

Writing Style as Cultural Change

Camilo García-Jimeno and Sahar Parsa

REVISED
July 28, 2025

WP 2024-23

<https://doi.org/10.21033/wp-2024-23>



FEDERAL RESERVE BANK *of* CHICAGO

*Working papers are not edited, and all opinions are the responsibility of the author(s). The views expressed do not necessarily reflect the views of the Federal Reserve Bank of Chicago or the Federal Reserve System.

Writing Style as Cultural Change*

Camilo García-Jimeno^{†1} and Sahar Parsa^{‡2}

¹Federal Reserve Bank of Chicago

²New York University

July 28, 2025

Abstract

We study how network structure and peer influence shape cultural change by tracking writing styles in economic theory from 1970 to 2019. We estimate a discrete-choice model where individual preferences, peer effects, and co-author bargaining drive gendered pronoun choice, leveraging variation in each author’s feasible co-author network as a source of exclusion restrictions. Estimation reveals a profession of conformists with strong peer influence: when an author’s peers shift from 20 percent to 70 percent feminine-only pronoun use, the author’s odds of adopting the feminine form more than double. Counterfactual simulations show that absent external societal trends, the early masculine norm would have persisted; however, once those pressures appeared, peer influence magnified their impact boosting long-run stylistic diversity. We also find that homophily in co-authorship sustained writing style diversity by allowing minority preferences to express freely. Demographic shifts — entry of women and new cohorts — did not initiate cultural change but accelerated it once under way by amplifying peer effects. Cultural change depended less on demographic turnover than on how newcomers rewired the network structure, turning peer influence from an initial drag into the engine of norm evolution.

Keywords: cultural change, language, gender, social norms, co-authorships, social networks

JEL Codes: D71, D83, D85, J16, Z1

*We especially thank Doris Pan, whose research assistance has been extraordinary. We also thank Ryan Perry, Liam Puknys, Aryan Safi, and Fangzhou Xie for their help at various stages, and Frank DiTraglia, Diego Jimenez-Hernandez, Martin Osborne, Debraj Ray, Martin Rotemberg, Ariel Rubinstein, Judit Temesvary, and participants at the UPenn micro seminar, the Northwestern Kellogg Political Economy seminar, the University of Glasgow applied micro seminar, the University of Calgary 2024 Empirical Microeconomics Workshop, the 2024 Petralia Political Economy Conference, the 2024 Fed System Econometrics Conference, the SEA Mini-Conference 2024, and the 2025 Stanford Network Economics Conference for their suggestions. Finally, we thank Professor Yang Feng for sharing his code with us. The views expressed in this article are those of the authors and do not necessarily reflect the position of the Federal Reserve Bank of Chicago or of the Federal Reserve System.

[†]camilo.garcia-jimeno@chi.frb.org, 230 S. LaSalle Street, Chicago, IL.

[‡]sahar.parsa@nyu.edu, 710 19th West 4th street, New York, NY 10012.

1 Introduction

Cultural change rarely proceeds in straight lines. Languages, dress codes, religious rituals, or shared beliefs can transform overnight in some societies yet persist for generations in others. What produces this mosaic of rapid shifts and stubborn continuity? While prior work highlights the role of peer influence and network structure in shaping the diffusion of norms, behaviors, and beliefs, empirical evidence tracing these dynamics remains limited. This paper offers micro-level evidence of how social ties within a professional network shape the process of cultural diffusion.

We study how cultural change takes hold and spreads by analyzing gendered pronoun use in academic economic theory papers from 1970 to 2019. This setting offers two advantages. First, gender typically plays no substantive role in the content, allowing authors discretion in assigning gender to abstract agents in formal models.¹ Because the content does not dictate gender, pronoun choice by economic theorists offers a unique measure of cultural expression capturing both individual beliefs and preferences (Baron, 1986) and social influences such as professional norms and expectations. Second, we construct a rich 50-year panel dataset of authors, their publications, and professional networks in a field cohesive enough for many to know each other directly, yet large enough to support indirect ties. This setting helps us identify peer influence through authors who shape others' behavior only indirectly.

The writing style of economics papers has undergone a steady transformation over the past 50 years (see Figure 1). In 1970, 80 percent of economic theory papers used exclusively masculine third-person pronouns (*he*, *him*, *himself*) for generic agents. By 2019, that share had dropped to 20 percent. Meanwhile, alternative pronoun forms — exclusively feminine (*she*, *her*, *herself*), exclusively plural (*they*, *them*, *themselves*), or mixed (*combining multiple forms*)² — emerged and spread at different times and rates. Plural and mixed forms began rising in the mid-1970s, while the feminine form gained traction around the 1990s. By 2019, masculine, feminine, and mixed forms each accounted for roughly 20 percent of published theory papers, with plural forms re-surging in the 2010s.

We model the individual choices that drive the cultural change in Figure 1 using a discrete choice framework with three novel components. First, we allow each author's preferences to depend on two sources: an idiosyncratic component constant across all their publications and a time-varying social component that reflects the behavior of their professional network.

¹For example, in principal-agent models, authors often use masculine pronouns for principals and feminine pronouns for agents. Stevenson and Zlotnick (2018) document similar patterns in economics textbooks, where gender is arbitrarily assigned to fictional characters. This contrasts with other fields where the topic often determines pronoun usage.

²This includes grammatical forms such as *he/she* or using different pronoun forms for different antecedents.

We construct this social component as the product of (i) the distribution of pronoun choices among an author’s past co-authors and cited authors, and (ii) a time-invariant, author-specific coefficient that captures their responsiveness to peer behavior. Second, this peer effect coefficient can be positive or negative, allowing for *conformists*, who move toward their peers’ past choices, and *contrarians*, who move away from them. Third, we model decisions in co-authored papers as the result of an implicit bargain between co-authors, with bargaining weights based on pairwise characteristics — such as differences in seniority, citation counts, or productivity — that reflect their relative influence. Coauthored papers constitute about 60 percent of our sample. By incorporating co-authorship explicitly, we avoid mis-attributing changes in pronoun use to an individual author when they may reflect a coauthor’s preference.

A well-known debate between Martin Osborne and Ariel Rubinstein in the preface to their game theory textbook ([Osborne and Rubinstein, 1994](#)) illustrates these components and shows how co-authorship can involve compromise between opposing preferences. Rubinstein advocated for the masculine pronoun, *he*, which he considered neutral, arguing that alternatives would be distracting: “...in academic material it is not useful to wave [language] as a flag.” Osborne disagreed: “...no language is ‘neutral’... The use of ‘he’ to refer to a generic individual... has its origins in sexist attitudes and promotes such attitudes... My preference is to use ‘she’ for all individuals.”

Estimating how professional peers influence writing style presents two empirical challenges. First, unobserved time-invariant author traits — such as ideological orientation or stylistic preferences — likely correlate with peers’ behavior. Although we observe many authors multiple times, we cannot difference out these traits in our nonlinear model. Unaccounted for, these latent effects act as nuisance parameters and bias our peer influence estimates. To address this challenge, we proxy for each author’s unobserved, time-invariant preference by assigning them to a latent community based on global co-authorship patterns.

Collaboration patterns in academia exhibit well-documented homophily on observable characteristics. If authors also sort on unobservable traits, then the residual structure of the co-authorship network, after accounting for observables, can reveal these latent dimensions. We formalize this idea using the stochastic block model (SBM), a canonical model in the community detection literature within network science ([Karrer and Newman, 2010](#); [Newman, 2018](#)). SBM recovers unobserved group membership — interpreted in our setting as shared latent traits — under the assumption that co-authoring probabilities partially depend on group membership. We set the number of groups to two to capture the broad, unobserved distinction between Osborne-like and Rubinstein-like orientations toward writing style. We estimate the SBM using the likelihood-based approach developed by [Feng et al. \(2023\)](#), and

adapt it to operate over each author’s acquaintance network — the set of feasible coauthors — rather than over all possible theorist pairs (see [Fafchamps et al. \(2010\)](#) for a related idea).

Many author pairs — due to differences in field or cohort — are not plausible collaborators. Including them would distort the estimation of the latent group structure. We construct the acquaintance network introducing a vector embedding method, *author2vec*, which adapts *word2vec* from natural language processing ([Mikolov et al., 2013](#)). Just as *word2vec* embeds words based on co-occurrence, *author2vec* embeds authors based on patterns of co-authorship and citation. We define the acquaintance set for each author as a subset of peers sufficiently similar in the “academic space” that spans their embeddings, based on cosine similarity. We interpret these pairs as the set of feasible collaborators.

The stochastic block model reveals homophily in co-authorship by ethnicity, gender, sub-field, age, and citation count, but not by lifetime publication output. It partitions theorists into two communities, with the smaller group — including Martin Osborne — representing 44 percent of the sample. Osborne and Rubinstein fall into different communities, mirroring their well-known disagreement on pronoun usage. For simplicity, we refer to the smaller group as the Osborne-type and to the larger group as the Rubinstein-type. These labels serve as illustrative proxies for contrasting approaches to writing style: a relatively more innovative (Osborne-type) and a relatively more traditionalist (Rubinstein-type) approach, respectively. Over time, the share of Osborne-types increased modestly from 40 percent in 1970 to 45 percent in 2019. Both communities are gender-balanced. In the model, we allow these two types to differ both in their idiosyncratic preferences and in their peer responsiveness.

A second empirical challenge arises from the possibility of time-varying shocks correlated with changes in peers’ writing styles, even after accounting for time-invariant author preferences. To identify peer effects, we adopt a control function approach. We estimate it using instrumental variables, drawing on methods from [Jochmans \(2023\)](#) and [Johnsson and Moon \(2021\)](#). Specifically, we exploit exclusion restrictions implied by the acquaintance network. We use variation in the writing style of co-authors of an author’s co-authors and citees — outside the author’s professional sphere — as instruments for peer influence. These instruments strongly predict the writing style choices of an author’s direct peers.³

Our results show that professional peer influence plays a central and quantitatively large role in shaping writing-style choices, revealing a profession of conformists. A shift in peers’ pronoun use from 20 to 70 percent feminine more than doubles the odds that an average economist adopts the feminine form. But this peer responsiveness alone does not explain

³In [section 5.1](#), we explain why a reduced-form IV strategy may fail to identify a well-defined treatment effect in a network setting with heterogeneous peer effects, motivating our more structural approach.

the observed aggregate pattern of change. At the aggregate level, simulations help us isolate the roles of peer effects and external cultural trends, which we also estimate in the model: in the absence of peer influence, societal trends alone would have displaced the masculine form in favor of the plural, producing less stylistic diversity than we observe. Conversely, in the absence of societal change, peer effects reinforce the status quo – driving the masculine form to 85 percent by 2019. Only when both forces interact do we reproduce the observed timing and diversity of pronoun use.

Peer effects are heterogeneous: women and Osborne-type economists respond more to the behavior of their peers. When the share of peers using the feminine form rises from 20 to 70 percent, the probability of adoption increases from 14 to 33 percent for a Rubinstein-type man and from 21 to 48 percent for an Osborne-type woman. Serendipitously — we do not identify community membership from writing style — the latent communities reflect expected idiosyncratic preference patterns: Osborne-types place a stronger penalty on the masculine form, echoing Martin Osborne’s challenge to the status quo. We find stronger conformity among Osborne-types and women, however, consistent with psychological evidence on gender differences in peer responsiveness ([Bond and Smith, 1996](#); [Eagly, 1983](#)). In a world dominated by the masculine form, their conformism reinforces the status quo working against their own preference. Only after external influences pushing towards more innovative forms percolate in the profession does their conformism begin amplifying the more novel writing styles.

As an implication, homophily in co-authorship does not hamper the adoption of new writing styles — instead, it helps sustain stylistic diversity. Authors with nontraditional preferences collaborate more freely with like-minded peers, enabling expression, while in mixed teams bargaining dilutes their influence. In our simulations, eliminating homophily (forcing cross-type, cross-gender co-authorship) reduces the use of feminine forms by 4 percentage points in the long run. Homophily can protect innovation by giving cultural minorities room to express their preferences, even as it limits exposure to difference.

Economists writing in the 1970s and 80s initiated the writing-style revolution; later cohorts largely imitated and amplified these changes. Although men initiated the shift away from masculine forms, women — especially Osborne-type women — accelerated its spread through greater peer responsiveness, enabled by homophily in co-authorship. This amplification is how the entry of women into the profession contributed to cultural change: not by shifting the ideological makeup of the profession, but by strengthening the network effects that accelerated the uptake of new norms. The growing share of Osborne-types had a moderate effect in the same direction, suggesting that cultural change depends not just on who enters a profession, but critically, on how they reshape the dynamics of peer influence once inside. In a counterfactual where the gender and ideological composition remain frozen at

their 1970 levels, the share of papers using the feminine form in 2019 falls by 3 percentage points, while masculine usage rises by 4.

An extensive literature in network theory emphasizes how network structure drives diffusion of new behaviors, views, and ideas (Bikhchandani et al., 1992; Kandori, 1992; Karni and Schmeidler, 1990; Young, 2014). We contribute by quantifying the real-world role of peer effects in the diffusion of a behavior — pronoun choice in theoretical economics — through an observed network. Unlike prior studies using language (Fryer and Levitt, 2004; Goldin and Shim, 2004; Lieberman and Bell, 1992), pronoun choice is discretionary and irrelevant to a paper’s scientific value or publication prospects, allowing us to isolate peer influence from economic incentives or institutional constraints. We also document how the overwhelming conformism of theoretical economists initially reinforced the status quo by slowing adoption of novel writing styles, then magnified external influences to accelerate their diffusion, consistent with theoretical literature (Goeree and Yariv, 2015; Jackson and Yariv, 2005; Matsuyama, 1991). Homophily played a protective role: by enabling like-minded collaboration, it allowed cultural minorities to express non-traditional preferences more freely and enabled innovation to survive long enough to spread (Golub and Jackson, 2012; Jackson et al., 2025). While homophily is documented across professional domains (Beaman, 2013; Davies et al., 2022; Ductor and Prummer, 2024; Kjelsrud and Parsa, 2024; Zeltzer, 2020; Zhu, 2018), we show how it shapes cultural innovation diffusion.

We also provide micro-level evidence within a unified framework that can generate both cultural persistence (Fernández and Fogli, 2009; Guiso et al., 2008) and cultural change (Becker and Woessmann, 2008; Fernández et al., 2019): the same network structures and peer influence processes can hinder or foster cultural change. Whereas Giuliano and Nunn (2021) explain cultural persistence and change through the selection of traditionalists and innovators, we emphasize peer effects working through network structures, explaining how societies can innovate on some norms but not others, rather than just cross-societal tendencies to innovate. We contribute to classic diffusion dynamics across domains, underscoring the broad relevance of peer and bargaining mechanisms from technology adoption (Goolsbee and Klenow, 2002; Griliches, 1957) to collective action mobilization (García-Jimeno et al., 2022; Kelli and Gráda, 2000).

Finally, we contribute methodologically by adapting *word2vec* into a method to locate academics in an embedded academic space which we refer to as *author2vec*. We exploit variation in this academic space to define feasible network ties, which provides us with exclusion restrictions and enables dimension-reduction in community detection.

2 The economics profession in the last half century

Beyond writing styles, the economics profession has undergone two significant changes over the past half-century: a rising share of female academic economists and increased academic collaboration, both reflected in our network of economic theorists. Panel A of [Figure 2](#) shows the share of women publishing (in red) steadily increased from 2 percent of authors in 1970 to 20 percent by 2020, while the share of papers with at least one female author (in blue) rose even faster, from less than 0.5 percent to over 30 percent.⁴ In the presence of gender differences in preferences, this rising share of women could account for a large fraction of the shifts in pronoun use shown in [Figure 1](#). Yet, papers with female authors follow similar writing-style trends to those with only male authors (see [Figure A.11](#)).

Co-authorship has become the norm: the co-authoring share rose from below 50 percent of all papers in 1970 to nearly 90 percent in 2020, representing 68 percent of all publications in economic theory. Whether increased collaboration entrenches the status quo or promotes stylistic change depends on three considerations: the extent of preference-based homophily in co-authoring, underlying conformism in the population, and how teams resolve these sources of disagreement over writing style choice.

Another potential driver is the entry of new cohorts with different stylistic preferences from incumbents. However, partitioning the set of authors into ten-year cohorts reveals minimal cohort effects in both co-authorship practices (Panel B of [Figure 2](#)) and writing styles (see [Figure A.12](#)). Pronoun usage shifted almost in parallel across cohorts, with only small level gaps. In contrast, early adopters of mixed and feminine forms — Arrow, Baumol, Black, Bowles, Crawford, Thomas Romer, and Spence — come from the earliest cohort, suggesting that senior scholars, not newcomers, led the stylistic change.

The aggregate changes in writing styles do not simply reflect differences across authors; there is considerable within-author variation over their individual careers. We construct a transition matrix of pronoun-form switches by tracking each author’s consecutive papers. While the diagonal entries of Panel A in [Table 2](#) show expected persistence in individual writing styles — with the probability of using the same style as high as 50 percent for the masculine and plural forms — the off-diagonals reveal significant switching.⁵ This within-scholar variation will help us disentangle the roles of co-authorship, cohort-differences, and social influences as drivers of the evolution of writing styles.

⁴This share matches the share of women in economics at large today ([Chari and Goldsmith-Pinkham, 2017](#); [Kjelsrud and Parsa, 2024](#); [Lundberg and Stearns, 2019](#)).

⁵The first row of Panel B in [Table 2](#) reports the stationary distribution under the transition matrix from Panel A. This implied long-run distribution closely matches the observed distribution around the mid-2000s with a third of only-masculine and of only-plural, a fifth of mixed and 12 percent of only-feminine articles.

3 The data

In this section we describe the five key components of our data collection. The online appendix contains a more detailed description.

Selection of economic theory articles. To construct the sample of economic theory papers, we first collected the set of all papers and authors in 1764 Economics and Economics-adjacent journals going back to 1852 from two sources: *Jstor* and *Crossref*.⁶ We restricted the initial sample of 710,000 published papers through a multi-step process detailed in Appendix subsection 11.1, which used journal information and full texts to classify papers as theory or not. Our final dataset comprises 66,854 articles — published between 1970 and 2019 — by 29,302 unique authors. We assigned unique identifiers to each, building an author-level panel dataset. We excluded articles with four or more authors, and papers from authors who only ever solo-authored.⁷

Third-person pronouns. To measure the dependent variable — the gendered pronoun forms used to refer to economic agents in each paper — we need to distinguish third-person pronouns that refer to agents in models from uses for other reasons. The growing shares of female economists and co-authorship risk confounding our measures of pronoun use with increasing references to female authors or collaborators. We tackle this problem using a *co-reference resolution* model, a natural language processing (NLP) tool that links pronouns and noun phrases to their antecedents, detecting when different expressions refer to the same entity.⁸ Specifically, for each paper we extract every third-person pronoun alongside their context window, then apply AllenNLP’s state-of-the-art neural co-reference-resolution module to identify the noun phrase each pronoun refers to.

After mapping each third-person pronoun to nouns in every segment, we keep only instances that refer to nouns in a keyword list referring to economic agents (see subsection 11.2). This list includes only gender-neutral nouns like “individual,” “worker,” or “agent.” Figure A.18 presents the top-50 nouns by frequency across all papers. For example *agent*, the

⁶We obtained the *Jstor* data under a data user agreement for the project and the *Crossref* data using the defunct *Crossref* API: <https://www.crossref.org/education/retrieve-metadata/rest-api/>.

⁷Out of the 66,858 papers, 32 percent are single authored and 68 percent are coauthored. Among the coauthored set, 67 percent have two, 27 percent three and only 6 percent have four authors or more. Authors who never co-authored constitute isolated components of the network. Because in the first step of our empirical strategy we classify authors into two underlying types using information from co-authorship links, there is no information to classify isolated components of the network, and we must exclude them. They represent 10 percent of all authors.

⁸For example, in the sentence “The consumer maximizes her utility subject to a budget constraint”, a co-reference resolution model can recognize that “her” refers to the noun “consumer”.

most common referenced noun in our sample, constitutes 6.5 percent of all third-person pronoun mentions. After having identified the relevant pronouns, we obtained the counts of masculine, feminine, and plural pronouns in each paper.⁹ The distribution reveals a striking pattern of mass points at 100 percent masculine, 100 percent feminine, and 100 percent plural, with the remaining papers showing an even balance of forms. This pattern led us to classify articles into four groups: masculine-only, feminine-only, plural-only, and mixed.

Co-authoring and citations networks. The metadata for each paper in our sample include information on its authors. Based on these data we built a time-varying co-authoring network dataset encoding as edges the cumulative number of co-authorships between every pair of authors every year between 1970 and 2019. Using *Microsoft’s Academic Graph* (MAG), we did a similar exercise to build a time-varying citations network.¹⁰ In contrast to the co-authoring network, the citations network is directed, allowing us to distinguish between backward (i cites j) and forward (i is cited by j) citations.

Other covariates. We first assign sub-fields to authors by embedding selected theory-relevant JEL field descriptions — using OpenAI’s text-embedding-ada-002 model — and averaging GPT embeddings of each author’s paper titles and their cited papers, and then matching authors to their three nearest fields.¹¹ We then classify the ethnic origin of authors with *Namsor*, a commercial software tool that identifies the likely regions of origin of names. Lastly, we infer the authors’ genders using their first names using R’s *Genderize* package, a probabilistic classifier for first names. Fourth, we aggregate citation counts for each author across publications.¹²

Acquaintance network. Our social network includes 30 thousand economists spanning a half-century of research across multiple sub-fields.¹³ Many pairs neither coauthor nor cite one another and because they work in different eras or areas, would never have had an opportunity to interact. Consequently, each theorist effectively “knows” (personally or through

⁹While *Allen NLP* has an accuracy of at least 75 percent in standard English text, our manual checks suggest an error rate of close to zero at the paper level.

¹⁰The MAG was a large-scale bibliographic knowledge database produced by Microsoft Research—covering millions of papers, authors, institutions, and their citation links—discontinued at the end of 2021 (<https://www.microsoft.com/en-us/research/project/microsoft-academic-graph>). We have citation data from MAG for 90 percent of our papers (60131 out of 66854 papers). See the supplemental Appendix for more details on construction and coverage.

¹¹See <https://openai.com/blog/new-embedding-models-and-api-updatesforAPIdetails>. The Appendix reports the full list of JEL fields.

¹²The citations data come from two sources: MAG and *Crossref*. See the Appendix for more details.

¹³This includes economic theorists but also economists in other fields who published theoretical papers.

their work) only a small slice of the overall network. We construct an underlying network of “feasible professional interactions”, hereafter referred to as the *acquaintance network*. To assign acquaintance edges between pairs of economists, we exploit the time and “academic” dimensions in a two step process. We first construct an “academic” mapping to measure “academic” distance between any two theorists, introducing an algorithm we call *author2vec*. We then define the acquaintance edge. Identifying these feasible pairs provides us with the exclusion restrictions to recover peer effects.

Step 1: Author2vec, embedding authors in academic space. *Author2vec* adapts the architecture and intuition of *word2vec* (Mikolov et al., 2013) to map scholars instead of words. Like *word2vec* — an NLP algorithm which assigns high-dimensional vectors (embeddings) to words based on local word co-occurrence frequencies — *author2vec* exploits how frequently pairs of economists co-author or co-occur in citations across our entire corpus to learn 100-dimensional embeddings encoding their positions in a high-dimensional academic space.¹⁴ In our adaptation, we treat each article as a sentence and its authors and cited economists as the words in the sentence. The resulting embeddings locate each scholar in a high-dimensional “academic space”, where proximity reflects academic similarity. For instance, economists who rarely co-author, cite, or appear together in reference lists receive low similarity scores — encoding both (i) local interactions (direct co-occurrence within the same paper) and (ii) global structure (indirect ties formed through chains of collaborators or citations across the corpus). We measure proximity via cosine similarity — pairwise normalized dot products ranging from -1 to 1 — with higher scores indicating greater academic proximity.

