

How Cultural Diversity Drives Innovation

Surnames and Patents in U.S. History*

Max Posch

Jonathan Schulz

Joseph Henrich[†]

September 8, 2025

Abstract

This paper examines the impact of cultural diversity on innovation. Focusing on the United States from 1850 to 1940, we develop a novel surname-based measure of cultural diversity and combine this with patent data. Leveraging quasi-random variation in counties' surname compositions driven by historical immigration, we find that rising diversity increased both the quantity and quality of innovation within counties and for individual inventors. Examining mechanisms, we provide evidence suggesting that greater surname diversity accelerated innovation both by expanding the range of ideas, skills and perspectives available for recombination and by fostering the diverse social interactions that facilitate idea sharing.

Keywords: Innovation, cultural diversity, surnames, immigration, collective brain

JEL Classification: O33, R11, N92, J15, Z13

*Max Posch and Jonathan Schulz share co-first authorship. We thank Leonardo Bursztyn and four anonymous referees for comments that significantly improved the paper. We are also grateful to Ran Abramitzky, Alberto Alesina, Mike Andrews, Pablo Balan, Michael Clemens, Tyler Cowen, Klaus Desmet, Dan Fetter, Vicky Fouka, Martin Fiszbein, Oded Galor, Erik Hornung, Chad Jones, Ross Mattheis, Petra Moser, Nathan Nunn, Tzachi Raz, and Slava Savitskiy for valuable feedback and discussions. We thank seminar and workshop participants at Bonn, CERGE-EI, Cologne, Copenhagen, Louvain, Lund, the NBER Summer Institute, NYU, and UBC. We are grateful to Enrico Berkes for generous data sharing. Research reported in this publication was supported by the John Templeton Foundation under Award Number 62161. Edited by Leonardo Bursztyn.

[†]Posch: University of Exeter Business School (m.posch@exeter.ac.uk). Schulz: George Mason University (jschulz4@gmu.edu). Henrich: Harvard University (henrich@fas.harvard.edu).

It is hardly possible to overrate the value [...] of placing human beings in contact with persons dissimilar to themselves, and with modes of thought and action unlike those with which they are familiar. [...] Such communication has always been, and is peculiarly in the present age, one of the primary sources of progress.

John Stuart Mill — *Principles of Political Economy*

1 Introduction

At least since John Stuart Mill (1871), scholars from various disciplines have posited that social interactions among diverse individuals stimulate innovation and creativity, primarily because such interactions foster the exchange and recombination of ideas (e.g., Jacobs, 1969; Glaeser et al., 1992; Weitzman, 1998; Muthukrishna and Henrich, 2016; Jones, 2023). Indeed, research from the history of technology (Usher, 2013), economics (Acemoglu et al., 2016; Mokyr, 2002; Akcigit et al., 2013) and related fields (Youn et al., 2015) finds that many, if not most, innovations arise from the recombination of existing ideas. From this perspective, greater cultural diversity may fuel faster innovation by supplying a broader range of know-how and perspectives. Here, within the historical U.S. context, we test this hypothesis using a new fine-grained measure of cultural diversity and two patent-based measures of innovation.

To conceptualize cultural diversity, drawing on the field of cultural evolution (Boyd and Richerson, 1985; Henrich, 2016), we define ‘culture’ as information stored in people’s heads that they acquired through social learning—by observing, listening to and copying others (Bisin and Verdier, 2001; Giuliano and Nunn, 2021). This definition includes beliefs, values, interests, heuristics, motivations and ways of thinking as well as some forms of human capital, such as know-how, skills and tacit knowledge (Nunn, 2021). These are all ‘culture’ because they are learned from other people and thus may persist over generations. Following this approach, we consider cultural traits to be a kind of ‘information,’ so ‘culturally-transmitted information’ refers to traits acquired from others.

Such information tends to cluster in (extended) family networks across generations because individuals often learn from older relatives, particularly from same-sex relatives such as fathers, grandfathers and uncles (for males). Sons in the U.S., for example, are far more likely to patent if their fathers patented and they typically invent in the same or a closely related technology subclass as their fathers (Bell et al., 2019). Similar patterns emerge from studies of the intergenerational transmission of entrepreneurship (Lindquist et al., 2015).¹

Because surnames, like many cultural traits in males, are transmitted patrilineally, they offer a natural lens through which to identify fine-grained clusters of culturally-transmitted information. Exploiting this parallel, we explore the unique contribution of surname-based cultural diversity in driving innovation and compare it with more conventional diversity measures, such as those based on ancestry or birth country (Ottaviano and Peri, 2006; Ager and Brückner, 2013; Alesina et al., 2016; Docquier et al., 2020; Fulford et al., 2020).

To construct our measure of cultural diversity, we collect all surnames reported in the full-count U.S. Census data from 1850 to 1940 and compute the entropic diversity of surnames across counties, which are presumed to be the primary locations for social interactions during this period. Counties might not capture all social interactions, especially in today’s interconnected world, but we will show that they provide a reasonable approximation within the pre-1950 historical context.

We link county-level surname diversity to two patent-based indicators of innovation. First, we connect it to the universe of patents filed in the United States, as recorded in the Comprehensive Universe of U.S. Patents (Berkes, 2018). Second, we integrate this patent data with an indicator for breakthrough patents created by Kelly et al. (2021). This measure identifies patents that are both highly novel and notably impactful. Breakthrough

¹Ample evidence across diverse societies reveal a central role for older same-sex relatives, particularly parents, in the process of cultural transmission (Henrich and Broesch, 2011; Hewlett, 2016; Schniter et al., 2023).

patents, or what we term simply 'breakthroughs,' are characterized by their relatively lower textual similarity to earlier patents and higher similarity to later ones.

To estimate the causal effect of surname diversity on innovation, we adapt an instrumental variable (IV) strategy from [Burchardi et al. \(2019\)](#) to surnames. We isolate quasi-random variation in U.S. counties' surname compositions arising from the interactions of two historical forces: (i) the staggered arrival of immigrants with different surnames and (ii) the changing attractiveness of counties to immigrants over time. Importantly, immigration does not always increase surname diversity: its impact depends on how common or rare the incoming surnames already are in a county. For example, if many immigrants with the surname "Corleone" arrive, it might increase diversity if "Corleones" are rare, or decrease it if they are already common within that county.

Using data from 1900 to 1940, we estimate the effect of surname diversity on our patent-based innovation measures at four levels of analysis: (i) counties, (ii) surname-counties (e.g., all "Smiths" living in Brooklyn), (iii) inventors and (iv) patents. Across the first three levels, we find that greater surname diversity within counties leads to faster innovation based on both all patents (innovation quantity) and just breakthroughs (innovation quality). All specifications include year or state-year fixed effects and, when appropriate, county, county-surname or inventor fixed effects. The latter two, respectively, absorb any unobserved and time-invariant characteristics of surname groups or individual inventors—such as their abilities, interests or knowledge—that may potentially affect the estimates. To account for scale effects, we hold population size constant using either county population sizes or an instrumental variable approach based on immigration flows. To address remaining identification concerns, we include county-specific linear trends in some specifications, which allow us to estimate the impact of surname diversity shocks on changes from the *growth rate* of local patenting rather than on absolute levels.

The estimated effects on the quantity and quality of innovation are large and economically relevant. At the county level, for example, a one standard deviation increase in

surname diversity, roughly corresponding to the growth in surname diversity in New York County between 1870 to 1940, raises patent rates by 107% and breakthrough patent rates by 80% relative to their respective sample means. These estimates correspond to 1.76 times and 0.76 times the estimated effect of population size, respectively. At the individual inventor level, the same one standard deviation increase in surname diversity boosts inventors' productivity by roughly half a patent (a 20% increase over the sample mean) and increases the probability of achieving a breakthrough by 10 percentage points over a five-year period. Over a 50-year career, stretching from 1870 to 1920, these estimates imply that the same inventor would obtain five additional patents and one breakthrough if they lived in New York County instead of 300 miles further north in Clinton County, NY. To confirm the stability of these patterns before 1900, we interact surname diversity with each decade in a county-level OLS analysis. The results suggest a persistent relationship between surname diversity and innovation that extends well back into the 19th century.

Since our IV strategy relies on historical immigration patterns to isolate quasi-random variation in surname diversity, a potential concern is that immigration might have a direct effect on innovation that is not channeled through surname diversity. We mitigate this concern in two ways. First, we orthogonalize our push-pull instruments by each county's history of immigration. Second, we control for the number of recent immigrants. We find that, although immigration is positively associated with innovation (confirming prior work, e.g., [Sequeira et al., 2020](#); [Arkolakis et al., 2020](#); [Terry et al., 2024](#)), including this as a control has minimal influence on our estimates of the impact of surname diversity on innovation. These findings reinforce our interpretation that surname diversity, rather than immigration itself, drives our results and bolsters confidence in our empirical strategy.

At all three levels of analysis, our results hold across a broad range of robustness checks, including placebo tests of dynamic changes (e.g., regressing prior-period innovation on current surname diversity), separate analyses of major U.S. regions (e.g., the South and West), log-transformations of the dependent variables, and alternative formulations

of surname diversity, including using weights based on cultural distance. Further, by including a squared term, we also find no clear evidence for a concave relationship between surname diversity and innovation. Crucially, our findings are robust when including more conventional measures of cultural diversity in these specifications, such as those based on birth country, ancestral originating country and race. These checks demonstrate robustness against a variety of challenges.

At the level of the individual patent, our fourth level of analysis, we show that greater surname diversity among co-inventors predicts both increased technological complexity for a patent, based on the number of technology codes, and a higher likelihood of becoming a breakthrough patent. This specification includes fixed effects for inventor team size, counties, state-periods, patent-technology-periods and, in its most stringent version, county-periods. To establish causality, we employ our IV for county-level surname diversity to extract quasi-random variation in patent-level surname diversity. The first-stage result confirms that greater surname diversity at the county level leads to greater surname diversity at the patent level. The second-stage result reveals large, positive effects: a one standard deviation increase in patent-level surname diversity adds approximately four technology codes (a 150% increase relative to the sample mean) to a patent and increases the likelihood of becoming a breakthrough by roughly 90 percentage points, most of which is at the extensive margin (going from one-inventor to two-co-inventor patents). By linking greater county-level surname diversity directly to more social interaction among diverse co-inventors at the level of the patent, these findings open the door to our discussion of mechanisms.

Our mechanism section considers how and why cultural diversity influences innovation. Beginning with the stylized fact that many, if not most, innovations represent recombinations of existing ideas, cultural diversity can shape rates of innovative recombination both directly, by expanding the set of ideas, perspectives and domains of know-how that can be recombined—the ‘informational channel’—and indirectly, by influencing the willingness

of individuals to interact, exchange information and potentially collaborate—the ‘social-behavioral channel’. Conceptually, both the informational and social-behavioral channels must be important since neither a population of diverse minds that never interact nor a group of cognitive clones who freely interact but all share the same information will generate recombinations (Henrich and Muthukrishna, 2024). Below, we offer evidence suggesting that rising surname diversity both enriches the local informational diversity—the raw materials for recombinative innovation—and fosters more social interaction among diverse minds.

Empirically, considering the informational channel, we find that surnames cluster with educational levels, occupations, ancestry and the technological categories of inventions (i.e., patent technology codes). Therefore, counties and inventor teams will tend to have greater informational diversity when they have higher surname diversity. Consistent with this, using our county-level specifications, we show that greater surname diversity leads to greater occupational and patent technology code diversity. These findings offer *prima facie* evidence that rising surname diversity also increases the information available for recombination. However, keep in mind that beyond these observables surnames likely also capture unobserved dimensions of informational diversity, including distinct interests, beliefs, thinking styles, preferred metaphors, know-how, tacit knowledge, attentional biases, intuitions and more. Like occupations and technology codes, a large body of research across several disciplines shows how individuals acquire many cognitive and motivational traits by learning from same-sex parents, family members and their local social networks (Bell et al., 2019; Christakis and Fowler, 2009; Clark, 2014; Güell et al., 2015; Henrich, 2016; Barone and Mocetti, 2021; Schniter et al., 2023).

For the social-behavioral channel, there are two alternative hypotheses. Research on ethnic fractionalization (Alesina and La Ferrara, 2005) and the ‘paradox of diversity’ (Schimmelpfennig et al., 2022) suggests that greater cultural diversity might lead people to withdraw from unfamiliar or less comfortable social exchanges, preferring interactions

with their extended families and co-ethnics. This effect would result in less diverse social interactions and potentially less innovation. Alternatively, as a population diversifies with newcomers who bring distinct expertise, interests and skills, heterogeneous social interactions may increase for two reasons: (1) individuals may find it harder to satisfy their needs solely within their small, homogeneous and relatively decreasing groups, and (2) the spread of groups with distinct information may create more opportunities for valuable positive-sum exchanges. Consistent with research on contact theory (e.g., [Allport, 1954](#); [Bursztyn et al., 2024](#)) and the impact of market integration on prosocial behavior ([Henrich et al., 2010](#); [Enke, 2022](#); [Agneman and Chevrot-Bianco, 2023](#); [Rustagi, 2024](#)), the potential for mutually beneficial engagement with outsiders may foster greater openness to social interactions with diverse individuals, which in turn may spur innovative recombination. Moreover, surname diversity (or the lack of it) partially captures the extent to which large family groups are present. Research on family ties connects intensive kin networks with lower trust in strangers and increased moral parochialism, both of which can inhibit interactions and collaborations among diverse individuals ([Alesina and Giuliano, 2014](#); [Schulz et al., 2019](#); [Enke, 2019](#); [Henrich, 2020](#)).

