. We estimate the counterfactual probability as

$$\Pi_{fo}^{g,\text{Geo}} = \frac{\exp(\lambda_{fo}) (T_{fo}^g a_c^g \tau_{fo}^{g,\text{UBV}})^\theta}{\sum_{f'o'} \exp(\lambda_{f'o'}) (T_{f'o'}^g a_{c'}^g \tau_{f'o'}^{g,\text{UBV}})^\theta}, \quad (19)$$

which measures the talent to job allocation when the world is flat, in which the matching is determined by other factors such as task complementarity (T_{fo}^g), sorting in amenity a_c^g , the unobserved factor $\tau_{fo}^{g,\text{UBV}}$, and firm-occupation-specific factor λ_{fo} .

6.2 Geography and the Matching to High-Paid Jobs

Using the model-predicted and counterfactual probabilities, we first estimate the effects of geography on the matching to high-paid jobs. We define high-paid jobs as those in the top 5% of the entry-level wage distribution in our sample. For university g , the effect of geography on the probability of matching to top 5% paid jobs is

$$\text{ProbTop}_g^{\text{Geo}} = \sum_{fo} \left(\Pi_{fo}^{g,\text{Benchmark}} - \Pi_{fo}^{g,\text{Geo}} \right) \times \mathbb{1}_{fo}^{\text{Top5\%}}. \quad (20)$$

Here, $\mathbb{1}_{fo}^{\text{Top5\%}}$ is an indicator that equals one if the salary of a job (firm and occupation) is in the top 5%. Figure 5 displays the probability of matching to top 5% job using blue circles and the effect of geography on this probability, $\text{ProbTop}_g^{\text{Geo}}$, using triangles. We display 30 universities: the fifteen with the highest geographic effect and the fifteen with the lowest.³¹

The blue circles in Figure 5 illustrate systematic variation across universities in this probability of interest. These differences likely reflect a combination of factors, including university reputation, educational quality, student ability, and geography. Among the listed universities, the highest probabilities are observed at Harvard (34.6%), Carnegie Mellon (34.5%), Georgetown (32.8%), UC Berkeley (29.8%), and Columbia (28.9%). In contrast, the lowest values are found in institutions such as San Diego State (8.8%) and Fordham (9.0%).

The red triangles show that geography substantially increases the probability of matching to high-paid jobs for universities located in the Bay Area and New York City. This result is unsurprising, given that the Bay Area and NYC together host nearly half of the top 5% highest-paid entry-level jobs according to Glassdoor. Notably, geographic proximity increases the probability of matching by 2.3 percentage points (ppts) for UC

³¹We select these 30 universities among those that have at least an 8% probability of matching to top 5% jobs.

Santa Cruz, 2.1 ppts for UC Berkeley, 2.0 ppts for Santa Clara University, 1.8 ppts for San Jose State, and 1.4 ppts for Columbia.

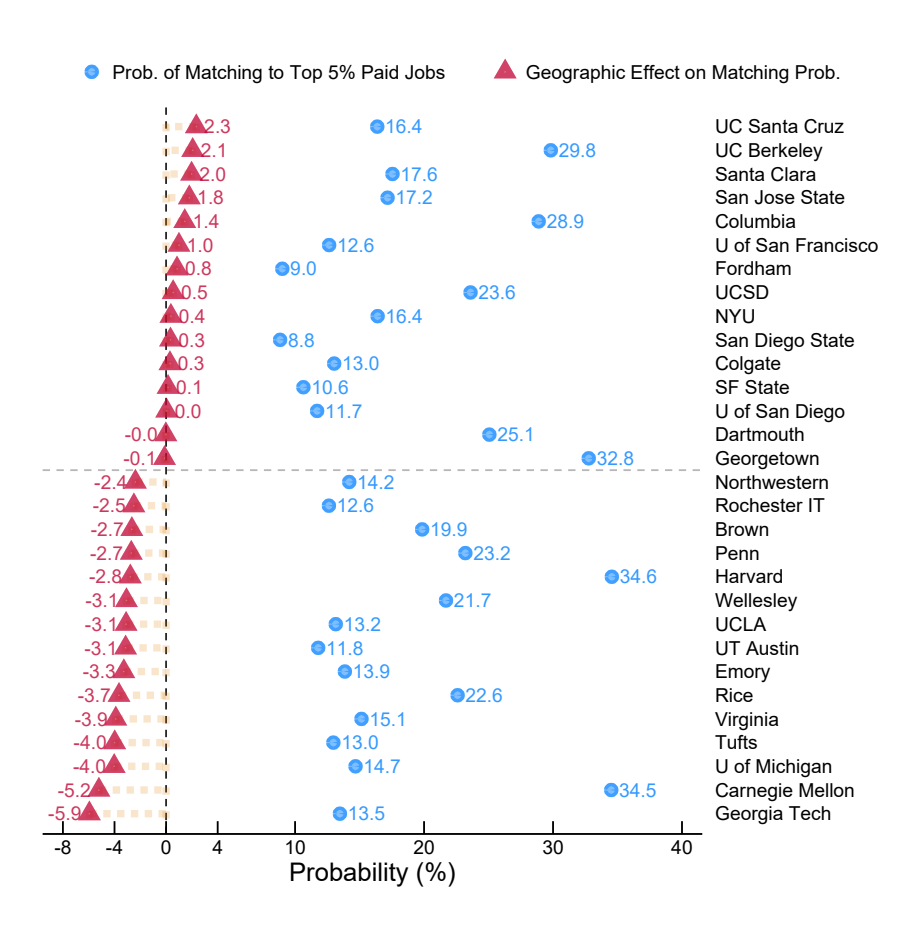


Figure 5: Geography and the Probability of Matching to Top 5% of Paid Jobs

Notes: The x-axis represents the probability in percentage points. The blue circles plot the overall probability of such matches. The red triangle plots the value of $\text{ProbTop}_g^{\text{Geo}}$, capturing the effect of geography on this probability. We display 30 universities: the fifteen with the highest probabilities and the fifteen with the lowest (most negative).

In contrast, geographic disadvantage reduces the probability most strongly for Georgia Tech (5.9 ppts), followed by Carnegie Mellon (5.2 ppts), and by 4.0 ppts for University of Michigan and Tufts. Graduates from other elite universities—including Harvard, Northwestern, Brown, and Penn—also appear to face geographic disadvantages in matching to high-paid jobs. Between the two extremes, geographic differences account for up to an 8.2 ppt gap in the probability of obtaining a top-paid job.³² In a recent study, Chetty, Deming and Friedman (2023) show that attending an Ivy-Plus college triples the likelihood of working in a prestigious firm, relative to attending an

³²This value is based on the difference between UC Santa Cruz (+2.3 ppts) and Georgia Tech (−5.9 ppts).

average flagship public university. Complementing this literature, we find a sizable geographic premium associated with university location.

6.3 The Geographic Premium in Salary

Next, we quantify the extent to which geographic factors affect the entry salaries of fresh graduates. To this end, we compute

$$\text{Wage}_g^{\text{Geo}} = \sum_{fo} (\Pi_{fo}^{g, \text{Benchmark}} - \Pi_{fo}^{g, \text{Geo}}) \times \text{wage}_{fo}. \quad (21)$$

wage_{fo} denotes the Glassdoor wage posted for firm f and occupation o . $\Pi_{fo}^{g, \text{Benchmark}}$ and $\Pi_{fo}^{g, \text{Geo}}$ are the matching probabilities given in Equation (17) and (19). By taking wage_{fo} unchanged, our approach is partial equilibrium, and the estimated wage effects from geography are driven entirely by changes in the matching probability, keeping other sorting channels—job content, amenities, and unobserved factors—unchanged.

Figure 6 plots the average annual salary (blue circles) and the wage premium associated with geography (red triangles); the latter is referred to as the geographic premium. We report results for 30 universities: the fifteen with the highest premiums and the fifteen with the lowest (most negative) premiums.

Universities differ systematically in the average entry salary earned by their fresh graduates. Our emphasis is that there is systematic variation in geographic premium across universities. The largest ones are observed for Columbia University and UC Berkeley, amounting to \$3.1K and \$2.4K in annual salary, respectively. Other universities in the Bay Area or New York City, such as UC Santa Cruz (\$1.9K), Santa Clara (\$1.2K), and NYU (\$1.1K), also rank among the highest.

By contrast, the most notable geographic penalties occur for Carnegie Mellon University (−\$6.7K). This large penalty reflects the university’s specialization in training STEM graduates while being located farther from San Jose and Seattle, where many large U.S. technology firms are concentrated and wages are among the highest. Despite this geographic disadvantage, Carnegie Mellon graduates earn an average annual salary of \$96.3K—higher than most universities displayed in Figure 6. For similar reasons, we also find large location penalties for Georgia Tech (−\$6.7K), Brown (−\$5.0K), Emory University (−\$4.6K), and the University of Pennsylvania (−\$4.2K).

The extent to which geographic factors affect the wages of college graduates depends on several elements: the geographic clustering of U.S. industries and firms, the skills imparted by universities that give graduates a comparative advantage in certain

industries or jobs, and the geographic barriers that limit graduate mobility. To analyze mechanisms that drive variation, Appendix Figure C.4 plots the geographic premium (y-axis) against the mover premium (x-axis) across cities. We measure the mover premium as the difference in average wages between movers and stayers among fresh graduates (2016–2018), which can be estimated directly from our sample. A negative mover premium implies that, on average, local stayers earn higher wages than those who migrate.

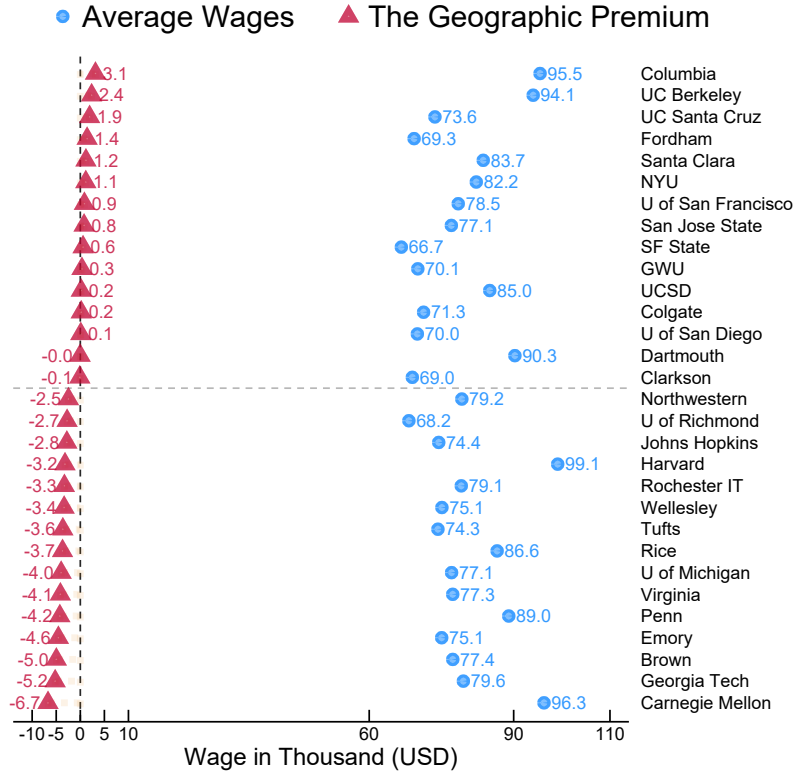


Figure 6: Average Wages and the Geographic Premium

Notes: The x-axis represents the average wage in thousands. The blue circles plot the average wage earned by fresh graduates from each university. The red triangle plots the value of $\text{Wage}_g^{\text{Geo}}$, capturing the effect of geography on average earnings. We display 30 universities: the fifteen with the highest geographic premium and the fifteen with the lowest (most negative).

We find a strong negative correlation: a simple OLS regression yields a coefficient of -0.07 (s.e. = 0.01) and an R^2 of 0.51. It appears that a single variable—the mover premium—explains more than half of the variation in the geographic premium, underscoring the central role of university location and graduates’ geographic mobility in shaping job outcomes.

The geographic premiums we discussed so far are expressed in nominal terms. How-

ever, cities differ substantially in their cost of living (Moretti, 2013, Albert and Monras, 2022), and part of the wages firms offer likely reflects local living costs, especially in expensive metropolitan areas. To account for this, we use the U.S. Bureau of Economic Analysis (BEA) Regional Price Parities (RPPs) for 2018. The RPPs provide relative price levels across MSAs, covering major expenditure categories such as housing, transportation, and food, and are constructed using data from the Bureau of Labor Statistics' Consumer Price Index, housing price data, and other regional sources. In a recent study, Diamond and Moretti (2024) demonstrate that the BEA index is strongly correlated with an alternative price index constructed using detailed consumption data.

We normalize the average value of the RPPs across all cities to one. In 2018, the normalized RPP was as high as 1.26 in San Francisco and 1.20 in New York City, compared to 1.00 in Pittsburgh (PA), and as low as 0.88 in Jackson (TN). We deflate Glassdoor wages using each city's RPP and re-estimate equation (21). Appendix Figure C.5 shows that adjusting for regional price variation reduces disparities in premiums across universities to some extent; however, sizable premiums persist for universities located in large metropolitan areas such as San Jose, San Francisco, and New York.