Panel A of Figure 3 illustrates the cosine similarity scores of Ariel Rubinstein and Martin Osborne (green nodes) alongside their respective ten closest economists (yellow), with larger nodes marking co-authors and tan nodes marking co-authors outside the closest set. Edge lengths correspond to academic distance; dashed edges denote citations; select nodes display cosine-similarity scores. The scores capture meaningful variation in academic proximity necessary to construct the acquaintance network: one of Osborne’s top ten neighbors is his co-author; three of Rubinstein’s are co-authors; and both share citations with many of their closest neighbors. Despite collaborating, Osborne and Rubinstein score only 0.12 in cosine similarity, reflecting their distinct research trajectories and no overlap between their respective sets of closest economists. Moreover, Osborne’s average proximity to his network exceeds Rubinstein’s, implying Rubinstein’s ties span greater academic diversity.

Step 2: Building the Acquaintance Network. Armed with the academic similarity scores,

¹⁴Given the size of our corpora, 67 thousand papers with 30 thousand unique authors, we set the embedding dimensionality to 100 which falls within the established guidelines — rich enough to capture varied collaboration patterns without over-parameterizing. See Appendix subsection 11.3 for details.

we construct an “acquaintance set” $Q_n(i)$ for each author i — the pool of scholars with whom i could co-author. Intuitively, $Q_n(i)$ includes i ’s co-authors and any scholar sufficiently close in academic space to be considered potential co-authors, provided their active years overlap. Let $L_n(j)$ be the n most similar authors to each author j by cosine score (including j himself), let $Y(i)$ be i ’s “active” years — from three years before his first publication to five years after his last — and let $A(i)$ be the set of i ’s co-authors (including i himself). Then

$$Q_n(i) = \{k : k \in L_n(j), \forall j \in A(i), Y(k) \cap Y(i) \neq \emptyset\}.$$

Our benchmark estimates use $n = 10$, with alternative specifications $n = \{5, 20\}$. The collection of these sets defines the acquaintance network, an underlying network on top of which co-authorships may form. Panel B of [Figure 3](#) plots distributions of academic similarities between Ariel Rubinstein and all other theorists: non-acquaintances (purple) concentrate near 0, acquaintances (blue/red/green) peak around 0.25 and exceed 0.30, with co-authors above the 75th percentile of acquaintances and the 99th percentile of non-acquaintances. Three of his ten closest neighbors (Eliaz Kfir, Michael Richter, and Yuval Salant) are co-authors. Similar patterns hold for most economists — average cosine similarity equals 0.53 among co-authors, 0.41 among acquaintances, and 0.03 among non-acquaintances. [Table 1](#) reports increased homophily on observables when moving from all author pairs to acquaintances and then to co-authors. Finally, [Figure A.13](#) plots pairwise log-degree scatter-plots across the co-authorship, citations, and acquaintance networks. Acquaintance degree varies widely at a given co-authorship or citation degree — especially for low-degree authors — showing that acquaintance ties capture information beyond collaborations or citations.

4 Model of writing style

We model third-person pronoun choice as a discrete-choice problem driven by (i) an author’s fixed idiosyncratic preference, (ii) social influence from evolving co-author and citation ties, and (iii) bargaining weights in co-authored papers. Below, we formalize each element (defining our network-based influence measure, the peer-effect heterogeneity, and the aggregation rule) before turning to estimation.

Let $a(i, j)t$ denote an article written by authors i and j published in year t , with $i = j$ for single-authored papers. The author(s) of each paper choose among writing styles — masculine-only, feminine-only, plural-only, or mixed usage — denoted $\rho \in \{m, f, p, x\}$. Their

payoff from style ρ on paper $a(i, j)t$ is

$$u_{a(i,j)t}(\rho) = \alpha_\rho + \omega(\mathbf{z}_{ij})[\beta_i r_{it}^\rho + \delta_i^\rho] + (1 - \omega(\mathbf{z}_{ij}))[\beta_j r_{jt}^\rho + \delta_j^\rho] + \epsilon_{a(ij)t}^\rho, \quad (1)$$

Three central components drive the payoffs. First, to accommodate pronoun choice in co-authored papers, we model the paper-level payoff as a weighted average of the authors' individual payoffs. The bargaining weight $\omega(\mathbf{z}_{ij}) \in [0, 1]$ captures author i 's relative influence and depends on pairwise differences in covariates $\mathbf{z}_{ij}$ such as cohort and citation counts, with $\omega(\mathbf{0}) = 1/2$ when authors are identical.¹⁵ We assume constant bargaining weights across choices, as relative influence should not depend on pronoun form. While (1) aggregates over two authors' payoffs for simplicity, we allow up to three-author papers in the estimation, and the framework extends naturally to larger teams.

Second, author-level payoffs include time-invariant author-specific preferences over pronoun styles, denoted δ_i^ρ . These can capture (unobserved) values or beliefs — such as views on gendered language — that shape an author's intrinsic preferences. Despite the panel nature of our data, we cannot simply difference them out in this nonlinear setting.¹⁶ To make estimation of the fixed effects δ_i^ρ tractable, we assume authors fall into one of two latent ideological types inspired by the Osborne–Rubinstein debate over gendered pronouns: Osborne-like (who favored moving away from the masculine form) or Rubinstein-like (who defended the masculine form). This modeling assumption — not an observed classification — reduces dimensionality. Authors draw their preferences from one of two type-specific distributions with means δ_O^ρ and δ_R^ρ , respectively. Defining O_i as a dummy for Osborne-type authors, we write:

$$\delta_i^\rho = \delta_O^\rho(1 - O_i) + \delta_R^\rho O_i + \xi_i^\rho, \quad \mathbb{E}[\xi_i^\rho] = 0 \quad (2)$$

Third and most importantly, the author level payoffs include a time-varying social influence component, $\beta_i r_{it}^\rho$, which captures how peer behavior shapes an author's stylistic choices. This term varies over time, pronoun choice ρ , and author i . r_{it}^ρ is the share of author i 's professional network that has used pronoun style ρ up to time t as the citation-weighted

¹⁵ $\mathbf{z}_{ij}$ include a dummy for shared ethnicity, gender, and subfield indicators; differences in cohort, lifetime citations, productivity (publication count); and the product of log productivities.

¹⁶ A standard approach in discrete choice panel settings eliminates fixed effects by conditioning on sufficient statistics (Chamberlain, 1980), such as an agent's total pronoun-style counts. We do not adopt this strategy for two reasons. First, for the nearly two-thirds of coauthored papers, each observation depends on two authors' fixed effects. Second, conditioning on sufficient statistics prevents us from recovering these fixed effects, which allows us to decompose the evolution of writing-style norms into the contributions of peer influence, idiosyncratic preferences, and co-authorship.

share of prior papers that used style ρ by i 's past coauthors and cited authors. Specifically,

$$r_{it}^\rho = \frac{\sum_{k \in A_i(t)} \sum_{a \in \{a(k, \cdot): \tau: \tau < t\}} \omega_k \mathbf{1}\{\text{Paper } a \text{ uses } \rho\} + \sum_{j \in \{j \text{ cited in } a(i, \cdot)t\}} \sum_{a \in \{a(j, \cdot): \tau: \tau < t\}} \omega_j \mathbf{1}\{\text{Paper } a \text{ uses } \rho\}}{\sum_{k \in A_i(t)} \sum_{a \in \{a(k, \cdot): \tau: \tau < t\}} \omega_k + \sum_{j \in \{j \text{ cited in } a(i, \cdot)t\}} \sum_{a \in \{a(j, \cdot): \tau: \tau < t\}} \omega_j}, \quad (3)$$

where $A_i(t)$ denotes the set of all co-authors of author i up to time t , and $C_i(t)$ is the set of authors cited by i at time t . We weight each peer's contribution by their citation prominence ω_ℓ , though we also undertook robustness checks using uniform weights, treating all peers as equally influential regardless of citation count.¹⁷ r_{it}^ρ varies across authors, publications, and pronoun forms as networks evolve and style prevalence shifts. Appendix Figure A.15 presents the distribution of r_{it}^ρ in our sample. The parameter β_i captures each author's individual sensitivity to peer influence. We allow for heterogeneity in peer responsiveness where β_i can take on both positive or negative values: authors with $\beta_i > 0$ are conformists, who follow their peers regardless of what they are choosing; those with $\beta_i < 0$ are contrarians, who move away from their peers regardless of what they are choosing. These can be seen as author-specific psychological predispositions, and thus common across choices.

We model β_i as a draw from a normal distribution conditional on author characteristics $\mathbf{w}_i$, including gender and idiosyncratic preference type (O_i) to capture the fact that the composition of the economics profession may have shifted — ideologically and demographically:

$$\beta_i | \mathbf{w}_i \sim \mathcal{N}(\mu(\mathbf{w}_i), \sigma(\mathbf{w}_i))$$

This structure lets us estimate the share of conformists — those with $\beta_i > 0$ — within any subgroup defined by $\mathbf{w}_i$.

The remaining components of the model are summarized as follows. The α_ρ are choice-specific intercepts. We maintain that $\mathbb{E}[r_{it}^\rho \alpha_\rho] = 0$, so there are no unobserved choice attributes correlated with the (endogenous) peer-influence regressor. This differs from residential-location or product-demand models, where amenities or quality give every choice an intrinsic payoff (e.g., Bayer and Timmins (2007); Nevo (2003)). In our setting, pronoun form has no intrinsic value. The contribution of a theory paper is unaffected by whether it uses masculine, feminine, plural, or mixed-gender pronouns.¹⁸ Finally, $\epsilon_{a(i,j)t}^\rho = \varphi_t^\rho + \tilde{\epsilon}_{a(i,j)t}^\rho$ represents

¹⁷Citation prominence is defined relative to the citation prominence of all other peers relevant for (i, t) :

$$\omega_\ell = \frac{1 + \text{Citations of } \ell}{\sum_{\text{All } j, k} (1 + \text{Citations of } \ell)}.$$

¹⁸One might object that pronoun choice affects readability, creating intrinsic value differences across styles.

time-varying unobservables. The φ_t^ρ reflect broad societal trends in the popularity of style ρ external to the professional network which we absorb with time fixed effects. The residual $\tilde{\epsilon}_{a(i,j)t}^\rho$ captures paper-specific shocks that may be correlated with the peer-influence variables $(r_{it}^\rho, r_{jt}^\rho)$: $\mathbb{E}[(r_{it}^\rho, r_{jt}^\rho)\tilde{\epsilon}_{a(i,j)t}^\rho] \neq \mathbf{0}$. For example, auto-correlation in the $\tilde{\epsilon}_{a(i,j)t}^\rho$'s will generate dependence with the social influences, r_{it}^ρ , through network effects. For example, a past shock on i induces him to choose a particular writing style; his conformist peers will subsequently mimic his choice; their choices now influence author i at time t . Moreover, because peer networks may reflect ideological homophily or past contagion effects, we expect the social influence term r_{it}^ρ to be dependent with intrinsic preferences δ_i^ρ . In the next section we introduce our estimation strategy and how we deal with these issues.

5 Estimation Strategy and Identification

To estimate the discrete choice model based on the preferences in (1) we must address two main econometric challenges. First, idiosyncratic time-varying unobservables, $\tilde{\epsilon}_{a(i,j)t}^\rho$, may be dependent with social influences through the network structure. Second, idiosyncratic preferences, δ_i^ρ , are unobserved and may also be dependent with social influences. In this section we address these two concerns.

5.1 Identification: Control Function and Acquaintance Network

We address the dependence between social influences and time-varying unobservables with a control function approach (Jochmans, 2023; Johnsson and Moon, 2021) leveraging the acquaintance network (defined in section 3). We instrument peer influence r_{it}^ρ using pronoun choices of authors *outside* i 's acquaintance set, who coauthored with, or were cited by, i 's coauthors or citees. To ensure excludability we only consider choices of authors outside i 's *direct network*, defined as all coauthors, citees, and most importantly, anyone in i 's acquaintance set. If i 's peers respond to peer influence, then the stylistic choices of these second-degree peers induce relevant time-varying variation in r_{it}^ρ .

We exploit the panel structure and construct instruments based on first differences of indirect peer exposure across two consecutive publications by author i : $\Delta z_{it}^\rho \equiv z_{it}^\rho - z_{it-1}^\rho$,

For instance, the mixed form — assigning different genders to different players — may help exposition. However, the growing use of feminine-only and plural-only forms suggests many authors do not perceive such benefits. If readability gains exist, they are likely second-order relative to the social dynamics we study. A second potential channel is editorial filtering: journals may implicitly reward or penalize certain styles, or authors may believe they do. We test this empirically and find no evidence of explicit editorial guidelines.

where

$$z_{it}^\rho = \frac{\sum_{k \in A_i(t) \cup C_i(t)} \sum_{P_i(k,t)} \omega_k \mathbf{1}\{\text{Paper } a(\ell, m)\tau \text{ uses } \rho\}}{\sum_{k \in A_i(t) \cup C_i(t)} \sum_{P_i(k,t)} \omega_k}, \quad (4)$$

$C_i(t)$ is the set of i 's citees up to time t , and $P_i(k, t) = \{a(\ell, m)\tau : \tau < t \text{ and } \ell \in A_k(t) \cap Q_i^C \cap C_i(t)^C, m \in Q_i^C \cap C_i(t)^C\}$ denotes the set of articles by authors outside of author i 's acquaintance set that have not been cited by author i , but who are past co-authors of one of his past co-authors or citees, k .¹⁹ This isolates changes in the writing style usage of newly acquired second-degree authors not directly connected to i . Their choices are excludable — they affect i 's style choice only through their influence on i 's peers.

In Figure 4 we illustrate the variation in Δz_{it}^ρ using Debraj Ray's professional network in 1993 (left) and 1994 (right). Green nodes denote Ray's acquaintances, with name labels for past coauthors; pink nodes represent coauthors of his coauthors who fall outside his acquaintance set. For example, in 1993 Douglas Bernheim had 4 co-authors outside Ray's acquaintance set. By 1994, Ray formed new co-authorships — including one with Kalyan Chatterjee, whose three past coauthors were also outside Ray's acquaintance set. These newly formed second-degree links contribute to the first-difference variation in Ray's instrument in 1994.²⁰ As an additional example, Appendix Figure A.14 illustrates the instrumental variables variation for Drew Fudenberg between 1992 and 1993.

We implement our control function approach by estimating a fractional response multinomial logit reduced form regression (Mullahy, 2011), given the fractions r_{it}^ρ add up to 1 across ρ :

$$\mathbb{E}[r_{it}^\rho | \Delta \mathbf{z}_{it}] = \frac{\exp(\Delta \mathbf{z}_{it}' \boldsymbol{\pi}^\rho)}{1 + \sum_{\rho \in \{m, f, x\}} \exp(\Delta \mathbf{z}_{it}' \boldsymbol{\pi}^\rho)}, \quad (5)$$

where $\Delta \mathbf{z}_{it} = (\Delta z_{it}^m, \Delta z_{it}^f, \Delta z_{it}^x)$. These conditional mean functions capture the share of variation in i 's peer pronoun usage induced by changes in the choices of those not directly connected to i but connected to i 's coauthors or citees. Under the identifying assumption

¹⁹If none of author i 's co-authors or citees have co-authors outside of his acquaintance set at time t , i.e., $\bigcup_{k \in A_i(t) \cup C_i(t)} P_i(k, t) = \emptyset$, we define $z_{it}^\rho = 1/4$ for all ρ , the maximum entropy multinomial distribution among four choices.

²⁰See Jochmans (2023) for a related approach using network distance to establish instrument exogeneity. Unlike the cross-sectional IV strategy of Bramoulle et al. (2009), which uses second-degree neighbors' *covariates* regardless of overlapping paths, our panel structure allows us to exploit time variation in second-degree exposure via past peer *choices*. Johnsson and Moon (2021) also propose using a control function to recover peer effects, but do so in a cross-sectional network setting with endogenous network links instead.

that $\mathbb{E}[\Delta z_{it}^\rho \tilde{\epsilon}_{a(i,j)t}] = 0$, the residuals from the fractional response model, $\eta_{it}^\rho = r_{it}^\rho - \widehat{\mathbb{E}}[r_{it}^\rho | \Delta \mathbf{z}_{it}]$, contain the endogenous variation in r_{it}^ρ which we include as a regressor in (1). For coauthored papers, we include both η_{it}^ρ and η_{jt}^ρ to account for the influence of both authors’ networks.

Table 3 reports estimates of π^ρ from (5), using plural-only as the reference category. The benchmark specification (top panel) includes both coauthors and citees in the network. Across all three columns — masculine-only, feminine-only, and mixed forms — a consistent pattern emerges: higher use of the masculine or feminine form by non-acquaintance peers raises the likelihood that an author’s peers adopt the same form, while greater use of the mixed form reduces the probability of adopting the masculine form. The middle and bottom panels, which restrict the network to coauthors only and citees only, yield similar qualitative patterns, with slightly larger coefficients in the citations-only case. Because fractional logit estimates are often not directly interpretable, Table 4 also presents linear probability models to assess robustness. We regress each pronoun form share on the level of the corresponding instrument (instead of the first difference) from (4), including author fixed effects to isolate within-author variation. The results mirror those in Table 3: the instrument for a given pronoun form consistently and significantly predicts the corresponding peer share.

Discussion: What about linear IV? Why not instrument peers’ average choices in a reduced form IV setting? In a network setting with heterogeneous peer effects, the standard monotonicity requirement for the first stage may not hold and the IV estimand may lack a causal interpretation. In our setting, we allow for both conformists and contrarians: Consider economists i, j, k , where j is i ’s co-author and k is j ’s co-author but is not in i ’s acquaintance set. If j is a conformist, he becomes more likely to adopt a writing style as it gains popularity among peers like k . Thus, i is a *complier* because j ’s behavior moves with k ’s. If j is a contrarian, he becomes more likely to adopt a style as it becomes less popular among peers like k . Thus, i is a *defier* because j ’s behavior moves against k ’s. As the treatment effects literature shows, IV does not recover a well-defined causal parameter for any subgroup in the presence of defiers (Angrist et al., 1996; Dahl et al., 2023).

5.2 Coauthorship Formation Model and Latent Preference Types

We turn to the second econometric challenge, namely time-invariant author preferences that correlate with social influences: These latent preferences must be accounted for in the model, yet remain unobserved. Supported by strong empirical evidence of homophily in co-authorship networks, we propose using the co-authorship network structure to identify latent preferences. A large literature documents sorting along observable dimensions such as gender, ethnicity, field, productivity, and social proximity (Besancenot et al., 2017; Duc-

tor et al., 2023; Ductor and Prummer, 2023; Fafchamps et al., 2010; Freeman and Huang, 2014; Önder et al., 2021), especially among economists. Collaboration patterns also exhibit “small world” features (Goyal et al., 2006; Newman, 2001).²¹ These facts suggest that authors may sort not only on observables, but also on unobservables — such as the latent preferences (e.g., values or beliefs) captured by δ_i^ρ in our model. If homophily operates on these dimensions, we should observe clustering: traditionalist authors ($O_i = 0$) coauthor disproportionately with other traditionalists, and innovative authors ($O_i = 1$) with other innovative ones. Our key insight is that network structure contains information about unobserved ideological similarity. In a setting where coauthorship reflects both observable and unobservable traits, if two authors differ on observables but still collaborate, we infer similarity on latent preferences. Conversely, if two authors are similar on observables but never coauthor, we infer ideological distance. These patterns of collaboration — or their absence — help us recover meaningful variation in unobserved author types. In practice, we estimate a homophily-based co-authorship model that assigns authors to one of two latent groups — those aligned with Rubinstein’s preferences and those with Osborne’s — adapting the *community detection* problem from Network Science.²²

5.2.1 The community detection problem

The *community detection* problem focuses on identifying unobserved types within a network (Karrer and Newman, 2010; Newman, 2001, 2018). We estimate a covariates-adjusted Stochastic Block Model (SBM) following Feng et al. (2023) that controls for observable homophily while inferring latent communities. The SBM, a workhorse inference-based community detection model, assumes a fixed number of communities (here, two) and models links between node pairs as Poisson draws conditional on pairwise covariates and community-membership parameters. Such model suits settings similar to ours with sparse (co-authorship) networks, and repeated coauthorship.²³ We recover an estimate of O_i for all authors with at

²¹Several mechanisms may drive these patterns, including gender differences in risk preferences (Lindenlaub and Prummer, 2020), asymmetric credit for joint work (Sarsons et al., 2021), and signaling concerns (Onucich and Ray, 2021). Co-authorship has grown markedly (Anderson and Richards-Shubik, 2022; Hammermesh, 2013; Kuld and O’Hagan, 2018; McDowell and Melvin, 1983), author age has risen (Hammermesh, 2015), and “small world” structures persist, with a few prolific economists linking otherwise distant peers (Goyal et al., 2006).

²²Economists have developed various methods to model network formation with unobserved link drivers. Some approaches sidestep estimating these unobserved effects altogether (Fafchamps et al., 2010; Graham, 2017), while others exploit additional structure (dePaula et al., 2018; Islam et al., 2022).

²³The SBM abstracts from timing and models the intensive margin of co-authorship. Since we aim to recover time-invariant author traits, this static specification is appropriate and avoids modeling the complex dynamics of collaboration over time.

least one coauthor.²⁴

Acquaintance network-adjusted SBM: We further adapt the covariates-adjusted SBM to restrict potential links to pairs within each other’s acquaintance sets. This serves two purposes. First, it mitigates the explosion in potential dyads as the number of authors grows. Second, including infeasible dyads would dilute estimates of homophily. For example, if ethnic similarity increases collaboration but partnerships are typically local, treating all distant same-ethnicity pairs as feasible would understate the effect. Our acquaintance network captures professionally “nearby” pairs; restricting the SBM to this feasible edge set ensures it focuses on links that could realistically form.

In practice, let each of the n economic theorists have an unobserved (to us) type $\tau_i \in \{\ell, c\}$, with π_ℓ denoting the share of type ℓ and $\pi_c = 1 - \pi_\ell$ the share of type c . Conditional on types, the number of co-authorships y_{ij} between $i \in Q(j)$ and $j \in Q(i)$ is Poisson distributed:

$$y_{ij} \sim \mathcal{P}(\omega_{\tau_i \tau_j} e^{\mathbf{x}'_{ij} \boldsymbol{\gamma}}), \quad \Omega = \begin{pmatrix} \omega_{\ell\ell} & \omega_{\ell c} \\ \omega_{\ell c} & \omega_{cc} \end{pmatrix}$$

where Ω captures the degree of type-based homophily in the network formation technology, when the diagonal elements exceed off-diagonal ones.²⁵ To account for homophily on observables, we include the following pairwise covariates in $\mathbf{x}_{ij}$: indicators for same sex and same ethnicity; the number of shared sub-fields; differences in age, log citations, and log productivity; and the log product of the two authors’ productivities.²⁶

The joint likelihood of observing co-authoring matrix $\mathbf{Y}$ and an assignment of types $\boldsymbol{\tau} = (\tau_1, \tau_2, \dots, \tau_n)$ (vector of latent types) is

$$\begin{aligned} \mathcal{L}(\mathbf{Y}, \boldsymbol{\tau} | \Omega, \boldsymbol{\gamma}, \pi, \mathbf{X}) &= \mathbb{P}(\mathbf{Y} | \boldsymbol{\tau}, \Omega, \boldsymbol{\gamma}, \pi, \mathbf{X}) \mathbb{P}(\boldsymbol{\tau} | \Omega, \boldsymbol{\gamma}, \pi, \mathbf{X}) \\ &\propto \prod_{i=1}^n \left[\prod_{j \in Q(i)} \left(\omega_{\tau_i \tau_j} e^{\mathbf{x}'_{ij} \boldsymbol{\gamma}} \right)^{y_{ij}} \exp(-\omega_{\tau_i \tau_j} e^{\mathbf{x}'_{ij} \boldsymbol{\gamma}}) \right] \pi_{\tau_i}, \end{aligned} \quad (6)$$

Solving the community detection problem involves maximizing (6) jointly over type shares π_ℓ , homophily parameters $\boldsymbol{\gamma}$ and Ω , and community assignments $\boldsymbol{\tau}$. Appendix [subsection 11.6](#) outlines the procedure, which follows [Feng et al. \(2023\)](#). The key insight is that the MLEs for π_ℓ and Ω have closed-form solutions given $\boldsymbol{\tau}$ and $\boldsymbol{\gamma}$, allowing us to compute a profile likelihood over just $\boldsymbol{\tau}$ and $\boldsymbol{\gamma}$.