To map out the social-behavioral channel, we construct measures of the strength of family ties and residential segregation for both surnames and ancestral-countries. These variables serve as proxy measures for the likelihood of social interactions between culturally distinct groups. The family ties measure reflects the tendency of individuals to form strongly homophilic networks with relatives ([Raz, 2025](#)). The measures of residential segregation assess the frequency of adjacent neighbors with either the same surname or ancestral-country relative to what is expected based on random assortment. Counties where people more freely assort provide greater opportunities for culturally diverse individuals to connect and share information.

Using our county-level specifications, we find that greater surname diversity weakens family ties and reduces residential segregation for both surnames and ancestral countries.

Interestingly, this holds even when controlling for ancestral-country and birth-country diversity. Together with our patent-level results, these findings suggest that increased fine-grained cultural diversity fosters more social interactions among heterogeneous individuals, thereby stimulating innovation.

This interpretation is further supported by an analysis of the impact of surrounding counties on innovation—that is, spillovers across county borders. We find that local proximity dominates the process, as the diversity of surrounding counties has a limited impact on the innovation within focal counties.

In sum, our analyses provide evidence of how rising cultural diversity fueled U.S. innovation, as captured by patents and breakthroughs, across four levels of analysis—counties, county surnames, inventors, and patents—from the mid-19th to the mid-20th centuries. Our exploration of mechanisms suggest that the fine-grained cultural diversity measured using surnames likely affects innovation through both informational and social-behavioral mechanisms. These results support the hypothesis that social interactions among diverse individuals foster innovation through recombination.

1.1 Contributions and Related Literature

We offer four main contributions. First, methodologically, we introduce a new fine-grained measure of cultural diversity based on surnames, demonstrate its relationship to informational diversity (e.g., occupations and patent technology codes) and compare it with more conventional cultural diversity measures in explaining innovation. Second, building on the push-pull method from the immigration literature ([Burchardi et al., 2019](#)), we develop instrumental variables for surname diversity as well as county population size and ancestral-country diversity. Third, we show how rising cultural diversity powerfully fueled innovation in the historical United States during the country’s rise to the technological frontier. Fourth, our analysis of mechanisms highlights the role of diverse social interactions—even at the level of the inventor and patent—and explores the influence of

rising surname diversity on the likelihood of diverse social interactions.

These contributions relate to several lines of work in economics. To start, our findings complement important research on the impact of population diversity on innovation and economic growth ([Ottaviano and Peri, 2006](#); [Kerr, 2008](#); [Ashraf and Galor, 2013](#); [Ager and Brückner, 2013](#); [Alesina et al., 2016](#); [AlShebli et al., 2018](#); [Arbatli et al., 2020](#); [Docquier et al., 2020](#); [Fulford et al., 2020](#); [Ashraf et al., 2021](#); [Fiszbein, 2022](#); [Galor et al., 2024](#)). In our conceptual framework, population diversity can fuel innovation if it captures differences in the information that individuals bring to social interactions, leading to the recombination of ideas and more novel creations.

Our results enrich this literature on population diversity and innovation in two ways. First, we demonstrate the importance of a particular form of cultural diversity in driving innovation during the era when the U.S. rose to global dominance. Second, by comparing surname diversity with both ancestral- and birth-country diversity, we highlight the larger impact of surname diversity and the distinct role of the social-behavioral channel. These results suggest that the impact of particular forms of diversity may depend on how they operates via the informational and social-behavior channels. Although certain forms of diversity may increase the potential for innovation via the informational channel, they can reduce or entirely short-circuit these effects if they suppress productive interactions among diverse individuals.

Our paper informs work connecting innovation to cities, population density, agglomeration, transportation and geographic proximity ([Jacobs, 1985](#); [Feldman and Audretsch, 1999](#); [Carlino et al., 2007](#); [Agrawal et al., 2008](#); [Kerr, 2010](#); [Glaeser, 2011](#); [Packalen and Bhattacharya, 2015](#); [Akcigit et al., 2017](#); [Berkes and Gaetani, 2021](#); [Atkin et al., 2022](#); [Andersson et al., 2023](#); [Chiopris, 2024](#)) as well as an emerging literature examining how social institutions, organizations and psychological dispositions (e.g., trust) spur innovation by facilitating diverse social interactions ([Algan and Cahuc, 2014](#); [de la Croix et al., 2018](#); [Nguyen, 2025](#); [Buggle, 2020](#); [Cinnirella et al., 2022](#); [Atkin et al., 2022](#); [Andrews,](#)

2023; Andrews and Lensing, 2024). In our theoretical account, these factors increase the likelihood of diverse individuals meeting and potentially exchanging ideas, leading to new recombinations.

Our paper contributes to these research lines in two ways. First, while we focus on cultural diversity, our results indicate that population size independently fuels innovation alongside diversity (we report the coefficients on population size throughout). Second, consistent with this literature’s emphasis on the potential for face-to-face meetings, our exploration of the social-behavioral mechanism suggests that greater surname diversity affects innovation partly through its impact on the potential for social interaction among diverse minds. Rising surname diversity at the county level leads to greater diversity at the patent level, resulting in patents that are more complex and more likely to be breakthroughs.

Our use of the quasi-random variation in surname diversity induced by immigration naturally connects our research to the literature examining the impact of immigration on innovation (Moser et al., 2014; Abramitzky and Boustan, 2017; Moser and San, 2020; Sequeira et al., 2020; Tabellini, 2020; Arkolakis et al., 2020; Abramitzky et al., 2023; Lissoni and Miguelez, 2024; Terry et al., 2024). In our conceptual framework, immigration can powerfully drive innovation when migrants bring relevant information and can socially interact in their new communities.

We contribute to the literature on immigration and innovation by showing that immigration operates on innovation partly through its influence on population size and cultural diversity. However, our results also suggest that something else is afoot in addition to these mechanisms: perhaps immigrants are self-selected on traits like individualism or trust (Beck Knudsen, 2024; Gorodnichenko and Roland, 2016) or the experience of moving itself makes individuals more social or creative (Maddux et al., 2010; Bernstein et al., 2022).

Finally, our focus on surname diversity connects our paper to the literature on kinship and family ties within economics (Alesina and Giuliano, 2010, 2014, 2015; Buonanno and

Vanin, 2017; Schulz et al., 2019; Enke, 2019; Henrich, 2020; Schulz, 2022; Bahrami-Rad et al., 2024; Ghosh et al., 2023). We contribute to this literature by showing how rising surname diversity, holding ancestral-country diversity constant, leads to less residential segregation, weaker family ties and patents with more diverse inventors.

Our presentation rolls out as follows: Section 2 describes the data and the operationalization of our key variables. Section 3 lays out the construction of the instruments for cultural diversity and population size, highlighting our identification assumption. Section 4 presents the estimates of the causal effect of county-level surname diversity on innovation at three nested levels of analysis—county, county-surname and inventor—and reports robustness checks across these levels. Section 5 presents our patent-level results, and Section 6 sheds light on potential mechanisms by offering evidence for both the informational and social-behavioral channels. Section 7 concludes with a review and consideration of policy implications.

2 Data and Measurement

We collect data on individuals’ surnames, locations, migration and patents. Summary statistics are presented in Tables B1 and B2, and additional details are in Appendix A.

2.1 Surname diversity

To measure surname diversity across counties and over time, we draw on the surname data provided by the full-count Integrated Public Use Microdata Series (Ruggles et al., 2021). We use the nine waves from 1850 to 1940 which contain the variable `namelast` of all individuals and county identifiers. Appendix A details how we clean the surname variable and harmonize county boundaries. Following Burchardi et al. (2019), we also obtain the variables `age` and `yrimmig` (the year of immigration) to estimate surname entropy for the mid-decades 1895, 1905, 1915, and 1925 by removing all individuals who were born or

immigrated after the mid-decade. The average surname population share within counties is 0.045% (s.d. = 0.131), the 25th percentile is 0.003%, and the 75th percentile is 0.040%. Figure B1 depicts the distribution.

We operationalize surname diversity using Shannon entropy. To motivate this measure and conceptually link it to recombinative innovations, consider a subpopulation consisting of N individuals partitioned into K surname groups, each of size N_k such that $\sum_{k=1}^{K} N_k = N$. Each group carries unique knowledge, labeled as $s_k \in \{a, b, c, d, \dots\}$ and $s_k \neq s_h$ for all $k \neq h$. When individuals from different surname groups meet, the likelihood of recombinative innovation increases. Information theory (Shannon, 1948) tells us that the average informational content (or the innovation potential) of such a population in which people randomly meet is

$$E = - \sum p_k \log_2 p_k \quad (1)$$

where $p_k = \frac{n_k}{N}$ is the probability that a person with group affiliation k is drawn and $\log_2 p_k$ is the informational content embedded in this individual (expressed in bits).

Shannon entropy is a central concept in information theory and is widely used in many scientific disciplines. The term $-\log_2 p_k$ is the self-information of subgroup k and captures the level of surprise (or the informational content of a specific outcome). The negative log reflects that rarely-encountered groups carry more surprise (or more information) compared to more frequent encounters. To arrive at Shannon entropy, the self-information is weighted by the probability of its occurrence and summed over all possible outcomes. For example, if the population only consists of one group k , the outcome of a draw is entirely predictable, resulting in an entropy of 0, and thus, no recombinations can arise through social interaction. In contrast, entropy for a population with a fixed number of groups is maximized if all groups are of equal size. In such a population, the average individual engaging in a random social interactions is most likely to observe someone different from themselves. A random draw will thus have more informational content (in

expectation), which is reflected by higher entropy.

In economics, a fractionalization index is usually used to capture diversity. For our purposes, however, Shannon’s entropy offers the conceptually clearest way to represent information diversity and has favorable mathematical properties (Carcassi et al., 2021). In particular, the fractionalization underweights the importance of rare surnames (p_k vs. $\log_2 p_k$), i.e., under the assumption that surnames carry unique pieces of information, rare surnames are more “valuable”—they carry a higher expected surprise. Empirically, the two measures are highly correlated in our setting ($\rho = 0.87$, Table B5).

Figure B2 displays surname diversity across U.S. counties in the year 1940, presenting both the raw data and the residual variation that is orthogonal to log population size.² Clear geographical patterns emerge from the data. Counties in California and most of the Northeast exhibit high surname diversity, whereas Utah and the Southern states show considerably lower diversity, indicating greater homogeneity in surnames. Some of these differences had persisted for nearly a century by 1940, as illustrated in Figure B4. The figure also indicates that surname diversity increased across all US census regions between 1850 and 1940. Table B3 reports surname diversity over time of the counties with the highest and lowest surname diversity in 1940.

Finally, we estimate the most likely ancestral origin country for each surname, denoted as $o(k)$, to account for ancestral-country specific factors in our empirical design and develop a measure of ancestral-country diversity. We achieve this by drawing on the universe of immigrants to the U.S. from 1880 to 1930 and identifying the ancestral country for each surname as the one where the majority of immigrants bearing that name were born. Our methodology is detailed in Appendix A. We achieve an accuracy rate of 68% among immigrants whose country of birth is recorded, representing 4% of the overall population. As reported in Table B6, there is a moderate correlation between ancestral-country diversity and surname diversity, with conditional correlation coefficients ranging

²Figure B3 illustrates the variation in surname diversity conditional on log county population from 1850 to 1930.

from 0.26 to 0.41.

2.2 Innovation

To measure innovation, we rely on patent data from the Comprehensive Universe of U.S. Patents (CUSP) data set compiled by [Berkes \(2018\)](#).³ For each patent, the data set provides inventor names and location of residence, filing and issuing years of patents, and the U.S. Patent and Trademark Office technology classifications. Using this data, we construct two kinds of innovation outcomes, one based on the number of patents and the second based on 'breakthrough' patents, which we call 'breakthroughs'. The latter was developed by [Kelly et al. \(2021\)](#) and identifies patents that are both highly novel and impactful based on their low textual similarity to earlier patents but high similarity to later ones. We use this measure of patent novelty and impact rather than the number of citations an individual patent has received because the U.S. Patent and Trademark Office did not consistently begin to record patent citations until after 1947.