6.4 The Bay-Area Premium

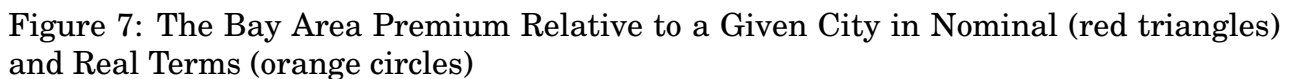
The evidence has shown that universities in the Bay Area occupy many slots among the top of the list, which include UC Berkeley (Top 20), UC Santa Cruz (Top 201-1000), Santa Clara (outside Top 1000), University of San Francisco (outside Top 1000), San Jose State (outside Top 1000), San Francisco State (outside Top 1000) are all among the top of the list in Figure 6.

Motivated by this fact, this section estimates the wage premium of studying in the Bay Area (aggregating over Ω_{Bay} , all universities in San Francisco and San Jose), relative to other cities. Specifically, we estimate the geographic premium for the Bay Area as

$$\frac{1}{N_{\text{Bay}}} \sum_{g \in \Omega_{\text{Bay}}} \sum_{fo} (\Pi_{fo}^{g, \text{Benchmark}} - \Pi_{fo}^{g, \text{Geo}}) \times \text{wage}_{fo}, \quad (22)$$

where N_{Bay} denotes the number of universities located in the Bay Area that appear in our sample. The term measures the average geographic premium of attending a university in the Bay Area (averaged across all local universities). Similarly, we can use Equation (22) to estimate the geographic premium for any U.S. city c . Taking the difference allows us to measure how geography differentially affects the entry-level salaries of graduates from the Bay Area relative to those from city c (referred to as the

Figure 7 displays the Bay Area premium relative to a given city. Again, we report results for 30 cities: the fifteen with the highest premiums (above the dashed line) and the fifteen with the lowest (below the dashed line). The largest geographic disadvantage is observed in Charlottesville, VA, with an average nominal wage penalty of \$7.6K. Midwest or southwest cities that host major state universities—such as El Paso (TX), Tucson (AZ), Lafayette (LA), Ann Arbor (MI), Greenville (SC), Bloomington (IN), Lansing (MI), Columbus (OH), Pittsburgh (PA), and Fargo (ND)—also appear on this list. Unsurprisingly, cities such as New York, Bellingham (WA), and San Diego appear at the bottom of Figure 7. Notably, New York City has a value of \$1.0K, implying that the NYC wage premium is \$1.0K lower than the Bay Area premium.



The orange circles represent real terms adjusted based on BEA RPPs. We see that

the Bay Area premium persists, though to a lesser extent, after accounting for the high cost of living in the Bay Area.

To place the magnitude of the Bay Area premium in context, we compare it with the wage premium associated with university ranking tiers. In nominal terms, attending a Top 20 university carries a premium of \$15K, while attending a Top 21–200 university yields a \$7K premium, relative to graduates from institutions outside the Top 1000.³³ In real terms, the corresponding premiums are \$18K for Top 20 universities and \$6K for Top 21–200 universities.

Using Midwest or Southwest university towns as the benchmark, the Bay Area premium—ranging from \$4.8K to \$6.2K—is sizable: it exceeds half the premium of attending a Top 21–200 university (in both nominal and real terms), amounts to roughly one-third of the nominal premium for the Top 20 universities, and equals about one-fifth of the real premium for the Top 20 universities.

7 Conclusions

This paper constructs a rich individual-level database by combining detailed information from LinkedIn profiles, job postings from Burning Glass Technology, and Glassdoor data to study geographic mobility and job search of fresh college graduates and analyze how these outcomes differ across universities of varying rankings and locations.

Using our dataset, we find that although college graduates are highly mobile outside their university city, most relocate within the same state or to nearby cities. Mobility patterns vary substantially with university ranking and the local economic opportunities available in the area where the university is located. We also show that movers tend to seek better job opportunities and a higher quality of life, consistent with compensating differentials that offset migration costs.

Estimating a model of college graduates sorting into firms and occupations, we find evidence of positive sorting: graduates from higher-ranked universities are more likely to match into cognitively intensive jobs and to locate in high-amenity cities. Interestingly, we find that geographic distance significantly reduces the probability of job matching for most universities, but it has no effect for the most elite institutions. Using the estimated model, we further quantify how geography shapes the labor market outcomes for each university. Our estimates reveal a sizable geographic premium in both nominal and real terms. Although prior research has highlighted the significant

³³Graduates from Top 20 universities earn an average salary of \$81K, compared with \$63K for those from Top 21–200, and \$56K for those from outside the Top 1000.

premium of attending elite universities, our findings suggest that attending a university geographically proximate to large or high-paying firms also offers a considerable premium.

The retention rate for local graduates is notably low for high-ranking universities and for cities that lack high-paying jobs. This trend raises important questions about how cities can retain their college-educated talent. Investing in higher education is one of the major public expenditures in the U.S. and has long been viewed as a key driver of long-term economic growth. Our results suggest that policies aimed at attracting productive firms may be more effective in fostering growth and development. By creating an environment that supports high-quality job opportunities, cities can become more attractive to graduates. Such policies may include offering tax incentives to firms, improving infrastructure, and fostering a business-friendly regulatory environment. We hope to explore these questions further in future research.

Table 3: Estimates of Equation (8)

	(1)	(2)	(3)	(4)	(5)
Cognitive $\times$ Top20	0.199*** (0.043)	0.189*** (0.042)	0.166*** (0.037)	0.168*** (0.037)	0.170*** (0.037)
Cognitive $\times$ Top21-200	0.073*** (0.020)	0.074*** (0.020)	0.085*** (0.019)	0.088*** (0.020)	0.086*** (0.019)
Cognitive $\times$ Top201-1000	0.029* (0.015)	0.031** (0.015)	0.040*** (0.014)	0.039*** (0.014)	0.038*** (0.014)
Social $\times$ Top20	0.065* (0.038)	0.057 (0.037)	0.045 (0.032)	0.045 (0.032)	0.043 (0.032)
Social $\times$ Top21-200	0.014 (0.020)	0.014 (0.020)	0.012 (0.017)	0.011 (0.018)	0.013 (0.018)
Social $\times$ Top201-1000	-0.021 (0.018)	-0.022 (0.018)	-0.014 (0.016)	-0.012 (0.016)	-0.011 (0.016)
Routine $\times$ Top20	-0.051 (0.046)	-0.051 (0.045)	-0.017 (0.036)	-0.017 (0.036)	-0.017 (0.037)
Routine $\times$ Top21-200	0.015 (0.017)	0.017 (0.017)	0.023 (0.014)	0.021 (0.015)	0.022 (0.015)
Routine $\times$ Top201-1000	0.018 (0.014)	0.020 (0.014)	0.018 (0.012)	0.016 (0.012)	0.017 (0.012)
Manual $\times$ Top20	-0.015 (0.032)	-0.016 (0.032)	-0.006 (0.028)	-0.007 (0.028)	-0.008 (0.028)
Manual $\times$ Top21-200	-0.026 (0.017)	-0.026 (0.018)	-0.043*** (0.016)	-0.043** (0.017)	-0.040** (0.017)
Manual $\times$ Top201-1000	-0.044*** (0.015)	-0.044*** (0.015)	-0.048*** (0.014)	-0.046*** (0.014)	-0.044*** (0.014)
Amenity $\times$ Top20		0.091** (0.036)	0.148*** (0.033)	0.156*** (0.033)	0.156*** (0.033)
Amenity $\times$ Top21-200		-0.003 (0.013)	0.031** (0.013)	0.032** (0.013)	0.032** (0.013)
Amenity $\times$ Top201-1000		-0.028** (0.013)	0.005 (0.010)	0.004 (0.010)	0.003 (0.010)
Ldist			-0.067*** (0.011)	-0.053** (0.021)	-0.052** (0.021)
Ldist $\times$ Top20			0.061** (0.024)	0.057** (0.029)	0.057** (0.029)
Ldist $\times$ Top21-200			0.016 (0.021)	0.016 (0.022)	0.016 (0.022)
Ldist $\times$ Top201-1000			-0.028** (0.013)	-0.023 (0.014)	-0.023 (0.014)
CZ			0.448*** (0.030)	0.428*** (0.090)	0.431*** (0.090)
CZ $\times$ Top20			0.177* (0.101)	0.163 (0.107)	0.161 (0.107)
CZ $\times$ Top21-200			0.148*** (0.056)	0.115* (0.064)	0.115* (0.064)
CZ $\times$ Top201-1000			0.087* (0.045)	0.066 (0.044)	0.066 (0.044)

Table 3: Estimates of Equation (8) (Continued)

	(1)	(2)	(3)	(4)	(5)
State			0.086*** (0.026)	-0.026 (0.069)	-0.026 (0.068)
State $\times$ Top20			0.197** (0.097)	0.246** (0.099)	0.248** (0.099)
State $\times$ Top21-200			0.146*** (0.043)	0.181*** (0.049)	0.180*** (0.049)
State $\times$ Top201-1000			0.035 (0.035)	0.048 (0.037)	0.047 (0.037)
Ldist $\times$ In-State				-0.006 (0.036)	-0.006 (0.036)
CZ $\times$ In-State				-0.211 (0.146)	-0.213 (0.146)
State $\times$ In-State				0.292** (0.117)	0.290** (0.117)
Ldist $\times$ Public				-0.019 (0.021)	-0.020 (0.021)
CZ $\times$ Public				0.234*** (0.083)	0.232*** (0.083)
State $\times$ Public				-0.131** (0.066)	-0.129* (0.066)
Cognitive $\times$ Public					0.019 (0.018)
Social $\times$ Public					-0.018 (0.018)
Routine $\times$ Public					-0.011 (0.017)
Manual $\times$ Public					-0.019 (0.015)
Observations	84,904	84,360	84,360	82,463	82,463
Adjusted R^2	0.68	0.67	0.74	0.75	0.75

Notes: Columns (1)-(3) report the estimated coefficients for equation (8). Column (1) has the regressors as the interaction of ranking dummies and BGT tasks; Column (2) adds amenities, and Column (3) adds geographic variables. Column (4) adds interactions of public universities and the share of in-state student enrollment. Column (5) adds interactions between public-university status and task measures. The models are estimated using OLS. The full sample covers 22 occupations, 25492 firms, and 266 universities. Standard errors are clustered at the firm-occupation level.

Table 4: Estimates of Equation (9)

	(1)	(2)	(3)	(4)	(5)
Cognitive $\times$ Top20	0.112*** (0.032)	0.110*** (0.032)	0.086*** (0.027)	0.086*** (0.028)	0.085*** (0.028)
Cognitive $\times$ Top21-200	0.029* (0.016)	0.030* (0.016)	0.038*** (0.015)	0.040*** (0.015)	0.040*** (0.015)
Cognitive $\times$ Top201-1000	0.015 (0.013)	0.016 (0.013)	0.020* (0.012)	0.021* (0.012)	0.021* (0.012)
Social $\times$ Top20	0.082** (0.033)	0.078** (0.032)	0.056** (0.027)	0.058** (0.027)	0.057** (0.027)
Social $\times$ Top21-200	0.018 (0.015)	0.019 (0.015)	0.013 (0.014)	0.014 (0.014)	0.015 (0.014)
Social $\times$ Top201-1000	-0.015 (0.014)	-0.016 (0.014)	-0.010 (0.012)	-0.008 (0.013)	-0.008 (0.013)
Routine $\times$ Top20	-0.058* (0.033)	-0.053 (0.033)	-0.030 (0.027)	-0.031 (0.027)	-0.031 (0.027)
Routine $\times$ Top21-200	0.000 (0.014)	0.001 (0.014)	0.001 (0.012)	-0.001 (0.012)	-0.002 (0.013)
Routine $\times$ Top201-1000	0.016 (0.011)	0.017 (0.011)	0.014 (0.010)	0.011 (0.010)	0.011 (0.011)
Manual $\times$ Top20	0.024 (0.028)	0.023 (0.028)	0.034 (0.025)	0.036 (0.025)	0.037 (0.026)
Manual $\times$ Top21-200	0.011 (0.014)	0.010 (0.014)	-0.003 (0.014)	-0.003 (0.014)	-0.004 (0.014)
Manual $\times$ Top201-1000	-0.014 (0.012)	-0.014 (0.012)	-0.013 (0.011)	-0.011 (0.011)	-0.012 (0.011)
Amenity $\times$ Top20		0.087** (0.034)	0.146*** (0.030)	0.154*** (0.031)	0.154*** (0.031)
Amenity $\times$ Top21-200		-0.005 (0.013)	0.031** (0.013)	0.034*** (0.012)	0.034*** (0.012)
Amenity $\times$ Top201-1000		-0.024* (0.013)	0.009 (0.010)	0.008 (0.010)	0.008 (0.010)
Ldist			-0.066*** (0.011)	-0.035* (0.020)	-0.035* (0.019)
Ldist $\times$ Top20			0.077*** (0.024)	0.062** (0.028)	0.062** (0.028)
Ldist $\times$ Top21-200			0.016 (0.019)	0.014 (0.020)	0.014 (0.020)
Ldist $\times$ Top201-1000			-0.022* (0.013)	-0.019 (0.014)	-0.019 (0.014)
CZ			0.448*** (0.030)	0.465*** (0.089)	0.465*** (0.089)
CZ $\times$ Top20			0.184* (0.099)	0.147 (0.105)	0.147 (0.105)
CZ $\times$ Top21-200			0.132** (0.053)	0.092 (0.061)	0.092 (0.061)
CZ $\times$ Top201-1000			0.084** (0.043)	0.060 (0.043)	0.060 (0.043)

Table 4: Estimates of Equation (9) (Continued)

	(1)	(2)	(3)	(4)	(5)
State			0.066*** (0.026)	-0.020 (0.068)	-0.020 (0.068)
State $\times$ Top20			0.277*** (0.102)	0.309*** (0.107)	0.310*** (0.107)
State $\times$ Top21-200			0.165*** (0.043)	0.196*** (0.049)	0.196*** (0.049)
State $\times$ Top201-1000			0.067* (0.034)	0.075** (0.038)	0.075** (0.038)
Ldist $\times$ In-State				-0.032 (0.033)	-0.032 (0.033)
CZ $\times$ In-State				-0.258* (0.138)	-0.258* (0.138)
State $\times$ In-State				0.236** (0.113)	0.236** (0.113)
Ldist $\times$ Public				-0.016 (0.020)	-0.016 (0.020)
CZ $\times$ Public				0.234*** (0.079)	0.235*** (0.079)
State $\times$ Public				-0.114* (0.067)	-0.114* (0.066)
Cognitive $\times$ Public					-0.004 (0.014)
Social $\times$ Public					-0.005 (0.015)
Routine $\times$ Public					0.002 (0.014)
Manual $\times$ Public					0.013 (0.013)
Observations	84,281	83,729	83,729	81,828	81,828
Adjusted R^2	0.75	0.75	0.81	0.81	0.81

Notes: Columns (1)-(3) report the estimated coefficients for equation (8). Column (1) includes the interaction of ranking dummies and BGT tasks; Column (2) adds amenities, and Column (3) adds geographic variables. Column (4) adds interactions of public universities and the share of in-state students enrollment. Column (5) adds interactions between public-university status and task measures. All models are OLS estimates. The full sample covers 22 occupations, 25492 firms, and 266 universities. Standard errors are clustered by university and reported in parentheses.