²⁴Because the SBM infers communities from global co-authorships, it cannot classify isolated authors, which we exclude from estimation (11% of the original network; 87% of them published only one paper).

²⁵The model accommodates single-authored papers as “self-edges.”

²⁶[Newman \(2018\)](#) shows that including this last covariate is akin to a SBM with “degree-correction”, which accommodates networks with high dispersion of their degree distribution.

5.2.2 Estimation results from the co-authoring model

Table 5 presents the estimates. Column 1 presents our community detection estimates under our benchmark acquaintance definition that sets $n = 10$ — i.e., seeding the acquaintance sets with the ten closest economists to each co-author. All results remain robust to the tighter (five closest) and looser (twenty closest) acquaintance sets (in Columns 2 and 3). The top panel reports coefficients on pairwise observables (γ): Except for the pairwise difference in productivities, co-authorship is significantly more likely among authors who share ethnicity, gender, or subfields, and less likely when age or citation gaps are large. These results confirm strong homophily on observables in economic theory collaborations. The bottom panel reports estimates of implied homophily along the unobserved type dimension, Ω , informed by the relative frequencies of observed co-authorships under the optimal community assignment τ . Conditional on observables, co-authorships are nine times higher between two ℓ types than between an ℓ and a c , and three times higher between two c types than between an ℓ and a c .

Turning to the community assignments, τ , we classify 56 percent of authors into one group and 44 percent into the other. Across specifications, Rubinstein and Osborne are always assigned to different groups, allowing us to label them: the Osborne type (relatively more innovative) and the Rubinstein type (relatively more traditionalist). In the next section, we show that Osborne-type authors are indeed less likely to use masculine forms. The Rubinstein group is larger, and the classification is stable: 85 percent of authors remain in the same group across all specifications. Although both communities are similar in aggregate size, the profession’s post-1970 expansion and rising female share could have shifted their composition across entry cohorts. Community shares vary modestly by entry cohort. Figure 5 shows that the Osborne share rose from 39% for authors entering in the 1970s–80s to 46% in the 1990s and later. The tilt toward the Osborne type is therefore modest and cannot, by itself, account for the pronounced shift in writing styles. As shown in Figure A.16, the two groups have similar distributions of gender, ethnicity, field, productivity, and citations.

Finally, we test whether community labels reflect institutional homophily. Although we lack systematic affiliation data for all authors, we collected information for faculty at top-ranked economics departments (1990–present). We regressed each author’s Osborne-group indicator on university dummies. Figure 6 shows a funnel plot with the 39 coefficients against their p-values: only three are significant at the 5 percent level, and 34 are smaller than 0.10 in magnitude. This suggests that community assignment is not driven by shared institutional ties, at least for this subset of theorists.

5.3 Estimation of the Choice Model

Armed with our community assignment, we set $O_i = 1$ for all authors classified in Osborne's community. Together with the control function estimates ([subsection 5.1](#)), we can write $u_{a(ij)t}^\rho$ from (1) as $u_{a(ij)t} = V_{a(ij)t}^\rho + \nu_{a(ij)t}$, where $V_{a(ij)t}^\rho(\beta_i, \beta_j) \equiv$

$$\tilde{\alpha}^\rho + \varphi_t^\rho + \omega(\mathbf{z}_{ij})[\beta_i r_{it}^\rho + \delta^\rho O_i] + (1 - \omega(\mathbf{z}_{ij}))[\beta_j r_{jt}^\rho + \delta^\rho O_j] + \lambda^\rho[\hat{\eta}_{it}^\rho + \hat{\eta}_{jt}^\rho], \quad (7)$$

$\tilde{\alpha}^\rho \equiv \alpha^\rho + \delta_R^\rho$, $\delta^\rho \equiv \delta_O^\rho - \delta_R^\rho$, and $\nu_{a(ij)t}$ are independent of $(r_{it}^\rho, r_{jt}^\rho, O_i, O_j)$ given $(\hat{\eta}_{it}^\rho, \hat{\eta}_{jt}^\rho)$ and type-1 extreme value distributed. As a functional form for the bargaining weights we use

$$\omega(\mathbf{z}_{ij}) = \frac{1}{1 + \exp(-\boldsymbol{\kappa}'\mathbf{z}_{ij})}$$

Under (7), and collecting in vector $\boldsymbol{\theta}$ all parameters, the unconditional likelihood of observing writing style $p_{a(ij)t} = \rho$ for paper $a(ij)t$ averages over the distribution of peer effects for each author conditional on their vector of characteristics $\mathbf{w}_i$:

$$\mathbb{P}(p_{a(ij)t} = \rho | \mathbf{w}_i, \mathbf{w}_j, \mathbf{z}_{ij}; \boldsymbol{\theta}) = \int \int \frac{\exp\left(V_{a(ij)t}^\rho(\beta_i, \beta_j)\right)}{1 + \sum_{s \in \{m, f, x\}} \exp\left(V_{a(ij)t}^s(\beta_i, \beta_j)\right)} d\Phi(\beta_i | \mathbf{w}_i) d\Phi(\beta_j | \mathbf{w}_j).$$

We think of the distribution of peer-effect heterogeneity as capturing traits that are possibly stationary in the overall population. Over the last fifty years, however, the Economics profession has grown in size. For example, while we see 1,620 economists from the 1970s cohort, we see 4,970 from the 1990s cohort, and 11,317 from the 2010 cohort. The profession also has shifted its sex composition towards women. Because the new entrants or women as a whole could differ in their preferences relative to incumbents, we allow $\mathbf{w}_i = (\text{woman}_i, O_i)$ to include the authors' sex and community assignment dummies.

The likelihood for the writing style choices across all articles, $\mathbf{P}$, is thus

$$L(\boldsymbol{\theta} | \mathbf{P}, \mathbf{W}, \mathbf{Z}) = \prod_a \prod_{\rho \in \{m, f, x, p\}} \mathbb{P}(p_{a(ij)t} = \rho | \mathbf{w}_i, \mathbf{w}_j, \mathbf{z}_{ij})^{\mathbf{1}\{p_{a(ij)t} = \rho\}}. \quad (8)$$

We use maximum simulated likelihood to estimate $\boldsymbol{\theta}$.²⁷ The vector of parameters includes the three pronoun specific intercepts $\tilde{\alpha}^\rho$, the three sets of time effects φ_t^ρ (in practice we include time effects for groups of 5 years with the exception of 1970-1974 and 1975-1979 for which we include a single time effect²⁸), the three coefficients δ^ρ on the Osborne community dummy,

²⁷See Appendix [subsection 11.7](#) for additional details about the estimator.

²⁸In 1970-74 no papers used the only feminine choice, so the time effect for that period is unidentified.

the three coefficients λ^p on the control function, five coefficients on the pairwise covariates κ on the bargaining weight function, and six coefficients governing the distribution of peer effect heterogeneity:²⁹ $\mu_{\mathbf{w}}$, $\sigma_{\mathbf{w}}$.

6 Findings and Simulation Exercises

Our main estimates use acquaintance sets $n = 10$, and incorporate both past co-authorship and citation ties as sources of peer influence.

6.1 Parameter estimates

The time effects φ_t^p capture shifts in the relative popularity of writing styles over time, reflecting broader societal changes outside the profession. The estimates (and confidence intervals) in Figure 7 reveal two main patterns: substantial changes in the relative preference for writing styles across cohorts, and a secular decline in the dispersion of those preferences. Among articles published in 1970–1979, the masculine form dominated, with likelihood ratios of 440 to 1, 22 to 1, and 1.1 to 1 over the feminine, plural and mixed forms. In contrast, the most recent articles show no preference among the masculine, feminine, and mixed forms, but favor the plural form, with a likelihood ratio of 2.7 to 1.³⁰

Panel A of Table 6 reports a subset of our parameter estimates. First, consider the *latent preference estimates*, δ^p . Relative to the plural form, Osborne-community authors are less likely to choose the masculine form (-0.36 , s.e. = 0.04), equally likely to choose the feminine form (-0.02 , s.e. = 0.13), and more likely to choose the mixed form (0.88 , s.e. = 0.07). This pattern confirms that the authors we identified as being part of Osborne’s community do share affinity with his writing style preferences, validating the interpretation of the community labels. The coefficients on the *control functions*, λ^p , are statistically significant across all writing styles, underscoring the importance of addressing endogeneity in peer influence. Turning to the *bargaining weight parameters*, κ , larger age gaps tilt decisions toward the older co-author, while larger gaps in citations and productivity favor the less cited and less productive co-author (conditional on their age difference). Sex and ethnicity differences do not affect bargaining weights.

In Panel B of Table 6 we turn to *peer effect estimates* β_i . The panel reports four sex-by-community group-level means and standard deviations for the estimated peer effect distributions. Figure 8 plots the implied densities. We find strong and precisely estimated local

²⁹The number of women in the profession is small, so we do not estimate separate variances for each sex.

³⁰For the oldest articles, $\exp(0.1)/\exp(-6) = 440$, $\exp(0.1)/\exp(-3) = 22$, and $\exp(0.1)/\exp(0) = 1.1$; for the most recent, $\exp(0)/\exp(-1) = 2.7$.

professional peer effects in writing style choices. Three patterns emerge. First, the estimates rule out the presence of contrarian economic theorists: in all groups, the peer effect distributions place positive mass only on positive values of β_i . Second, peer effects show only moderate heterogeneity overall — especially among Rubinstein-community economists, 95 percent of whom fall between 1.5 and 1.6. Third, Osborne-community theorists exhibit stronger average peer effects, particularly women, whose mean peer effect is 1.92 — significantly higher than that of all other groups. As a consequence, their strong conformism undermines their idiosyncratic preferences and reinforces the status quo when the masculine forms dominate. Together, these findings portray an overwhelmingly conformist profession, with innovative economists responding more strongly to peer behavior. To assess magnitudes, consider two examples. A Rubinstein-community man whose peers shift from using 80 percent masculine-20 percent feminine pronouns to 30 percent masculine-70 percent feminine becomes 46 percent less likely to choose the masculine form and 2.2 times more likely to choose the feminine form. An Osborne-community woman facing the same shift becomes 38 percent less likely to choose the masculine form and 2.6 times more likely to choose the feminine form.³¹ These estimates imply that the growth of the Osborne group (see Figure 5) and the rising share of women in the profession (see Figure 2) have together amplified the aggregate strength of peer influence over time.

6.2 Peer influence, homophily, and demographic change

To shed more light on the role of peer effects in writing norms over the last fifty years, we simulate the model under alternative counterfactual scenarios with a focus on the roles of conformism and homophily in co-authorship. We then discuss cohort effects, which have featured prominently in the cultural change literature. We conclude with robustness checks.

Quantifying the role of peer influences. We begin with the baseline simulation, shown in Panel A of Figure 9, which holds fixed the set of articles and professional network links, simulates pronoun form choices for each paper, and updates the peer influence variables r_{it}^p each year to determine choices in subsequent periods. The simulation starts from the observed distribution of choices in 1970–1974. Panel A of Figure 9 reproduces the historical path of pronoun choices in Figure 1 with remarkable accuracy. It captures both the decline of the masculine form and the rise of plural and feminine forms — matching the timing and shape of the trends in Figure 1. The only notable discrepancy is a slightly earlier rise in the simulated popularity of the mixed form, which takes off more gradually in the data.

³¹For men, $\exp(1.54 \times 0.3) / \exp(1.54 \times 0.8) = 0.46$, and $\exp(1.54 \times 0.7) / \exp(1.54 \times 0.2) = 2.15$; for women, $\exp(1.92 \times 0.3) / \exp(1.92 \times 0.8) = 0.38$, and $\exp(1.92 \times 0.7) / \exp(1.92 \times 0.2) = 2.6$.

We next ask: What would the trajectory have looked like if the main force for change came from the local peer influences within the profession? In this counterfactual we hold the broader cultural environment constant — proxied by the time effects φ_t^ρ frozen at their 1970 values (Figure 7) — and simulate how professional interactions alone shape writing style trajectories. Panel B of Figure 9 plots the resulting aggregate pronoun use over time. Absent external trends, peer influence reinforces the dominant norm: masculine-only usage rises from 65 to 80 percent, while plural use falls from 33 to 18 percent. Rather than driving change, conformist interactions entrench the status quo. The last column of Table 7 reports the percentage point differences between this scenario and the baseline simulation in 2019.

We then flip the experiment: What if societal trends evolve as observed, but peer influence plays no role? To explore this, we turn off peer effects ($\beta_i = 0$ for all authors), while preserving the full sequence of estimated time effects φ_t^ρ . The results, shown in Panel C of Figure 9, stand in sharp contrast to the previous scenario. Without peer reinforcement, stylistic change happens faster: plural-only usage rises steeply, surpassing 55 percent by the late 1980s — about 10 p.p. above the baseline at that point. The masculine form declines more rapidly as well, while mixed and feminine forms largely track the baseline. By 2019, the plural share is 7 p.p. higher than in the baseline, and the masculine share 4 points lower. Meanwhile, mixed and feminine forms lag behind. Both societal trends and these local professional peer effects are necessary to reproduce the realized trajectory.

These last two experiments underscore the dual role of peer effects, either reinforcing prevailing norms or fracturing them. To quantify how local peer influence and societal trends shape the long-run diversity of writing norms, we use entropy — an information theory measure summarizing the dispersion of a distribution. Entropy takes its highest (normalized) value of 1 when all four styles (masculine, feminine, mixed, plural) are used equally, and falls as usage becomes more concentrated in one form — values near 0 indicate dominance by a single form. Figure 10 plots the average entropy by the end of our period (2017–2019) under 33 simulation scenarios. We vary both the external influences by setting time effects to $\pi\varphi_t^\rho + (1 - \pi)\varphi_{1970}^\rho$ for $\pi \in \{0, 0.1, \dots, 1\}$ and the peer effects across three scenarios (none, our estimated values, and a high-conformity benchmark where everyone adopts Osborne-community women’s β_i).³²

Across all peer-effect regimes, entropy rises with π : stronger external influences promote

³²Using the estimated choice probabilities for each article $(p_a^m, p_a^f, p_a^x, p_a^p)_t$, we compute the article’s entropy as $E_{a,t} = -\sum_{\rho \in \{m,f,x,p\}} p_a^\rho \log(p_a^\rho)$, and average over all articles published in 2017–2019. We further normalize by the maximum possible entropy ($\log(4)$):

$$H = \frac{\sum_{a,t \in \{2017-2019\}} E_{a,t}}{\sum_{a,t \in \{2017-2019\}} 1} \frac{1}{\log(4)} \in [0, 1].$$

greater diversity in writing styles. But peer effects shape how diversity emerges. Under weak societal influence (π near 0), peer conformity suppresses innovation, with stronger peer conformity entrenching the dominant norm, leading to significantly lower diversity. As external influences become more salient, however, the professional networks subject to stronger peer influences increase their long-run entropy faster than those with weaker peer effects. The gap across peer effect scenarios closes at $\pi = 1$ (the estimated external trends). The scenario under the estimated peer effects (red line) produces the highest diversity — more than societal influence could achieve on its own. In short, peer effects do not just reinforce the status quo; they also intensify change once it begins to take hold.

Homophily, co-authorship and diversity in writing styles. We now turn to the role of academic collaboration in shaping writing norms. Co-authorship now dominates theory papers — from under 50 percent in 1970 to nearly 90 percent in 2020. Whether this rise fosters innovation or entrenches tradition depends on conformity, homophily, and bargaining inside teams. In practice, co-authorship is highly assortative with 82 percent of papers written by authors from the same community, and 88 percent featuring same-sex teams (slightly less than the 92 percent that would obtain under random matching).³³ To isolate homophily, we simulate a scenario that forces every collaboration to pair opposite communities and opposite sexes (panel F, Figure 9). By 2019 masculine-only usage is 4 p.p. higher and feminine-only 4 p.p. lower than in the baseline — diverse teams slow, rather than speed, innovation. Why would more diverse collaborations yield less writing style innovation? The answer lies in how bargaining weights compare to population shares. Consider the two extreme scenarios of full homophily and full heterophily. Under full homophily, collaborators stick with like-minded partners. Co-authorship gives all types the opportunity to express their preferences within paper, leaving voices undiluted. The aggregate behavior across papers reflects the population shares of the respective communities. Under full heterophily, however, every paper becomes a mixed team: pronoun choices now reflect the bargaining weights within each team. Whether one group’s voice is amplified or diminished under homophily versus heterophily depends on how population-level aggregation compares to within-team bargaining dynamics. In our setting, the baseline is highly homophilous but not fully so. Nevertheless, the same mechanism applies: the more innovative types (Osborne-Men and Osborne-Women) lose influence when their voices must compete within mixed teams rather than being aggregated at the population level in homogeneous teams.

³³Because our community labels are identified from co-authorship itself, it is unsurprising for community homophily to be high.

Compositional changes and cohort effects. Finally, we assess the role of demographic shifts in the evolution of writing styles. Over the past fifty years, as the economics profession grew in size, it changed in composition with the rising share of female economists, and the gradual replacement of older cohorts by new ones. A large literature attributes cultural change to cohort effects (in contrast to period effects): newcomers with different beliefs and preferences reshape prevailing practices. To what extent can these changes account for the observed evolution in writing styles? But earlier findings cast doubt on this explanation: these shifts coincided with a modest increase in Osborne-community representation — from 41 percent in 1970 to 48 percent in 2020. Still, compositional change could matter interacted with the peer dynamics: Osborne-community and women in particular, are more conformists. To explore this mechanism, we simulate a scenario that holds demographics at their 1970 levels (2 percent women, 41 percent Osborne) throughout. Panel E of [Figure 9](#) and column 3 of [Table 7](#) show that the masculine-only usage is 4 p.p. higher, feminine-only 3 p.p. lower, and mixed 1 p.p. lower by 2019 compared to the baseline. Demographic turnover, therefore, played a limited role; the high degree of conformism we observe dampens the influence of new entrants. Lasting change requires not just new voices, but network ties that can amplify them through peer influence.

6.3 Robustness: Choice-specific unobservables

Our model assumes that the value of writing-style choices is purely social, obviating the need to deal with choice-specific unobservables. To conclude, we explore several plausible sources of underlying, non-social variation in the perceived value of writing styles.

Beliefs about journal editors’ preferences. If authors believe that editors favor certain styles, such perceptions could shape writing choices. For example, authors may believe that male and female editors differ in stylistic preferences and act accordingly. To test this, we assembled editorial board data from 1970 to the present for the top five general-interest journals and five leading theory journals. We then estimated linear probability models of masculine-only usage at the article level, regressed on the average number of women on the board in the three years prior to publication. Results appear in [Table 8](#). Column 2 includes author fixed effects to isolate within-author variation. Overall, the share of female editors does not predict gendered writing style. Columns 5 and 6 present results separately for papers with and without a female author: we find no effect for men, but a positive and significant effect for papers authored by women, even with fixed effects. Female authors are

more likely to use masculine-only forms when more women serve as editors.³⁴

Expectations of conformity by un-tenured professors. If writing styles are perceived to matter for publication (and career concerns), un-tenured economists may respond to such perceptions through their choices. For instance, early-career scholars might perceive the profession as expecting more traditional stylistic norms — prompting them to adopt more conservative forms to signal conformity. We estimate linear probability models of gendered pronoun choice on a dummy variable equal to 1 for articles with at least an author in the first six years of their academic career. Authors at an early stage of their career are more likely to choose the plural writing style (relative to all other three styles), both with and without author fixed effects (top row of [Table 9](#), columns 5–6).

Differential changes in women’s preferences. While we find small gender gaps in preferences, it is possible that women’s ability to express their preferences varies over time. Prior work shows women may only assert distinct preferences once they represent a critical mass within a network (e.g., [Owen and Temesvary \(2018\)](#) on bank boards). In our setting, the small share of female theorists early on may have limited their ability to assert their stylistic preferences. As more women entered the profession, their preferences may have become more visible and influential. To test for evolving preferences, in [Table A.12](#), we examine whether the distribution of Osborne-community membership differs across female cohorts. We find no meaningful differences across cohorts.

Signaling preferences across degrees of journal prestige. Authors may associate different writing styles with journals of varying prestige. We test this by regressing pronoun choice on either the log journal ranking or a top-five-journal dummy. As reported in [Table 9](#), articles in “top five” economics journals are less likely to use plural forms (column 6) and more likely to use mixed forms (column 8) even after including author fixed effects.

Underlying complementarities between sub-fields and writing styles. Writing styles may vary systematically across sub-fields. For example, contract theory’s principal-agent models may more often adopt mixed pronoun usage. More abstract sub-fields may favor the plural form. We estimate linear probability models of style choice on subfield dummies and present the results in [Table 10](#). Even-numbered columns include author fixed

³⁴This finding may seem counterintuitive. However, [Kosnik \(2022\)](#) documents more negative sentiment in articles authored by women, particularly in prestigious journals. She argues this is driven by career concerns, because papers with more negative writing styles tend to receive more citations.

effects. We find that authors are more likely to use mixed forms and less likely to use plural forms in papers classified under Collective Decision-Making, Game Theory, Information Economics, and Welfare Economics.

Motivated by these findings, we re-estimate our writing-style model including four additional shifters of choice-specific payoffs: (i) a dummy for at least one female author, (ii) a dummy for at least one un-tenured author, (iii) the journal’s ranking, and (iv) a dummy for articles in the sub-fields Collective decision-making, Game theory, Information economics, or Welfare economics. Our main results remain unchanged.

7 Conclusions

We have shown how changing writing style norms in economic theory — traced through gendered pronoun choices from 1970 to 2019 — stem not only from external societal shifts but crucially from local peer dynamics. By combining a discrete-choice model of pronoun use with exclusion restrictions drawn from a feasible co-author network and accounting for latent author preferences via community detection, we identify significant conformist peer effects: early on, they entrenched the masculine norm; later, they magnified societal pressures and sustained long-run stylistic diversity.

Our quantitative simulations demonstrate that women’s and young economists’ entry accelerated this transformation; homophily in co-authorship preserved minority preferences and fostered stylistic diversity. These findings carry three broader implications. First, understanding cultural transformation requires going beyond accounting for the demographic characteristics of the social network under study; attention to network structure and its evolution are key. Cultural change hinges not just on demographic turnover but also on how newcomers reshape the professional network structure and the bargaining dynamics embedded in collaboration. Second, diversity initiatives focused on demographic composition may miss crucial network effects — the impact of new entrants depends critically on how they alter influence patterns within existing professional ties. Third, our finding that homophily can protect innovation — by giving cultural minorities space to express preferences — challenges the conventional wisdom about heterogeneous teams leading to cultural innovation.