We construct innovation outcome variables at four levels of aggregation: (i) county, (ii) surname-county, (iii) inventor, and (iv) patent. The county-level outcomes measure the number of (breakthrough) patents filed by inventors residing in county i during period t , normalized by the county population.⁴ For patents filed by multiple inventors, possibly

³Although patents have been widely used in economics and other disciplines to study innovation, important concerns remain ([Griliches, 1990](#); [Moser, 2013](#); [Lerner and Seru, 2022](#)). These include the fact that many innovations are not patented, industries have variable patenting tendencies, types of inventions have different patentability, and increased patenting in a specific technology category could inhibit innovation rates. Seeking to address these concerns, prior work has demonstrated that results using patents per capita parallel those using alternative measures, including using patent citations ([Terry et al., 2024](#); [Acemoglu et al., 2016](#)), patents with novel technology codes ([Lerner and Wulf, 2007](#)), the presence of 'creative' ([Gomez-Lievano et al., 2017](#)) or 'supercreative' ([Bettencourt et al., 2007](#)) occupations, exhibits and prizes at World Fairs ([Dowey, 2017](#); [Moser, 2013](#); [Squicciarini and Voigtländer, 2015](#)) and economic productivity (e.g., [Alesina et al., 2016](#); [Sequeira et al., 2020](#); [Terry et al., 2024](#)). Overall, while far from perfect, current evidence suggests that patents offer a valuable proxy for innovative activity.

⁴In our main analysis of the causal impact of surname diversity on innovation, we follow [Terry et al. \(2024\)](#) in normalizing patent outcomes by start-of-panel population values, i.e., population in 1900. This choice accounts for the likelihood that population growth, potentially endogenous due to innovative regions attracting more people, could impact the results. In the OLS estimation of the relationship between diversity and patents back until 1850, we normalize the outcomes by contemporaneous county population.

from different counties, we divide the patent count by the number of inventors. The county-surname-level outcomes measure the same quantities but are aggregated and normalized at the level of surnames within counties. The inventor-level outcomes measure the number of patents and breakthroughs filed by inventor j during period t . Finally, the patent-level outcomes are an indicator equal to one if a patent was a breakthrough patent and zero otherwise, as well as the number of technology codes on the patent. The latter measures the complexity of the patent and has been interpreted as a metric of recombination. Details on our patent-based measures can be found in Appendix A.

3 Constructing an Instrument for Surname Diversity

We will estimate the effect of surname diversity on innovation at four levels: (i) counties, (ii) surname-counties, (iii) inventors, and (iv) patents. For example, we will use the following county-level equation:

$$Y_i^t = \beta_1 \text{Surname diversity}_i^t + \beta_2 \text{Population}_i^t + \alpha_{s(i),t} + \alpha_i + \varepsilon_i^t \quad (2)$$

where $\text{Surname diversity}_i^t$ denotes the surname diversity in county i in period t . Y_i^t represents the outcome of interest, typically the number of patents or breakthrough patents filed in county i in the five-year period starting in year t , normalized by the county's population. The coefficient β_1 is our main interest. Population_i^t denotes county-level population, which we include to control for scale effects. The parameters $\alpha_{s(i),t}$ and α_i are state-period and county fixed effects, respectively. These fixed effects enable the estimation of β_1 from changes within the same county over time, while controlling for both persistent and time-varying differences across states. The error term is denoted by ε_i^t . In our most restrictive model specifications, we include county-specific linear trends, represented by $\alpha_i \times t$. These trends control for any inherent county-specific trend in Y_i^t , allowing us to focus only on the variation in the growth rate of patenting over time within

each county.

The main concern with the OLS estimate of β_1 is that reverse causality or unobserved factors that co-determine surname diversity and innovation might bias the estimates. For instance, a highly innovative county may attract a more diverse set of migrants, which would increase surname diversity (Abramitzky and Boustan, 2017). Similarly, highly-skilled individuals, who may be more prone to innovate, might prefer more diverse counties. In both instances, a positive correlation between surname diversity and innovation could be observed even if no causal relationship exists. On the other hand, some historical evidence suggests that immigrants tended to settle in lower-income (i.e., less innovative) areas during this period, as discussed by Sequeira et al. (2020). This possibility would bias our estimates downward.

Therefore, we employ an instrumental variable (IV) strategy to estimate the causal effect of surname diversity on innovation. We observe that migration, along with births, deaths and marriages, plays a significant role in changing counties' surname diversity. Thus, we leverage an IV approach from the immigration literature that allows us to isolate quasi-exogenous variation in surname diversity and to estimate the local average treatment effect (LATE) of immigration-induced changes in surname diversity on innovation as opposed to changes stemming from births, deaths, marriage, or domestic migration.

It is important to recognize that immigration does not uniformly increase surname diversity: its impact depends on the existing local surname distribution. For example, if families with the surname "Smith" immigrate to a county with few or no residents named "Smith," surname diversity increases. Conversely, if they move to a county where the "Smiths" are already common, diversity decreases. Table B4 underscores this observation by listing counties with the highest and lowest immigrant shares in 1900. Notably, counties with minimal immigrant presence (less than 1%), such as Harris County, Georgia (surname diversity: 8.55) and Jackson County, Texas (8.37), exhibit surname diversity comparable to counties with high immigrant shares (over 50%), like Logan County, North Dakota

(7.5) and Webb County, Texas (8.56). These findings suggest that the mere presence of immigrants does not necessarily drive surname diversity. Consistent with this, controlling for immigration does not significantly alter our estimates (see Section 4.4).

To construct the instrument for surname diversity, we adapt the historical push-pull approach of Burchardi et al. (2019) to predict the number of residents with surname k in county i in period t , $N_{k,i}^t$. This approach isolates quasi-random variation in surname stocks across counties by predicting these stocks using historical migration patterns. The instrument for surname diversity is then computed by calculating entropy based on these predicted surname stocks.

To understand how the push-pull interaction terms affect surname stocks and diversity, consider a historical example adapted from Bursztyn et al. (2024). Beginning in the 1910s, many "Garcias" (Mexican individuals) immigrated due to the Mexican Revolution, representing a substantial "push". Meanwhile, Detroit (Wayne County, MI) attracted large numbers of immigrants because of the booming automobile industry—a large "pull" for Detroit.⁵ The push-pull interaction increased the stock of Garcias in Detroit from the early 1910s (Garcia push 1915 $\times$ Detroit pull 1915), which contributed to greater surname diversity, as relatively few Garcias lived in Detroit before the 1910s. In contrast, after World War I, the "Johnsons," "Smiths," "Andersons," "Browns," and "Millers" once again became the most common immigrants, a "push" from British surnames. The push-pull interaction thus increased the stock of these surnames in Detroit from the early 1920s (e.g., Brown push 1925 $\times$ Detroit pull 1925). However, these increases did not necessarily contribute to greater diversity, as many "Johnsons," "Smiths," "Andersons," "Browns," and "Millers" were already present.

Step 1: Predicting Surname Stocks Burchardi et al. (2019) show that a reduced form

⁵Figure B5 shows how the timing of migrations varies for common surnames, highlighting the spike in "Garcias" and other Mexican surnames during the 1910s and 1920s. Figure B6 indicates that US counties became attractive to immigrants at various points in time; for instance, Detroit (Wayne, MI) gained appeal starting in the mid-1910s.

model of immigration driven by push (such as economic or political conditions in the immigrants' origin countries) and pull shocks (like economic opportunities in the destination counties), combined with a leave-out strategy, can identify variation in ancestry shares that is plausibly exogenous to local factors and ancestral country-destination county specific factors.

We tailor this approach to surnames. Formally, our model is given by:

$$N_{k,i}^t = \delta_{o(k),i} + \delta_{k,r(i)} + \sum_{\tau=1880}^{t-1} b^\tau \underbrace{I_{k,-r(i)}^\tau}_{\text{Push}} \underbrace{\frac{I_{i,-o(k)}^\tau}{I_{-o(k)}^\tau}}_{\text{Pull}} + \sum_{\tau=1880}^{t-1} d^\tau \frac{I_i^\tau}{I^\tau} + u_{i,k}^t, \quad (3)$$

where $N_{k,i}^t$ denotes the stock of people (in thousands) with surname k residing in county i in year t ; $I_{k,-r(i)}^\tau$ is the push factor, given by the total number of immigrants (in thousands) with surname k who arrive in the U.S. in the period ending in year τ (1880, 1895, 1900, 1905, 1910, 1915, 1920, 1925, 1930) and settle *outside* the region containing county i . The variable $\frac{I_{i,-o(k)}^\tau}{I_{-o(k)}^\tau}$ is the pull factor, which captures the relative attractiveness of a specific county i during the period ending in τ . It is given by the share of immigrants settling in county i , excluding immigrants who likely are from the same origin country as those with surname k . $\delta_{o(k),i}$, denotes origin country-county fixed effects, removing any time-invariant factors that make specific counties more attractive to immigrants from the origin country that is most strongly associated with individuals with surname k (e.g., Swedish immigrants tended to settle in counties in Minnesota). $\delta_{k,r(i)}$ are surname-region fixed effects. They remove time-invariant unobserved factors that may make specific census regions (i.e., one of the nine census divisions defined by the Census Bureau) more attractive to migrants with certain surnames. Our surname setting enables us to also include $\sum_{\tau=1880}^{t-1} d^\tau \frac{I_i^\tau}{I^\tau}$, i.e., the relative share of migrants who settle in a county in each period τ . This term captures the time-varying relative attractiveness of a county in the past. It isolates the push-pull instruments from county-level conditions that drew migrants in each period τ up to $t-1$, which may still affect innovation in period t .

Like [Burchardi et al. \(2019\)](#), we adopt a leave-out strategy. That is, we exclude immigrants settling in the census region r where county i is located from the push factor, denoted by $-r(i)$, and we remove immigrants with surnames associated with the same origin country o as those with surname k from the pull factor, denoted by $-o(k)$. This approach ensures that our estimates are not driven by confounding factors that vary at the origin country-destination county level such as geographic similarity of the two regions. For instance, consider migrants from Denmark, a leading dairy producer at the time, who settled in areas suitable for dairy farming. If these Danes and their descendants continue to innovate in the dairy industry, and this innovation pulls more Danes to dairy counties in numbers large enough to affect aggregate flows from Denmark, then failing to exclude Danish migrants when constructing the pull factor for, say, the "Jensens" may confound the analysis. By removing immigrants from the same originating country when forming the pull factor, we address concerns that such confounding factors bias our estimates.

From this estimation, we derive predicted surname stocks

$$\hat{N}_{k,i}^t = \sum_{\tau=1880}^{t-1} \hat{b}^{\tau} \left(I_{k,-r(i)}^{\tau} \frac{I_{i,-o(k)}^{\tau}}{I_{-o(k)}^{\tau}} \right)^{\perp} \quad (4)$$

where $\hat{b}^{\tau}$ are the coefficients estimated from equation (3) and $\left(I_{k,-r(i)}^{\tau} \frac{I_{i,-o(k)}^{\tau}}{I_{-o(k)}^{\tau}} \right)^{\perp}$ are the residuals of a regression of the push-pull interactions on all the controls in (3), which ensures that the instruments are orthogonal to these controls.

Step 2: Constructing the Instrument for Surname Diversity We calculate the instrument for surname diversity using the predicted surname stocks $\hat{N}_{k,i}^t$:

$$\widehat{\text{Surname diversity}}_i^t = - \sum_k \left(\frac{\hat{N}_{k,i}^t}{\sum_k \hat{N}_{k,i}^t} \log_2 \left(\frac{\hat{N}_{k,i}^t}{\sum_k \hat{N}_{k,i}^t} \right) \right). \quad (5)$$

The data allow us to construct the instrument for surname diversity for eight periods: 1900, 1905, 1910, 1915, 1920, 1925, 1930 and 1940. These periods will be part of our main causal analysis.

Equation (2) controls for the population size of counties. Therefore, we also construct an instrument for population size by summing up the predicted stocks in each county for each period:

$$\widehat{\text{Population}}_i^t = \sum_k \hat{N}_{k,i}^t. \quad (6)$$

Identifying Assumption The identifying assumption is given by:

$$I_{k,-r(i)}^\tau \frac{I_{i,-o(k)}^\tau}{I_{-o(k)}^\tau} \perp \varepsilon_i^t \quad \forall k, \tau < t. \quad (7)$$

It requires that any factor affecting a temporary change in a county's innovation in t (ε_i^t) does not systematically correlate with earlier immigration of individuals with a given surname k to other regions within the US ($I_{k,-r(i)}$) interacted with immigrants from different origin countries than k simultaneously settling in that destination county i ($I_{i,-o(k)}/I_{-o(k)}$). If this condition is satisfied, the predicted stocks of surnames used to instrument for surname diversity are exogenous to innovation, and so is the instrument for surname diversity.⁶

⁶Our approach differs from the standard shift-share approach discussed in the literature. Typically, shift-share instruments are used as excluded instruments in a two-stage least squares regression, with identification achieved either by a large number of uncorrelated shifts (Borusyak et al., 2022) or by exogenous shares (Goldsmith-Pinkham et al., 2020). In contrast, we use a series of shift-share terms to predict stocks of surnames across U.S. counties. We leverage how the timing of historic waves of immigration (push) coincided with the timing of county-level attractiveness to immigrants (pull)—an arrangement that can be viewed as a series of historical quasi-experiments. We then use these predicted stocks to construct predicted surname diversity, which serves as an instrument for realized surname diversity in a new regression equation with a different set of controls. Despite these differences, our approach shares key features with the “many uncorrelated shifts” approach. For example, the number of surnames is very large (e.g., 454,566 in the 1940 regression, Columns 15–16 of Table 1). The fixed effects in our specification ensure that the push-pull instruments leverage only the residual variation in the shifts after absorbing time-invariant county-origin country and surname-destination region factors. The leave-out strategy removes immigration from a surname's origin-country, which help mitigate issues related to clustered (by origin country) and serially correlated shifts. Finally, including the full history of county immigration share controls approximates controlling for the sum of shares.