Table 5: OLS Estimates of Equation (8) with Additional Characteristics

	(1)	(2)	(3)	(4)
Cognitive $\times$ Top20	0.168*** (0.037)	0.185*** (0.039)	0.147*** (0.043)	0.119*** (0.037)
Cognitive $\times$ Top21-200	0.085*** (0.019)	0.091*** (0.020)	0.071*** (0.022)	0.047*** (0.017)
Cognitive $\times$ Top201-1000	0.039*** (0.014)	0.041*** (0.014)	0.032** (0.015)	0.009 (0.011)
Social $\times$ Top20	0.043 (0.032)	0.029 (0.035)	0.054 (0.038)	0.035 (0.032)
Social $\times$ Top21-200	0.012 (0.017)	0.000 (0.019)	0.015 (0.021)	0.009 (0.016)
Social $\times$ Top201-1000	-0.014 (0.016)	-0.017 (0.017)	-0.010 (0.017)	-0.004 (0.013)
Routine $\times$ Top20	-0.015 (0.036)	-0.018 (0.040)	-0.016 (0.042)	-0.005 (0.037)
Routine $\times$ Top21-200	0.022 (0.014)	0.023 (0.017)	0.024 (0.018)	0.031** (0.014)
Routine $\times$ Top201-1000	0.018 (0.012)	0.018 (0.013)	0.019 (0.013)	0.022** (0.010)
Manual $\times$ Top20	-0.007 (0.028)	-0.025 (0.029)	-0.026 (0.031)	-0.024 (0.026)
Manual $\times$ Top21-200	-0.043*** (0.016)	-0.059*** (0.017)	-0.060*** (0.018)	-0.049*** (0.013)
Manual $\times$ Top201-1000	-0.048*** (0.014)	-0.053*** (0.014)	-0.053*** (0.015)	-0.042*** (0.010)
Amenity $\times$ Top20	0.135*** (0.033)	0.172*** (0.035)	0.130*** (0.035)	0.099*** (0.033)
Amenity $\times$ Top21-200	0.034*** (0.013)	0.040*** (0.014)	0.007 (0.016)	-0.009 (0.014)
Amenity $\times$ Top201-1000	0.004 (0.010)	0.001 (0.012)	-0.021 (0.014)	-0.032*** (0.012)
Ldist	-0.055*** (0.020)	-0.031 (0.025)	-0.056** (0.027)	-0.025 (0.024)
Ldist $\times$ Top20	0.055** (0.028)	0.064** (0.029)	0.043 (0.027)	0.025 (0.025)
Ldist $\times$ Top21-200	0.018 (0.020)	0.012 (0.020)	-0.002 (0.018)	-0.013 (0.016)
Ldist $\times$ Top201-1000	-0.021 (0.013)	-0.023* (0.013)	-0.030** (0.013)	-0.027** (0.011)
CZ	0.406*** (0.088)	0.312*** (0.116)	0.388*** (0.119)	0.155 (0.104)
CZ $\times$ Top20	0.243** (0.111)	0.332*** (0.123)	0.344*** (0.126)	0.305*** (0.115)
CZ $\times$ Top21-200	0.202*** (0.059)	0.270*** (0.066)	0.260*** (0.066)	0.186*** (0.058)
CZ $\times$ Top201-1000	0.099** (0.043)	0.113** (0.046)	0.107** (0.046)	0.060 (0.038)
State	-0.029 (0.067)	0.051 (0.086)	-0.065 (0.087)	0.058 (0.078)
State $\times$ Top20	0.208** (0.104)	0.237** (0.110)	0.163 (0.111)	0.108 (0.102)
State $\times$ Top21-200	0.157*** (0.044)	0.141*** (0.046)	0.100** (0.048)	0.063 (0.044)
State $\times$ Top201-1000	0.041 (0.034)	0.034 (0.036)	0.014 (0.036)	0.004 (0.033)

Table 5: OLS Estimates of Equation (8) with Additional Characteristics (Continued)

	(1)	(2)	(3)	(4)
Cognitive $\times$ Most Selective		-0.025 (0.032)	-0.028 (0.031)	-0.024 (0.028)
Cognitive $\times$ Highly Selective		0.038 (0.024)	0.029 (0.023)	0.025 (0.018)
Cognitive $\times$ Selective		0.014 (0.015)	0.019 (0.015)	0.003 (0.011)
Social $\times$ Most Selective		-0.010 (0.029)	-0.006 (0.028)	0.001 (0.022)
Social $\times$ Highly Selective		0.005 (0.023)	0.013 (0.023)	0.006 (0.018)
Social $\times$ Selective		-0.022 (0.016)	-0.023 (0.016)	-0.008 (0.012)
Routine $\times$ Most Selective		0.011 (0.028)	0.014 (0.027)	0.009 (0.023)
Routine $\times$ Highly Selective		-0.022 (0.020)	-0.021 (0.020)	-0.012 (0.016)
Routine $\times$ Selective		-0.004 (0.013)	-0.005 (0.013)	-0.007 (0.010)
Manual $\times$ Most Selective		-0.011 (0.023)	-0.010 (0.023)	-0.019 (0.020)
Manual $\times$ Highly Selective		0.008 (0.020)	0.009 (0.020)	0.000 (0.014)
Manual $\times$ Selective		-0.032** (0.014)	-0.032** (0.014)	-0.013 (0.010)
Amenity $\times$ Most Selective		-0.035* (0.020)	-0.048** (0.020)	-0.037** (0.017)
Amenity $\times$ Highly Selective		-0.015 (0.023)	-0.032 (0.023)	-0.017 (0.018)
Amenity $\times$ Selective		0.011 (0.010)	-0.002 (0.011)	0.006 (0.010)
Ldist $\times$ Most Selective		-0.039* (0.021)	-0.049** (0.022)	-0.053*** (0.020)
Ldist $\times$ Highly Selective		0.024 (0.024)	0.011 (0.025)	-0.007 (0.023)
Ldist $\times$ Selective		-0.019 (0.027)	-0.017 (0.027)	-0.028 (0.023)
CZ $\times$ Most Selective		0.035 (0.102)	0.013 (0.099)	0.063 (0.088)
CZ $\times$ Highly Selective		0.004 (0.086)	-0.016 (0.086)	0.022 (0.073)
CZ $\times$ Selective		0.168 (0.107)	0.144 (0.107)	0.140 (0.092)
State $\times$ Most Selective		-0.128 (0.079)	-0.146* (0.076)	-0.146** (0.068)
State $\times$ Highly Selective		0.056 (0.071)	0.034 (0.073)	0.006 (0.066)
State $\times$ Selective		-0.036 (0.082)	-0.019 (0.081)	-0.057 (0.076)

Table 5: OLS Estimates of Equation (8) with Additional Characteristics (Continued)

	(1)	(2)	(3)	(4)
Ldist $\times$ Public	-0.029** (0.014)	-0.025 (0.026)	-0.025 (0.026)	-0.002 (0.022)
CZ $\times$ Public	0.192*** (0.054)	0.078 (0.103)	0.085 (0.102)	0.041 (0.087)
State $\times$ Public	-0.124*** (0.036)	-0.105 (0.076)	-0.107 (0.076)	-0.055 (0.070)
Ldist $\times$ In-State	-0.001 (0.025)	-0.020 (0.027)	-0.018 (0.027)	-0.026 (0.024)
CZ $\times$ In-State	-0.127 (0.118)	-0.073 (0.138)	-0.103 (0.134)	0.029 (0.114)
State $\times$ In-State	0.266*** (0.086)	0.185* (0.099)	0.207** (0.097)	0.056 (0.082)
Cognitive $\times$ Stem			0.224*** (0.067)	0.192*** (0.052)
Social $\times$ Stem			-0.143** (0.066)	-0.110** (0.049)
Routine $\times$ Stem			-0.017 (0.059)	-0.035 (0.044)
Manual $\times$ Stem			0.007 (0.049)	0.003 (0.040)
Ldist $\times$ Stem			0.135** (0.055)	0.107** (0.052)
CZ $\times$ Stem			-0.157 (0.190)	0.038 (0.173)
State $\times$ Stem			0.490*** (0.158)	0.292** (0.145)
Amenity $\times$ Stem			0.148*** (0.047)	0.135*** (0.041)
Network				1.675*** (0.048)
Observations	84,360	84,360	84,274	84,274
Adjusted R^2	0.74	0.75	0.75	0.79

Notes: This table reports the estimated coefficients for equation (8). All columns are OLS estimates. The full sample covers 22 occupations, 25492 firms, and 266 universities. Standard errors are clustered by university and reported in parentheses.

Table 6: Estimates of Equation (8) with Disaggregated Groups

	(1)	(2)	(3)	(4)
Cognitive $\times$ Top20	0.199*** (0.043)	0.189*** (0.042)	0.167*** (0.037)	0.170*** (0.037)
Cognitive $\times$ Top21-100	0.077*** (0.023)	0.077*** (0.024)	0.084*** (0.022)	0.087*** (0.023)
Cognitive $\times$ Top101-200	0.075** (0.036)	0.082** (0.036)	0.084*** (0.032)	0.089*** (0.033)
Cognitive $\times$ Top201-500	0.047** (0.020)	0.048** (0.020)	0.058*** (0.018)	0.058*** (0.018)
Cognitive $\times$ Top501-1000	0.017 (0.017)	0.021 (0.018)	0.016 (0.015)	0.016 (0.015)
Social $\times$ Top20	0.060 (0.038)	0.052 (0.038)	0.043 (0.032)	0.044 (0.032)
Social $\times$ Top21-100	-0.016 (0.022)	-0.018 (0.022)	-0.008 (0.019)	-0.008 (0.020)
Social $\times$ Top101-200	0.080*** (0.031)	0.080*** (0.031)	0.058** (0.025)	0.058** (0.026)
Social $\times$ Top201-500	-0.029 (0.024)	-0.029 (0.024)	-0.017 (0.020)	-0.016 (0.021)
Social $\times$ Top501-1000	-0.031* (0.019)	-0.034* (0.019)	-0.009 (0.016)	-0.007 (0.016)
Routine $\times$ Top20	-0.052 (0.046)	-0.052 (0.046)	-0.015 (0.037)	-0.016 (0.037)
Routine $\times$ Top21-100	0.041** (0.018)	0.041** (0.018)	0.035** (0.016)	0.031* (0.016)
Routine $\times$ Top101-200	-0.041 (0.034)	-0.036 (0.034)	-0.006 (0.026)	-0.005 (0.027)
Routine $\times$ Top201-500	0.018 (0.018)	0.018 (0.018)	0.019 (0.016)	0.018 (0.016)
Routine $\times$ Top501-1000	0.016 (0.015)	0.020 (0.015)	0.016 (0.013)	0.013 (0.013)
Manual $\times$ Top20	-0.010 (0.032)	-0.012 (0.032)	-0.007 (0.028)	-0.007 (0.028)
Manual $\times$ Top21-100	-0.022 (0.021)	-0.022 (0.021)	-0.036* (0.019)	-0.035* (0.020)
Manual $\times$ Top101-200	-0.040* (0.023)	-0.038 (0.023)	-0.058*** (0.021)	-0.059*** (0.022)
Manual $\times$ Top201-500	-0.060*** (0.019)	-0.059*** (0.019)	-0.061*** (0.017)	-0.058*** (0.017)
Manual $\times$ Top501-1000	-0.030* (0.017)	-0.029* (0.017)	-0.032** (0.015)	-0.030* (0.015)
Amenity $\times$ Top20		0.348*** (0.127)	0.519*** (0.114)	0.547*** (0.116)
Amenity $\times$ Top21-100		0.106* (0.060)	0.137** (0.056)	0.137** (0.055)
Amenity $\times$ Top101-200		-0.348*** (0.080)	0.062 (0.056)	0.089 (0.058)
Amenity $\times$ Top201-500		-0.014 (0.066)	0.049 (0.051)	0.054 (0.051)
Amenity $\times$ Top501-1000		-0.171*** (0.066)	-0.023 (0.048)	-0.040 (0.048)
Ldist $\times$ Public				-0.015 (0.021)
CZ $\times$ Public				0.228*** (0.081)
State $\times$ Public				-0.131** (0.066)