References

ANDERSON, K. A. AND S. RICHARDS-SHUBIK (2022): “Collaborative Production in Science: An Empirical Analysis of Coauthorships in Economics,” *The Review of Economics and Statistics*, 104, 1241–1255.

- ANGRIST, J., G. IMBENS, AND D. RUBIN (1996): “Identification of Causal Effects Using Instrumental Variables,” *Journal of the American Statistical Association*, 91, 444–472.
- BARON, D. (1986): *Grammar and Gender*, New Haven, CT: Yale University Press.
- BAYER, P. AND C. TIMMINS (2007): “Estimating Equilibrium Models of Sorting Across Locations,” *The Economic Journal*, 117, 353–374.
- BEAMAN, L. A. (2013): “Social networks and the dynamics of labor market outcomes: Evidence from refugees resettled in the US,” *Review of Economic Studies*, 80, 128–161.
- BECKER, S. O. AND L. WOESSMANN (2008): “Luther and the Girls: Religious Denomination and the Female Education Gap in Nineteenth-Century Prussia,” *Scandinavian Journal of Economics*, 110, 777–805.
- BESANCENOT, D., K. HUYNH, AND F. SERRANITO (2017): “Co-authorship and research productivity in economics: Assessing the assortative matching hypothesis,” *Economic Modelling*, 66, 61–80.
- BIKHCHANDANI, S., D. HIRSHLEIFER, AND I. WELCH (1992): “A Theory of Fads, Fashion, Custom, and Cultural Change as Informational Cascades,” *Journal of Political Economy*, 100, 992–1026.
- BOND, R. AND P. B. SMITH (1996): “Culture and conformity: A meta-analysis of studies using Asch’s (1952b, 1956) line judgment task,” *Psychological Bulletin*, 119, 111–137.
- BRAMOULLE, Y., H. DJEBBARI, AND B. FORTIN (2009): “Identification of Peer Effects through Social Networks,” *The Journal of Econometrics*, 150, 41–55.
- CHAMBERLAIN, G. (1980): “Analysis of Covariance with Qualitative Data,” *Review of Economic Studies*, 47, 225–238.
- CHARI, A. AND P. GOLDSMITH-PINKHAM (2017): “Gender Representation in Economics across Topics and Time: Evidence from the NBER Summer Institute,” NBER Working Paper No. 23953.
- DAHL, C., M. HUBER, AND G. MELLACE (2023): “It is Never too LATE: A New Look at Local Average Treatment Effects with or without Defiers,” *The Econometrics Journal*, 26, 378–404.
- DAVIES, R., D. CLARKE, K. KARBOWNIK, A. LIM, AND K. SHIBAZAKI (2022): “Gender representation in economics across topics and time: Evidence from the NBER,” .

- DEPAULA, A., S. RICHARDS-SHUBIK, AND E. TAMER (2018): “Identifying Preferences in Networks with Bounded Degree,” *Econometrica*, 86, 263–288.
- DUCTOR, L., S. GOYAL, AND A. PRUMMER (2023): “Gender and Collaboration,” *The Review of Economics and Statistics*, 105, 1366–1378.
- DUCTOR, L. AND A. PRUMMER (2023): “Gender Homophily, Collaboration, and Output,” Working Paper, Universidad de Granada.
- (2024): “Gender homophily, collaboration, and output,” *Journal of Economic Behavior & Organization*, 221, 477–492.
- EAGLY, A. H. (1983): “Gender and social influence: A social role interpretation,” *Psychological Review*, 94, 378–404.
- FAFCHAMPS, M., S. GOYAL, AND MARCO VAN DER LEIJ (2010): “Matching and Network Effects,” *Journal of the European Economic Association*, 8, 203–231.
- FENG, Y., S. HUANG, AND J. SUN (2023): “Pairwise Covariates-Adjusted Block Model for Community Detection,” ArXiv.
- FERNÁNDEZ, R., P. PARSA, AND M. VIARENGO (2019): “The AIDS Epidemic and Shock-Driven Shifts in Consumption Patterns,” *Journal of Economic Behavior & Organization*, 162, 123–138.
- FERNÁNDEZ, R. AND A. FOGLI (2009): “Culture: An Empirical Investigation of Beliefs, Work, and Fertility,” *American Economic Journal: Macroeconomics*, 1, 146–177.
- FREEMAN, R. B. AND W. HUANG (2014): “Collaborating with People like me: Ethnic co-authorship within the US,” NBER Working Paper No. 19905.
- FRYER, R. AND S. LEVITT (2004): “The Causes and Consequences of Distinctively Black Names,” *Quarterly Journal of Economics*, 119, 767–805.
- GARCÍA-JIMENO, C., A. IGLESIAS, AND P. YILDIRIM (2022): “Information Networks and Collective Action: Evidence from the Women’s Temperance Crusade,” *American Economic Review*, 112, 41–80.
- GIULIANO, P. AND N. NUNN (2021): “Understanding Cultural Persistence and Change,” *Review of Economic Studies*, 88, 1541–1581.
- GOEREE, J. AND L. YARIV (2015): “Conformity in the Lab,” *Journal of the Economic Science Association*, 1, 15–28.

- GOLDIN, C. AND M. SHIM (2004): “Making a name: Women’s Surnames at Marriage and Beyond,” *Journal of Economic Perspectives*, 18, 143–160.
- GOLUB, B. AND M. JACKSON (2012): “How Homophily Affects the Speed of Learning and Best Response Dynamics,” *Quarterly Journal of Economics*, 127, 1287–1338.
- GOOLSBEE, A. AND P. KLENOW (2002): “Evidence on Learning and Network Externalities in the Diffusion of Home Computers,” *Journal of Law and Economics*, 45, 317–343.
- GOYAL, S., M. VAN DEL LEIJ, AND J. L. MORAGA-GONZALEZ (2006): “Identification of Peer Effects through Social Networks,” *The Journal of Political Economy*, 114, 403–412.
- GRAHAM, B. (2017): “An Econometric Model of Network Formation with Degree Heterogeneity,” *Econometrica*, 85, 1033–1063.
- GRILICHES, Z. (1957): “Hybrid Corn: An Exploration in the Economics of Technological Change,” *Econometrica*, 25, 501–522.
- GUIO, L., P. SAPIENZA, AND L. ZINGALES (2008): “Long-Term Persistence,” Working Paper 14278, National Bureau of Economic Research.
- HAMMERMESH, D. (2013): “Six Decades of Top Economics Publishing: Who and How?” *Journal of Economic Literature*, 51, 162–172.
- (2015): “Age, Cohort, and Coauthorship,” NBER Working Paper No. 20938.
- ISLAM, A., J. S. KIM, S. SONG, AND Y. ZENOU (2022): “Network Formation with Unobserved Homophily: Identification and Consistent Estimation,” Monash University.
- JACKSON, M., S. NEI, E. SNOWBERG, AND L. YARIV (2025): “The Dynamics of Networks and Homophily,” Working paper.
- JACKSON, M. AND L. YARIV (2005): “Diffusion on Social Networks,” *Economie Publique*, 16, 3–16.
- JOCHMANS, K. (2023): “Peer Effects and Endogenous Social Interactions,” *Journal of Econometrics*, 235, 1203–1214.
- JOHNSSON, I. AND H. MOON (2021): “Estimation of Peer Effects in Endogenous Social Networks: Control Function Approach,” *The Review of Economics and Statistics*, 103, 328–345.

- KANDORI, M. (1992): “Social Norms and Community Enforcement,” *Review of Economic Studies*, 59, 63–80.
- KARNI, E. AND D. SCHMEIDLER (1990): “Fixed Preferences and Changing Tastes,” *American Economic Association Papers and Proceedings*, 80, 262–267.
- KARRER, B. AND M. NEWMAN (2010): “Stochastic Block Models and Community Structure in Networks,” ArXiv.
- KELLI, M. AND C. O. GRÁDA (2000): “Market Contagion: Evidence from the Panics of 1854 and 1857,” *American Economic Review*, 90, 1110–1124.
- KJELSRUD, A. AND S. PARSA (2024): “Mentorship and the Gender Gap in Academia,” *Available at SSRN 5193388*.
- KOSNIK, L.-R. (2022): “Who Are the More Dismal Economists? Gender and Language in Academic Economics Research,” *American Economic Association P&P*, 112, 1–5.
- KULD, L. AND J. O’HAGAN (2018): “Rise of multi-authored papers in economics: Demise of the ‘lone star’ and why?” *Scientometrics*, 114, 1207–1225.
- LIEBERSON, S. AND E. BELL (1992): “Children’s First Names: An Empirical Study of Social Taste,” *American Journal of Sociology*, 98, 511–554.
- LINDENLAUB, I. AND A. PRUMMER (2020): “Network Structure and Performance,” *Economic Journal*, 131, 851–898.
- LUNDBERG, S. AND J. STEARNS (2019): “Women in economics: Stalled progress,” *Journal of Economic Perspectives*, 33, 3–22.
- MATSUYAMA, K. (1991): “Custom versus Fashion: Hysteresis and Limit Cycles in a Random Matching Game,” Working Paper, Northwestern U.
- MCDOWELL, J. M. AND M. MELVIN (1983): “The Determinants of Co-Authorship: An Analysis of the Economics Literature,” *Review of Economics and Statistics*, 65, 155–160.
- MIKOLOV, T., K. CHEN, G. CORRADO, AND J. DEAN (2013): “Efficient Estimation of Word Representations in Vector Space,” ArXiv Working Paper 1301.3781v3.
- MULLAHY, J. (2011): “Multivariate fractional regression estimation of econometric share models,” Working Paper 16354, University of Wisconsin-Madison.

- NEVO, A. (2003): “Measuring Marker Power in the Ready-to-Eat Cereal Industry,” *Econometrica*, 69, 307–342.
- NEWMAN, M. (2001): “The Structure of Scientific Collaboration Networks,” *Proceedings of the National Academy of Sciences*, 98, 404–409.
- (2018): *Networks*, Oxford, UK: Oxford U Press.
- ONUCICH, P. AND D. RAY (2021): “Signaling and Discrimination in Collaborative Projects,” *American Economic Review*, 113, 210–252.
- OSBORNE, M. AND A. RUBINSTEIN (1994): *A Course in Game Theory*, Cambridge, MA: MIT Press.
- OWEN, A. AND J. TEMESVARY (2018): “The Performance Effects of Gender Diversity on Bank Boards,” *Journal of Banking and Finance*, 90, 50–63.
- SARSONS, H., K. GÖRKHANI, E. REUBEN, AND A. SCHRAM (2021): “Gender Differences in Recognition for Group Work,” *Journal of Political Economy*, 129, 101–147.
- STEVENSON, B. AND H. ZLOTNICK (2018): “Representations of Men and Women in Introductory Economics Textbooks,” *American Economic Association P&P*, 108, 180–185.
- YOUNG, P. (2014): “The Evolution of Social Norms,” WP No. 726, Oxford University.
- ZELTZER, L. (2020): “Network Effects in Referral Hiring: The Role of Social Connections in Job Placements,” *Labour Economics*, 67, 101–115.
- ZHU, L. (2018): “Gender and social networks in job search,” *Labour Economics*, 54, 55–65.
- ÖNDER, A., S. SCHWEITZER, AND H. YILMAZKUDAY (2021): “Specialization, field distance, and quality in economists’ collaborations,” *Journal of Informetrics*, 15, 1751–1577.

8 Tables

	(1)	(2)	(3)	(4)
	Co-authors	Acquaintances	Non-coauthors	All
Same ethnicity	0.38 (0.49)	0.23 (0.42)	0.16 (0.36)	0.16 (0.36)
Same sex	0.77 (0.42)	0.76 (0.43)	0.71 (0.45)	0.71 (0.45)
Common fields	1.41 (0.98)	1.00 (0.91)	0.29 (0.56)	0.29 (0.56)
Age difference	9.20 (8.83)	10.89 (9.16)	13.67 (10.72)	13.67 (10.72)
Citations difference	4,720 (12,576)	5,060 (12,335)	2,457 (7,244)	2,457 (7,244)
Productivity difference	12.51 (15.99)	11.91 (14.92)	5.07 (8.65)	5.07 (8.65)
Log productivity product	3.53 (1.94)	3.42 (1.71)	1.70 (1.38)	1.70 (1.38)
Pairs	50,778	748,023	429,238,173	429,288,951

Table 1: Pairwise Characteristics. The table reports means and standard deviations (in parenthesis) for pairwise characteristics across pairs of economists. Columns (1), (2), and (3) restrict the set to pairs who: co-authored with each other; are in each other’s acquaintance sets; never co-authored with each other. Column (4) includes all pairs of economists in the professional network.

<i>Panel A</i> Transition matrix for all sequences of pairs of articles				
<i>From/To</i>	Masculine (1)	Feminine (2)	Plural (3)	Mixed (4)
Masculine	0.52	0.06	0.24	0.18
Feminine	0.18	0.31	0.24	0.26
Plural	0.28	0.09	0.49	0.14
Mixed	0.28	0.14	0.20	0.38

<i>Panel B</i> Implied stationary distributions				
	Masculine (1)	Feminine (2)	Plural (3)	Mixed (4)
Overall	0.35	0.12	0.31	0.22
Only single-authored	0.43	0.09	0.29	0.19
Only 70s cohort	0.51	0.04	0.29	0.16
Only 80s cohort	0.39	0.08	0.33	0.20
Only 90s cohort	0.33	0.12	0.31	0.24
Only 00s cohort	0.29	0.16	0.30	0.25
Only 10s cohort	0.26	0.19	0.31	0.24

Table 2: Transition matrix and stationary distributions. Panel A presents the implied transition matrix across all sequential pairs of articles. Panel B presents the implied stationary distribution based on Panel A (overall), and based on transition matrices that restrict attention to sequential single-authored pairs of articles, and sequential pairs of articles by author cohorts. The corresponding transition matrices for the single-authored and cohort groups appear in [Table A.11](#).

Fractional Multinomial Response Models			
Social network: Co-authors and cited			
Dep var: Share of articles by author i 's social network using writing style			
	<u>Masculine</u>	<u>Feminine</u>	<u>Mixed</u>
	(1)	(2)	(3)
Δz_{it}^m	1.76	0.08	1.72
	(0.06)	(0.06)	(0.09)
Δz_{it}^f	2.12	6.67	4.65
	(0.04)	(0.14)	(0.08)
Δz_{it}^x	-0.15	2.81	1.27
	(0.08)	(0.11)	(0.04)
Obs.	68,837		
Social network: Only co-authors			
Dep var: Share of articles by author i 's social network using writing style			
	<u>Masculine</u>	<u>Feminine</u>	<u>Mixed</u>
	(4)	(5)	(6)
Δz_{it}^m	1.01	0.57	1.04
	(0.06)	(0.03)	(0.04)
Δz_{it}^f	1.55	2.77	2.49
	(0.05)	(0.07)	(0.06)
Δz_{it}^x	0.59	1.40	0.97
	(0.07)	(0.06)	(0.03)
Obs.	68,837		
Social network: Only cited			
Dep var: Share of articles by author i 's social network using writing style			
	<u>Masculine</u>	<u>Feminine</u>	<u>Mixed</u>
	(7)	(8)	(9)
Δz_{it}^m	2.10	0.37	1.84
	(0.06)	(0.07)	(0.09)
Δz_{it}^f	2.22	8.22	5.44
	(0.04)	(0.13)	(0.09)
Δz_{it}^x	-0.76	3.21	1.31
	(0.07)	(0.11)	(0.05)
Obs.	68,837		

Table 3: Control Function Models of Pronoun Choice. The table presents coefficient estimates of the fractional multinomial choice conditional mean equations. The explanatory regressors measure the change in (weighted) average pronoun choice of peers of a given author's peers who are not his acquaintances. The baseline category is the plural form. The top panel considers co-authors and citees as peers. The middle panel considers only co-authors as peers. The bottom panel considers only citees as peers.

Social network: Co-authors and cited				
	<u>Masculine</u>	<u>Feminine</u>	<u>Mixture</u>	<u>Plural</u>
	(1)	(2)	(3)	(4)
z_{it}^m	0.53 (0.01)	-0.01 (0.00)	0.03 (0.00)	-0.55 (0.01)
z_{it}^f	-0.33 (0.02)	0.66 (0.01)	0.05 (0.01)	-0.39 (0.01)
z_{it}^x	-0.03 (0.01)	-0.00 (0.01)	0.64 (0.01)	-0.61 (0.01)
Authors FEs	Y	Y	Y	Y
R^2	0.62	0.62	0.60	0.48
F-statistic	578	382	337	122
Social network: Only co-authors				
	<u>Masculine</u>	<u>Feminine</u>	<u>Mixture</u>	<u>Plural</u>
	(5)	(6)	(7)	(8)
z_{it}^m	0.40 (0.01)	-0.01 (0.01)	0.05 (0.01)	-0.44 (0.01)
z_{it}^f	-0.12 (0.02)	0.48 (0.01)	0.07 (0.01)	-0.43 (0.01)
z_{it}^x	-0.03 (0.01)	0.05 (0.01)	0.48 (0.01)	-0.50 (0.01)
Authors FEs	Y	Y	Y	Y
R^2	0.40	0.37	0.39	0.41
F-statistic	102	139	72	46
Social network: Only cited				
	<u>Masculine</u>	<u>Feminine</u>	<u>Mixture</u>	<u>Plural</u>
	(9)	(10)	(11)	(12)
z_{it}^m	0.54 (0.01)	-0.00 (0.00)	0.04 (0.00)	-0.57 (0.01)
z_{it}^f	-0.48 (0.02)	0.83 (0.01)	0.09 (0.01)	-0.44 (0.01)
z_{it}^x	-0.03 (0.01)	-0.01 (0.00)	0.71 (0.01)	-0.66 (0.01)
Authors FEs	Y	Y	Y	Y
R^2	0.68	0.77	0.66	0.49
F-statistic	1003	1130	670	143
Obs.	84,434	84,434	84,434	84,434

Table 4: Robustness: Linear Models for Pronoun Choice. The table presents coefficient estimates of the within-author panel linear regression models for the four pronoun form shares. The explanatory regressors measure the (weighted) average pronoun choice of peers of a given author's peers who are not his acquaintances. The baseline category is the plural form. The top panel considers co-authors and citees as peers. The middle panel considers only co-authors as peers. The bottom panel considers only citees as peers.

Pairwise Covariate	Acquaintance Set Definition		
	$Q_{10}(i)$	$Q_5(i)$	$Q_{20}(i)$
$\underline{\gamma}$			
Same ethnicity	1.03 (0.10)	0.92 (0.08)	1.13 (0.12)
Same sex	0.24 (0.12)	0.23 (0.10)	0.25 (0.14)
Common fields	0.76 (0.05)	0.67 (0.04)	0.85 (0.06)
Cohort difference	-1.98 (0.51)	-1.64 (0.44)	-2.29 (0.61)
Citations difference	-4.86 (0.19)	-4.65 (0.17)	-4.96 (0.23)
Productivity difference	0.34 (0.49)	0.21 (0.41)	0.36 (0.58)
Log Productivity Product	0.49 (0.02)	0.47 (0.02)	0.51 (0.03)
$\underline{\Omega}$			
$\omega_{\ell\ell}$	0.18 (0.01)	0.34 (0.02)	0.09 (0.01)
$\omega_{\ell c}$	0.02 (0.01)	0.05 (0.01)	0.01 (0.01)
ω_{cc}	0.06 (0.02)	0.11 (0.03)	0.03 (0.02)
Rubinstein-type share	0.56	0.55	0.57

Table 5: Community Detection Estimates. The table presents maximum likelihood estimates of the covariates-adjusted stochastic block model for community detection (Feng et al., 2023). The first column presents results under the ten-closest acquaintance set definition. The second column presents results under the five-closest acquaintance set definition. The third column presents results under the 20-closest acquaintance set definition. The model is estimated on the 29,302 authors who co-authored at least once.

Panel A: Parameters			
	Masculine	Feminine	Mixed
	(1)	(2)	(3)
α (Intercepts)	-1.04 (0.06)	-0.74 (0.04)	-0.60 (0.04)
δ (Osborne-type dummy)	-0.36 (0.04)	0.02 (0.13)	0.88 (0.07)
λ (Control function)	0.13 (0.06)	0.28 (0.03)	0.14 (0.03)
ω (Bargaining power)			
Age diff.		1.10 (0.62)	
Citations diff.		-1.19 (0.89)	
Productivity diff.		-2.30 (0.59)	
Sex diff.		0.01 (0.17)	
Same ethnicity		-0.14 (0.11)	
Panel B: Peer effect heterogeneity			
<i>Community</i>	Sex	μ	σ
Osborne	Men	1.63 (0.16)	0.03 [0.003 , 0.31]
	Women	1.92 (0.17)	
Rubinstein	Men	1.54 (0.14)	0.02 [0.002 , 0.15]
	Women	1.55 (0.16)	
Observations		56,239	

Table 6: Parameter Estimates of the Writing Style Model. The table presents the parameter estimates from the multinomial choice model that considers both (weighted) past co-authors and past citees as peer influences, estimated using simulated maximum likelihood. The corresponding time effects are reported in Figure 7. The first four rows report choice-specific parameters. The parameters on the bargaining weights and the peer effect distributions are common across choices. The table reports standard errors for all parameter estimates except for the standard deviations of the peer effect distributions. For those we report confidence intervals that rely on the delta-method.

Percentage points difference relative to baseline share in 2019						
	Baseline (1)	Freeze to 1970 (2)	No Peer Effects (3)	Own cohort peers (4)	Freeze 1970 types (5)	No Homophily (6)
Masculine	0.19	0.60	-0.04	-0.00	0.04	0.04
Feminine	0.23	-0.23	-0.01	-0.01	-0.03	-0.04
Mixed	0.26	-0.24	-0.01	0.02	-0.01	0.00
Plural	0.32	-0.14	0.07	-0.01	0.00	-0.00

Table 7: Simulated end-line writing style shares under alternative scenarios. The table presents the difference in the average distributions of pronoun form shares by 2019 relative to the baseline simulation using the estimated parameters from Table 6, under alternative scenarios. The first column presents the end point distribution under the baseline stationary simulation. The second column freezes the external trends to 1970 ($\varphi_t^p = \varphi_{1970}^p$). The third column supposes no peer effects ($\beta_i = 0$). The fourth column restricts peer influences to exist only between members of the same cohort. The fifth column freezes the sex and community-type distributions to their 1970-73 averages. The last column supposes no homophily in co-authoring ($\omega = 1$). In all simulations, the observed co-authorships are held fixed.

	Dependent variable: Only masculine pronouns dummy					
	Overall		With female author(s)		With only male authors	
	(1)	(2)	(3)	(4)	(5)	(6)
Exposure	0.16 (0.07)	-0.07 (0.12)	0.75 (0.20)	1.59 (0.61)	0.08 (0.08)	-0.15 (0.12)
Year FEs	Y	Y	Y	Y	Y	Y
Journal FEs	Y	Y	Y	Y	Y	Y
Author FEs	N	Y	N	Y	N	Y
Obs.	10,918	6,804	1,465	519	9,453	6,013

Table 8: Exposure to Female Editors. The table presents linear probability models at the article level estimated by OLS, on the sub-sample of articles published by the authors from our theorists professional network in one of ten major Economics journals (top-5 general interest and the top 5 economic theory journals based on *RePEC*'s rankings of September 2023). The exposure variable is the average number of female editors of in the board of a journal, over the three years period prior to an article's publication date. Besides year and journal fixed effects, odd-numbered columns also include the date of first publication, the total number of publications, the total number of citations, the ethnicity, and the community assignment (Osborne/Rubinstein) of each author. Columns (1) and (2) include all articles in any of the ten journals. Even-numbered columns include author fixed effects instead. Columns (3) and (4) only include articles with at least one female author. Columns (5) and (6) only include articles with both male authors.