Predictive Power of the Instrument Table 1 reports the estimates of equation (3) for each period from 1900 to 1940. Following Burchardi et al. (2019), to facilitate the interpretation of coefficients as the marginal effect of migrations in that period, we sequentially orthogonalize each instrument with respect to the previous instruments. The push-pull instruments predict surname stocks well, with R^2 values ranging from 0.58 to 0.62 in columns without any controls. The F -statistic for the Wald test of the joint nullity of the push-pull instruments is highly significant in all periods. The coefficients on the instruments hardly change when we add the fixed effects and past immigration controls. The results indicate that push-pull instruments from most periods are significant predictors of the stock of surnames in subsequent years. For example, the estimates reported in Column 16 suggest that most of the instruments from 1880 to 1930 are significant predictors of the stock of surnames in 1940.⁷

Using the predicted surname stocks across counties and periods derived from estimating equation (3), we construct the instruments for surname diversity, equation (5), and population, equation (6). Figure B9 provides maps displaying this instrument, $\widehat{\text{Surname diversity}}_i^t$ in equation (5), for each 5-year period from 1900 to 1940, net of county and state-period fixed effects. Figure B10 displays the same variation but additionally nets out county-specific trends. The maps show substantial variation in the instrument over time and across counties.

4 County-level Surname Diversity and Innovation

Using the quasi-exogenous variation captured by our instrument, we provide estimates of the causal effect of county-level surname diversity on innovation at three levels—counties, surname-counties and inventors—and then assess the robustness of the results. We begin with our county-level findings.

⁷Qualitatively, our results parallel those of Burchardi et al. (2019), who estimate the push-pull factor at the level of originating countries (not surnames).

4.1 County-level Estimates

Table 2 reports our main county-level findings. Our baseline IV estimates of β_1 in equation (2) are in Columns 2 and 5 of Panel C, respectively. They estimate the impact of surname diversity in year t on the number of patents and breakthrough patents filed over the subsequent five years relative to other counties in the same state. The coefficients are positive and statistically significant for both patents and breakthrough patents. They suggest that a one standard deviation increase in a county's surname diversity increases the flow of patents by 1.13 per 1,000 people, a 107% increase with respect to the sample mean (1.05). The estimated effect on breakthrough patents is 0.104 per 1,000 people, an 80% increase with respect to the sample mean (0.13). The difference in surname diversity between average entropy in 1940 and Loving, TX, the county with the lowest surname diversity (see Table B3), implies $1.287 \times (9.71 - 6.01) = 4.762$ per 1,000 residents more patents and $0.115 \times (9.71 - 6.01) = 0.426$ per 1,000 residents more breakthroughs, corresponding to roughly 129 additional patents and 12 additional breakthroughs.⁸ When we compare these coefficients to those of population size, we find that the coefficient of surname diversity is about 1.76 times that of population size for patents and about 0.76 times for breakthrough patents. This suggests that the causal effects of scale (based on county population size) and diversity (based on surnames) are of the same order of magnitude.⁹

Panels A–C show our OLS, reduced-form and IV estimates, respectively. Panel D reports the first stage of the IV. Across specifications, the IV estimates are stable and similar to the OLS estimates, suggesting biases cancel out. These estimates hold even when we control for county-specific time trends, which allows us to focus on variation in the

⁸The coefficients are taken from unreported estimates of regressions that use unstandardized surname diversity as independent variables.

⁹We obtain similar magnitudes when estimating Poisson regressions. For patents, the estimated coefficient is 0.7641, which implies a treatment effect of $\exp(0.7641) - 1 = 1.15$, suggesting a 115% increase; for breakthrough patents, the estimated coefficient is 0.7801, implying a treatment effect of $\exp(0.7801) - 1 = 1.18$, or a 118% increase.

growth rate of patents within the same counties over time. Similarly, across specifications, the first stage has a strong F -statistic for both surname diversity and population size.

To visualize these effects, Figure 1 plots the reduced-form relationships between the instrument for surname diversity and both patent outcomes. They show that these relationships are strong, approximately linear and not driven by a small set of observations. Figure B11 shows that these relationships hold—even sharpen—when we additionally control for county-specific trends. Figure B12 depicts the first-stage relationship for our IV.

While data for our IV approach is only available from 1900 to 1940, Figure B13 presents OLS estimates for each decade extending back to 1850, the earliest census wave for which surname data are available. By regressing both patents and breakthroughs per capita on surname diversity (interacted with decadal dummies) while controlling for population size and including fixed effects for counties and state-periods, the figure shows the stability of this relationship across time. Beginning in the mid-19th century, we consistently find large, positive, and remarkably stable coefficients over time for both patents and breakthroughs.

4.2 County-surname-level Estimates

A potential concern with interpreting our findings so far is that surname-specific traits, such as abilities, interests, or knowledge, might drive innovation rather than the diversity of these traits. [Clark \(2014\)](#) and [Barone and Mocetti \(2021\)](#), for example, find that rare surnames are proxies for the vertical transmission of traits, and these traits might affect innovation.

We address this concern by estimating the following county-surname-level specification:

$$Y_{i,k}^t = \beta_1 \text{Surname diversity}_i^t + \beta_2 \text{Population}_i^t + \alpha_{i,k} + \alpha_{s(i),t} + \alpha_{k,t} + \varepsilon_{i,k}^t \quad (8)$$

where, as before, $\text{Surname diversity}_i^t$ and Population_i^t are county i 's surname diversity

and population size in year t . $Y_{i,k}^t$ is now the number of patents or breakthroughs filed by people with surname k residing in county i in the five-year period starting in t (in terms of 1,000 residents with the same surname in the year 1900). For example, 18,351 individuals with the surname ‘Johnson’ resided in Cook County (IL) in 1900 and filed 69 patents and 1 breakthrough patent between 1900 and 1904. Therefore, while the surname diversity remains defined at the county-period level, the innovation outcomes vary at the surname-county-period level.¹⁰

Crucially, we can now include county-surname fixed effects as well as surname-period fixed effects (denoted by $\alpha_{i,k}$ and $\alpha_{k,t}$, respectively), which absorb any traits specific to individuals with a particular surname in a given county or time period (e.g., traits specific to all ‘Johnsons’ in Cook County or in 1940). The remaining parameters and variables are as in equations (2) and the coefficient of interest remains β_1 . Standard errors are clustered on states and each observation is weighted by its relative share of the county’s population, making these estimates comparable to the county-level estimates in Table 2.

Table 3 reports our main surname-county level estimates. We find positive and statistically significant effects with effect sizes comparable to those in Table 2. In our most stringent specifications (columns 4 and 8), the estimates for β_1 indicate that a one standard deviation increase in surname diversity results in a 133% increase in patents (relative to the sample mean) and a 91% increase in breakthroughs. By comparison, a standard deviation increase in the size of a county’s population results in increases of 145% and 170% for patents and breakthroughs, respectively. Crucially, the surname-fixed effects ensure that the estimates are orthogonal to any unobserved surname-specific characteristics. Intuitively, this suggests that the “Winklers” get more innovative when they are in more culturally diverse places.¹¹

¹⁰Here, we switch the unit of analysis, resulting in an increase in the number of observations. For each county-period, the number of observations corresponds to the number of surnames in that county during that time period. We focus on surname-county observations that are balanced over the full panel from 1900 to 1940. See Appendix A for details on how we constructed the sample.

¹¹In many county-surname specifications, the reduced form results for breakthrough patents are imprecisely estimated, while the IV estimates are strong. The binscatter plot of the relationship, based on the fixed

4.3 Inventor-level Estimates

To shed light on the impact of surname diversity on innovation at the individual level, we extend our analysis to the inventor level. The patent data for this era does not include unique inventor IDs. This poses challenges related to identifying inventors due to inconsistencies in name strings and the presence of common names such as "John Smith." Since we only observe the name and the location of residence, this is particularly difficult for mobile inventors. Despite these obstacles, we have assigned unique inventor IDs for approximately 47% of the patents filed between 1900 and 1944 based on name similarities and their county of residence (see Appendix A for details). This effectively restricts the sample to non-moving inventors. While data on mobile inventors would be useful, these data still allow us to study how the productivity of inventors who remain in the same county changes over time as their county's diversity changes.

Using this dataset, we estimate the following inventor-level specification:

$$Y_j^t = \beta_1 \text{Surname diversity}_{i(j)}^t + \beta_2 \text{Population}_{i(j)}^t + \alpha_j + \alpha_{s(i),t} + \varepsilon_j^t, \quad (9)$$

where Y_j^t is the number of patents and breakthroughs filed by inventor j residing in county i during the five-year period starting in t . The specification includes inventor fixed effects (α_j) and state-period fixed effects ($\alpha_{s(i),t}$), which absorb unobserved inventor-specific traits and time-varying factors at the state level. By including inventor fixed effects, we ensure that our results are not influenced by any time-invariant personal characteristics of the inventors. As a result, the regression relies solely on variation arising from immigration-induced, quasi-random changes in the surname diversity of the inventor's county of residence. Standard errors are clustered at the state level.

Table 4 presents our main inventor-level results. Across specifications, we find positive effects specification in Column 6 and shown in Figure B14, together with our robustness test that drops counties with very high surname diversity (Table B34), suggests that this lack of significance is at least partly due to a small number of outliers at the top of the predicted surname diversity distribution.

effects of surname diversity on both patents and breakthrough patents. The IV estimates in Panel C indicate that a one standard deviation increase in county-level surname diversity leads inventors to file about half an additional patent during the five-year period starting in t , on average. This increase represents a 20% rise relative to the sample mean (2.24). The effect on breakthrough patents is around one-tenth of an additional breakthrough, roughly a 40% increase relative to the sample mean (0.26)—although the estimate lacks precision.¹² The estimates are very similar when we additionally control for population size, which interestingly does not affect patenting at the inventor level. The first stage has sufficiently strong F -statistics for both surname diversity and population size.

4.4 Robustness and Extensions

We test the robustness of our findings using nine different approaches at all three levels of analysis. First, to assess the impact of immigration on our findings, we explicitly control for it. Second, as a placebo check, we regress past innovation on future surname diversity. Third, to assess the role of broad geographic heterogeneity, we conduct our main analyses separately in each of the four major U.S. Census regions. Fourth, we consider alternative forms for both our innovation outcomes and key predictor variables. For innovation, we assess the impact of using log transformations. For surname diversity, we evaluate five alternative formulations. Fifth, we probe the robustness of our results by dropping the controls from our zero stage. Sixth, we drop the most surname diverse counties from the sample. Seventh, we explore the role of non-linear effects in surname diversity. Eighth, we assess the potential impact of schooling on our findings. Finally, we examine robustness to incorporating more conventional measures of cultural diversity.

¹²The wide confidence interval is driven by the most surname-diverse counties. Excluding the top 5–10 percent of counties in terms of surname diversity leaves the IV coefficient on breakthrough patents almost unchanged but cuts its standard error by roughly one-half and raises the first-stage F from 12 (in Column 4) to above 142, making the estimate statistically significant (see Table B35).

4.4.1 Sensitivity to Immigration

Because our instrument is constructed from immigration data, a potential concern is that immigration, operating through channels not related to either population size or surname diversity, may influence patenting. Previous research suggests that immigrants significantly fuel innovation due to factors like higher skills, entrepreneurial spirit or unique patentable knowledge (Moser et al., 2014; Abramitzky and Boustan, 2017; Sequeira et al., 2020; Terry et al., 2024).

We address this concern in two ways. Fundamentally, we orthogonalized the push-pull instruments in equation (3) using counties' historical immigration patterns. Thus, addressing this concern lies at the core of our identification strategy. Here, to confirm the effectiveness of this approach, we additionally control for the number of recent immigrants arriving in each county during each period. Our instrument is based on changes in surname diversity induced by immigration, which depends on the existing surname distribution in a county. Thus, this approach allows us to separate the effects of recent immigration from those of surname diversity.