Table 6: Estimates of Equation (8) with Disaggregated Groups (Continued)

	(1)	(2)	(3)	(4)
Ldist			-0.067*** (0.011)	-0.054** (0.021)
Ldist $\times$ Top20			0.061** (0.024)	0.058** (0.029)
Ldist $\times$ Top21-100			0.044* (0.025)	0.045* (0.026)
Ldist $\times$ Top101-200			-0.048** (0.022)	-0.048** (0.024)
Ldist $\times$ Top201-500			-0.033** (0.015)	-0.028* (0.016)
Ldist $\times$ Top501-1000			-0.017 (0.015)	-0.013 (0.015)
CZ			0.447*** (0.030)	0.428*** (0.091)
CZ $\times$ Top20			0.180* (0.101)	0.167 (0.108)
CZ $\times$ Top21-100			0.187*** (0.065)	0.163** (0.071)
CZ $\times$ Top101-200			0.058 (0.090)	0.013 (0.098)
CZ $\times$ Top201-500			0.124** (0.057)	0.095* (0.057)
CZ $\times$ Top501-1000			0.043 (0.054)	0.032 (0.053)
State			0.086*** (0.026)	-0.036 (0.069)
State $\times$ Top20			0.197** (0.097)	0.250** (0.099)
State $\times$ Top21-100			0.160*** (0.051)	0.193*** (0.055)
State $\times$ Top101-200			0.116* (0.065)	0.165** (0.071)
State $\times$ Top201-500			0.017 (0.045)	0.035 (0.047)
State $\times$ Top501-1000			0.060 (0.038)	0.067* (0.041)
Ldist $\times$ In-State				-0.009 (0.035)
CZ $\times$ In-State				-0.206 (0.147)
State $\times$ In-State				0.306*** (0.116)
Observations	86,010	85,462	84,360	82,463
Adjusted R^2	0.67	0.67	0.74	0.75

Notes: Columns (1)-(3) report the estimated coefficients for equation (8). Column (1) includes the interaction of ranking dummies and BGT tasks; Column (2) adds amenities, and Column (3) adds geographic variables. Column (4) adds interactions of public universities and the share of in-state students enrollment. All models are OLS estimates. Standard errors are clustered by university and reported in parentheses.

Table 7: Estimates of Equation (8) using US News Rankings

	(1)	(2)	(3)	(4)
Cognitive $\times$ Top20	0.181*** (0.052)	0.165*** (0.050)	0.121*** (0.042)	0.122*** (0.042)
Cognitive $\times$ Top200	0.065*** (0.020)	0.066*** (0.020)	0.063*** (0.017)	0.064*** (0.018)
Social $\times$ Top20	0.057 (0.046)	0.053 (0.046)	0.043 (0.036)	0.042 (0.036)
Social $\times$ Top200	-0.000 (0.020)	-0.000 (0.020)	-0.004 (0.016)	-0.004 (0.016)
Routine $\times$ Top20	-0.059 (0.057)	-0.066 (0.056)	-0.030 (0.043)	-0.029 (0.042)
Routine $\times$ Top200	-0.007 (0.019)	-0.006 (0.019)	-0.004 (0.015)	-0.006 (0.016)
Manual $\times$ Top20	-0.009 (0.038)	-0.009 (0.038)	0.015 (0.032)	0.015 (0.032)
Manual $\times$ Top200	0.018 (0.017)	0.019 (0.017)	0.013 (0.015)	0.013 (0.015)
Amenity $\times$ Top20		0.088** (0.037)	0.144*** (0.035)	0.155*** (0.035)
Amenity $\times$ Top200		-0.007 (0.014)	0.035*** (0.013)	0.036*** (0.013)
Ldist			-0.106*** (0.010)	-0.080*** (0.020)
Ldist $\times$ Top20			0.074*** (0.021)	0.062*** (0.023)
Ldist $\times$ Top200			0.070*** (0.021)	0.066*** (0.022)
CZ			0.516*** (0.036)	0.449*** (0.122)
CZ $\times$ Top20			0.143 (0.132)	0.127 (0.141)
CZ $\times$ Top200			0.097 (0.062)	0.089 (0.069)
State			0.113*** (0.025)	0.090 (0.081)
State $\times$ Top20			0.117 (0.150)	0.142 (0.150)
State $\times$ Top200			0.167*** (0.044)	0.179*** (0.045)
Ldist $\times$ Public				-0.037 (0.028)
CZ $\times$ Public				0.217* (0.127)
State $\times$ Public				-0.125 (0.086)
Ldist $\times$ In-State				0.001 (0.046)
CZ $\times$ In-State				-0.158 (0.202)
State $\times$ In-State				0.177 (0.143)
Observations	57,788	57,479	57,479	56,007
Adjusted R^2	0.67	0.67	0.75	0.75

Notes: Columns (1)-(3) report the estimated coefficients for equation (8). Column (1) includes the interaction of ranking dummies and BGT tasks; Column (2) adds amenities, and Column (3) adds geographic variables. Column (4) adds interactions of public universities and the share of in-state student enrollment. All models are OLS estimates. Standard errors are clustered by university and reported in parentheses.

References

- Abowd, J. M., Kramarz, F. and Margolis, D. N. (1999), ‘High wage workers and high wage firms’, *Econometrica* **67**(2), 251–333.
- Acemoglu, D. and Autor, D. (2011), Skills, tasks and technologies: Implications for employment and earnings, in ‘Handbook of labor economics’, Vol. 4, Elsevier, pp. 1043–1171.
- Albert, C. and Monras, J. (2022), ‘Immigration and spatial equilibrium: the role of expenditures in the country of origin’, *American Economic Review* **112**(11), 3763–3802.
- Amanzadeh, N., Kermani, A. and McQuade, T. (2024), Return migration and human capital flows, Technical report, National Bureau of Economic Research.
- Atalay, E., Phongthiengtham, P., Sotelo, S. and Tannenbaum, D. (2018), ‘New technologies and the labor market’, *Journal of Monetary Economics* **97**, 48–67.
- Atalay, E., Phongthiengtham, P., Sotelo, S. and Tannenbaum, D. (2020), ‘The evolution of work in the united states’, *American Economic Journal: Applied Economics* **12**, 1–36.
- Atalay, E., Sotelo, S. and Tannenbaum, D. (2024), ‘The geography of job tasks’, *Journal of Labor Economics* **42**(4), 979–1008.
- Autor, D., Dorn, D., Hanson, G. H., Jones, M. R. and Setzler, B. (2025), Places versus people: the ins and outs of labor market adjustment to globalization, Technical report, National Bureau of Economic Research.
- Auxier, B., Anderson, M. et al. (2021), ‘Social media use in 2021’, *Pew Research Center* **1**(1), 1–4.
- Barron’s Educational Series, I., ed. (2017), *Barron’s Profiles of American Colleges 2017*, 33rd edn, Barron’s Educational Series, Hauppauge, NY.
- Batra, H., Michaud, A. and Mongey, S. (2023), Online job posts contain very little wage information, Technical report, National Bureau of Economic Research.
- Becker, G. S. (1973), ‘A theory of marriage: Part i’, *Journal of Political Economy* **81**(4), 813–846.
- Bonhomme, S., Holzheu, K., Lamadon, T., Manresa, E., Mogstad, M. and Setzler, B. (2023), ‘How much should we trust estimates of firm effects and worker sorting?’, *Journal of Labor Economics* **41**(2), 291–322.
- Bound, J. and Holzer, H. J. (2000), ‘Demand shifts, population adjustments, and labor market outcomes during the 1980s’, *Journal of labor Economics* **18**(1), 20–54.
- Braxton, J. C. and Taska, B. (2023), ‘Technological change and the consequences of job

- loss', *American Economic Review* **113**(2), 279–316.
- Bryan, G. and Morten, M. (2019), 'The aggregate productivity effects of internal migration: Evidence from indonesia', *Journal of Political Economy* **127**(5), 2229–2268.
- Burstein, A., Morales, E. and Vogel, J. (2019), 'Changes in between-group inequality: computers, occupations, and international trade', *American Economic Journal: Macroeconomics* **11**(2), 348–400.
- Cadena, B. C. and Kovak, B. K. (2016), 'Immigrants equilibrate local labor markets: Evidence from the great recession', *American Economic Journal: Applied Economics* **8**(1), 257–90.
- Card, D. (2001), 'Immigrant inflows, native outflows, and the local labor market impacts of higher immigration', *Journal of Labor Economics* **19**(1), 22–64.
- Card, D., Heining, J. and Kline, P. (2013), 'Workplace heterogeneity and the rise of west german wage inequality', *The Quarterly journal of economics* **128**(3), 967–1015.
- Chetty, R., Deming, D. J. and Friedman, J. N. (2023), Diversifying society's leaders? the determinants and causal effects of admission to highly selective private colleges, Technical report, National Bureau of Economic Research.
- Chetty, R., Friedman, J. N., Saez, E., Turner, N. and Yagan, D. (2020), 'Income segregation and intergenerational mobility across colleges in the united states', *The Quarterly Journal of Economics* **135**(3), 1567–1633.
- Costinot, A. and Vogel, J. (2010), 'Matching and inequality in the world economy', *Journal of Political Economy* **118**(4), 747–786.
- Costinot, A. and Vogel, J. (2015), 'Beyond ricardo: Assignment models in international trade', *Annual Review of Economics* **7**(1), 31–62.
- Dale, S. B. and Krueger, A. B. (2002), 'Estimating the payoff to attending a more selective college: An application of selection on observables and unobservables', *The Quarterly Journal of Economics* **117**(4), 1491–1527.
- Deming, D. J. (2017), 'The growing importance of social skills in the labor market', *The Quarterly Journal of Economics* **132**(4), 1593–1640.
- Deming, D. J. and Noray, K. (2020), 'Earnings dynamics, changing job skills, and stem careers', *The Quarterly Journal of Economics* **135**(4), 1965–2005.
- Deming, D. and Kahn, L. B. (2018), 'Skill requirements across firms and labor markets: Evidence from job postings for professionals', *Journal of Labor Economics* **36**(S1), S337–S369.
- Diamond, R. (2016), 'The determinants and welfare implications of us workers' diverging location choices by skill: 1980–2000', *American Economic Review* **106**(3), 479–524.

- Diamond, R. and Gaubert, C. (2022), ‘Spatial sorting and inequality’, *Annual Review of Economics* **14**, 795–819.
- Diamond, R. and Moretti, E. (2024), ‘Where is standard of living the highest? local prices and the geography of consumption’.
- Dorn, D., Schoner, F., Seebacher, M., Simon, L. and Woessmann, L. (2025), ‘Multidimensional skills on linkedin profiles: Measuring human capital and the gender skill gap’.
- Eeckhout, J. (2018), ‘Sorting in the labor market’, *Annual Review of Economics* **10**(1), 1–29.
- Eeckhout, J. and Kircher, P. (2011), ‘Identifying sorting—in theory’, *The Review of Economic Studies* **78**(3), 872–906.
- Greenland, A., Lopresti, J. and McHenry, P. (2019), ‘Import competition and internal migration’, *Review of Economics and Statistics* **101**(1), 44–59.
- Hall, B. H., Jaffe, A. B. and Trajtenberg, M. (2001), ‘The nber patent citation data file: Lessons, insights and methodological tools’.
- Hanson, G. H. and Liu, C. (2016), High-skilled immigration and the comparative advantage of foreign-born workers across us occupations, *in* ‘Talent Flows in the Global Economy’, University of Chicago Press.
- Hanson, G. H. and Liu, C. (2021), Immigration and occupational comparative advantage, Technical report, National Bureau of Economic Research.
- Hanushek, E. A. and Kimko, D. D. (2000), ‘Schooling, labor-force quality, and the growth of nations’, *American economic review* pp. 1184–1208.
- Hershbein, B. and Kahn, L. B. (2018), ‘Do recessions accelerate routine-biased technological change? evidence from vacancy postings’, *American Economic Review* **108**(7), 1737–72.
- Kennan, J. and Walker, J. R. (2011), ‘The effect of expected income on individual migration decisions’, *Econometrica* **79**(1), 211–251.
- Khanna, G. and Morales, N. (2025), ‘The IT boom and other unintended consequences of chasing the american dream’.
- Lagakos, D. and Waugh, M. E. (2013), ‘Selection, agriculture, and cross-country productivity differences’, *American Economic Review* **103**(2), 948–80.
- Lopes de Melo, R. (2018), ‘Firm wage differentials and labor market sorting: Reconciling theory and evidence’, *Journal of Political Economy* **126**(1), 313–346.
- Martellini, P., Schoellman, T. and Sockin, J. (2024), ‘The global distribution of college graduate quality’, *Journal of Political Economy* **132**(2), 434–483.
- Moretti, E. (2013), ‘Real wage inequality’, *American Economic Journal: Applied Eco-*

- nomics* **5**(1), 65–103.
- Munshi, K. (2003), ‘Networks in the modern economy: Mexican migrants in the us labor market’, *The Quarterly Journal of Economics* **118**(2), 549–599.
- Pellegrina, H. S. and Sotelo, S. (2021), Migration, specialization, and trade: Evidence from brazil’s march to the west, Technical report, National Bureau of Economic Research.
- Ruggles, S., Alexander, J. T., Genadek, K., Goeken, R., Schroeder, M. B., Sobek, M. et al. (2010), ‘Integrated public use microdata series: Version 5.0 [machine-readable database]’, *Minneapolis: University of Minnesota* **42**.
- Song, J., Price, D. J., Guvenen, F., Bloom, N. and Von Wachter, T. (2019), ‘Firming up inequality’, *The Quarterly journal of economics* **134**(1), 1–50.
- Spitz-Oener, A. (2006), ‘Technical change, job tasks, and rising educational demands: Looking outside the wage structure’, *Journal of labor economics* **24**(2), 235–270.
- Tombe, T. and Zhu, X. (2019), ‘Trade, migration, and productivity: A quantitative analysis of china’, *American Economic Review* **109**(5), 1843–72.
- Yamaguchi, S. (2012), ‘Tasks and heterogeneous human capital’, *Journal of Labor Economics* **30**(1), 1–53.
- Zimmerman, S. D. (2019), ‘Elite colleges and upward mobility to top jobs and top incomes’, *American Economic Review* **109**(1), 1–47.