Dep. Var.	Masculine		Feminine		Plural		Mixed	
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
First 6 Years	-0.022 (0.004)	-0.011 (0.007)	0.018 (0.002)	-0.005 (0.004)	0.022 (0.004)	0.023 (0.007)	-0.018 (0.003)	-0.007 (0.006)
Log(Rank)	0.006 (0.002)	0.002 (0.002)	-0.001 (0.001)	0.001 (0.001)	0.016 (0.001)	0.008 (0.002)	-0.022 (0.001)	-0.011 (0.002)
Top 5 Journal	0.019 (0.006)	0.006 (0.008)	0.005 (0.003)	-0.002 (0.004)	-0.073 (0.006)	-0.020 (0.007)	0.049 (0.005)	0.016 (0.007)
Year FEs	Y	Y	Y	Y	Y	Y	Y	Y
Author FEs	N	Y	N	Y	N	Y	N	Y
Obs.	66,533	48,632	66,533	48,632	66,533	48,632	66,533	48,632

Table 9: Alternative Drivers of Writing Style Choices. The table presents coefficient estimates from linear probability models at the article level, separately regressing dummy variables for each type of writing style on three different variables. In the first row we report results for models that include a dummy variable equal to 1 if at least one author is, at the time of publishing the paper, at most 6 years since his first publication, as a proxy for the tenure track period. In the second row we report results for models that include the log rank of the journal where the article was published, based on the most recent ranking here: www.researchbite.com. It combines an h-index, an impact score, and the SJR score. In the third row we report results for models that include a dummy variable equal to 1 if the journal where the article was published is either *Econometrica*, *The Review of Economic Studies*, *The Journal of Political Economy*, *The American Economic Review*, or *The Quarterly Journal of Economics*. Odd columns present results for models without author fixed effects. Even columns present results for models with author fixed effects instead.

	<u>Masculine</u>		<u>Feminine</u>		<u>Plural</u>		<u>Mixed</u>	
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
Analysis of Collective Decision-Making	0.038 (0.011)	0.018 (0.017)	0.053 (0.009)	-0.001 (0.012)	-0.234 (0.010)	-0.067 (0.014)	0.144 (0.011)	0.050 (0.017)
Distribution	0.001 (0.027)	0.006 (0.038)	0.044 (0.021)	0.023 (0.027)	-0.012 (0.029)	-0.030 (0.039)	-0.033 (0.020)	0.002 (0.034)
Financial Economics	0.027 (0.007)	0.005 (0.011)	-0.012 (0.004)	0.001 (0.006)	-0.035 (0.007)	-0.004 (0.010)	0.020 (0.006)	-0.003 (0.009)
Game Theory	0.072 (0.008)	-0.012 (0.013)	0.037 (0.006)	-0.005 (0.009)	-0.232 (0.007)	-0.030 (0.011)	0.123 (0.008)	0.047 (0.012)
General Equilibrium	0.090 (0.016)	0.022 (0.021)	0.019 (0.011)	0.006 (0.014)	-0.087 (0.016)	-0.011 (0.020)	-0.023 (0.012)	-0.018 (0.017)
Household Behavior and Family Economics	-0.053 (0.027)	-0.057 (0.046)	0.016 (0.023)	-0.046 (0.035)	-0.096 (0.031)	-0.042 (0.045)	0.133 (0.030)	0.145 (0.044)
Information, Knowledge, and Uncertainty	0.052 (0.009)	0.001 (0.014)	0.015 (0.007)	-0.003 (0.010)	-0.239 (0.008)	-0.064 (0.011)	0.172 (0.009)	0.066 (0.014)
Market Structure, Pricing, and Design	-0.005 (0.008)	0.012 (0.013)	0.012 (0.006)	-0.004 (0.009)	-0.062 (0.009)	-0.010 (0.013)	0.055 (0.007)	0.001 (0.012)
Micro-Based Behavioral Economics	-0.032 (0.023)	0.010 (0.041)	0.046 (0.020)	0.010 (0.032)	-0.147 (0.024)	-0.052 (0.034)	0.132 (0.025)	0.032 (0.046)
Production and Organizations	-0.020 (0.017)	0.055 (0.027)	-0.032 (0.011)	-0.018 (0.015)	-0.040 (0.018)	-0.038 (0.024)	0.091 (0.016)	0.002 (0.025)
Welfare Economics	0.026 (0.012)	-0.003 (0.017)	0.067 (0.010)	0.021 (0.012)	-0.177 (0.011)	-0.048 (0.016)	0.084 (0.012)	0.030 (0.016)
Year FEs	Y	Y	Y	Y	Y	Y	Y	Y
Author FEs	N	Y	N	Y	N	Y	N	Y
Obs.	66,533	48,637	66,533	48,637	66,533	48,637	66,533	48,637

Table 10: Differences in pronoun style choice by sub-fields. The table presents estimates from linear probability models at the article level, separately regressing dummy variables for each type of writing style on sub-field indicators. Odd columns present results for models without author fixed effects. Even columns present results for models with author fixed effects instead.

9 Figures

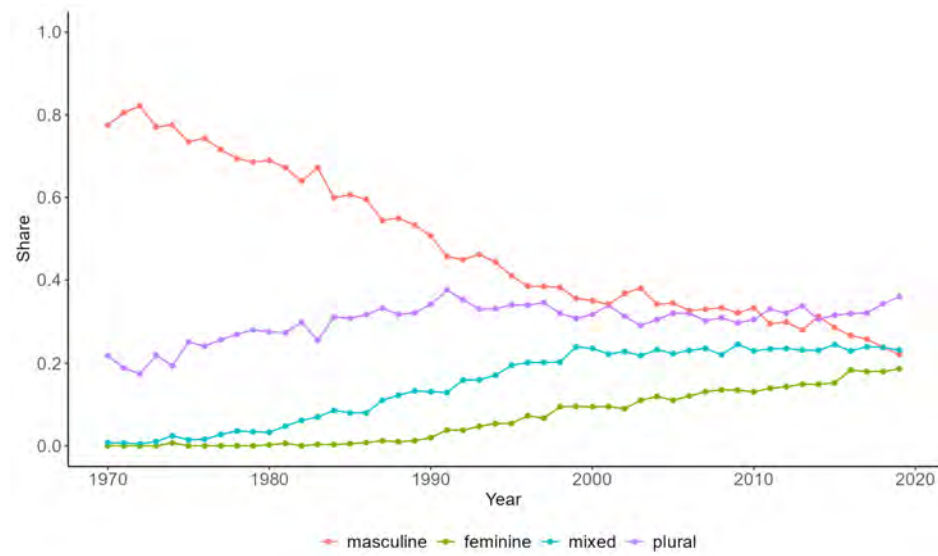
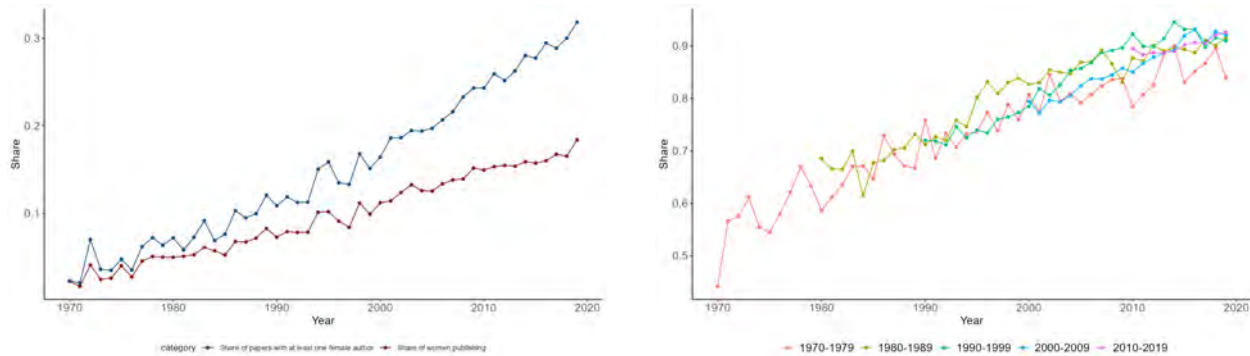


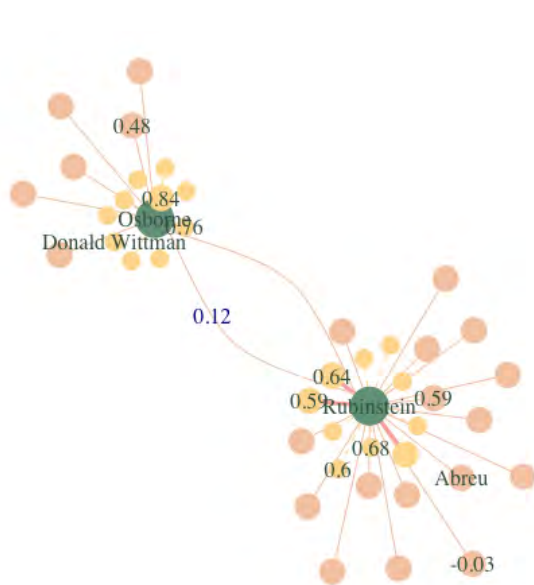
Figure 1: Distribution of pronoun use over time in theory papers, 1970-2019.



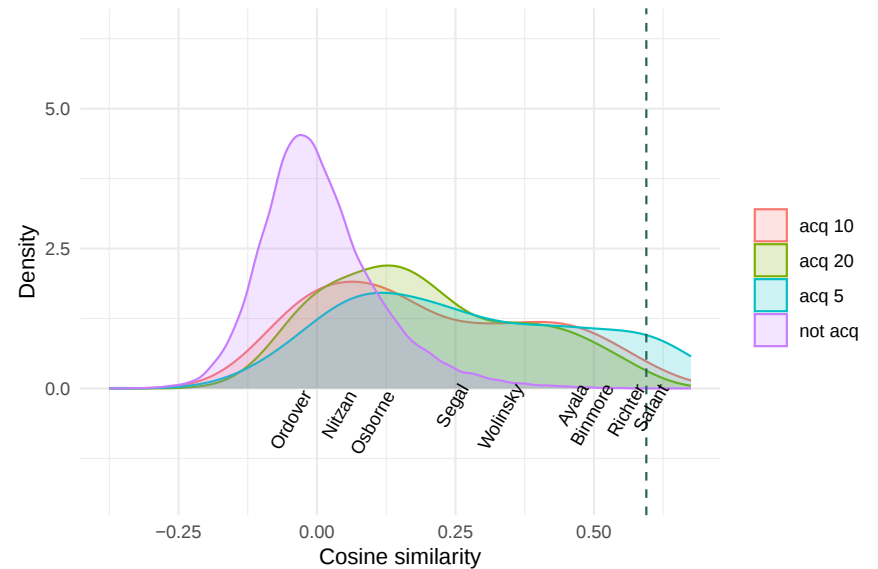
(a) Participation of women in the economic theory.

(b) Share of co-authored papers, by cohort.

Figure 2: Long-term change in the economics profession.

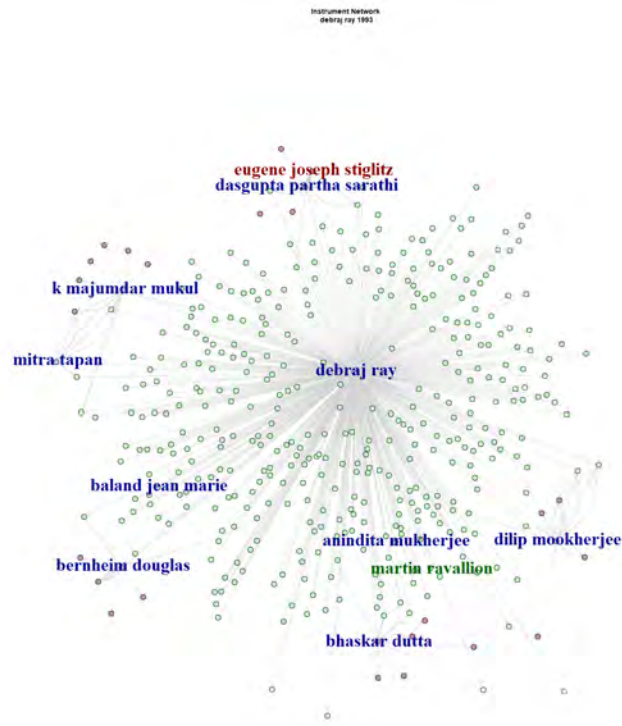


(a) Osborne and Rubinstein's local co-author network.

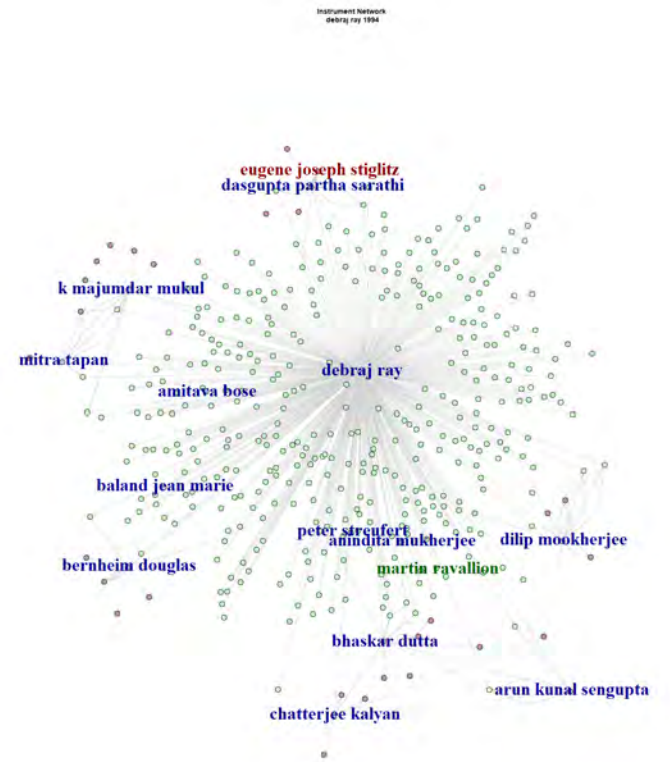


(b) Distribution of academic cosine similarity between Ariel Rubinstein and all other economists.

Figure 3: Illustration: Ariel Rubinstein's and Martin Osborne's local peer network, and distribution of Ariel Rubinstein's academic similarities. In panel (a), solid edges represent co-authorships. Dashed edges represent acquaintances who are not co-authors. Yellow circles represent each author's ten closest authors in academic cosine similarity. The lengthier edges represent longer distances. In panel (b), A subset of Rubinstein's co-authors are marked along the x axis by their names; the density of his non-acquaintances appears in pink; the densities of his acquaintance sets appear in blue ($n = 5$), red ($n = 10$), and green ($n = 20$). The vertical dashed line represents the location of Rubinstein's tenth most similar author.



(a) Debraj Ray's network, 1993



(b) Debraj Ray's network, 1994

Figure 4: Example network and instrumental variables variation. The figure illustrates the instrumental variables variation induced by co-authors of co-authors who are not acquaintances of an author, for the case of Debraj Ray in 1993 and 1994. His co-authors appear in blue, his acquaintances appear in green, and non-acquaintances appear in pink.

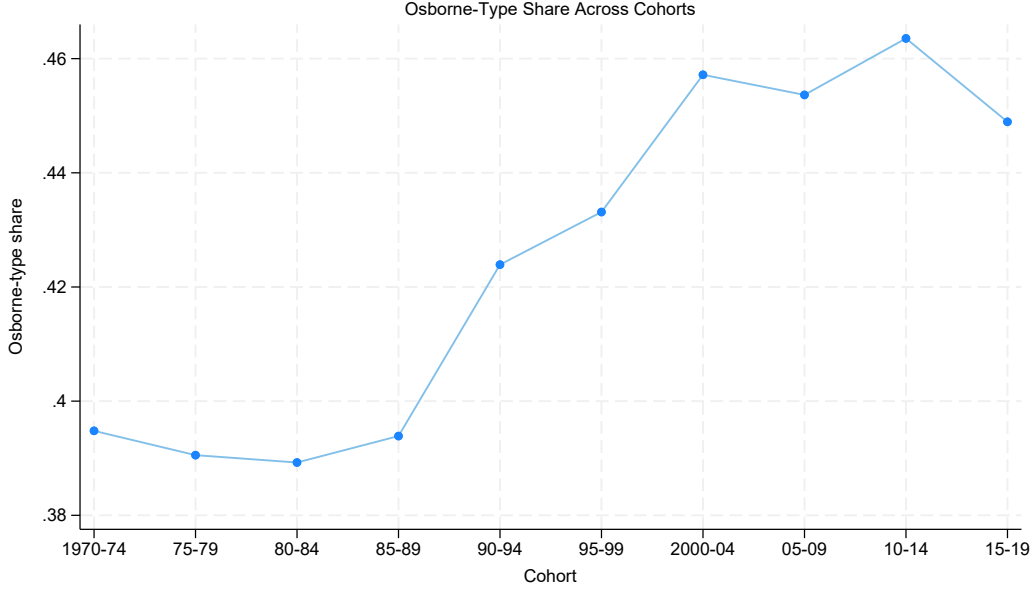


Figure 5: Osborne Type Share across Cohorts. Share of authors assigned to Osborne's community, by 5-year cohorts of economists based on the community detection estimates based on the ten-closest acquaintance set definition.

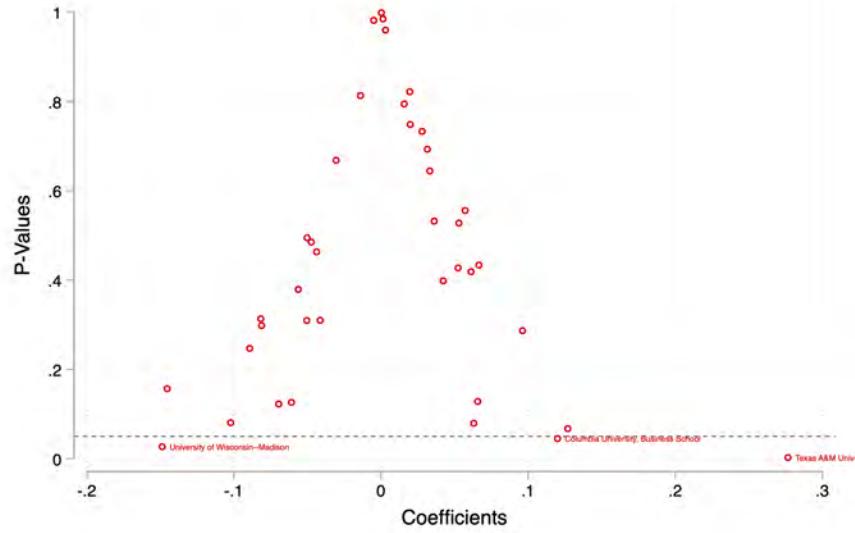


Figure 6: University affiliations and the Osborne-type dummy. Distribution of coefficient sizes and p-values by university to predict the Osborne-type dummy in a regression of 1,868 unique authors in 39 academic departments and 2,592 authors-x-department of the form: $\text{Osborne type dummy}_i = a + \beta \text{University } j \text{ dummy}_i + \epsilon_i$. The dashed line represents a p-value of 0.05.

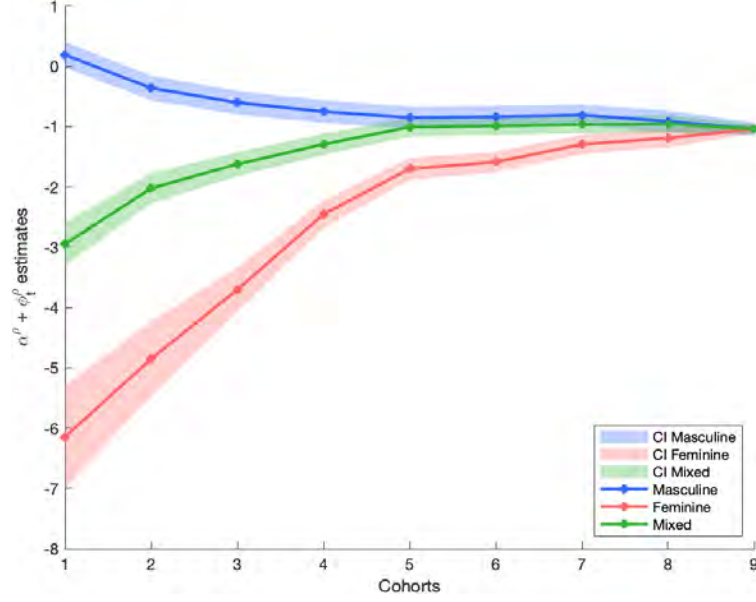


Figure 7: Time effects. Point estimates and 95% confidence bands for the writing-style-specific time-effect coefficients φ_t^ρ (plus the corresponding intercept α^ρ). We define time periods based on publication years, and group them as follows: 1970-79, 1980-84, 1985-89, 1990-94, 1995-99, 2000-04, 2005-09, 2010-14, 2015-19.

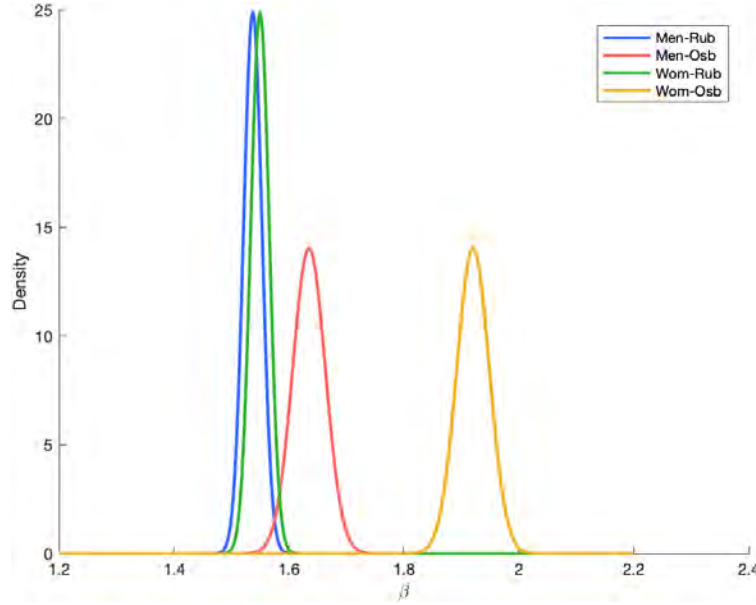
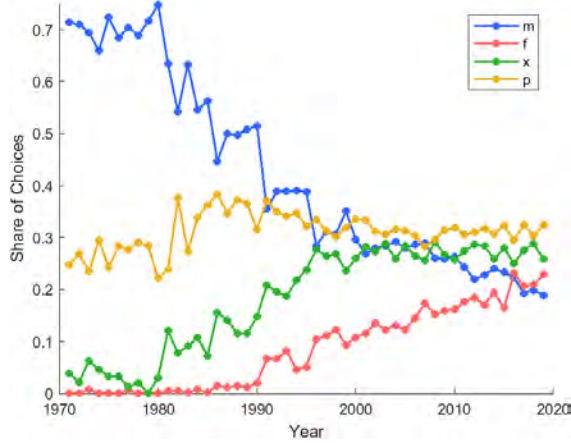
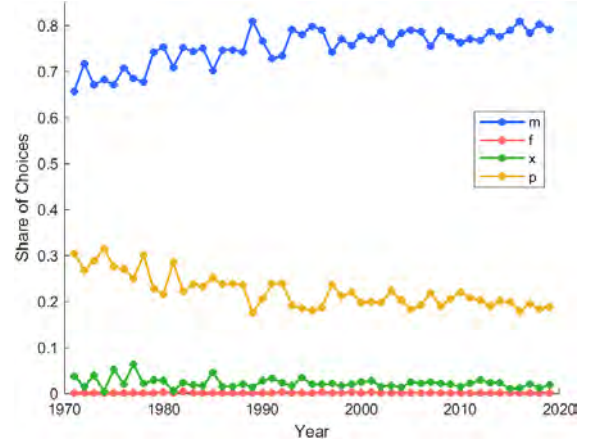


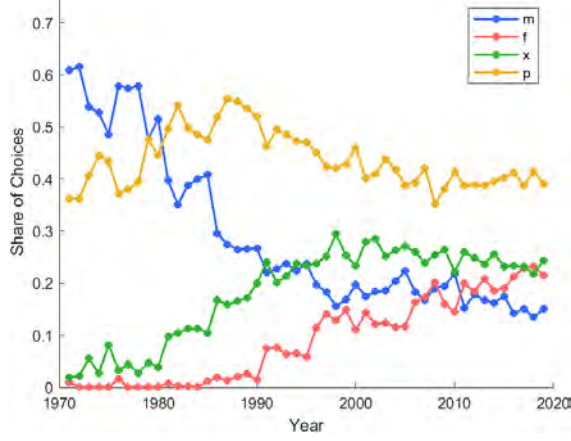
Figure 8: Peer effects. Estimated Normal distributions of peer effect heterogeneity for Osborne and Rubinstein-type men and women economists. $\mu(\text{Rub, men}) = 1.54$ (std.err = 0.14); $\mu(\text{Rub, women}) = 1.55$ (std.err = 0.16); $\mu(\text{Osb, men}) = 1.63$ (std.err = 0.16); $\mu(\text{Osb, women}) = 1.92$ (std.err = 0.17); $\sigma(\text{Rub}) = 0.02$ (confidence interval = [0.002, 0.15]); $\sigma(\text{Osb}) = 0.03$ (confidence interval = [0.003, 0.31]).