Table B7 presents these results for our county-level specification (2).¹³ In line with previous work, we find a positive relationship between immigration and innovation across specifications. Crucially, for our purposes, we also find that controlling for immigration has little impact on the estimates for surname diversity. For example, the IV estimate for patents per 1,000 people is 1.087 without controlling for immigration (Column 1 of Panel C), compared to 1.063 with this control (Column 2). The estimates for breakthrough patents are 0.083 and 0.079 in Columns 5 and 6, respectively. We observe similar patterns when controlling for immigration in the county-surname (Table B8) and inventor-level analysis (Table B9). Taken together, these findings suggest that the orthogonalization of our instrumental variable to the influence of immigration worked as expected and

¹³We exclude data from the period ending in 1940 due to the lack of immigration year information necessary for calculating recent immigration.

confirms that our estimates are not confounded by immigration-related factors other than surname diversity.

4.4.2 Placebo Test and Dynamic Effects

Tables B10, B11 and B12 show the dynamic effects of changes in surname diversity on our innovation outcomes at the county, county-surname and inventor level, respectively. Two findings stand out. First, looking back, we do not observe a positive effect of current surname diversity in period t on past innovation in the previous period, $t - 5$ (Columns 1 and 5). This lack of evidence for reverse causality strengthens our confidence in the identification strategy. Second, looking forward, we find that surname diversity impacts innovation not only in period t (0 to 5 years after we observe surnames) but also in period $t + 5$ (5 to 10 years after our observation of surnames) at both the county and county-surname levels.

4.4.3 Estimates for Major U.S. Regions

Tables B13, B14 and B15 report the estimates for the four major U.S. census regions: the Midwest, Northeast, South, and West. We observe that across specifications, the region-specific estimates are uniformly positive and mostly statistically significant. At the county level, the range of IV coefficients for patents across regions (Panel C, Column 1) spans from 0.79 to 1.74, and for breakthrough patents (Column 3), from 0.07 to 0.19. The larger coefficients for the West are particularly interesting, possibly reflecting the region's unique historical and immigration context. The estimates for the South further refute the idea that our results are driven by high immigration flows since the South experienced relatively low rates of immigration during this period.

4.4.4 Alternative Forms of Key Variables

We consider the sensitivity of our results to alternative forms of measuring both innovation and surname diversity. For innovation, we justify our decisions to normalize (using population in 1900) and winzorize our patent-based measures on statistical grounds, as detailed in Section 2 and Appendix A. Nevertheless, we examine the effects of using log-transformations (without winsorization) for patents and breakthroughs to address the right skewness present in these measures, and we also consider the impact of removing the normalization. To avoid dropping observations with zeroes, we add 1 to the patent variables. Tables B16 and B17 show the results for log population-normalized outcomes and Tables B18, B19 and B20 report the estimates for log non-normalized patents and breakthroughs. In sum, our results are robust to these transformations of the dependent variables.¹⁴

For surname diversity, Section 2 explains our choice of an entropic diversity measure based on conceptual, statistical and parsimony considerations. Nevertheless, it is important to explore how critical this choice is to our results. Conceptually, diversity is composed of three parts: (1) ‘surname variety,’ or a function of the number of distinct surnames in a county, (2) ‘dispersion,’ which indicates how individuals are distributed across surnames (the opposite of ‘balance’ or ‘concentration’) and (3) ‘disparity,’ which measures the differences between surnames (e.g., [Stirling, 2007](#)). Our measure combines ‘variety’ and ‘dispersion’ but does not include a metric to capture the ‘disparity’ or cultural distance among surnames. That is, our entropy measure treats each surname as an independent category equidistant from other surnames. This means that our approach does not account for the fact that the Smiths and Browns might share more information due to a shared cultural background than the Smiths and Garcias.

To address this, we create a distance-weighted measure of surname diversity by in-

¹⁴For log-transformed breakthroughs, the coefficients at the county-surname level are positive but sometimes imprecisely estimated.

corporating a proxy for cultural distance based on genetic distance (following [Spolaore and Wacziarg, 2009](#)). As explained in Appendix A, we assigned each surname to an ancestral-country and then used the genetic distance between these countries and the U.S. to assign a d value to each surname. Specifically, we compute $-\sum dp \log p$, where d is the genetic distance, our proxy for cultural distance. The resulting weighted measure of surname diversity is highly correlated with our baseline entropy measure ($\rho \approx 0.87$).

Table B21 reports the estimates for our county-level equation (2) using this distance-weighted entropy measure as the main endogenous variable. We find that the distance-weighted entropy measure increases both of our innovation measures across specifications (Columns 1–4 and 6–9). However, when we include our more parsimonious baseline entropic measure of surname diversity (Columns 5 and 10), the coefficients for distance-weighted diversity shrink, lose significance and even switch directions (in 5 out of 6 estimates). We observe similar patterns at the county-surname and inventor level (Tables B22 and B23).

These results suggest that adding a measure of disparity, at least at the level of the ancestral country, does not improve our results. Next, we break our surname diversity measure down into its two components: variety and dispersion. Table B24, B25 and B26 report OLS comparisons between our baseline surname diversity and its components. At the county level, we find that when our conceptually-motivated entropy measure of surname diversity is included, the coefficients on two components are small, mostly poorly estimated and sometimes go in the wrong direction. We find broadly similar patterns at the county-surname and inventor levels. Taken together, these results indicate that our entropic surname diversity—the interactions of variety and dispersion—captures the relevant relationship between surname diversity and innovation in an effective way. Neither of its components is the dominant driver of the results, suggesting that a high degree of variety is much more effective in fostering innovation when paired with a high degree dispersion.

Finally, while we prefer our entropic measure for the reasons outlined above, we also report results from using a more traditional measure of diversity in economics: a fractionalization index, calculated as $1 - \sum p$. Tables B27, B28 and B29 demonstrate that the county and inventor-level results are robust to using surname fractionalization rather than surname entropy.

4.4.5 Excluding Controls from Zero-stage Estimation

A potential concern with our empirical strategy arises from the inclusion of controls and fixed effects in the zero-stage regression used to predict surname stocks. Specifically, while our push-pull interaction terms are orthogonalized with respect to these controls, the estimated coefficients b^τ from equation (3) could still be functions of the included controls and fixed effects. This dependency may introduce endogeneity into our predicted surname stocks $N_{k,i}^t$, as these stocks are computed using the estimated coefficients b^τ . Consequently, the exogeneity of our instrument—the predicted surname entropy measure—might be compromised if the coefficients capture effects related to the controls rather than solely the exogenous variation from the push-pull factors.

To explore the importance of this concern, we conduct a robustness check by re-estimating the zero-stage regression without any controls or fixed effects. By removing these additional variables, we ensure that the estimated coefficients b^τ are solely determined by the push-pull interactions. This approach isolates the variation in surname stocks that is attributable only to the shocks captured by the push and pull factors.

The odd-numbered columns in Table 1 report these estimates. As expected, the R^2 statistics drop by about 0.10. However, the coefficients on the push-pull instruments remain stable. We then recalculate the predicted surname stocks $N_{k,i}^t$ using these new coefficients and reconstruct our predicted surname entropy measure. Re-estimating our county-, surname-county and inventor-level specifications with this alternative instrument, we find that the results remain consistent with our original findings (Tables B30, B31 and

B32). The estimated effects are similar in magnitude and statistical significance, suggesting that our conclusions are not driven by the inclusion of controls in the zero-stage regression.

4.4.6 Dropping Most Surname Diverse Counties, Rare and Frequent Surnames

A potential concern may be that outliers drive our results. While all regressions rely on within-state variation, one might be concerned that the results largely reflect the effects of large, diverse urban areas. To assess this possibility, we replicate our estimates in subsamples that drop the 5%, 10% and 20% most diverse counties (measured in 1900). These estimates are reported in Tables B33, B34 and B35. We find that the estimates in these subsamples are very similar to our baseline ones.

A related concern may be that locally rare and frequent surnames may affect our results. To tackle this concern, we recompute the entropy and population measures after excluding the top and bottom 5% of surnames in each county. The results, presented in Tables B36, B37 and B38 are very similar to our main findings, suggesting that our findings are not unduly driven by the inclusion of locally rare and frequent surnames.

4.4.7 Surname Diversity Squared

Some scholars have suggested that cultural diversity may bear a concave relationship with innovation because greater cultural diversity, at a certain point, may induce less social interaction ([Schimmelpfennig et al., 2022](#)), thus reducing innovation indirectly via the social-behavioral channel. To address this potential non-linearity, we add a squared term for surname diversity to our regressions at the county, county-surname and inventor level. As reported in Tables B39, B40 and B41, we do not find clear evidence of any non-linear effects of surname diversity on innovation.

4.4.8 Education

Given that Americans experienced rising levels of formal schooling over our study period, and that education is linked to innovation, we consider how including education and educational diversity in our analysis impacts our key effects. Indeed, while the cultural diversity captured by surname diversity provides a rich source of ideas, more schooling may facilitate the transformation of these ideas into patentable innovations. Notably, both [Alesina et al. \(2016\)](#) and [Terry et al. \(2024\)](#) show that skill levels of immigrants influence rates of innovation.

To assess the importance of formal education in our setting, we estimate specifications based on the cross-section of U.S. counties in the year 1940—the only census wave in our sample that reports years of schooling. Tables B42 and B43 present the estimates for specifications that include years-of-schooling diversity and average years of schooling, respectively. As a benchmark, we first regress our two measures of innovation on surname diversity and population size (Table B42, Columns 1 and 4). Then, in Columns 2 and 5, we regress these innovation variables on years of schooling diversity, and in Columns 3 and 6, we include both diversity measures.

We find that years-of-schooling diversity is significantly related to our innovation outcomes; however, the coefficients are about half the size to those on surname diversity. When both diversities are included, the coefficients on surname diversity remain stable, while the ones on educational diversity shrink by another half.

Next, turning to Table B43, we regress our patent outcomes on average years of schooling in Columns 1 and 4. In Columns 2 and 5, we add surname diversity, and in Columns 3 and 6 we examine the interaction between surname diversity and years of schooling.

Across all specifications, we find a significant effect of surname diversity on innovation. Schooling is also related to patenting; however, its inclusion affects the coefficients on surname diversity to a limited degree. Notably, the interaction between the two independent variables is highly significant. These findings suggest that surname diversity fuels

patenting, especially among more highly educated individuals. Indeed, the magnitude of the relationship between schooling and innovation almost doubles with an increase in surname diversity by one standard deviation.

4.4.9 Surname vs. Other Measures of Cultural Diversity

How does our surname diversity measure compare to more conventional indicators of cultural diversity, which have typically been based on birth-countries or even broader categories such as continental origins (e.g., "Asians")? One might argue that surnames are merely a noisy proxy for the cultural diversity captured by ancestral origin countries. To examine this, we conduct two "horse race" analyses that directly compare surname diversity with alternative diversity measures.

To this end, we draw on the Census data to construct measures of diversity based on birthplace, race, and occupation. We include occupational diversity because, although occupation is not traditionally considered a form of culture, individuals' occupational choices are influenced by learning from same-sex parents and broader community networks. Additionally, as mentioned in Section 2, we infer the ancestral country for each surname using Census data from 1850 to 1940 to predict the most likely country of origin. This completes our set of four alternative diversity measures.

Table B6 presents the conditional correlations between surname diversity and these alternative measures, controlling for population and various geographical and temporal fixed effects. For instance, ancestral-country diversity is moderately correlated with surname diversity, with conditional correlation coefficients ranging from 0.26 to 0.41, depending on the fixed effects and population control variables included.

To assess the predictive power of surname diversity relative to these alternative measures, we estimate a series of OLS regressions, the results of which are reported in Table 5. Panels A and B show the estimates for patents and breakthroughs, respectively. Replicating Columns 3 and 7 from Table 2, Column 1 establishes a baseline model that includes our

county-level surname diversity measure and county population size.

Birth- and Ancestral-country Diversity In Columns 2 and 4, we replace surname diversity with birth-country diversity and ancestral-country diversity, respectively. The standardized coefficients on these alternative diversity measures are positive and statistically significant, though notably smaller than the coefficient on surname diversity in the baseline model (Column 1). This positive relationship aligns with prior research (Ottaviano and Peri, 2006; Ager and Brückner, 2013; Alesina et al., 2016; Docquier et al., 2020; Fulford et al., 2020). In Columns 3 and 5, we reintroduce surname diversity into the models. When surname diversity is included, the relationships between innovation and both birthplace and ancestral-country diversity shrink dramatically and cannot be distinguished from zero at conventional levels of precision. By contrast, the coefficient on surname diversity remains large and robust, comparable to that in the baseline model. If surname diversity were merely a noisy proxy for ancestral or birth-country diversity, we would expect its coefficient to attenuate when these alternative measures are included; yet, this is not what we observe.

Racial Diversity Columns 6 and 7 incorporate racial diversity into the analysis. The inclusion of racial diversity does not materially affect the coefficients on surname diversity, which remain stable and statistically significant.

Occupational Diversity In Column 8, we examine the relationship between occupational diversity and innovation. Consistent with our expectations, occupational diversity exhibits a strong positive relationship with innovation, although the coefficient is approximately half the magnitude of that on surname diversity in the baseline model. In Column 9, we add surname diversity to the regression. The coefficient on surname diversity remains substantial and statistically significant, with only a slight reduction compared to the baseline. This finding is consistent with evidence that occupation is one among several

traits transmitted across generations.