Online Appendix

College Graduates in the Labor Market: Geographic Mobility and Sorting into Firms and Occupations

Chen Liu
NUS

Sibo Liu
HKBU

Yifan Zhang
CUHK

A Data Appendix

A.1 Data Description

The LinkedIn Data. Our primary data are purchased from Revelio Labs, which compiles information from publicly available LinkedIn profiles and other sources. LinkedIn is recognized as a major online platform for professional networking, where individuals voluntarily provide their work experiences and educational backgrounds for job search and career development purposes. As of 2023, LinkedIn has more than 1 billion registered members across over 200 countries and territories.³⁴ LinkedIn user profiles essentially function as self-reported resumes, containing detailed information on individuals' educational and employment histories. These include the universities they attended, the degrees and fields of study they pursued, their employers, job titles, and the dates during which they held these positions. In the United States, a majority of college graduates use LinkedIn (Auxier, Anderson et al., 2021). Our version of the LinkedIn data from Revelio Labs has only recently begun to be used in economics research. A few recent studies have employed it to examine the returns to international migration (Amanzadeh, Kermani and McQuade, 2024), Indian engineering migrants to the United States (Khanna and Morales, 2025), and the gender skill gap (Dorn, Schoner, Seebacher, Simon and Woessmann, 2025).

Because university names, employer names, and job titles are all self-reported in LinkedIn profiles, we undertake extensive work to harmonize and standardize this information across our three sources (LinkedIn, Burning Glass, and Glassdoor). Employer names are cleaned to ensure consistency across three platforms. The database for LinkedIn (Revelio Labs) provides OCCSOC occupation codes for user self-reported job titles using internal text-based algorithms. For job titles in Glassdoor, we employ a large language model (ChatGPT 4o) to obtain OCCSOC occupation codes. University names are also systematically standardized and linked to external sources of institutional data. To capture university location, we assign each institution a city and state based on its primary campus location, and then map these to geographic codes: commuting zone (CZ) and state codes.

Finally, LinkedIn profiles include detailed geographic information about individuals' current jobs. Crucially, this data specifies the location of each job—not just the employer's headquarters—at multiple levels, including state, metropolitan area, city, and even street address. This granularity enables us to analyze graduates' geographic mobility based on the actual location of their employment. Specifically, we use this job location information to identify where graduates are employed and compare it with the

³⁴For more details, please refer to the [Wikipedia page](#) of LinkedIn.

location of their alma mater. This comparison allows us to measure geographic mobility relative to where graduates attended college.

Since our focus is on job matching among fresh college graduates, we restrict the sample to individuals who obtained their bachelor’s degrees (as their highest degree) between 2016 and 2018, are currently employed by a U.S. firm, and received their college education at a U.S. institution. Our objective is to measure individuals’ first “primary” job immediately after graduation. We use firm and occupation information from their 2018 job record. For those reporting multiple jobs in their profile, we select the position in which they had the longest tenure. We exclude individuals currently working as interns and those still enrolled in graduate programs (master’s or doctoral).

The Burning Glass Data. The data contains the universal job posting data collected by Burning Glass Technology (BGT), and was first used in [Hershbein and Kahn \(2018\)](#). We use four BGT task variables commonly used in the literature: cognitive, social, routine, and manual. Following [Spitz-Oener \(2006\)](#), [Atalay, Phongthiengtham, Sotelo and Tannenbaum \(2020\)](#), and [Deming and Kahn \(2018\)](#), we measure these task variables from job advertisements based on keywords.

The BGT data have already processed job titles and mapped them to OCCSOC codes. Location information for vacancies is also standardized and available at multiple geographic levels (MSA, CZ, and state codes), which we can directly use. Employer names, however, are not fully cleaned. We harmonize and standardize these to create a crosswalk with LinkedIn and Glassdoor (see Appendix A.2).

We measure task intensity using Burning Glass Technologies (BGT) job advertisements, focusing on cognitive, social, routine, and manual tasks as captured by the skill requirements listed in job ads. Specifically, BGT reports whether a given skill (from thousands of listed skills) is required for each job title. For each task category, we count the number of required skills and compute task measures as percentile rankings across all postings. We then take the average across postings that share the same firm name and occupation code. These BGT task variables are thus measured at the firm and occupation levels, and each variable is standardized to range between 0 and 1.

For cognitive tasks, we based our selection on two sets of words. The first set includes those related to [Spitz-Oener \(2006\)](#). The second set incorporates words that refer to advanced computer software or skills. The selected words are as follows:

- problem solving, research, analytical, critical thinking, math, statistics, development
- Microsoft C#, Microsoft SQL, Microsoft server, social media platforms, virtual private networking, Microsoft visual C++, C (programming language), statistics, statistical software, software development, simulation software, scripting, sql databases and programming, neuroscience, machine learning, mathematics, aerospace engineering, application programming interface (API), application development, automation engineering, big data, C and C++, cache (computing), chemical engineering, cloud computing, computer hardware, data analysis, data mining, data science

For social tasks, we adopt the keywords based on [Deming and Kahn \(2018\)](#). These words are

- communication, teamwork, collaboration, negotiation, presentation, supervisory, leadership, management, mentoring, staff

For routine tasks, we based our selection on two sets of words. The first set includes words directly related to the administrative nature of jobs. The second set includes software commonly used by admin-

istrative staff.

- budget, accounting, cost, account management, admin, billing, administration, education administration
- Microsoft, spreadsheet, Photoshop, Google Docs, Google Maps, Google Drive, Google apps, Macintosh OS, YouTube, Facebook, payroll, accounting and finance software, administrative support
- human resource management systems, human resources software, identity management, record keeping

Manual tasks use the following keywords:

- customer, client, service, physical abilities, repair, cleaning, sales

The Glassdoor Data. Our data source for salary is collected from Glassdoor. Glassdoor is an online platform where workers can review employers, report their earnings, and search for jobs. To encourage participation, Glassdoor uses a “give-to-get” model: users who submit an employer review or salary report gain access to others’ anonymous data. Users share a wide range of information, including their compensation details such as base pay, bonuses, currency, and job-related information such as years of work experience, employment status, job title, location, and employer. To keep access, users must update their data annually if they have not submitted a recent review or salary report. For our research purposes, the dataset includes rich employer-employee matches with detailed information about individual workers.

We obtain a snapshot of Glassdoor data collected between September and October 2024, which includes detailed wage information by firm, occupation, location, and years of experience. The original sample contains 183,257 companies and 10.5 million wage records, of which 5.04 million are at the employer–title level and 5.45 million at the employer–title–location level. For use in our study, we match job titles from Glassdoor to OCCSOC occupation codes, company names to those in Burning Glass and LinkedIn, and job locations to commuting zone (CZ) codes.

A.2 Data Processing

Step 1: Standardizing Employer’s Names across Three Sources. Our analysis draws on firm-level information from Burning Glass Technologies (BGT), LinkedIn, and Glassdoor. Since our empirical analysis relies heavily on task variables from BGT, we standardize firm names across all three sources and focus primarily on the set of firms that appear in both LinkedIn and the BGT database. The original BGT records are at the firm-occupation-location-time level and are described in detail in Section A2. Details on the standardization and matching procedures are provided below.

1. **BGT firm list.** We retrieve the original firm names from BGT, which include 462,295 unique, non-standardized entries. Since the names in BGT are extracted from job postings and often vary in how company names are displayed, multiple records may correspond to the same firm.

Following the standardization method used in [Hall, Jaffe and Trajtenberg \(2001\)](#), we employ a series of cleaning routines for organization names and create two versions of cleaned firm names. The “standard name” retains basic firm-related words, such as “GROUP” or “INC”, while the “stem

name" is a shorter version that contains only the company's name that has been simplified to its most basic form.³⁵ Examples of the standardized results are provided in Table A.1. We use the BGT "stem name" as the tracking benchmark to match firm names across BGT, LinkedIn, and other databases. Firm names are cleaned and standardized using source-specific methods, with different matching techniques applied to align them to the BGT "stem name".

Table A.1: Examples of BGT Name Standardization

BGT name	Standard name	Stem name
Belimo	BELIMO	BELIMO
Belimo Air Controls	BELIMO AIR CONTROLS	BELIMO AIR CONTROLS
Belimo Air Controls Incorporated	BELIMO AIR CONTROLS INC	BELIMO AIR CONTROLS
Belimo Aircontrols Usa Incorporated	BELIMO AIRCONTROLS USA INC	BELIMO AIRCONTROLS USA
Belimo Americas	BELIMO AMERICAS	BELIMO AMERICAS
Tesla	TESLA	TESLA
Tesla Gigafactory	TESLA GIGAFACTORY	TESLA GIGAFACTORY
Tesla Incorporated	TESLA INC	TESLA
Tesla Motors	TESLA MOTORS	TESLA MOTORS

Notes: This table provides examples of how BGT names are standardized and stemmed.

After standardization, some firms may share the same standard name or stem name. For example, there are 455,238 unique standard names and 423,149 unique stem names out of 462,295 original BGT firm names. As shown in the table above, "TESLA", "TESLA MOTORS", and "TESLA GIGAFACTORY" essentially refer to the same firm, which we collapse into a single entity in the fuzzy matching procedure.

2. **Matching with LinkedIn firm list.** To align with other data sources and our research design, we filter a list of LinkedIn users who had active job experience in 2018. We then extract all job records for those users across all time periods. This approach yields a sample of 196,167,307 job records from 5,899,025 U.S. firm names, each assigned a unique ID by the data vendor. As the data are self-reported, multiple records may correspond to the same firm. To reconcile firm names with those in the BGT dataset for research purposes, we perform a similar standardization process to obtain both standardized and stemmed firm names. As before, the "stem name" is used to match firms across the BGT and LinkedIn datasets. Examples are provided in Table A.2.

Table A.2: Examples of LinkedIn Name Standardization

LinkedIn name	Standard name	Stem name
BJ TERRONI CO	BJ TERRONI CO	BJ TERRONI
B.J. Terroni Co., Inc.	BJ TERRONI CO INC	BJ TERRONI
B.J. Terroni Company, Inc	BJ TERRONI CO INC	BJ TERRONI
BJ Terroni Company, Inc	BJ TERRONI CO INC	BJ TERRONI
JPMorgan Chase	JP MORGAN CHASE	JP MORGAN CHASE
JP Morgan Chase	JP MORGAN CHASE	JP MORGAN CHASE
J.P. Morgan Chase	JP MORGAN CHASE	JP MORGAN CHASE
JPMorgan Chase Bank	JP MORGAN CHASE BANK	JP MORGAN CHASE BANK
JPMorgan Chase Bank, N.A.	JP MORGAN CHASE BANK NA	JP MORGAN CHASE BANK NA

Notes: This table provides examples of how company names found on LinkedIn are standardized and stemmed for analysis.

Again, after standardization, some firms may share the same standard name or stem name. For example, there are 5,699,863 unique standard names out of 5,899,025 unique original LinkedIn US firm names. As shown in the table above, "JP MORGAN CHASE", "JP MORGAN CHASE

³⁵The algorithm used can be downloaded [here](#).

BANK”, and “JP MORGAN CHASE BANK NA” essentially refer to the same firm, which we collapse into a single entity in the fuzzy matching procedure. The number of records in LinkedIn firms is larger than that in the BGT firm list. We use the shorter BGT firm list as the master file and perform a left join with the matched LinkedIn firms. To this end, we carry out a three-step matching procedure as follows.

- Exact matching using standard names. Initially, records are compared using the standard names from both lists. If the standard names are identical, the records are considered a match.
- Exact matching using stem names. For records not yet matched, the next step uses "stem names." Similar to the first step, but using a potentially simplified or base form of the names (e.g., removing suffixes or prefixes). Matches are made when stem names match exactly between the two lists.
- Fuzzy matching. For records that remain unmatched after exact matching attempts, a fuzzy matching technique is applied. This process is restricted to name pairs where the first six characters are identical. Fuzzy matching algorithms generate similarity scores based on standard names and based on stem names. Records are retained if both generated similarity scores are above 0.85.³⁶

Through the procedure described above, we identify 51,363 disambiguated firm names that link LinkedIn records with BGT data. These disambiguated employers are associated with 75,037,667 jobs (firm–occupation–location), accounting for 38.3% of the original sample. Non-matched records can be attributed to three main factors: (1) we currently consider only LinkedIn positions from 2018, and therefore the LinkedIn sample does not fully overlap with the coverage of BGT; (2) the matching scheme may not capture all potential firm matches, as we adopt a conservative approach that prioritizes the accuracy of matched records; and (3) all firm names with stemmed name lengths shorter than three characters are excluded from the analysis. Among all matched records, 92% are based on exact matches using either standardized or stemmed names, while the remaining 8% rely on fuzzy matching. Examples of matched records are provided in Table A.3.