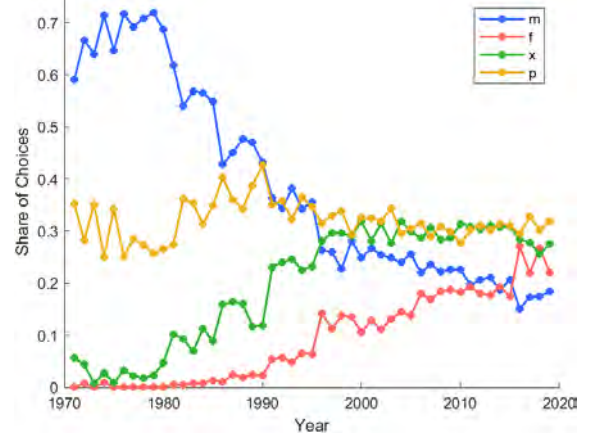
(a) Model Fit: Baseline simulation



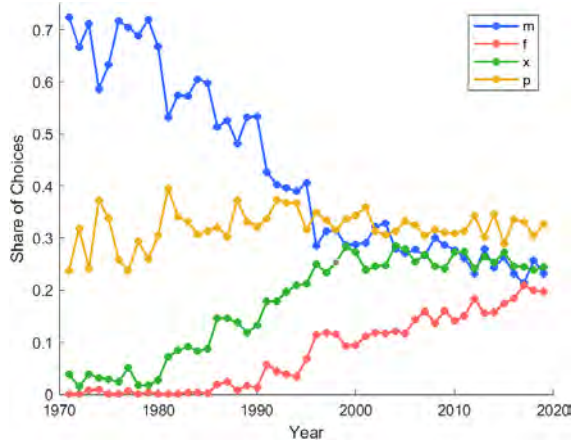
(b) External influences frozen: $\varphi_t^p = \varphi_{1970}^p$



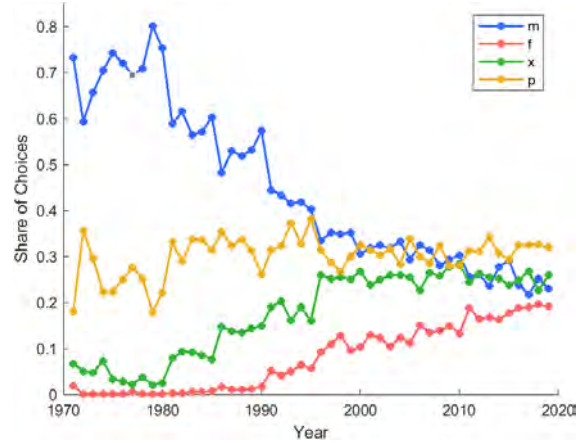
(c) No peer effects: $\beta^p = 0$



(d) Peer influences only from own cohort



(e) Type composition frozen



(f) No homophily

Figure 9: Simulated distribution of writing style choices over time: Alternative scenarios. The figures plot the time evolution of the aggregate distribution writing styles from simulated choices based on the estimated parameters from Table 6 and alternative assumptions. As starting values for the peer influences, the simulation uses the observed average 1970-1974 distribution of choices.

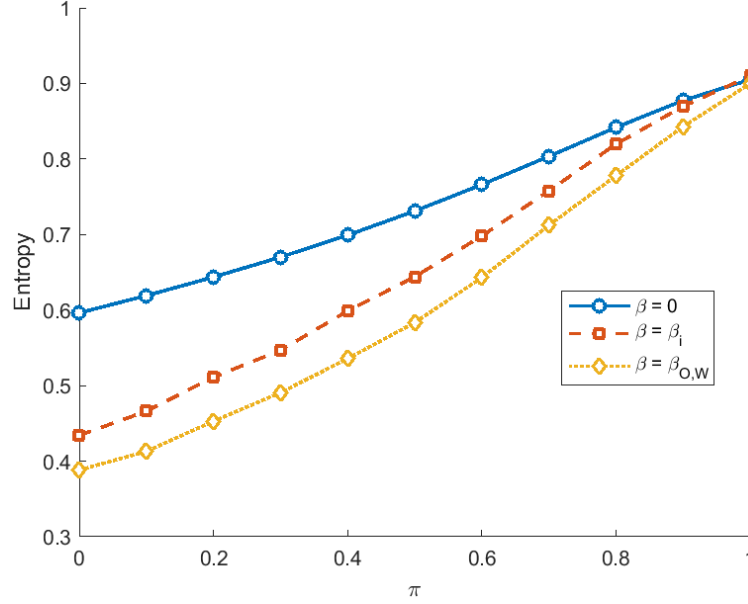


Figure 10: Entropy Simulations. The figure plots the entropy of the distribution of choice probabilities averaged over 2017-2019 under alternative scenarios for the societal trend strength (horizontal axis) and peer effects strength (colored curves). The red curve corresponds to simulations using the estimates peer effects. The blue curve corresponds to simulations where peer effects are shut down. The yellow curve corresponds to simulations where peer effects for all economists are as strong as the estimated mean peer effects on Osborne-community women.

10 Supplemental Appendix I (Online): Additional Tables and Figures

<i>From/To</i>	<u>Masculine</u> (1)	<u>Feminine</u> (2)	<u>Plural</u> (3)	<u>Mixed</u> (4)
<i>Single authored</i>				
Masculine	0.63	0.03	0.21	0.13
Feminine	0.14	0.43	0.20	0.23
Plural	0.32	0.07	0.50	0.11
Mixed	0.28	0.11	0.19	0.42
<i>70s cohort</i>				
Masculine	0.65	0.02	0.22	0.10
Feminine	0.22	0.27	0.21	0.30
Plural	0.41	0.04	0.44	0.12
Mixed	0.33	0.07	0.24	0.35
<i>80s cohort</i>				
Masculine	0.54	0.04	0.25	0.17
Feminine	0.19	0.27	0.27	0.27
Plural	0.31	0.07	0.49	0.13
Mixed	0.31	0.11	0.22	0.36
<i>90s cohort</i>				
Masculine	0.48	0.07	0.24	0.20
Feminine	0.19	0.27	0.25	0.29
Plural	0.25	0.10	0.50	0.15
Mixed	0.28	0.14	0.20	0.38
<i>00s cohort</i>				
Masculine	0.46	0.09	0.23	0.22
Feminine	0.17	0.33	0.24	0.25
Plural	0.23	0.13	0.50	0.15
Mixed	0.26	0.17	0.18	0.39
<i>10s cohort</i>				
Masculine	0.41	0.12	0.24	0.23
Feminine	0.16	0.38	0.22	0.24
Plural	0.21	0.13	0.52	0.14
Mixed	0.24	0.19	0.18	0.39

Table A.11: Sub-group transition matrices. Transition matrices for single-authored to single-authored papers, and for different cohorts of authors corresponding to Panel B of [Table 2](#).

	Osborne dummy
Woman	-0.021 (0.053)
Woman $\times$ 1980	-0.050 (0.061)
Woman $\times$ 1990	-0.009 (0.057)
Woman $\times$ 2000	-0.001 (0.055)
Woman $\times$ 2010	0.003 (0.055)
Obs.	29302

Table A.12: The table presents the coefficients and standard errors from a cross-sectional linear regression at the author level, of the Osborne dummy on a dummy for whether the author is a woman, and interactions of it with cohort dummies.

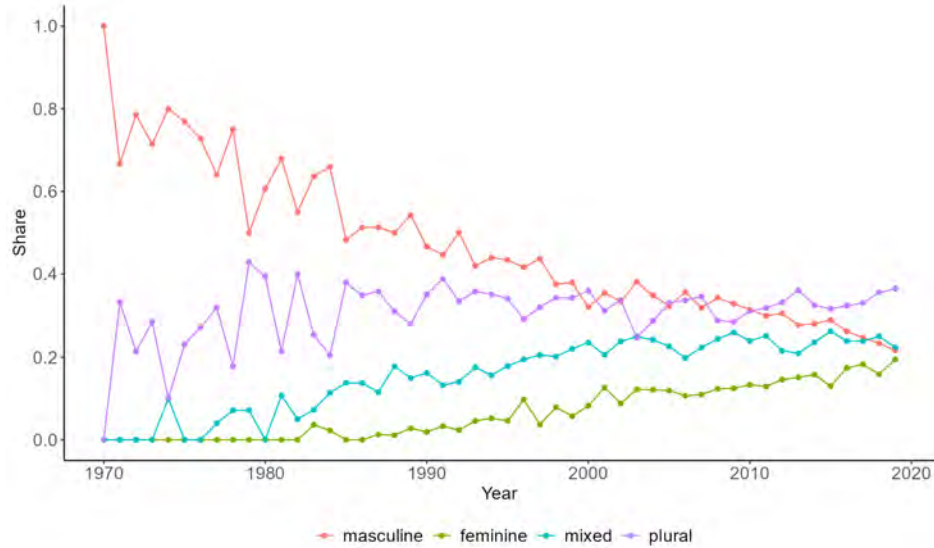
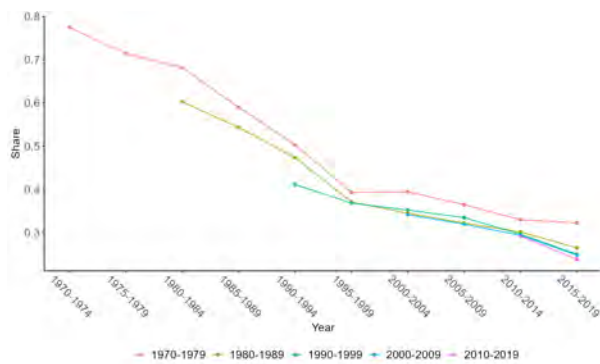
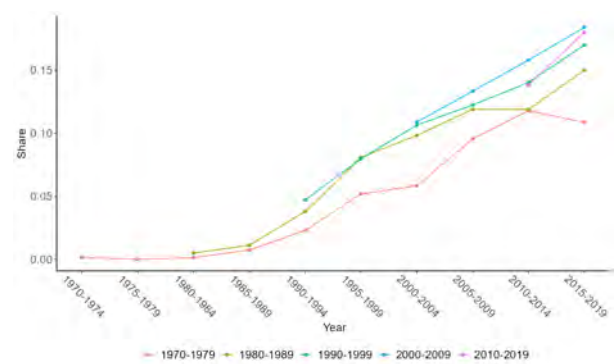


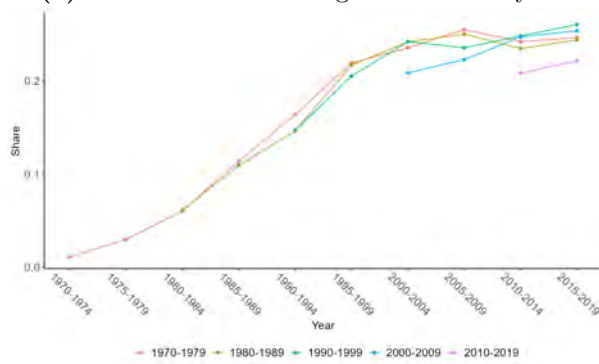
Figure A.11: Distribution of pronoun use over time for papers authored by women, 1970-2019.



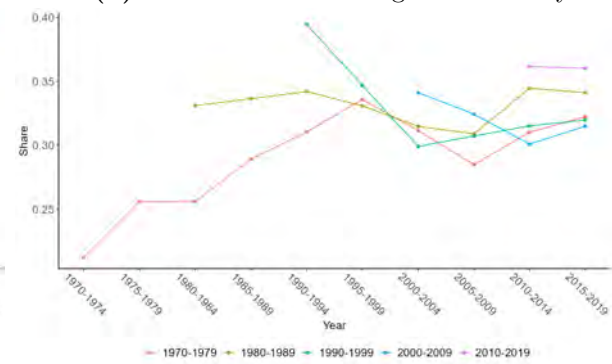
(a) Share of authors using masculine only.



(b) Share of authors using feminine only.

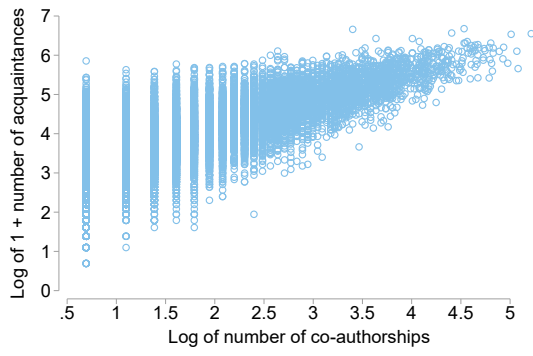


(c) Share of authors using mixed.

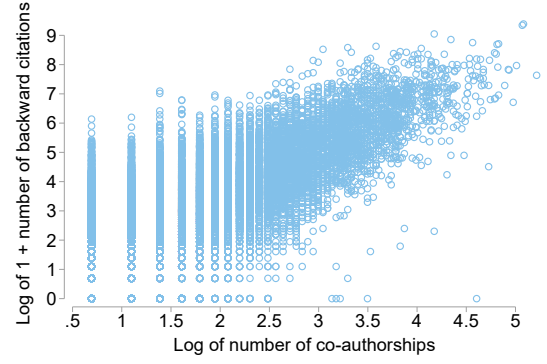


(d) Share of authors using plural only.

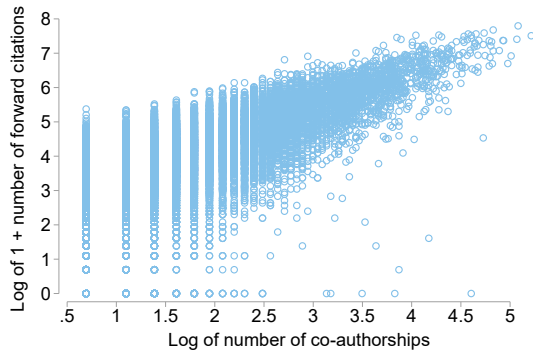
Figure A.12: Distribution of pronoun use over time, by cohorts.



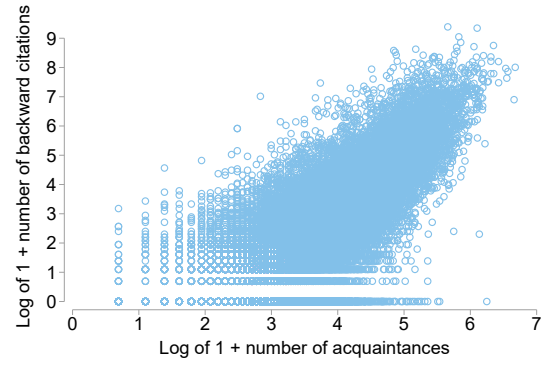
(a)



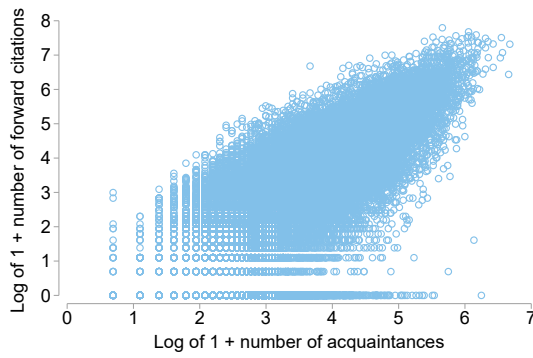
(b)



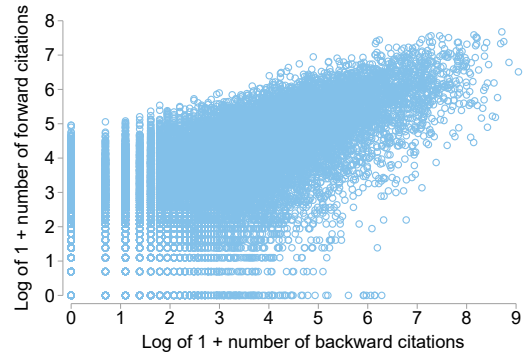
(c)



(d)

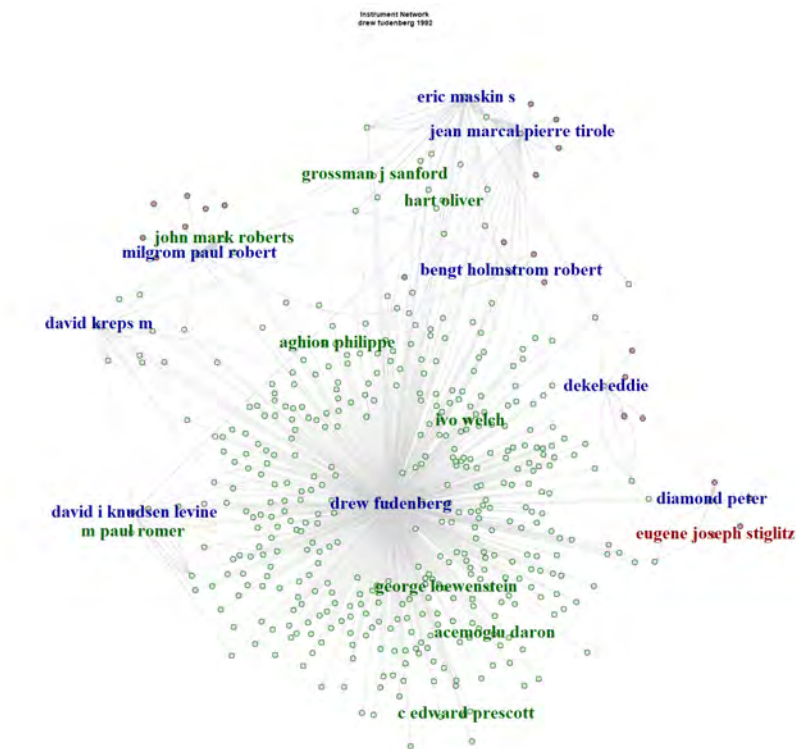


(e)

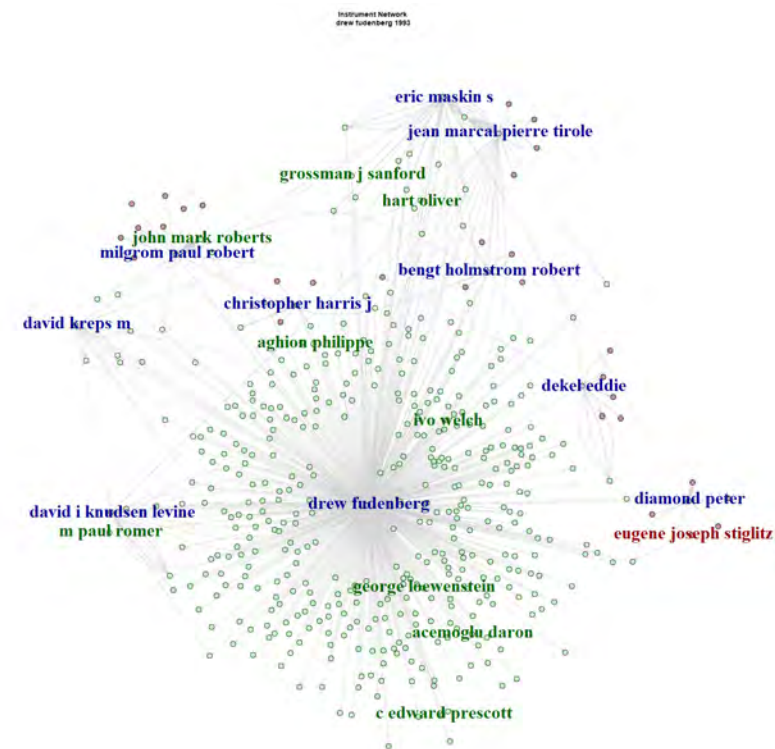


(f)

Figure A.13: Degree distributions across networks.



(a) Drew Fudenberg's network 1992



(b) Drew Fudenberg's network 1993

Figure A.14: Example network and instrumental variables variation. The figure illustrates the instrumental variables variation induced by co-authors of co-authors who are not acquaintances of an author, for the case of Drew Fudenberg in 1992 and 1993. His co-authors appear in blue, his acquaintances appear in green, and non-acquaintances appear in pink.