Immigrant Shares by Birth Country We also consider the population shares of immigrants from each of the 59 birth countries to account for the role of immigration from specific origins. By including these immigrant shares, we aim to control for the potential impact that immigrants from particular countries may have on innovation outcomes. When we include these shares in the regression models in Column 10, we find that the coefficients on surname diversity remain virtually unaffected.

Combined Diversity Measures Finally, Columns 11 and 12 include all diversity measures simultaneously in the model. The coefficients on surname diversity remain stable and significant, similar to that in the baseline model. In contrast, the coefficients on birth country, ancestral-country, and racial diversity become statistically insignificant or negative. Occupational diversity continues to show positive relationships with innovation, remaining significant except in Column 11 but loses significance in Column 12, where county-specific linear trends are included. However, the coefficients on occupational diversity are substantially smaller—ranging from about 40% to 10%—than those on surname diversity.

IV Estimates for Ancestral-Country Diversity To extend our analysis beyond the correlations observed in Table 5, we construct an IV for ancestral-country diversity using the same methodology applied to surname diversity, but substituting inferred ancestral countries for surnames. This approach yields an IV for ancestral-country diversity that parallels our instrument for surname diversity, allowing us to employ a similar IV strategy.

To generate predicted stocks of individuals from each ancestral country, we estimate a variant of equation (3), where the dependent variable is the number of individuals from a specific ancestral country residing in a given U.S. county. The push and pull factors are redefined at the ancestral country–county level. We adjust the leave-out strategy by

excluding immigrants from countries neighboring the ancestral country from the pull factor (see Appendix A for details). The estimates from these zero-stage regressions are reported in Table B44. Consistent with our findings for surnames, the push-pull instruments are highly significant predictors of ancestral-country stocks.

Using these predicted ancestral-country stocks across counties and time periods, we construct the IV for ancestral-country diversity. We then re-estimate equation (2), incorporating both surname diversity and ancestral-country diversity as independent variables.

Table B45 presents the results of these IV regressions. The coefficients on surname diversity remain large, stable, and precisely estimated, with magnitudes comparable to those reported in Table 2. In contrast, the coefficients on ancestral-country diversity are small, imprecisely estimated, or even negative in some specifications.

These results support the OLS horse race findings in Table 5, suggesting that surname diversity is not simply an imprecise proxy for ancestral-country diversity. Below, in our discussion of mechanisms, we will delve deeper into the pathways through which these different types of cultural diversity affect innovation.

5 Patent-level Surname Diversity and Innovation

We extend our analysis from the impact of county-level surname diversity on innovation to estimate the effect of patent-level surname diversity on patent quality and complexity.¹⁵

We estimate:

$$Y_p^t = \beta \text{ Surname diversity of patent } p^t + \alpha_{n(p)} + \alpha_{i(p)} + \alpha_{s(p),t} + \alpha_{f(p),t} + \varepsilon_p^t, \quad (10)$$

where Surname diversity of patent_{*p*}^{*t*} is the diversity of surnames of the inventors who filed the patent *p* in the five-year period starting in *t*. *Y_p^t* represents either the breakthrough

¹⁵In our sample, 162,424 patents (about 11%) have more than one patentee, and 11,963 of these have a single surname (Table B46). These percentages remained fairly stable over the historical period we consider (Figure B15).

indicator or the number of technology codes assigned to the patent. The latter has been used as a measure of the complexity of a patent and the degree of recombination (Strumsky et al., 2011; Akcigit et al., 2013; Fiszbein, 2022; Youn et al., 2015).¹⁶ The specification includes fixed effects for inventor team size ($\alpha_{n(p)}$) because larger teams often have higher surname diversity and tackle more complex and impactful projects. It includes county ($\alpha_{i(p)}$) and state-period ($\alpha_{s(p),t}$) fixed effects to absorb geographic and time-varying differences in team diversity, patent quality and complexity. It also includes fixed effects for patent technology fields interacted with period fixed effects ($\alpha_{f(p),t}$).¹⁷ Different fields may have varying team compositions, and patents in these fields can differ in impact and complexity, which may change over time. In our most stringent specifications, we include county-period fixed effects ($\alpha_{i(p),t}$), absorbing any county-specific, time-varying unobserved factors. If multiple inventors reside in different counties, the patent is assigned to each county and weighted by $\frac{1}{n}$, where n is the number of co-inventors. Standard errors are clustered at the county level.

To identify the causal effect of patent-level surname diversity on patent outcomes, we instrument for patent-level surname diversity using our IV for county-level surname diversity. This approach allows us to estimate the local treatment effect of *patent* surname diversity on the quality and complexity of innovation for patents whose surname diversity is influenced by county-level variation.

Table 6 reports these estimates. Panels A (OLS), B (reduced-form) and C (IV) indicate that greater surname diversity at the patent level increases the likelihood of a patent becoming a breakthrough and having more technology codes, suggesting these are more complex recombinations. The IV effect size is large, implying that a one-standard deviation increase in surname diversity at the patent level delivers roughly a 91 percentage point increase in the likelihood of breakthroughs and an increase of roughly four additional technology codes per patent (the mean number of codes per patents is 2.49), according

¹⁶There are 159,402 patent technology codes in our sample from 1900 to 1944.

¹⁷These are six broad technological fields: chemical, computers, drugs, electrical, mechanical or other.

to Columns 2 and 5. We further confirm that the OLS estimates remain robust when we incorporate county-period fixed effects, thus comparing patents within the same county and five year period (Columns 3 and 6).

Panel D shows the first stage of the IV. The estimates suggest that a one standard deviation increase in predicted county-level surname diversity increases the diversity of co-inventor surnames at the patent level by around 0.01 standard deviations. This establishes that higher surname diversity in a county generates more diverse interactions among patentees, at least in terms of their surnames. This relatively small effect arises, in part, from the fact that most patents during this period have only a single inventor.

Notably, the IV estimates here are substantially larger than the OLS estimates. Part of this gap might be due to the moderate strength of the IV. With an F -statistic of 9–10, the IV is at the lower bound of what is commonly accepted as a sufficiently strong instrument. Additionally, we suspect a downward bias affects the OLS estimates due to the cumulative nature of innovation. It is well established that patents build on prior patents, and particularly on breakthroughs ([Acemoglu et al., 2016](#); [Youn et al., 2015](#); [Akcigit et al., 2013](#)). Breakthroughs often integrate ideas from diverse domains, leading to ‘refinements’ and ‘reuse patents’ among similarly diverse or even the same sets of co-inventors who build on the insights found in the breakthrough. This creates reverse causality running from breakthroughs to both patent surname diversity and technology codes. Such refinements and reuse patents rarely become breakthroughs themselves and tend to focus on specific domains of the breakthrough, thus, the downward bias for both outcomes. Our IV taps a source of plausibly exogenous variation in surname diversity, removing this bias and delivering larger estimates. Without the IV, the observed correlation between patent-level surname diversity and breakthrough likelihood and patent complexity seems small.

One might worry that our entropic measure of surname diversity might not apply well to inventor teams. To address this, Table B47 reports the relationship between county-level surname diversity and three alternative measures of patent-level diversity,

separately capturing variety ($\log$ number of distinct surnames among patentees, $\log N$) and dispersion ($(-\sum p \log p)/\log N$) as well as our measure that includes disparity (genetic distance-weighted surname entropy: $-\sum dp \log p$). The IV estimates in Panel C indicate that shifts in county-level surname diversity result in similar changes across these alternatives. When we include team size fixed effects (even-numbered columns), the estimates imply that surname diversity decreases the likelihood that inventors on a same-sized team share the same surname.¹⁸

Finally, we investigate whether our main patent-level findings replicate when using genetic distance-weighted surname diversity and ancestral-country diversity. First, our estimates hardly change when using patent- and county-level genetic distance-weighted surname diversity, although the instrument has become weaker (Table B49). Second, Panel C of Table B50 shows that county-level ancestral-country diversity is not meaningfully correlated with patent-level outcomes. We also do not observe a relationship with the ancestral-country diversity of patents, and even a negative association with the surname diversity of patents (Panel D), suggesting that greater ancestral-country diversity within counties does not lead to more diverse collaborations among inventors. Lastly, Panel A reports that ancestral-country diversity is positively correlated with patent-level outcomes, implying that once a collaboration is established, patents filed by more diverse teams, in terms of ancestral backgrounds, tend to be more complex and are more likely to be high impact. However, this correlation disappears for the production of breakthroughs in Panel B when we add patent-level surname diversity.

Taken together, these results establish an important set of links: they connect surname diversity at the county-level to more diverse interaction at the patent-level and then connect the more diverse teams to more complex inventions that are more likely to become

¹⁸A comparison with the estimates reported in the even- and odd-numbered columns of Tables B47 and B48, showing the results with and without team size fixed effects for all patents and those with 2 or more co-inventors, respectively, suggests that shifts in county-level diversity may (1) increase the number of 2-person teams (with different surnames) but (2) for larger teams (2+ co-inventors), primarily operate through greater team diversity (not size).

breakthroughs. This opens the door to probing the the underlying mechanisms more deeply.

6 Mechanisms

Our hypothesis posits that greater surname diversity increases the production of both patents and breakthrough patents because innovation often emerges through the recombination of ideas that occurs when diverse individuals interact. Above, we connected county-level surname diversity to patent-level diversity and to innovation quality. To further flesh out the underlying mechanisms, we use census and patent data to explore the extent to which surname diversity affects the informational component—different people must possess distinct ideas—and social-behavioral component—individuals must be willing to interact and share their ideas. Jointly, these two components should facilitate the creation of innovative recombinations.

6.1 Informational Mechanism

Our approach assumes that surnames in the U.S. (1850-1940) capture significant informational diversity within populations, which arises from social learning and the intergenerational transmission of information. Conceptually, this aligns with previous work on population diversity and growth that uses markers like country of birth or ancestry to measure such diversity (e.g., [Alesina et al., 2016](#); [Fulford et al., 2020](#)). Here, as proof of concept, we show that at least some forms of knowledge and skills cluster in surnames during our study period. As noted above, we argue that surnames also capture other informational differences (not observed here), such as those related to techniques, languages, perspectives, intuitions and interests.¹⁹

¹⁹Research from across the social sciences indicates that surnames should capture substantial and important informational differences. Three lines of evidence are relevant. First, across populations, people rely on different languages, thinking styles, decision heuristics, preferences, metaphors, attentional biases and ritual practices ([Nisbett, 2003](#); [Falk et al., 2018](#); [Henrich, 2020](#)). Work in cognitive science, for example, indicates

6.1.1 Surnames Cluster Information

To assess whether surnames capture relevant information in our context, we calculate Herfindahl-Hirschman Indices (HHIs) that capture how strongly surnames cluster in four domains: (1) education, (2) occupations, (3) birthplaces and (4) technology patent classes.²⁰ The construction of the concentration measure for each domain proceeds in two steps. First, using occupation as an example, we calculate HHIs for each surname across all occupations. This gives us a measure of how strongly each surname clusters in occupations. We normalize this measure such that a value of zero implies a uniform distribution of a surname across occupations and a value of one implies that a surname only occurs within a single occupation. Second, we average the surname-specific HHI across all surnames, weighted by the number of people with a given surname. This averaged index reveals the overall surname concentration across occupations based on the full population. Following this protocol, we construct the HHIs for the other informational domains and apply the same approach to groupings based on the birth country and ancestral country, replacing surnames with more conventional measures of cultural diversity.

Table 7 reports the surname concentration indices for the different domains, samples, and years in the first row. Here, we use only adult men (ages 25-60), but Table B51 reports the results when all individuals are included. In these tables, all surname concentration indices are well above zero, indicating that surnames are concentrated in years of schooling (Column 1), occupations (Columns 2 to 5), countries of birth (Columns 6 and 7), regions of

that speaking and thinking in different languages has consequences for people's perceptions, attention and reasoning (Blasi et al., 2022). So, culturally-evolved cognitive differences are common. Second, research on social mobility indicates that surnames reflect unique skills and socialization (Clark, 2014; Güell et al., 2015; Barone and Mocetti, 2021) and much cultural transmission flows from same-sex parents to their children (Hewlett, 2016; Schniter et al., 2023). Third, as explained above, analyses of U.S. patents and Swedish entrepreneurship reveal the significance of cultural transmission from both same-sex parents. The latter two points indicate that some of these cognitive differences will occur among family groupings.