Table A.3: Examples of BGT and LinkedIn Name Matching Methods

Matching method	BGT name	LinkedIn name
Exact standard/stem name	Blue Ridge Sales	blue ridge sales inc
Exact standard/stem name	Blue Rose Consulting Llc	blue rose consulting group, inc.
Exact standard/stem name	Blue Sky Property	blue sky property group
Fuzzy matching	Blue Streak Reprographic	blue streak reprographics
Fuzzy matching	Blue Water Automotive Systems	blue water automotive, inc
Fuzzy matching	Blue Willow Counseling	blue willow counseling ctr

Notes: This table illustrates examples of exact and fuzzy matching between BGT and LinkedIn company names.

3. **Matching with Glassdoor firm list.** The original Glassdoor sample consists of 5,448,727 wage records at the employer-title-location level and an additional 5,048,711 records at the employer-title level. These wage records are associated with 183,267 unique employer names. Using the

³⁶We use Stata function `matchit` to create the similarity scores. We experiment with alternative methods using different lengths of the initial characters as a starting point and different score cutoffs. The current approach yields similar results to other optimized combinations.

same matching procedures described above, we disambiguate these names and match them with firms in the BGT dataset. Specifically, we apply the standardization algorithm, match based on standardized and stemmed names, and employ the same fuzzy matching approach. As a result, we identify 12,175 disambiguated firm names that can be matched with our BGT sample. These matched employers account for 66.8% of the wage records in the Glassdoor dataset. Examples of matched records are provided in Table A.4.

Table A.4: Examples of BGT and Glassdoor Name Matching Methods

Matching method	BGT name	Glassdoor name
Exact standard/stem name	Bahwan Cybertek, Inc	Bahwan CyberTek
Exact standard/stem name	Bain Company Incorporated	Bain & Company
Exact standard/stem name	Worldwide Flight Services Incorporated	Worldwide Flight Services
Fuzzy matching	Zeeland Lumber Supply Company	Zeeland Lumber & Supply
Fuzzy matching	Marmic Fire Safety Company Incorporated	Marmic Fire & Safety
Fuzzy matching	Hard Rock Hotels And Casinos	Hard Rock Hotel & Casino

Notes: This table provides examples showing how company names from BGT are matched with Glassdoor names using exact and fuzzy matching methods.

Step 2: Standardizing OCCSOC Codes of Job Titles for Three Sources. Our study requires occupation information that can be linked to standard OCCSOC codes. The occupation titles associated with these codes are obtained from [IPUMS USA](#). In total, 487 distinct occupations are identified. We summarize how this occupation information is utilized across the three main datasets used in our analysis.

1. **LinkedIn job titles.** The LinkedIn data provider, Revelio Labs, assigns an OCCSOC code to each job title using an internal text-based algorithm. In the 2018 LinkedIn dataset, 385 distinct OCCSOC codes are identified. Occupations not captured by this process are primarily labor-intensive roles, which are typically underrepresented in LinkedIn data.
2. **Burning job titles.** The job posting data from Burning Glass Technologies (BGT) includes OCCSOC codes directly, which have been validated and used in the literature (e.g., [Braxton and Taska \(2023\)](#)).
3. **Glassdoor job titles.** Job titles in the Glassdoor wage data are highly non-standard and noisy. From the full sample, we observe 572,691 distinct job title descriptions. To standardize these, we utilize ChatGPT-4o to match each title to the closest OCCSOC code, resulting in 563,294 successful matches(match rate=98.4%). A summary of the results is provided in Table A.5.

Table A.5: Examples of Job Title Matching Between Glassdoor and OCCSOC

Glassdoor Job Title	Matched OCCSOC Code	Matched OCCSOC Title
Assistant Vice President Consultant Risk Tech	15-1199.09	Risk Management Specialists
Senior Systems Data Analyst	15-2041.01	Data Scientists
Associate Audiovisual Technician	27-4011.00	Audio and Video Technicians
Tax Servicing Specialist	13-2082.00	Tax Preparers
Debt Consolidation	13-2072.00	Loan Officers
Financial Controller	11-3031.00	Financial Managers

Notes: This table shows examples of how job titles from Glassdoor are matched to OCCSOC codes and standardized occupational titles.

Step 3: Standardizing LinkedIn University Name and Matching to World Ranking and IPEDS.

Our analysis requires university ranking and location. We perform the following procedures to obtain this information.

1. **LinkedIn university names.** LinkedIn provides user-generated education information, resulting in high variability. This includes differences in the type of degree, field of study, program details (if provided), and the university name. As shown in Table A.6, the self-reported university names are particularly noisy, as they often include various name variants, abbreviations, department names, campus or school designations, and entries in multiple languages. For example, column (1) shows how Harvard University is self-reported by LinkedIn users, and column (2) shows the example for Purdue University.

Table A.6: Examples of University Name Variations

Example 1 (Harvard)	Example 2 (Purdue)
Harvard	Purdue University
Harvard University	Purdue Global University
Harvard Law School	Purdue North Central
Harvard University Extension School	Purdue School of Engineering and Technology IUPUI
Harvard College	Purdue U Indiana U
Harvard Business School	Purdue College of Technology Columbus
HBS	Purdue University Calumet
Harvard University Kennedy School of Government	Purdue University Daniels School of Business
John F. Kennedy School of Government	Purdue University Global
Harvard T.H. Chan School of Public Health	Purdue University Krannert School of Management

Notes: This table presents examples of how the same institution can appear under different names, illustrating the need for name normalization in university data.

We filter degrees and associated university names for analysis using the following criteria:

- In our main sample, we first consider a list of users with a U.S.-based job and extract all education information for those users, obtaining 79,407,923 degree records.
- We retain degrees at the bachelor’s level and above—specifically, bachelor’s, master’s, and doctoral degrees.
- We keep only records associated with U.S. universities. Since the country information is missing for many institutions, we infer the university’s country based on the most frequent job location of its graduates.
- We drop records with a missing degree start year.
- We conduct preliminary cleaning by removing special symbols from university names.

The above procedure results in 31,696,635 degree records associated with 426,388 unique university name entries, which remain highly noisy.

2. **University rankings and a list of universities to consider.** The noisy nature of LinkedIn university names hinders our ability to disambiguate institutions and match them with external datasets. To address these challenges, we construct a list of major U.S. universities to restrict our sample to a meaningful subset of records for further data cleaning. This list is based on prominent university rankings, which naturally provide the ranking needed for additional analysis. Specifically, we combine the following three sources:

- **World University Rankings (WUR)**: WUR provides annual rankings for 2,000 global universities, among which we consider 358 universities located in the United States. For our study period, we use the 2019 edition of the WUR rankings.
- **US News³⁷**: The U.S. News ranking includes 2,145 universities, but only the Top 1,000 institutions are assigned a specific rank; universities beyond the Top 1,000 are unranked. Among these, 284 are U.S.-based universities, of which 197 are ranked.
- Universities considered in [Chetty et al. \(2020\)](#), including over 2000 US universities.

There is substantial overlap across the three university lists, although institution names often appear in varied forms. We merge the lists and standardize university names before matching them to the LinkedIn database, ultimately identifying 2,300 distinct U.S. universities. The composition of the final university list used in our analysis is summarized in Table A.7. Notably, 235 universities (10.2%) appear in all three lists, while the list from [Chetty et al. \(2020\)](#) alone covers 1,925 universities (83.7%).

Table A.7: Source of the University List

Sources	Freq	Percent
Chetty et al. (2020)	1,925	83.70
US News	16	0.70
US News, WUR	94	4.09
US News, WUR, Chetty et al. (2020)	235	10.22
WUR	30	1.30
Total	2,300	100.00

Notes: This table summarizes the sources used to determine universities analyzed in our study. WUR refers to the World University Rankings.

3. **Matching LinkedIn university names.** We implement a multi-step matching procedure between LinkedIn university entries and the standardized school list. This process begins with 426,388 unique university name entries from LinkedIn, as filtered in the preceding stage.

- **Exact and fuzzy matching.** To obtain matched records, we experiment with both exact and fuzzy matching approaches in an iterative fashion. We begin by matching records with exact character strings. We then compute similarity scores using a fuzzy matching algorithm.³⁸ Several heuristics prove useful in identifying matches. For example, when the first 12 characters of university names are identical and the similarity score exceeds 0.8, the records are generally correctly matched.
- **First-30-character matching.** University names with the first 30 characters identical are also found to match with high probability.
- **Manual verification.** We manually verify unmatched records where university names appear more than 2,000 times in the entire database.
- **ChatGPT-assisted matching.** As explained above, raw university names from LinkedIn often include abbreviations, department names, school names, program names, or even multiple languages. We utilize ChatGPT to assist in identifying the correct matched institutions in such cases.

³⁷We retrieve the ranking information from [US News](#) as of 3/13/2023.

³⁸We again use the Stata function `matchit` to generate similarity scores.

Through the matching procedure described above, we successfully linked 20,543 raw LinkedIn university name entries to our standardized list of U.S. universities, corresponding to 2,170 major institutions. These matched universities are associated with 26,691,378 degrees held by LinkedIn users working in 2018, representing 84.2% of the original filtered dataset (31,696,635 records).³⁹ A summary of the matched records is presented in Table A.8.

Table A.8: LinkedIn University Matching Methods

Match Method	Raw LinkedIn Universities Matched	Degrees Matched
Exact and fuzzy	2,684	8,639,110
First-30-character	11,320	50,746
Manual	119	742,885
ChatGPT-assisted	6,420	17,258,637
Total	20,543	26,691,378

Notes: This table provides the number of universities and degrees matched using various methods, including traditional exact/fuzzy matching, manual review, and ChatGPT-assisted matching.

Step 4: Standardizing the Geographic Information of Jobs and Universities. For our research purposes, we require city-level geographic information for the job locations listed in LinkedIn and Glassdoor, as well as for users' graduating universities.

1. **Job location in LinkedIn.** As explained above, we begin by extracting LinkedIn users with an active job position in 2018, and then retrieve all available job history associated with those users. The resulting dataset includes 3.89 million unique job location records. However, the data is noisy, as it includes geographic information at varying levels of granularity—such as state, county, city, and street levels. For our research purposes, we focus on extracting U.S. city-level information from the raw address field.
 - **Extract Cities.** We start with a list of 31,254 city names⁴⁰. We search the job location addresses using city names as keywords and require the state information to be consistent. We also manually correct various issues, such as abbreviations (e.g., "NYC" for New York City) and ambiguous city names that appear in multiple states. After cleaning and validation, we successfully identify 2,998,495 records with reliable city-level information. The remaining records are manually reviewed, found to be inaccurate or incomplete, and subsequently excluded. Examples are provided below in Table A.9.
 - **Match with CZ Codes.** The geographic coordinates of the extracted cities are obtained using the Google Geocoding API and are mapped to Commuting Zone (CZ) boundaries (based on the 1990 definition)⁴¹, as well as to Metropolitan and Micropolitan Statistical Areas. Among the 75,037,667 job positions matched with firm names from the Burning Glass database, 68,663,767 (91.5%) are successfully associated with a city name and thus assigned with a CZ code.

³⁹The 26,691,378 LinkedIn users include individuals who graduated in any year, regardless of the employer or occupation reported.

⁴⁰Source: [SimpleMaps](#).

⁴¹Source: [The Health Inequality Project](#).

Table A.9: Examples of Location Extraction and CZ Code Inference

Raw Location in LinkedIn	State	Extracted City	Inferred CZ Code
45 W. 111th Street, Chicago, IL 60627	Illinois	Chicago	16001
1356 Bellevue St. Green Bay, WI	Wisconsin	Green Bay	55004
5500 Cloverleaf Pkwy Cleveland	Ohio	Cleveland	39002
Shiawassee Michigan area.	Michigan	–	–
United States, MI, Orion Township	Michigan	–	–
Unit Number 4, DDA Local Shopping Centre, Hemkunt Co	Colorado	–	–
Central, Nebraska	Nebraska	–	–

Notes: This table shows examples of how raw LinkedIn location strings are parsed to extract U.S. states, cities, and commuting zone (CZ) codes.

2. **Job location in Glassdoor.** We also require city-level information for the Glassdoor wage records. Among 3,952,145 wage records with location information linked to employers matched with Burning Glass firms, there are 24,246 unique location descriptions. We apply the same matching approach described above. After this process, 8,913 records remain unmatched. Given that Glassdoor location data is relatively clean, we use the Google Geocoding API to extract precise geographic coordinates for these remaining records. In the final sample, we successfully obtained city and CZ information for 3,950,469 out of 3,952,145 wage records. Some examples are provided in Table A.10.

Table A.10: Examples of Location Extraction from Glassdoor and CZ Code Inference

Raw Location in Glassdoor	State	Extracted City	Inferred CZ Code
Austin, TX	Texas	Austin	31201
Bala Cynwyd, PA	Pennsylvania	Bala Cynwyd	19700
La Jolla, CA	California	La Jolla	38000
Whippany, NJ	New Jersey	Whippany	19600

Notes: This table shows how structured location data from Glassdoor job postings is parsed and linked to commuting zone (CZ) codes.

3. **University location in LinkedIn.** For each university, we extract geographic coordinates using the Google Geocoding API and map them to cities, Commuting Zone (CZ) boundaries, Metropolitan and Micropolitan Statistical Areas.

The data we processed includes college graduates from all U.S. institutions across all graduating years, covering 13.7 million US employees. For our analysis, we restrict the sample to LinkedIn users who: (1) received a bachelor’s degree (as their highest degree) between 2016 and 2018 from a U.S. institution ranked in the Top 2000 of the WUR; (2) were employed in the United States in 2018; (3) have employer names and occupational titles that can be clearly identified and matched to the BGT data; and (4) have identifiable employer geographic locations. To improve estimation precision, we further restrict the sample to U.S. universities with at least 100 LinkedIn users.