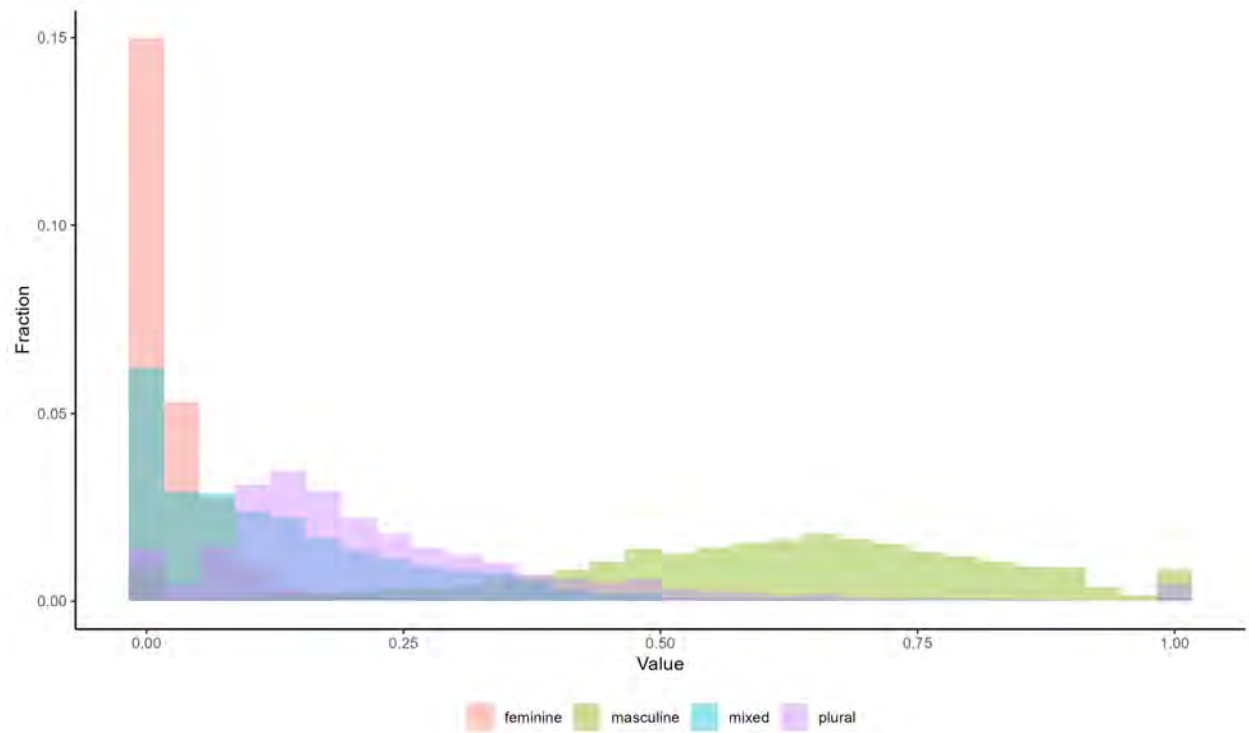
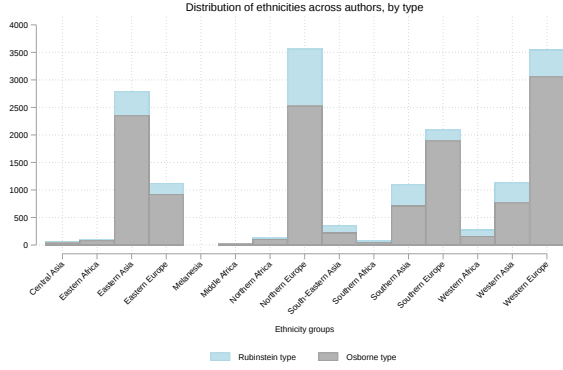
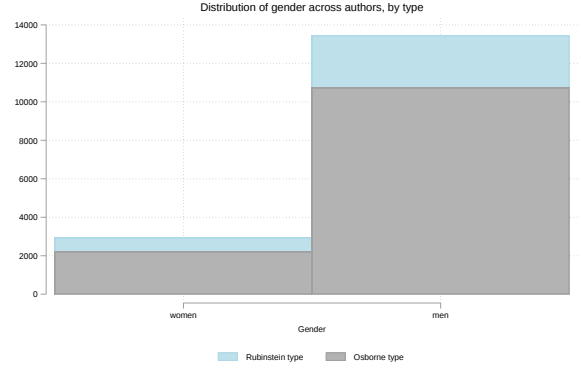


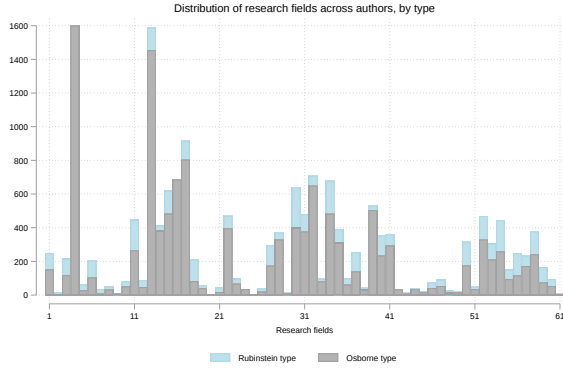
Figure A.15: Distribution of the endogenous regressors $r_{a(ij)t}^{\rho}$. The figure plots the distribution of the peer choice regressors for the four different writing styles.



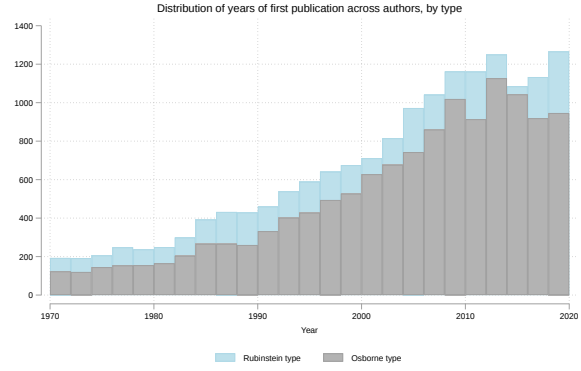
(a) Ethnicity



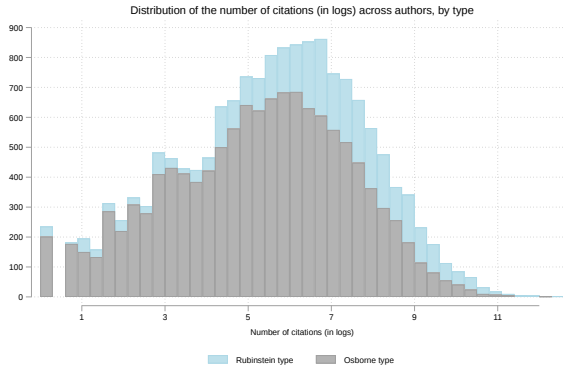
(b) Gender



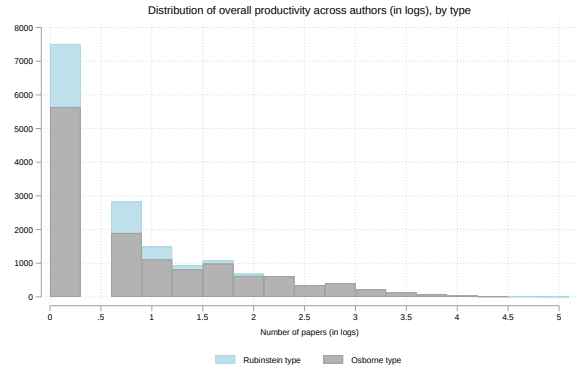
(c) Fields



(d) First publication



(e) Citations



(f) Productivity

Figure A.16: Distributions of author characteristics by community type assigned.

11 Supplemental Appendix II (Online): Methodological Details

11.1 Selection of the sample of articles and authors

We use several sources to put together the sets of articles and authors that underlie our study. From *Jstor* and *Crossref* we obtained the metadata and the full texts of a large set of papers from Economics and Economics-related academic journals. We obtained the *Jstor* data under a data user agreement for the project. We obtained the *Crossref* data using their API.³⁵

This resulted in 710 thousand articles. The set is over-inclusive, however. It contains papers in all fields of Economics, whereas our purpose is to put together a set of economic theory articles only. We implement a layered procedure to filter out articles unlikely to be theoretical, and to make sure we keep articles likely to be theoretical.

1. We exclude articles with corrupted metadata:

- Missing a title.
- Missing authors.
- Missing the articles' text. These are articles for which our *Crossref* API retrieval generated a line of metadata but no associated article text. We inspected the list titles of this set of articles, and found 849 that we clearly identified as economic theory papers. We proceeded to directly retrieve the text of these articles, and included them back.

2. We exclude any article whose metadata suggests it is not a standard academic paper. This includes a reference to any of the following labels:

"Note from the editor"	"Photograph"
"Meeting of the econometric society"	"Meetings of the econometric society"
"Accepted Manuscripts"	"List of members"
"Announcement"	"Announcements"
"Award"	"Awards"
"Front matter"	"Back Matter"
"Book review"	"Book reviews"
"Call for papers"	"Distinguished fellow"
"Referees"	"Editorial"
"Editor"	"Election of fellows"
"Errata"	"Erratum"
"Addendum"	"Correction:"
"Correction to:"	"Retracted Article"

³⁵See <https://www.crossref.org/education/retrieve-metadata/rest-api/>. We used the R package `crminer` to retrieve the data. This package is no longer maintained, and to our knowledge, Crossref discontinued its open-access full-text retrieval service as of December 2020 -after we accessed it-.

"Corrigendum"	"European meeting"
"Fellows"	"Foreward"
"In memoriam"	"Obituary"
"Report of the committee"	"Report on the adhoc committee"
"Report of the director"	"Report of the editor"
"Report of the managing editor"	"Report of the representative"
"Report of the secretary"	"Report of the treasurer"
"Submission"	"Report of the President"
"Thesis titles"	"Author index"
"Discussion"	"Preface"
"Foreword"	"Index"
"Comment"	"Contributors"
"Abstracts"	"Noticeboard"
"IMACS"	"Reply"
"Note"	"Rejoinder"
"Presidential address"	"Hardback"
"Hardcover"	"Paperback"
"Actuarial Vacancy"	"Secretary-Treasurer"
"Secretary/Treasurer"	"Treasurer"
"ISBN"	"pp\\\"."
"Conference"	"Symposium"
"Verlag"	"pages"
"Tribute"	"(Eds) "
"Listing Service"	"Content of Volume"
"Contents of Volume"	

3. We exclude all articles from academic journals that are either exclusively econometric or statistical, or from unrelated fields. Below is the list of journals whose articles we exclude:

"Econometric Theory"
 "Econometrics Journal"
 "Journal of Applied Econometrics"
 "Journal of Econometrics"
 "Physica A: Statistical Mechanics and its Applications"
 "Statistics & Probability Letters"
 "Stochastic Processes and their Applications"
 "Applied Energy"
 "Energy"
 "Resources and Energy"
 "Renewable Energy"
 "The Electricity Journal"
 "Marine Policy"
 "Computational Statistics & Data Analysis"
 "Mitigation and Adaptation Strategies for Global Change"

"Journal of Classification"
 "World Patent Information"
 "Social Indicators Research: An International and Interdisciplinary
 Journal for Quality-of-Life Measurement"
 "Journal of Multivariate Analysis"
 "Metrika: International Journal for Theoretical and Applied Statistics"
 "Statistical Papers"
 "Annals of the Institute of Statistical Mathematics"
 "Journal of the Royal Statistical Society Series A"
 "Journal of the Royal Statistical Society. Series C (Applied Statistics)"
 "Journal of Time Series Analysis"
 "Statistical Methods & Applications"
 "Applied Mathematics and Computation"
 "Mathematics and Computers in Simulation (MATCOM)"
 "Global Finance Journal"
 "Children and Youth Services Review"
 "European Journal of Operational Research"
 "Mathematical Methods of Operations Research"
 "Mathematics of Operations Research"

4. We directly included in our final set all articles from strictly economic theory journals:

"Journal of Economic Theory"
 "American Economic Journal: Microeconomics"
 "Economic Theory"
 "Games and Economic Behavior"
 "International Journal of Game Theory"
 "Games"
 "Journal of Public Economic Theory"

5. For all other articles which had not been filtered out at this stage, we implement an algorithm to classify them as likely theoretical. For this purpose, we constructed a list of microeconomics keywords and a list of econometrics keywords.

The list of microeconomics keywords is:

game, player, utility, coalition, equilibrium,
 equilibria, rational, preference, core, Bayesian,
 pricing, welfare, marginal cost, theoretic, induction,
 signalling, strategic, bargaining, proposal, dynamic,
 Markov, subgame, monopoly, duopoly, oligopoly, cooperation,
 free rid, punish, design, contract, first best, second best,
 model, theory, theories, theoretical, auction, bid,
 dominance, risk, payoff dominant, backward induction, Cournot,
 Stackelberg, Nash, Aumann, unique, existence, multiplicity,

pure, mixed, coordination, hawk, dove, battle of the sexes,
 battle of the sex, matching pennies, prisoner, efficient,
 efficiency, evolutionary, replicator, dynamics, stable,
 opponent, ambiguity aversion, strategies, payoffs,
 expected utility, common knowledge, match, beliefs, intuitive
 criterion, fixed point, delay, market design, zero-sum,
 n-person, linear programming, Marshallian, compensated variation,
 transitive, transitivity, club, Rules of thumb, rule of thumb,
 Shapley value, Axiom, Axiomatic, Normal form, Extensive form,
 Information set, Impossibility, Information structure,
 private information, asymmetric information, moral hazard,
 adverse selection, surplus, incentive constraint,
 participation constraint, transferable utility, quasi-linear

The list of econometrics keywords is:

estimator, instrument, asymptotic variance, regression,
 two stage least square, maximum likelihood,
 generalized method of moments, multiple test, delta method,
 continuous mapping theorem, measurement error, moment condition

- (a) We include any paper containing at least 250 microeconomics keywords and no econometrics keywords.
- (b) We include any paper satisfying all of the following criteria:
 - Contains the word *proof* in its text.
 - Contains at least ten microeconomics keywords.
 - Contains ten times more microeconomics keywords as econometrics keywords.
- (c) We then identify all authors from papers from (a) and (b), and among the remaining not-yet-included papers, we include those which satisfy both of the following conditions:
 - It includes authors from this list.
 - it has ten times more microeconomics theory keywords as econometrics specific keywords, or has zero econometrics specific keywords.

This concludes the first component of the selection of papers into our sample, and yields 70062 articles written by 48626 authors.

6. At this stage, some of these 48626 author names correspond to differing spellings of the name of the same underlying author. We implemented an algorithm to find the alternative spellings of the same author, to then collapse these alternative spellings into a single author. First, we compute the frequencies of each name component (e.g., a first name, a last name, etc.) among all author names. We also extract the initials of each full name. We then identify, for each author, his least common name component (we call it the rare component). For example, for *Jean Marcal Tirole*, *Marcal* is its rare

component, as its frequency is the smallest among the three components of this name. Next we split the sample of author full names into two sets. A set A of authors whose rare component is unique in the data set, and none other of the components of their names are a rare component of any other author, and a set B with its complement.

The uniqueness of at least one word in the names of authors in set A implies they are highly unlikely to have duplicates. Set A has 8137 authors. In contrast, authors in set B have a rare component that is not unique in the data set. For each author $i \in B$, we produce a list of potential duplicates $D(i) = \{j, k, \dots\} \subset A \cup B$ containing the author identifiers of each author sharing i 's rare component. We then compare the initials of i 's name to the initials of the names of every element of this potential match list to thin out these lists as follows: if j 's initials are not a subset, a super set, or identical to the initials of i , we exclude j from the list. If the resulting match list for i is empty, we consider i to have no duplicates and hence to be unique. We identify 18413 authors as unique in this step.

For authors i for whom this procedure yields non-empty potential match sets $D(i)$, we further make pairwise comparisons of each of the name components of i to each of the name components of j with overlapping initials. If there is not at least one identical pairing among all these comparisons, we exclude j from the list in an additional thinning step. If the resulting match list for i is empty, we consider i to have no duplicates and hence to be unique. We identify 6336 authors as unique in this step. This leaves us with $15740 = 48626 - 8137 - 18413 - 6336$ author names i with potential duplicates $D(i)$, with corresponding initials and at least one identical name component from set B .

We then move to compare them to their potential duplicates using information about their articles. To do this, we first take the titles of the articles of each author i , and retrieve *ChatGPT* embeddings for each title separately, $\mathbf{e}_{ia}$, and for the grouping of all the titles of the author's articles, $\tilde{\mathbf{e}}_i$. For each pair of potential duplicate authors we compute the cosine similarity between each pairing of their articles and find the highest of these cosine similarities, s_{ij}^{max} . For each pair of potential duplicate authors we compute the cosine similarity between their grouped-titles embedding, $\tilde{s}_{ij}$. We then apply the following rule:

- (a) If authors i and j share the same rare component (stronger signal), and $\min\{s_{ij}^{max}, \tilde{s}_{ij}\} \geq 0.8$, consider i and j to be the same author.
- (b) If authors i and j share a name component that is not the rare one for one of the authors (weaker signal), and $\min\{s_{ij}^{max}, \tilde{s}_{ij}\} \geq 0.9$, consider i and j to be the same author.
- (c) Otherwise, consider i and j to be unique distinct authors.

We chose the cutoff values for these rules by inspecting the sample to trade-off type 1 and type 2 errors as best as possible. In this way, we incorporate information from both the pair of authors' names and from the similarity in their articles, to assess whether they are actually the same individual.

For the remaining set of names i and potential duplicates $D(i)$ we find the most similar duplicate of i , $m_i = \operatorname{argmax}_{j \in D(i)} \tilde{s}_{ij}$. We then find the most similar author to m_i : m_{m_i} . If $m_i \neq m_{m_i}$, i.e., if the most similar duplicate of i does not have i as its most similar duplicate too, we consider them to be distinct authors unless $\tilde{s}_{im(i)} > 0.85$. Otherwise, we classify them as the same author. This final step is particularly useful for a handful of cases with a multiplicity of differing but closely similar name variations.

For the top 200 authors in our data set based on citations, we manually checked for alternative spellings of their names, and collapsed the duplicates accordingly. At this point we are left with 46655 unique authors.

7. We exclude articles missing their publication date, or with publication dates prior to 1970 or posterior to 2019. We also exclude articles that do not use any third person pronouns as described in [subsection 11.2](#), and articles that do not have at least one known author matched to it.
8. Finally, we excluded articles with four or more authors, and papers from authors who only ever solo-authored.³⁶

This concludes our construction of the sample of articles and authors, and yields 66,854 articles written by 29,302 unique authors.

11.2 Classification of the pronoun use style of articles: Allen NLP coreferencing

Our methodology demands that we classify the writing style of each article as it relates to the gender choices for its third person pronouns. We rely on the *Allen* natural language processing (NLP) package, a state-of-the-art neural network model.³⁷ For each paper, we identify every instance of one of the following third-person pronouns: Masculine:

he, him, his, himself.

Feminine:

she, her, hers, herself.

Plural:

they, them, their, theirs, themselves.

Mix:

he or she, him or her, his or her, himself or herself, he and she, him and her, his and her, himself and herself.

³⁶Authors who never co-authored constitute isolated components of the network. Because in the first step of our empirical method we classify authors into two underlying types using information from co-authorship links, there is no information to classify isolated components of the network, and we must exclude them.

³⁷See <https://demo.allennlp.org/coreference-resolution>.

Consider the extreme case in which 2 the revision node is reached almost certainly, i.e., $E-1$. In this situation 1 player 1 can " blackmail " 0 player 2 by choosing a strategy which makes 0 player 2 play the strategy that gives 1 player 1 the payoff higher than x . If the possibility of reaching 2 revision node is small, however, 1 player 1 should also take into account the possibility that 1 his blackmail can not affect 3 0 player 2's supgame strategy, and 3 1 his strategy must face z . In that case, the probability of which is $1-s$, player 1's payoff 5 decreases by at least some positive amount, say d ; recall that (f, f_2) is constructed so that 4 each player can not decrease 4 his opponent's payoff without decreasing 4 his own payoff by at least d . For sufficiently small ϵ , 5 this loss can not be compensated by the gain obtained through the revision of supgame strategy by 0 player 2.

Figure A.17: Allen NLP coreferencing example.

For each identified pronoun, we extract the sentence containing the pronoun, and the sentences preceding and following it. We then run the *Allen NLP* coreferencing model on this text segment. This model relates the pronoun to its corresponding noun within the segment. For example, if we feed it the sentence "John ate an apple and he liked it", *Allen NLP* will indicate that "he" refers to John, and that "it" refers to apple. Figure A.17 illustrates the form of the *Allen NLP* output, in a paragraph from an article in our sample.

After parsing every segment involving a pronoun in an article, we obtain a list of proper nouns and their corresponding coreferenced third-person pronouns in the article. *Allen NLP* is known to achieve at least a 75 percent accuracy in standard English text. At the paper level, our manual checks suggest an error rate of almost zero.

In a next step, we use a list of keyword economic agent nouns, to select the *Allen NLP* coreferenced nouns in each paper that correspond to economic agents the articles are referring to. We use the following list:

'individual', 'worker', 'agent', 'principal', 'loser',
'representative', '[pl]ayer', 'trader', 'competitor', 'winner',
'citizen', 'messenger', 'manufacturer', 'investor', 'bank',
'government', 'criminal', 'member', 'researcher', 'opponent',
'group', 'respondent', 'party', 'incumbent', 'buyer', 'legislator',
'officer', 'prisoner', 'insured', 'insurance', 'owner', 'lender',
'challenger', 'cooperator', 'employer', 'customer', 'participant',
'borrower', 'mover', 'recipient', 'household', 'innovator', 'leader',
'rival', 'follower', 'contestant', 'intermediaries', 'voter',
'dictator', 'ceo', 'monopolist', 'migrant', 'candidate', 'manager',
'peer', 'user', 'trustee', 'oligopolist', 'employee', 'firm',
'regulator', 'person', 'maker', 'auctioneer', 'type', 'intruder',
'outsider', 'insider', 'people', 'dealer', 'entrepreneur',
'policymaker', 'nature', 'negotiator', 'neighbo[r]', 'executive',
'physician', 'generation', 'child', 'parent', 'newcomer', 'friend',
'professional', 'retailer', 'resident', 'student', 'subject',
'seller', 'partner', 'bidder', '[c]onsumer', 'organization',
'those who', 'sender', 'receiver', 'stockholder', 'team', 'speculator',
'supplier', 'producer', 'labourer', 'laborer', 'landholder', 'farmer',
'developer', 'creditor', 'politician', 'planner', 'arbitrageur',

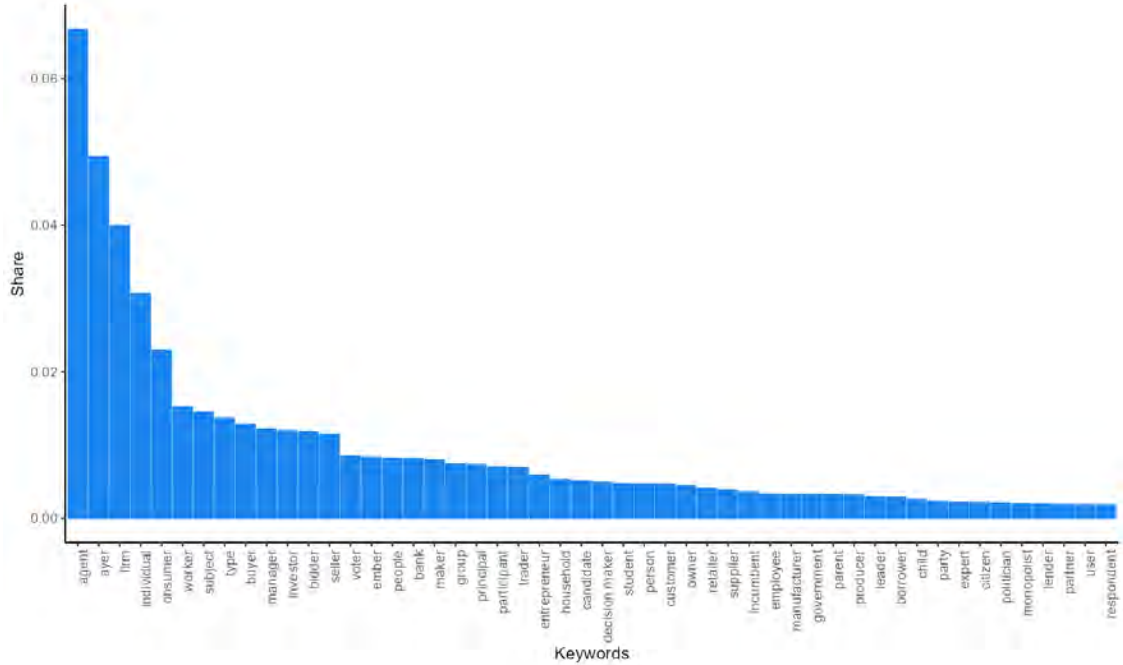


Figure A.18: Distribution of agent nouns used for co-referencing across articles: top 50.