²⁰There are 307 and 407 patent technology classes in our 1880-9 and 1949 samples, respectively, including categories like "Geometrical Instruments", "Stoves and Furnace", and "Chemistry: Electrical and Wave Energy."

birth within Germany (Column 8),²¹ and patent technology classes (Columns 9 and 10).²² To illustrate, two people with the same surname in 1880 have a roughly 17% (above chance) probability of holding the same occupation out of 249 possible occupations (Column 2). For example, "Smith" and "Johnson", the most frequent surnames in the 1880 census, show a distinct pattern. The Smiths outnumber the Johnsons by a factor of about 1.65 times. Yet, among individuals who reported "blacksmith" as their occupation, there are 2.46 times more Smiths than Johnsons. "Smith," of course, has a long-standing association with metal-working and this result illustrates this pattern in 1880. Similarly, two same-surname immigrants have about a 42% probability of having the same country of birth (Column 6) and, if they are German, a 27% probability of being from the same inner-German region (out of 32 regions). Finally, surnames cluster on educational attainment (Column 1) and on the fine-grained technology classes applied to patents (Columns 9 and 10). The latter finding is consistent with Bell et al.'s 2019 analysis of father-son patenting patterns in the U.S. at the dawn of the 21st century.

To put these concentration indices for surnames into context, we compare them to parallel HHI indices for country of birth (row 2), country of ancestral origin (row 3) and residence county (row 4). We calculate these in the same fashion as with surnames. Compared to country of birth and ancestry, surnames are about equally indicative of education and occupation; however, for patent technology classes, surnames are substantially more indicative.²³ In 1880, occupations are relatively more concentrated in counties than

²¹This domain is restricted to German regions because fine-grained subregional birthplace data are available for this country only.

²²The patent data set does not allow us to uniquely identify the universe of inventors. Hence, we cannot detect all inventors who file multiple patents in the same technology category, which could bias the concentration upwards. We still report this statistic because this bias is likely small, given the low level of regional clustering in this variable (row 2), where we would expect a similar upward bias if regional mobility among inventors is not very high.

²³A potential concern is that, although surnames may often be nested within coarser categories like country of origin, regional birthplace and race, there have been historical processes that muddy this hierarchical nesting. For example, many formerly enslaved Africans carry the European surnames of their enslavers (Cook et al., 2022). Surname diversity may thus underestimate the diversity stemming from African heritage. To address this, we construct a more fine-grained measure and check how it relates to our main indicator. This measure creates additional 'surname categories' based on race-surname combinations. For example, the number of white "Jacksons" enters the diversity indicator as a separate category from the number of black

in surnames, but this difference narrows markedly by 1940. By contrast, surnames are substantially more indicative of countries and regions of birth compared to immigrants' residence counties.

Building on the clustering of important informational types demonstrated in Table 7, Table 8 uses equation (2) to connect our county-level surname diversity measures to occupational and patent technology code diversity. Column 1 and 3 of Panel C show that a one standard deviation increase in surname diversity leads to roughly a one-third standard deviation increase in occupational and technology code diversity. This finding is robust to additionally controlling for ancestral-country (Columns 2 and 4) and birth-country diversity (Table B52). Taken together, these results supports the role of the informational channel in driving innovation. However, as noted previously and shown in Table 5, surname diversity captures more than just occupational diversity.

6.2 Social-Behavioral Mechanisms

Distinct forms of cultural diversity could potentially produce different indirect effects on innovation by either fostering or suppressing social interactions among diverse minds. As described above, some existing research leads us to suspect that greater diversity might lead to greater segmentation, fragmentation and polarization in a population (Alesina and La Ferrara, 2005). This would inhibit many forms of social interaction and informational exchanges, thereby reducing the chances for innovative recombinations (Andrews, 2023; Andrews and Lensing, 2024). Alternatively, as surname diversity increases and the relative sizes of informationally distinct surname groups shrink, the rate of diverse social interactions may increase for two reasons: (i) people may be less able to satisfy their needs

"Jacksons". Similarly, we calculate a surname diversity indicator that further differentiates along country of birth. Table B5 shows that the main surname diversity indicator in 1940 is almost perfectly correlated with those more fine-grained diversity measures. Furthermore, we obtain similarly high correlation coefficients between the main surname diversity indicator and indicators that are based on (i) phonetically uncorrected surnames, (ii) surnames of men only, (iii) surnames of household heads only, and (iv) surnames of whites only.

by interacting solely within their own small and relatively diminishing groups, and (ii) the spread of different groups with distinct information may create new avenues for mutually beneficial exchanges. Such fruitful engagements with outsiders may foster greater openness to social interactions with diverse others, which in turn (and often incidentally) may promote the recombination of ideas. Both hypotheses could apply to different forms of diversity, with large ethnic or racial groupings operating differently from smaller groups, such as families, lineages or clans.

Based on prior work, greater surname diversity may result in greater social interaction among diverse minds. This is because surname diversity partially captures the extent to which large family groups are present. A growing body of research suggests that strong family ties and large residential kinship groups influence both people’s psychology—leading to less trust and openness toward strangers—and their social behavior—resulting in less interaction with outgroup members ([Enke, 2019](#); [Alesina and Giuliano, 2014](#); [Schulz et al., 2019](#); [Schulz, 2022](#); [Bahrami-Rad et al., 2024](#)).

To better understand the potential indirect impacts of surname diversity on innovation, we construct three proxy variables for the likelihood of diverse social interactions at the county level: (i) the residential segregation of surnames, (ii) the residential segregation of ancestral-country groups and (iii) the strength of family ties.

The first two measures, calculated in parallel fashion, aim to capture the residential clustering for two types of groups, one based on surnames and the other based on ancestral countries. For each, we constructed a measure of next-door neighbor segregation closely following the methodology of [Logan and Parman \(2017\)](#). Constructing these measures involves two steps: first, quantifying the degree of segregation for each group (surname or ancestral country) in each county during specific periods; second, averaging these degrees across all groups within a county to obtain a county-level measure of segregation (see Appendix A). This gives us two measures of segregation: (1) for surnames, our measure captures the cluster of same-surname pairs at the street-level, indicating segregation *beyond*

what would be expected from random choices of residence; and (2) for ancestral countries, our measure captures the cluster of same-ancestral-country pairs at the street-level. Figure B7 depicts the geographic distributions of both types of residential segregation in 1940.

Crucially, these segregation measures are constructed to account for the distributions of surnames or ancestral countries at the county level. Positive values signify segregation beyond what would occur by chance, while negative values imply that neighbors are more diverse than what chance alone would predict, indicating a *preference for diversity*. Assessing the degree of segregation in relation to a county's pattern against that expected by chance allows us to capture people's settlement preferences. Thus, an association between surname diversity and the segregation of either surnames or ancestral countries is not merely driven by more diverse counties being necessarily less segregated. Rather, it reflects differences in fine-grained settlement preferences while holding background diversity constant. Both surname segregation and ancestral-country segregation measures are highly negatively correlated with surname diversity ($\rho = -0.71$ and $\rho = -0.54$, respectively), indicating that individuals in counties with low surname diversity tend to live in closer geographic proximity to others with the same surname and country of ancestry.

The third measure, the strength of family ties, captures the size, homogeneity, and stability of households as social units (Raz, 2025). It represents the first principal component of four county-level variables: (i) the divorce-to-marriage ratio, (ii) the share of elderly people living without a relative, (iii) the share of people living with at least one person who is not their relative, and (iv) the mean size of families (Appendix A). Higher values for the strength of family ties indicate a preference for surrounding oneself with relatives. Figure B8 illustrates the geographic distribution of the strength of family ties in 1940. This variation is fairly strongly negatively correlated with surname diversity ($\rho \approx -0.50$), suggesting that individuals in counties with low surname diversity tend to have strong kin-based networks.

To analyze the effect of surname diversity on the likelihood of diverse social interactions,

we re-estimate our county-level specification (2), replacing the dependent variable by our three proxy measures. Table 8 reports the results in Columns 5, 7 and 9. The IV estimate in Column 5 of Panel C indicates that a one standard deviation increase in surname diversity reduces the residential segregation of surname groups by 0.41 standard deviations (below what would be expected by chance alone). For the segregation of ancestral-country groups, the impact of surname diversity is even larger (Column 7, Panel C): a one standard deviation increase in surname diversity leads to a one standard deviation reduction in residential segregation of ancestral-country groups. Similarly, as shown in Column 9 of Panel C, a one standard deviation increase in surname diversity results in the strength of family ties weakening by 0.32 standard deviations. In all cases, the OLS and reduced-form estimates in Panels A and B are similar to the IV estimates. These results suggest that an increase in surname diversity makes diverse social interactions within counties more likely.

We also consider what happens when we add our measure of ancestral-country diversity to these specifications. These estimates are presented in Columns 6, 8 and 10 in Table 8. Crucially, we find that the coefficients on surname diversity remain negative, large, well-estimated and stable. However, in contrast, the results for ancestral-country diversity are not consistent across Panels A, B, and C. The IV estimates in Panel C are positive and well-estimated, suggesting that greater ancestral-country diversity may result in greater segregation and stronger family ties—precisely the opposite of the impact of surname diversity. However, the coefficients in the OLS and reduced-form regressions are small, sometimes poorly estimated, and sometimes flip signs. We observe similar patterns when we incorporate birth-country diversity instead of ancestral-country diversity (Table B52). Thus, unlike for surname diversity, we do not find consistent evidence that ancestral- or birth-country diversity have the same impacts on social interaction.

The above results expand upon our earlier patent-level findings in two ways. First, Table 6 shows that an increase in surname diversity at the county level results in an

increase in surname diversity at the patent level and a larger number of technology codes per patent. The former indicates that more surname-diverse counties spur more diverse interactions among inventors, while the latter suggests that such interactions encourage more complex recombinations. The increased likelihood of more diverse co-inventors is not surprising given the results presented above: growing surname diversity leads to more residential intermixing—both in terms of surnames and ancestral countries—and to weaker family ties, thus enabling inventors to interact more outside their family and ethnic group. Second, by contrast, our patent-level results do not suggest that greater ancestral-country diversity at the county level results in more heterogeneous interactions at the patent-level (based on both ancestral countries and surnames, Table B50). Consistent with this, when surname diversity is included in the county-level model (Table 8), we do not observe robust relationships between an increase in ancestral-country diversity and either family ties or residential segregation.

6.2.1 Local Proximity and Surname Diversity

Our hypothesis proposes that most innovations arise from the recombination of ideas, often sparked by social interactions among diverse minds. Thus, we expect that people acquire ideas from others they frequently interact with, often serendipitously, in their daily lives (Atkin et al., 2022). If true, physical proximity and local diversity will play a big role (Jaffe et al., 1993; Carlino and Kerr, 2015). To explore this, we study the impact of the surname diversity found in neighboring counties on local innovation. If local social interactions are the most important, then the impact of neighboring counties should be small or negligible relative to that of a focal county’s surname diversity. Thus, we compute the surname diversity of people residing in surrounding counties at successively further distances from each focal county. Specifically, for each county i at time t , we pool the individuals in surrounding counties within 100 miles, compute their surname diversity and construct a separate instrument for these individuals, excluding i itself. We repeat

this exercise for individuals between 100 and 200 miles and between 200 and 300 miles.²⁴

Table 9 presents the results. For patents per capita, as shown in Columns 1 to 3 of Panel C, the IV coefficients for surname diversity within the focal county are large and do not decrease in magnitude when we include the surname diversity of surrounding counties. The effect size of surname diversity in neighboring counties within a 100-mile radius is about half the size of the coefficient for own-county surname diversity, and the coefficients for surname diversities in more distant locals are small and poorly estimated. For breakthrough patents (Columns 4 to 6, Panel C), the IV coefficients for own-county surname diversity remain stable with the addition of surrounding counties. Here, the coefficients for surname diversity of surrounding counties at any distance are small and poorly estimated. Overall, consistent with the importance of the social-behavioral mechanism, these findings suggest that the causal link between surname diversity and innovation is geographically localized, suggesting that meeting and interacting face-to-face on a day-to-day basis may be an important ingredient for innovation during our study period. The results also offer empirical support for our choice of the county as our primary unit of analysis.

7 Conclusion

The United States rose to dominate global innovation during the period from 1850 to 1940. To better understand the factors driving this process, we consider the impact of America's changing cultural diversity on innovation. Our core hypothesis is that many, if not most, innovations arise from the recombination of existing ideas, approaches and techniques that come together when diverse minds meet, exchange ideas, and sometimes collaborate. To create a fine-grained proxy for this diversity, we use surnames drawn from the complete U.S. Census to calculate an entropic measure of variation. To measure innovation, we use

²⁴We use the [NBER's County Distance Database](#) to compute these areas for each county.

both patents and a text-based measure of breakthrough patents. For causal identification, we employ an instrumental variable approach that uses historical immigration flows (push-pull interactions) to extract plausibly exogenous shifts in surname diversity across space and time. Our IV analyses suggest a causal connection between surname diversity and innovation at four levels: the county, surname-county, inventor and patent. At the county level (1900–1940), we estimate that a standard deviation increase in surname diversity leads to roughly 107% more patents per capita and 80% more breakthrough patents per capita relative to their means, corresponding to 1.76 times and 0.76 times the effect of population size, respectively. These relationships are remarkably stable over time and likely extend well back into the 19th century. At the inventor level, a standard deviation increase in surname diversity raises inventors’ productivity by roughly five patents and one breakthrough patent over a 50-year career. At the patent level, a one-standard deviation increase in surname diversity delivers roughly a 91 percentage point increase in the likelihood of a patent becoming a breakthrough and roughly 150% more technology codes per patent.