We also restrict the analysis to 266 U.S. universities ranked in the WUR.⁴² Our final sample covers 266 universities, 25492 distinct firms, and a total of 244,632 LinkedIn users.⁴³

⁴²WUR ranks the top 2,000 universities globally, of which 348 are U.S. universities.

⁴³Throughout the paper, we define a firm as the combination of a company name and the location of its establishment. A firm enters our sample if at least one LinkedIn user is observed as employed there.

B Validating the Sample

Since BGT data has been widely used and validated in previous studies (Hershbein and Kahn, 2018, Atalay et al., 2024), we focus on validating our LinkedIn sample in various ways.

B.1 Validating the LinkedIn Sample

We first assess the spatial representativeness of the LinkedIn data by comparing it with the ACS. For both datasets, we restrict attention to workers whose highest degree is a college degree and who were employed in 2018. The left panel of Figure B.1 plots each commuting zone's share of national college graduate employment using LinkedIn data (y-axis) and ACS data (x-axis). The dashed line represents the 45-degree line. Most commuting zones lie close to this line, indicating a high degree of similarity in employment shares between the two datasets. Nonetheless, LinkedIn users appear to be overrepresented in a few major cities, such as New York, San Francisco, and Seattle, and underrepresented in others, such as Los Angeles, Newark. Across U.S. commuting zones, the correlation between city sizes implied by the two datasets is high: the correlation equals 0.95, and the OLS regression slope equals 1.12 (s.e. = 0.01).

We next examine the occupational representativeness of the LinkedIn data. The right panel of Figure B.1 plots the shares of employment across two-digit SOC occupations using LinkedIn data (y-axis) and ACS data (x-axis). Again, most points lie close to the 45-degree line. Perhaps as expected, LinkedIn users are disproportionately represented in high-skilled, business- and technology-oriented occupations such as business and finance, and computer and mathematics, while they are underrepresented in service and administrative occupations such as education, sales, and office administration. Once more, we find a strong correlation between the two datasets: the correlation equals 0.92, and the OLS regression slope equals 0.98 (s.e. = 0.09).

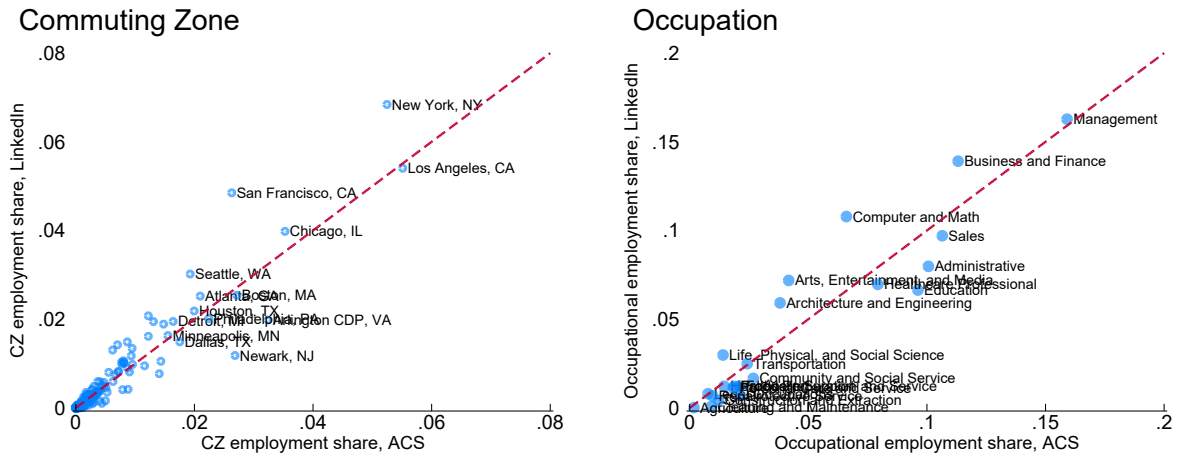


Figure B.1: Comparison of LinkedIn and IPUMS-ACS Data: Commuting Zones (left) and Occupations (right)

Our third validation exercise examines the extent to which LinkedIn data represent graduating class sizes across U.S. universities. Specifically, we compare the share of graduates from each university in

LinkedIn with that reported in the Integrated Postsecondary Education Data System (IPEDS), published by the National Center for Education Statistics. For both data, we use LinkedIn users who graduated between 2016 and 2018. Figure B.2 plots the national share of graduates by university, with LinkedIn data on the y-axis and IPEDS data on the x-axis. For comparability, shares are expressed as percentages. We find a strong positive relationship between the two measures. An Ordinary Least Squares (OLS) regression yields a coefficient of 1.76 (standard error = 0.038), with an R^2 of 0.73.

These results indicate that LinkedIn data are, to a considerable extent, representative of the U.S. college-educated labor force and graduating class sizes.

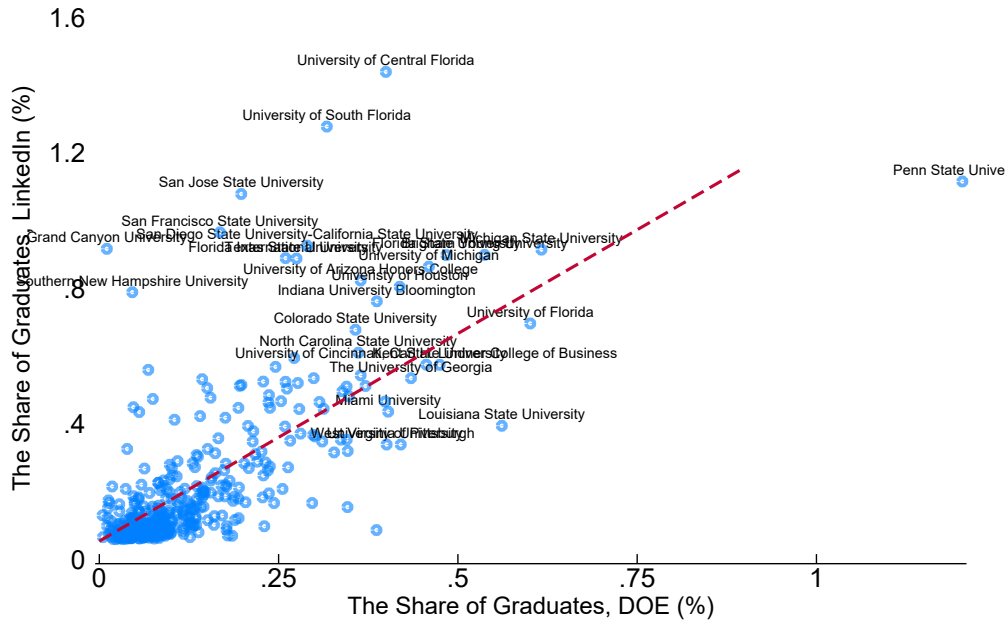


Figure B.2: Comparison of College Class Sizes: LinkedIn vs. IPEDS

B.2 Validating the Glassdoor Wages

Glassdoor wage data are available at the firm, city, and occupation levels. We validate these data in three ways. First, we compute the average Glassdoor wage by commuting zone and compare it with corresponding estimates from the ACS. In the ACS, wages are calculated for college graduates who are full-time workers (defined as working more than 35 hours per week and more than 40 weeks per year).

The left panel of Figure B.3 compares the average annual salaries from the two sources across 722 U.S. CZs. The dashed line represents the 45-degree line. The two measures are strongly correlated: a simple regression of Glassdoor wages on ACS wages yields a coefficient of 0.71 and an R^2 of 0.23. The right panel of Figure B.3 compares average annual salaries across occupations. Using 22 two-digit OCCSOC occupations, a simple regression produces a coefficient of 1.03 and an R^2 of 0.84.

Finally, Figure B.4 compares average annual salaries across CZ–occupation pairs, where a simple regression yields a coefficient of 0.87 and an R^2 of 0.55. Taken together, the evidence indicates that Glassdoor wage data captures much of the variation in wages across cities and occupations.

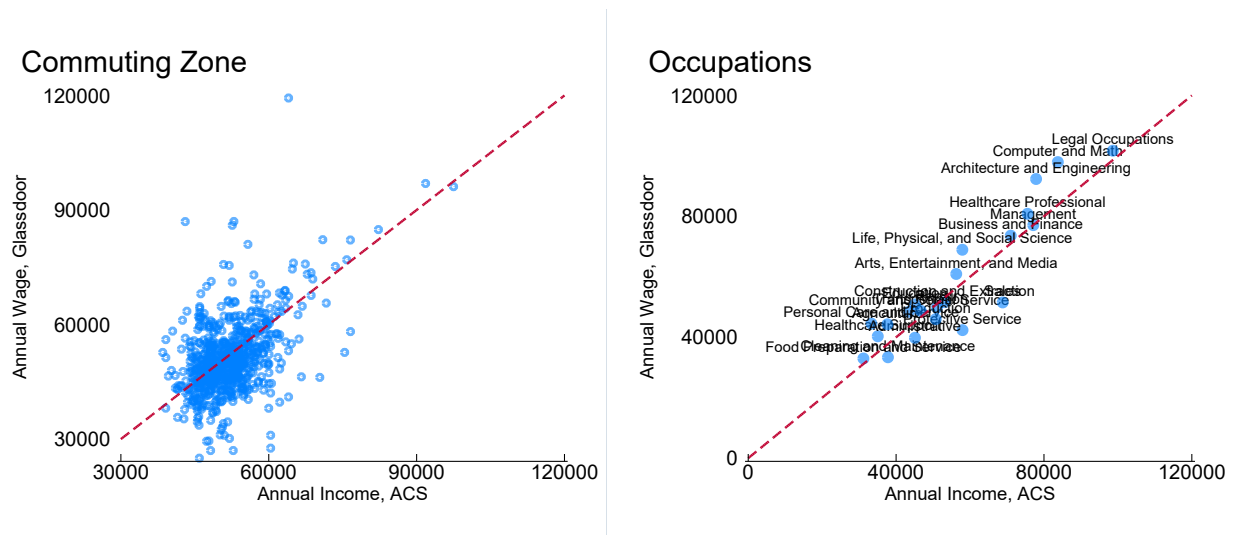


Figure B.3: Annual Wage by Commuting Zones (left) and Occupations (right): Glassdoor vs. ACS

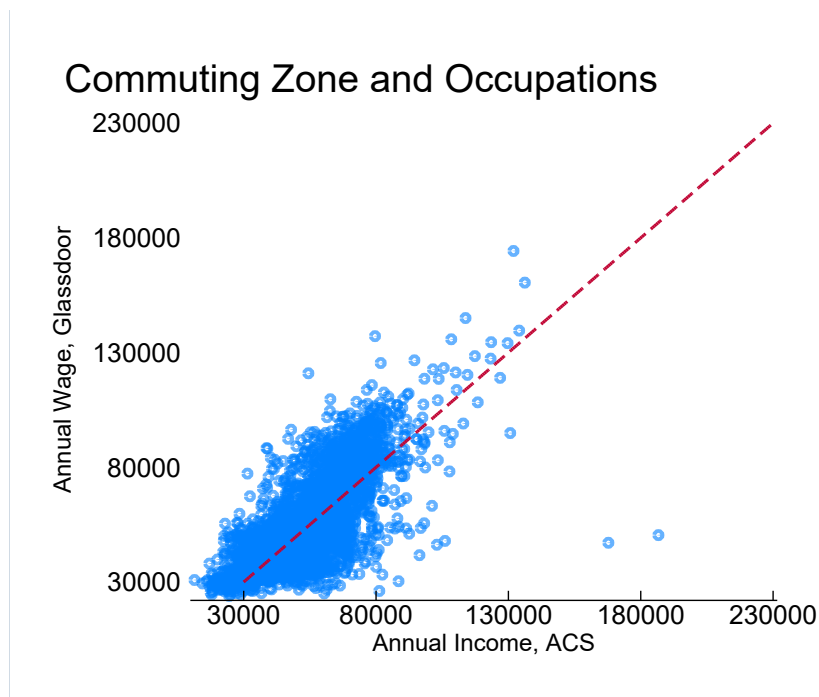


Figure B.4: Annual Wage by Commuting Zone-Occupation Pairs: Glassdoor vs. ACS

C Tables and Figures

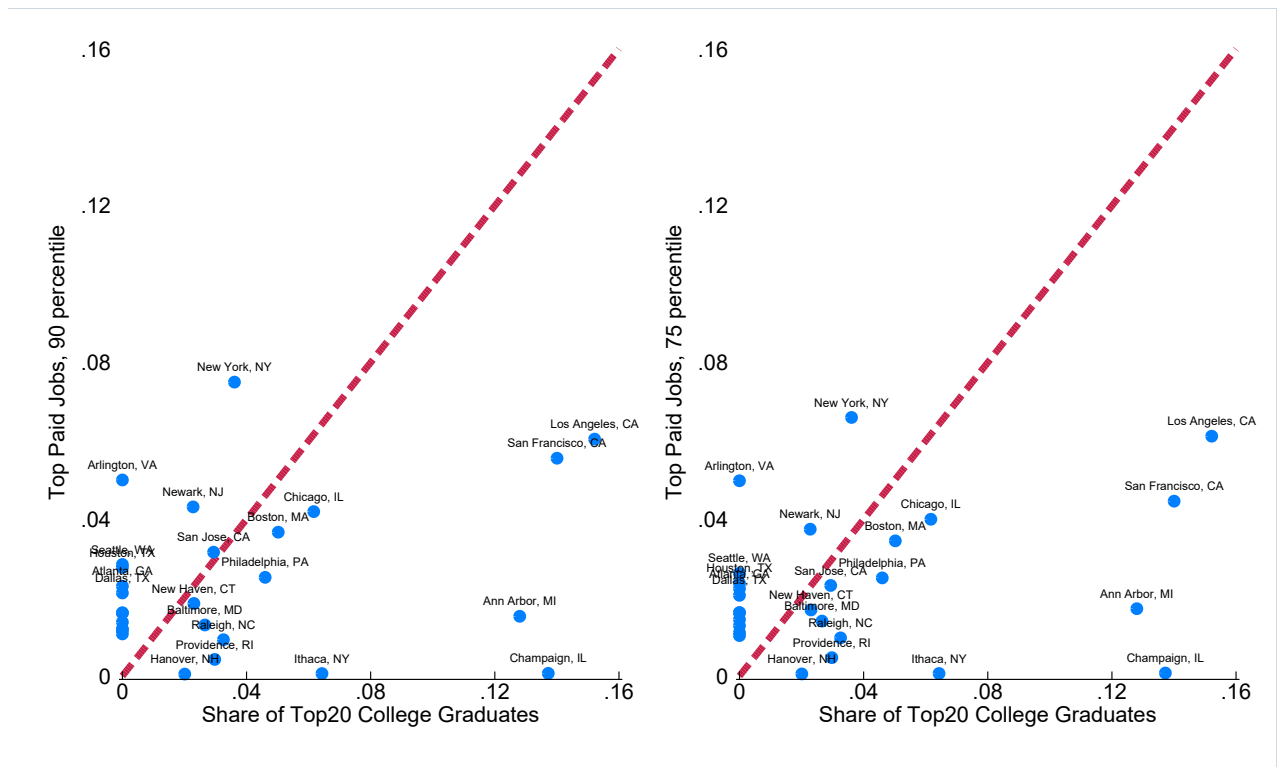


Figure C.1: Spatial Distribution of Top 20 university Graduates and Jobs: 90th Percentile (Left) and 75th Percentile (Right) as High-Paying Jobs