'committee', 'board', 'bargainer', 'herder', 'defendant', 'plaintiff', 'jury', 'jurist', 'juror', 'judge', 'colleague', 'faculty', 'scientist', 'analyst', 'applicant', 'baron', 'bureaucrat', 'contractor', 'decision - maker', 'decisionmaker', 'decisions makers', 'entrant', 'expert', 'landlord', 'merchant', 'mutant', 'offender', 'peasant', 'proposer', 'purchaser', 'responder', 'teacher', 'venture capitalist', 'tortfeasor', 'commuter', 'insurer'.

After identifying all instances of pronoun use referring to any of the agent nouns listed above, we count the number of times masculine, feminine, plural, or a combination, are used in each paper to refer to them. Figure A.18 presents the distribution of these agent nouns across the full sample of article texts, for the top 50 most frequently used agent nouns.

We classify an article as masculine if it only uses masculine pronouns. We classify an article as feminine if it only uses feminine pronouns. We classify an article as plural if it only uses plural pronouns. We classify an article as mixed if it uses a combination of more than one type of pronoun.

11.3 Measurement of the relative spatial location of authors: *Author2vec*

To identify a set of plausible coauthors for each author in our sample, we adapted the *Word2vec* algorithm to our setting. *Word2vec* is a widely used algorithm in computer science designed to capture semantic relationships between words based on their co-occurrence patterns in a body of text (corpus). It is based on the distributional hypothesis proposition

in linguistics, according to which words appearing in similar contexts tend to have similar meanings. Within a given corpus (e.g., the congressional record), it uses the relative frequencies with which pairs of words appear near each other (right before or after, within a few words of each other, etc.) to assign a high-dimensional vector of real numbers to each word –referred to as the word’s *embedding*–.³⁸ We denote word i ’s embedding by $\mathbf{e}_i$. An embedding contains cardinal information about the word’s meaning in relation to all other words in the corpus: words that are closer to each other in this vector space, say using a Euclidean distance norm, are deemed to be closer to each other in meaning, because the relative frequencies with which they appear near other words is similar.

Consider a word w_j in some sentence, and refer to it as the center word. Consider other words in the same sentence found at most m ³⁹ words away from w_j , and refer to them as context words. Denote this set as $M(j)$. *Word2vec* allows for each word j to have an embedding as center word, $\mathbf{e}_j^c$, and an embedding as context word $\mathbf{e}_j^o$. *Word2vec* defines the conditional probability of observing context word w_k given center word w_j using the *softmax* function as

$$\mathbb{P}(w_k|w_j) = \frac{\exp(\mathbf{e}_k^o \mathbf{e}_j^c)}{\sum_{\ell} \exp(\mathbf{e}_{\ell}^o \mathbf{e}_j^c)}$$

Making the dot product between context word w_k and center word w_j large relative to all other words in the corpus makes this probability high.

Word2vec chooses the collection of vectors $\{\mathbf{e}_j^o, \mathbf{e}_j^c\}_{j=1}^W$ for all words in the corpus that maximizes the joint likelihood of observing the actual context-center pairs:

$$L(\theta) = \prod_{j=1}^W \prod_{k \in M(j)} \mathbb{P}(w_k|w_j)$$

The solution to this problem minimizes the difference between the predicted conditional probabilities and the actual distribution of word pairings in the corpus. In a final step one can average the estimated center and context embeddings of each word to obtain a single embedding for the word.

Word2vec is, implicitly, a network-based model where words are nodes, and edges between words exist when two words are near each other in the corpus –how near being a parameter chosen by the researcher–. The idea we propose here is to rely on the same logic, applied to the social network of economists in our sample, to measure ‘academic similarity’ across authors. We call this algorithm *Author2vec*. Authors play the role of words, cross-citation relationships play the role of edges between them, and we compute an embedding vector for each author⁴⁰. Two authors with close embeddings will be authors who cite and are cited by similar subsets of other authors, in the same way that words with close embeddings are words that appear near similar subsets of other words. In this sense, such authors are nearby in ‘academic’ space, and we will rely on this academic distance to restrict the set of authors that could feasibly be co-authors of a given author.

³⁸Large language models such as *ChatGPT*, for example, rely on a corpus that may include all of the internet, and on embeddings of many thousands of dimensions.

³⁹ m is a radius chosen by the researcher. If $m = 1$, for example, we only consider the word directly preceding and the word directly succeeding w_j as context words.

⁴⁰In practice we allow for 100-dimensional embeddings for the authors.

To implement our *Author2vec* methodology we transform each article a in our data set into a vector $\mathbf{v}_a$ of author identifiers that includes identifiers for all authors that either co-authored the paper or that are cited in the paper. Each such vector is analogous to a sentence in standard *Word2vec*. The collection of all such vectors $\{\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_{N_a}\}$ constitutes our corpus. We define a pair of authors to be ‘near’ if they appear in the same article vector. We can then use the frequencies with which each author is ‘near’ every other author within our corpus of articles in exact analogy to how *Word2vec* uses the frequency with which a given word appears before or after (near) every other word within the corpus of text.

We rely on the *Microsoft Academic Graph* (MAG)⁴¹ and *Jstor* data sets to retrieve network-related information about the set of authors in our sample, including co-authorship relationships and forward and backward citation relationships.

11.4 Construction of the acquaintance sets

We rely on the author embeddings from our *Author2vec* methodology to compute the cosine similarity (dot product of two vectors divided by the product of their lengths) between each pair of authors in our sample, $s_{i,j}$, as a scalar measure of academic proximity⁴²:

$$s_{i,j} = \frac{\mathbf{e}_i' \mathbf{e}_j}{|\mathbf{e}_i| |\mathbf{e}_j|}$$

Our premise is that pairs of authors far from each other in this academic space are effectively unable to consider each other as potential co-authors. We compute an acquaintance set of potential co-authors for each author, $Q_n(i)$, as follows: we take the union of the n closest authors to author i , all co-authors of author i , and the n closest authors to each of i ’s co-authors. We then exclude from this set any author who does not overlap in his productive years—defined as the range of years between three years before the author’s first publication and five years after the author’s last publication—with author i . By construction, $Q_n(i)$ includes all authors who did co-author with i at some point and a number of other authors who did not, but who are close enough in academic space that it is likely i could have considered them as co-authors. Our benchmark estimates use acquaintance sets with $n = 10$, but we also set $n = 5$ or 20 in alternative specifications.

11.5 Measurement of covariates

11.5.1 Assignment of sub-fields for authors: *ChatGPT* embeddings

Co-authorship decisions are likely influenced, among other characteristics, by the overlap in the sub-fields of study of authors. We assign sub-fields of specialization to the authors in our sample as follows: first, we borrow the *Journal of Economic Literature* (JEL) fields classification, and select a subset of the JEL fields which we deem relevant in our context. The following is the list of JEL fields we use:

⁴¹See <https://www.microsoft.com/en-us/research/project/microsoft-academic-graph>.

⁴²Cosine similarity is the most commonly used distance measure in the network science-large language models literature.

- C6 Mathematical Methods • Programming Models • Mathematical and Simulation Modeling
- C7 Game Theory and Bargaining Theory
- C9 Design of Experiments
- D1 Household Behavior and Family Economics
- D2 Production and Organizations
- D3 Distribution
- D4 Market Structure, Pricing, and Design
- D5 General Equilibrium and Disequilibrium
- D6 Welfare Economics
- D7 Analysis of Collective Decision-Making
- D8 Information, Knowledge, and Uncertainty
- D9 Micro-Based Behavioral Economics
- E2 Consumption, Saving, Production, Investment, Labor Markets, and Informal Economy
- E3 Prices, Business Fluctuations, and Cycles
- E4 Money and Interest Rates
- E5 Monetary Policy, Central Banking, and the Supply of Money and Credit
- E6 Macroeconomic Policy, Macroeconomic Aspects of Public Finance, and General Outlook
- E7 Macro-Based Behavioral Economics
- F1 Trade
- F3 International Finance
- G1 General Financial Markets
- G2 Financial Institutions and Services
- G3 Corporate Finance and Governance
- G4 Behavioral Finance
- G5 Household Finance

- H1 Structure and Scope of Government
- H2 Taxation, Subsidies, and Revenue
- H3 Fiscal Policies and Behavior of Economic Agents
- H4 Publicly Provided Goods
- H5 National Government Expenditures and Related Policies
- H6 National Budget, Deficit, and Debt
- H7 State and Local Government • Intergovernmental Relations
- H8 Miscellaneous Issues
- I1 Health
- I2 Education and Research Institutions
- I3 Welfare, Well-Being, and Poverty
- J. Labor and Demographic Economics
- K. Law and Economics
- L1 Market Structure, Firm Strategy, and Market Performance
- O1 Economic Development
- O2 Development Planning and Policy
- O3 Innovation • Research and Development • Technological Change • Intellectual Property Rights
- O4 Economic Growth and Aggregate Productivity
- P. Political Economy and Comparative Economic Systems
- R. Urban, Rural, Regional, Real Estate, and Transportation Economics
- Z1 Cultural Economics • Economic Sociology • Economic Anthropology

We then retrieve the *ChatGPT* embedding corresponding to all the words in the description of each of these fields, including the text describing its subfields.⁴³ This gives us an

⁴³We retrieve ChatGPT-3 embeddings of 1536 dimensions, based on their text-embedding-ada-002 model. See <https://openai.com/blog/new-and-improved-embedding-model>. Because ChatGPT's embeddings are estimated for a large corpus of English text, they are ideal as measures of relative similarity between common-use words. One of the main advantages of LLM word embeddings is their cardinal nature, allowing arithmetic operations that preserve relative meanings. As an example often used in this literature, subtracting the embedding for the word *man* from the embedding for the word *king*, and then adding the embedding for the word *woman* yields an embedding that is remarkably close to the embedding for the word *queen*.

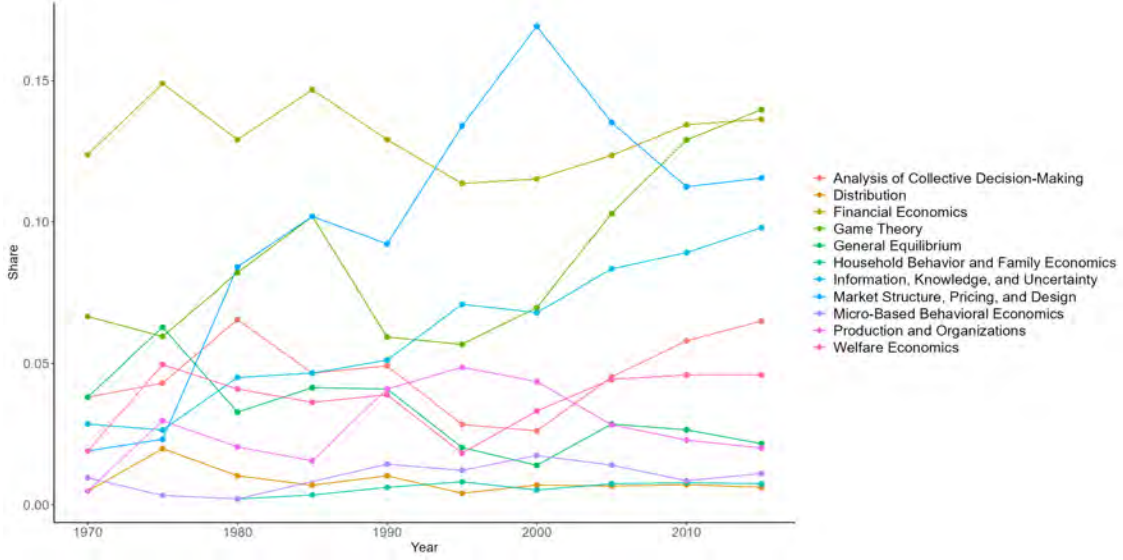


Figure A.19: Distribution sub-fields by 10-year cohorts.

embedding for each field j , $\bar{\mathbf{f}}_j$, with $j = 1, \dots, J$. In parallel, for each author i in our sample we create a collection K_i of the words in the titles of all of i 's articles, and the words in the titles of all papers cited in i 's articles. Next we retrieve the *ChatGPT* embedding for the collection of all words in K_i . This gives us an average embedding for author i , $\bar{\mathbf{g}}_i$. Next we compute cosine similarity distances between each author and each field,

$$\sigma_{i,j} = \frac{\bar{\mathbf{g}}_i' \bar{\mathbf{f}}_j}{|\bar{\mathbf{g}}_i| |\bar{\mathbf{f}}_j|}.$$

Finally, we assign to each author the three sub-fields with the smallest cosine similarity distances and use those to create dummy variables indicating sub-field membership.

In Figure A.19 we plot the distribution of sub-fields by 5-year cohorts of articles. Most fields have remained stable, with some exceptions: “Market structure, pricing and information” has grown steadily from 2 percent in 1970 to 12 percent today, and “Game theory” has grown from 7 percent in 1970 to 14 percent today.

11.5.2 Classification of the of ethnic origin of authors: *Namsor*

We rely on the authors' full names we obtained directly from the articles in our data set to assign an ethnic origin to each author. We do this using *Namsor*⁴⁴, a software tool specialized in identifying the likely regions of origin of proper names and last names from cultures all around the world. For each component of an author's name—first name, middle name, last name—Namsor reports a most likely origin at the sub-region level (e.g., Western Europe, South-east Asia, Middle East, etc.). As the ethnic origin of author i , we assign the modal sub-region reported by Namsor across all of the author's name components. For the small subset of cases with ties, we relied on *ChatGPT* prompts containing Namsor's guesses, and

⁴⁴See <https://namsor.app>.

retrieved *ChatGPT*’s best guess responses.

11.5.3 Classification of the gender of authors: Genderize package in R

We rely on the authors’ first names we obtained directly from the articles in our data set to assign a gender to each author. We do this using the *Genderize* package in R,⁴⁵ a software tool that has been trained on a large corpus of text as a probabilistic gender classifier for first names. We face one challenge: first and last names appear in no particular order. Sometimes first names appear before last names, and sometimes the other way around. Thus, we proceeded by genderizing each component of an author’s full name. For example, we asked the package to assign a gender to both “Debraj” and “Ray” separately. We then classified the authors as follows: if both components were assigned the same gender, we assigned that gender to the author. If there was a discrepancy across components, we identified the most popular of the components and assigned that gender to the author. We cross checked the quality of our gender assignment algorithm manually.

11.5.4 Computing citation counts of authors

We directly pulled estimated citation counts for each paper from the *Microsoft Academic Graph* (MAG) data set and from the *Crossref* dataset when the MAG information was unavailable. We then assigned to each author the sum of citations of the author’s articles.

11.5.5 Assignment of institutional affiliations of authors

For a subset 47 US institutions, we matched the theorists in our sample with their home department using a manually collected dataset. A department is included if it is in the top 50 list of the *RePEc* U.S. department rankings in 2013, 2014 and 2015.⁴⁶ The department level dataset covered all faculty members as well as their titles from 1995 to 2019 from two sources (department websites and course catalogues). We matched our sample of theorists to the faculty members in these departments using their names. This sums up to a total of 11,087 theorists with affiliation info.

11.6 Description of the methodology to estimate the community detection model based on [Feng et al. \(2023\)](#)

Taking logs from (6), we can express the log likelihood compactly as

$$\log \mathcal{L} = \sum_{t \in \{\ell, c\}} n_t(\boldsymbol{\tau}) \log(\pi_t) + \sum_{t, t' \in \{\ell, c\}} M_{tt'}(\boldsymbol{\tau}) \log(\omega_{tt'}) - \sum_{t, t' \in \{\ell, c\}} \omega_{tt'} B_{tt'}(\boldsymbol{\tau}, \boldsymbol{\gamma}) + \sum_{i=1}^n \sum_{j=i}^n q_{ij} y_{ij} \mathbf{x}'_{ij} \boldsymbol{\gamma} \quad (9)$$

⁴⁵See <https://www.rdocumentation.org/packages/genderizeR/versions/2.0.0>.

⁴⁶See <https://ideas.repec.org/top/old/1505/top.usecondept.html>.

where $q_{ij} = 1$ if $j \in Q(i)$,

$$n_t(\boldsymbol{\tau}) = \sum_{i=1}^n 1\{\tau_i = t\}$$

is the total number of type t authors under assignment $\boldsymbol{\tau}$,

$$M_{tt'}(\boldsymbol{\tau}) = \sum_{i=1}^n \sum_{j=i}^n q_{ij} y_{ij} 1\{\tau_i = t, \tau_j = t'\}$$

is the number of co-authorships between a type t and a type t' authors under assignment $\boldsymbol{\tau}$, and

$$B_{tt'}(\boldsymbol{\tau}, \boldsymbol{\gamma}) = \sum_{i=1}^n \sum_{j=i}^n q_{ij} e^{\mathbf{x}'_{ij} \boldsymbol{\gamma}} 1\{\tau_i = t, \tau_j = t'\}$$

is an aggregate of the covariate influence in co-authorship formation among type t and a type t' authors under assignment $\boldsymbol{\tau}$.

We can first take the FOC with respect to π_t and Ω . With respect to π_ℓ :

$$\begin{aligned} n_\ell(\boldsymbol{\tau}) \frac{1}{\pi_\ell} + (n - n_\ell(\boldsymbol{\tau})) \frac{1}{1 - \pi_\ell} (-1) &= 0 \\ \Rightarrow \\ \pi_\ell^{MLE} &= \frac{n_\ell(\boldsymbol{\tau})}{n} \end{aligned} \tag{10}$$

With respect to $\omega_{tt'}$,

$$\begin{aligned} \frac{M_{tt'}(\boldsymbol{\tau})}{\omega_{tt'}} - B_{tt'}(\boldsymbol{\tau}, \boldsymbol{\gamma}) &= 0 \\ \Rightarrow \\ \omega_{tt'}^{MLE} &= \frac{M_{tt'}(\boldsymbol{\tau})}{B_{tt'}(\boldsymbol{\tau}, \boldsymbol{\gamma})} \end{aligned} \tag{11}$$

Plugging back (10) and (11) into (9), we obtain the profile likelihood:

$$\begin{aligned} \log \mathcal{L}^* &= \sum_{t \in \{\ell, c\}} n_t(\boldsymbol{\tau}) \log \left(\frac{n_t(\boldsymbol{\tau})}{n} \right) + \sum_{t, t' \in \{\ell, c\}} M_{tt'}(\boldsymbol{\tau}) \log \left(\frac{M_{tt'}(\boldsymbol{\tau})}{B_{tt'}(\boldsymbol{\tau}, \boldsymbol{\gamma})} \right) \\ &\quad - \sum_{t, t' \in \{\ell, c\}} \frac{M_{tt'}(\boldsymbol{\tau})}{B_{tt'}(\boldsymbol{\tau}, \boldsymbol{\gamma})} B_{tt'}(\boldsymbol{\tau}, \boldsymbol{\gamma}) + \sum_{i=1}^n \sum_{j=i}^n q_{ij} y_{ij} \mathbf{x}'_{ij} \boldsymbol{\gamma} \end{aligned}$$

Notice that the third sum is a constant equal to the total number of co-authorships, so it does not depend on $\boldsymbol{\tau}$ or $\boldsymbol{\gamma}$.

Thus, maximizing $\log \mathcal{L}^*$ is equivalent to maximizing

$$\log \tilde{\mathcal{L}}^* = \sum_{t \in \{\ell, c\}} n_t(\boldsymbol{\tau}) \log \left(\frac{n_t(\boldsymbol{\tau})}{n} \right) + \sum_{t, t' \in \{\ell, c\}} M_{tt'}(\boldsymbol{\tau}) \log \left(\frac{M_{tt'}(\boldsymbol{\tau})}{B_{tt'}(\boldsymbol{\tau}, \boldsymbol{\gamma})} \right) + \sum_{i=1}^n \sum_{j=i}^n q_{ij} y_{ij} \mathbf{x}'_{ij} \boldsymbol{\gamma} \quad (12)$$

For a given ideological assignment $\tilde{\boldsymbol{\tau}}$, the terms of the form $n_t \log(n_t/n)$ and $M_{tt'} \log(M_{tt'})$ do not depend on $\boldsymbol{\gamma}$, so

$$\hat{\boldsymbol{\gamma}}(\tilde{\boldsymbol{\tau}}) = \operatorname{argmax}_{\boldsymbol{\gamma}} \left\{ \sum_{i=1}^n \sum_{j=i}^n q_{ij} y_{ij} \mathbf{x}'_{ij} \boldsymbol{\gamma} - \sum_{tt' \in \{\ell, c\}} M_{tt'}(\tilde{\boldsymbol{\tau}}) \log(B_{tt'}(\tilde{\boldsymbol{\tau}}, \boldsymbol{\gamma})) \right\}$$

This objective function is strictly concave in $\boldsymbol{\gamma}$, so it has a unique solution that can be easily found with a BFGS algorithm.

We can now plug in $\hat{\boldsymbol{\gamma}}(\tilde{\boldsymbol{\tau}})$ in (12):

$$\log \tilde{\mathcal{L}}^*(\tilde{\boldsymbol{\tau}}) = \sum_{t \in \{\ell, c\}} n_t(\tilde{\boldsymbol{\tau}}) \log \left(\frac{n_t(\tilde{\boldsymbol{\tau}})}{n} \right) + \sum_{t, t' \in \{\ell, c\}} M_{tt'}(\tilde{\boldsymbol{\tau}}) \log \left(\frac{M_{tt'}(\tilde{\boldsymbol{\tau}})}{B_{tt'}(\tilde{\boldsymbol{\tau}}, \hat{\boldsymbol{\gamma}}(\tilde{\boldsymbol{\tau}}))} \right) + \sum_{i=1}^n \sum_{j=i}^n q_{ij} y_{ij} \mathbf{x}'_{ij} \hat{\boldsymbol{\gamma}}(\tilde{\boldsymbol{\tau}}) \quad (13)$$

The space of possible vectors $\boldsymbol{\tau}$ is very large; there are 2^n possible vectors. [Feng et al. \(2023\)](#) propose an algorithm that works very well:

1. Pick an arbitrary $\tilde{\boldsymbol{\tau}}$, and find $\hat{\boldsymbol{\gamma}}(\tilde{\boldsymbol{\tau}})$.
2. Maximize (13) evaluated at $\hat{\boldsymbol{\gamma}}$ using an EM algorithm. For details on the EM algorithm, see [Feng et al. \(2023\)](#).
3. This yields an allocation $\tilde{\boldsymbol{\tau}}(\hat{\boldsymbol{\gamma}})$.
4. Iterate if desired, although in practice the first iteration will already deliver a very accurate allocation.

11.7 Description of the methodology to estimate the multinomial choice model through simulated maximum likelihood

We maximize (8) using the method of maximum simulated likelihood. This entails numerically simulating the double integral that averages over the distribution of peer effects conditional on $\mathbf{w}_i$. We simulate this integral with a discrete sum. The estimator takes the

form

$$\ln \hat{L}(\gamma) = \sum_{a=1}^N \sum_{\rho \in \{m, f, x, p\}} 1\{p_{a(ij)t} = \rho\} \times \\ \ln \left[\frac{1}{B_1} \frac{1}{B_2} \sum_{b_i(\mathbf{w}_i)=1}^{B_1} \sum_{b_j(\mathbf{w}_j)=1}^{B_2} G_\rho \left(V_{a(ij)t}^\rho(b_i, b_j) \right) \right],$$

where

$$G_\rho(v^\rho) = \frac{\exp(v^\rho)}{1 + \sum_{s \in \{m, f, x\}} \exp(v^s)},$$

and the $b_k(\mathbf{w}_k)$ are draws for the β coefficients for each author from normal distributions conditional on $\mathbf{w}_k$, and B_1, B_2 are the number of draws for approximating the integrals. For single-authored papers the integral is effectively one dimensional.