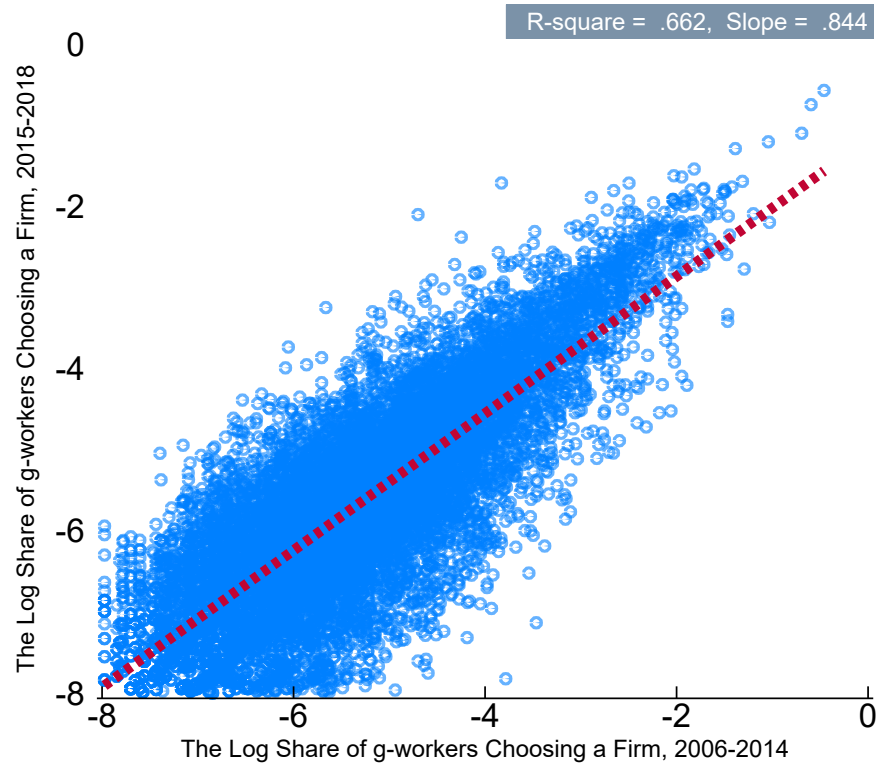


Figure C.2: The Log Share of g-worker Choosing a Firm, Recent Graduates (2016-2018) versus Former Graduates (2005-2014)

This figure plots the log share of graduates from a given college choosing a firm, for fresh graduates (2016-2018) against early cohorts who graduated before 2014. The dashed line plots the linear fit, which shows a slope coefficient of 0.844 (s.e. = 0.006) and an R^2 of 0.662 .

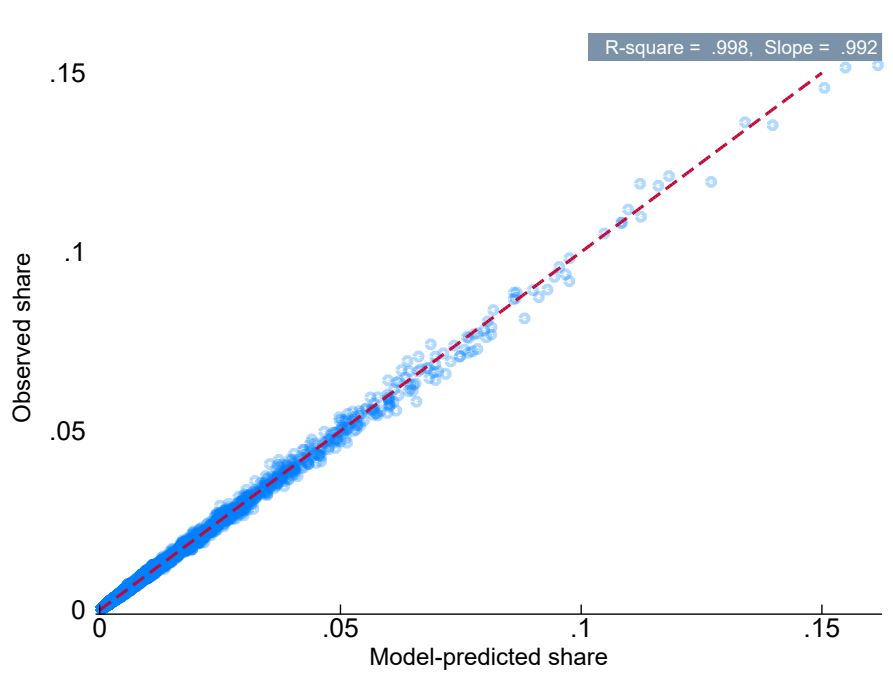


Figure C.3: The University-to-Job Allocation Shares, Data vs. Model-predicted Shares

This figure plots the log share of graduates from a given college choosing a firm, for fresh graduates (2016-2018), against early cohorts who graduated before 2014. The dashed line plots the linear fit, which shows a slope coefficient of 0.844 (s.e. = 0.006) and an R^2 of 0.662 .

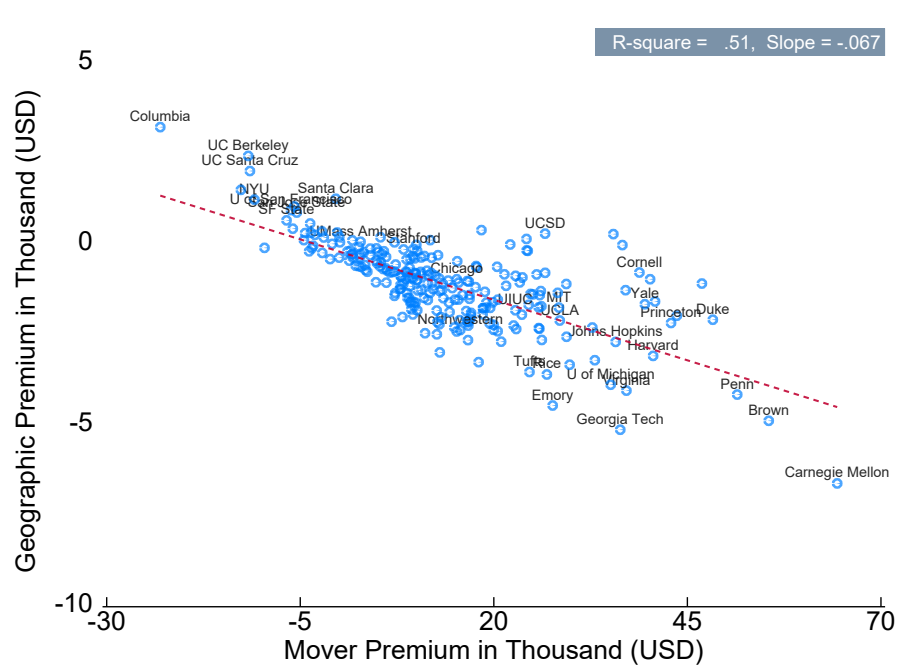


Figure C.4: Geographic Premium (y-axis) Against Mover Premium (x-axis)

This Figure plots the geographic premium (y-axis) against the mover premium (x-axis) across cities. We measure the mover premium as the difference in average wages between movers and stayers among fresh graduates (2016–2018), directly estimated from our sample. A negative mover premium implies that, on average, local stayers earn higher wages than those who migrate.

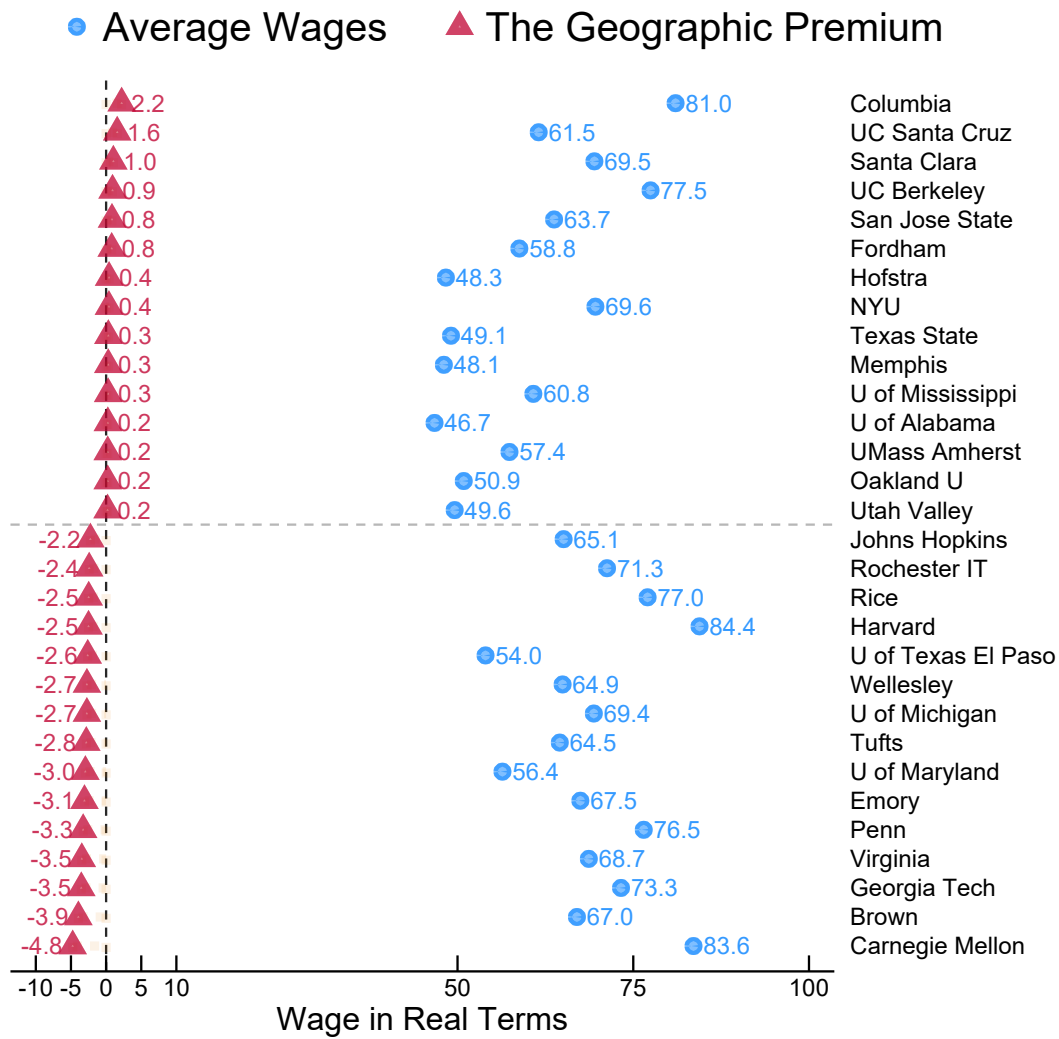


Figure C.5: Average Wages and Geographic Premium in Real Terms

Figure C.5 reports results for 30 universities: the fifteen with the highest premiums and the fifteen with the lowest (most negative) premiums in real terms. We select these 30 universities from those with an average annual salary of at least \$65K.

Table C.1: Admission Criteria by College Tiers

College Tier	Minimum SAT Scores	Minimum ACT Scores	Rankings in High school
Most selective	655	29	Top 10-20%
Highly selective	620	27	Top 20-35%
Very selective	573	24	Top 35-50%
Selective	500	21	Top 50-65%
Less selective	below 500	below 21	Top 65%
Nonselective	None	None	None

Notes: This table shows the SAT/ACT scores and the ranking in high school transcripts or class rank that are typically required by college admission. The information is based on Barron's Profiles of American universities ([Barron's Educational Series, 2017](#)).