To explore the causal pathways running from surname diversity to innovation, we close by offering evidence supporting the connection between surname diversity and (1) informational diversity, which opens the door directly to more innovative recombinations, and (2) the potential for heterogeneous social interaction, which increases the likelihood that diverse individuals will meet and interact. Focusing on the informational channel, we show that both surnames and ancestral countries cluster with education, occupations, originating countries and the technological categories of patents, but both birth- and ancestral-country diversity do not have an effect on innovation beyond the impact captured by surname diversity. We also show that greater surname diversity results in greater occupational and patent technology code diversity. Focusing on the social-behavioral channel, we demonstrate that greater surname diversity at the county level results in weaker family ties and less residential segregation among people with both different

surnames and ancestral countries, thus creating conditions conducive to more diverse social interactions.

Broadly, our results highlight the need for researchers to more deeply consider how, when and why diversity and immigration foster innovation. Our comparisons of surname diversity versus ancestral-country diversity suggest that the former, which captures more granular cultural differences, may be the more powerful stimulant for innovation. Indeed, our findings indicate that the latter influences innovation by capturing part of what surname diversity includes. Our analysis of mechanisms suggests that this difference is most likely due to how the two forms of diversity influence social interaction. Greater surname diversity results in weaker family ties, less residential segregation and more diversity at the patent level. By contrast, ancestral-country diversity does not stimulate more social interaction; indeed, there are hints that it may result in less interaction. The differences here likely arise from the parochialism and distrust that appear among large ethnic groups, which are not usually matched at similar levels among different families or smaller groupings. Such findings suggest there may be optimal forms of diversity that jointly maximize the richness of information and the potential for social interactions.

Our findings offer insights for designing policies aimed at enhancing innovation, suggesting that any factors that increase the likelihood of open-minded, face-to-face social interactions among people with diverse backgrounds can foster faster innovation. Policies related to immigration, housing, education, urban design (zoning, density) and transportation (among others), as well as local norms and beliefs that influence tolerance, cooperation and trust towards strangers, outsiders and foreigners, can play a role. Crucially, the most powerful forms of diversity come from having smaller and more uniformly distributed kinship groups or networks rather than large country or nationality-based groups. Specifically, our results suggest that policies that encourage family-based migration, resulting in clan-like clusters of relatives, may suppress innovation and that educational investments may have larger innovation payoffs when applied to local populations with more

fine-grained cultural diversity.

Data Availability

Data and code replicating the tables and figures in this article can be found in [Posch et al. \(2025\)](#) in the Harvard Dataverse, <https://doi.org/10.7910/DVN/EXE7CK>.

References

- Abramitzky, Ran, Philipp Ager, Leah Boustan, Elicor Cohen & Casper W. Hansen** (2023) “The effect of immigration restrictions on local labor markets: Lessons from the 1920s border closure”, *American Economic Journal: Applied Economics*, 15 (1), pp. 164–191.
- Abramitzky, Ran & Leah Boustan** (2017) “Immigration in American economic history”, *Journal of Economic Literature*, 55 (4), pp. 1311–45.
- Acemoglu, Daron, Ufuk Akcigit & William R. Kerr** (2016) “Innovation network”, *Proceedings of the National Academy of Sciences*, 113 (41), p. 11483–11488.
- Ager, Philipp & Markus Brückner** (2013) “Cultural diversity and economic growth: Evidence from the US during the age of mass migration”, *European Economic Review*, 64, pp. 76–97.
- Agneman, Gustav & Esther Chevrot-Bianco** (2023) “Market participation and moral decision-making: Experimental evidence from greenland”, *The Economic Journal*, 133 (650), p. 537–581.
- Agrawal, Ajay, Devesh Kapur & John McHale** (2008) “How do spatial and social proximity influence knowledge flows? Evidence from patent data”, *Journal of Urban Economics*, 64 (2), pp. 258–269.
- Akcigit, Ufuk, John Grigsby & Tom Nicholas** (2017) “Immigration and the rise of American ingenuity”, *American Economic Review*, 107 (5), pp. 327–331.
- Akcigit, Ufuk, William R Kerr & Tom Nicholas** (2013) “The mechanics of endogenous innovation and growth: Evidence from historical U.S. patents”.
- Alesina, Alberto F. & Paola Giuliano** (2015) “Culture and institutions”, *Journal of Economic Literature*, 53 (4), p. 898–944.
- Alesina, Alberto & Paola Giuliano** (2010) “The power of the family”, *Journal of Economic Growth*, 15 (2), p. 93–125.
- Alesina, Alberto & Paola Giuliano** (2014) “Chapter 4 - Family Ties”, Philippe Aghion & Steven N. Durlauf eds. *Handbook of Economic Growth*, 2, Elsevier, pp. 177–215.
- Alesina, Alberto, Johann Harnoss & Hillel Rapoport** (2016) “Birthplace diversity and economic prosperity”, *Journal of Economic Growth*, 21 (2), pp. 101–138.

- Alesina, Alberto & Eliana La Ferrara** (2005) “Ethnic Diversity and Economic Performance”, *Journal of Economic Literature*, 43 (3), pp. 762–800.
- Algan, Yann & Pierre Cahuc** (2014) “Chapter 2 - trust, growth, and well-being: New evidence and policy implications”, Philippe Aghion & Steven N. Durlauf eds. *Handbook of Economic Growth*, 2 of Handbook of Economic Growth, Elsevier, pp. 49–120.
- Allport, Gordon W.** (1954) *The Nature of Prejudice*, Oxford, England, Addison-Wesley.
- AlShebli, Bedoor K., Talal Rahwan & Wei Lee Woon** (2018) “The preeminence of ethnic diversity in scientific collaboration”, *Nature Communications*, 9 (1), p. 5163.
- Andersson, David, Thor Berger & Erik Prawitz** (2023) “Making a market: Infrastructure, integration, and the rise of innovation”, *The Review of Economics and Statistics*, 105 (2), p. 258–274.
- Andrews, Michael** (2023) “Bar Talk: Informal Social Interactions, Alcohol Prohibition, and Invention”.
- Andrews, Michael J. & Chelsea Lensing** (2024) “Cup of joe and knowledge flow: Coffee shops and invention”.
- Arbatli, Cemal Eren, Quamrul H Ashraf, Oded Galor & Marc Klemp** (2020) “Diversity and Conflict”, *Econometrica*, 88 (2), pp. 727–797.
- Arkolakis, Costas, Sun Kyoung Lee & Michael Peters** (2020) “European Immigrants and the United States’ Rise to the Technological Frontier”.
- Ashraf, Quamrul & Oded Galor** (2013) “The “Out of Africa” Hypothesis, Human Genetic Diversity, and Comparative Economic Development”, *American Economic Review*, 103 (1), pp. 1–46.
- Ashraf, Quamrul H., Oded Galor & Marc Klemp** (2021) “Chapter 22 - The Ancient Origins of the Wealth of Nations”, Alberto Bisin & Giovanni Federico eds. *The Handbook of Historical Economics*, Academic Press, pp. 675–717.
- Atkin, David, Keith Chen, M. & Anton Popov** (2022) “The returns to face-to-face interactions: Knowledge spillovers in silicon valley”, Working Paper 30147, National Bureau of Economic Research, Cambridge, MA.

- Bahrami-Rad, Duman, Jonathan Beauchamp, Joseph Henrich & Jonathan Schulz** (2024) “Kin-based institutions and economic development”, Working Paper 4200629, Social Science Research Network.
- Barone, Guglielmo & Sauro Mocetti** (2021) “Intergenerational mobility in the very long run: Florence 1427–2011”, *The Review of Economic Studies*, 88 (4), pp. 1863–1891.
- Beck Knudsen, Anne Sofie** (2024) “Those who stayed: Selection and cultural change in the age of mass migration”.
- Bell, Alexander, Raj Chetty, Xavier Jaravel, Neviana Petkova & John Van Reenen** (2019) “Who Becomes an Inventor in America? The Importance of Exposure to Innovation”, *Quarterly Journal of Economics*, 134 (2), pp. 647–713.
- Berkes, Enrico** (2018) “Comprehensive universe of US patents (CUSP): Data and facts”.
- Berkes, Enrico & Ruben Gaetani** (2021) “The Geography of Unconventional Innovation”, *The Economic Journal*, 131 (636), pp. 1466–1514.
- Bernstein, Shai, Rebecca Diamond, Abhisit Jiranaphawiboon, Timothy McQuade & Beatriz Pousada** (2022) “The contribution of high-skilled immigrants to innovation in the united states”, Working Paper 30797, National Bureau of Economic Research, Cambridge, MA.
- Bettencourt, L M A, J Lobo & D Strumsky** (2007) “Invention in the city: Increasing returns to patenting as a scaling function of metropolitan size”, *Research Policy*, 36 (1), pp. 107–120.
- Bisin, Alberto & Thierry Verdier** (2001) “The Economics of Cultural Transmission and the Dynamics of Preferences”, *Journal of Economic Theory*, 97 (2), pp. 298–319.
- Blasi, Damián E., Joseph Henrich, Evangelia Adamou, David Kemmerer & Asifa Majid** (2022) “Over-reliance on english hinders cognitive science”, *Trends in Cognitive Sciences*, 26 (12), p. 1153–1170.
- Borusyak, Kirill, Peter Hull & Xavier Jaravel** (2022) “Quasi-Experimental Shift-Share Research Designs”, *The Review of Economic Studies*, 89 (1), pp. 181–213.
- Boyd, Robert & Peter J Richerson** (1985) *Culture and the Evolutionary Process*, Chicago, IL, University of Chicago Press.

- Buggle, Johannes C** (2020) “Growing collectivism: Irrigation, group conformity and technological divergence”, *Journal of Economic Growth*, 25 (2), pp. 147–193.
- Buonanno, Paolo & Paolo Vanin** (2017) “Social Closure, Surnames and Crime”, *Journal of Economic Behavior & Organization*, 137, pp. 160–175.
- Burchardi, Konrad B, Thomas Chaney & Tarek A Hassan** (2019) “Migrants, Ancestors, and Foreign Investments”, *The Review of Economic Studies*, 86 (4), pp. 1448–1486.
- Bursztyn, Leonardo, Thomas Chaney, Tarek A. Hassan & Aakaash Rao** (2024) “The Immigrant Next Door”, *American Economic Review*, 114 (2), pp. 348–384.
- Carcassi, Gabriele, Christine A Aidala & Julian Barbour** (2021) “Variability as a better characterization of Shannon entropy”, *European Journal of Physics*, 42 (4), p. 045102.
- Carlino, G A, S Chatterjee & R M Hunt** (2007) “Urban density and the rate of invention”, *Journal of Urban Economics*, 61 (3), pp. 389–419.
- Carlino, Gerald & William R. Kerr** (2015) “Agglomeration and Innovation”, *Handbook of Regional and Urban Economics*, 5, Elsevier, pp. 349–404.
- Chen, Jiafeng & Jonathan Roth** (2023) “Logs with Zeros? Some Problems and Solutions”, *The Quarterly Journal of Economics*, p. qjad054.
- Chiopris, Caterina** (2024) “Spatial networks and the diffusion of ideas”.
- Christakis, N.A. & J.H. Fowler** (2009) *Connected: The Surprising Power of Our Social Networks and How They Shape Our Lives*, Little, Brown.
- Cinnirella, Francesco, Erik Hornung & Julius Koschnick** (2022) “Flow of ideas: Economic societies and the rise of useful knowledge”, *CESifo Working Paper*, 9836.
- Clark, Gregory** (2014) *The Son also Rises: a Surprising Look how Ancestry still Determines Social Outcomes*, Princeton, Princeton University Press.
- Cook, Lisa D., John Parman & Trevon Logan** (2022) “The antebellum roots of distinctively black names”, *Historical Methods: A Journal of Quantitative and Interdisciplinary History*, 55 (1), pp. 1–11.
- de la Croix, David, Matthias Doepke & Joel Mokyr** (2018) “Clans, guilds, and markets: Apprenticeship institutions and growth in the pre-industrial economy”, *The Quarterly Journal of Economics*, 133 (1), pp. 1–70.

- Docquier, Frédéric, Riccardo Turati, Jérôme Valette & Chrysovalantis Vasilakis** (2020) “Birthplace diversity and economic growth: Evidence from the US states in the Post-World War II period”, *Journal of Economic Geography*, 20 (2), pp. 321–354.
- Dowey, James** (2017) “Mind over matter: Access to knowledge and the British industrial revolution”, Ph.D. dissertation, Dissertation, The London School of Economics and Political Science.
- Enke, Benjamin** (2019) “Kinship, Cooperation, and the Evolution of Moral Systems*”, *The Quarterly Journal of Economics*, 134 (2), pp. 953–1019.
- Enke, Benjamin** (2022) “Market exposure and human morality”, *Nature Human Behaviour*, 7 (1), pp. 134–141.
- Falk, Armin, Anke Becker, Thomas Dohmen, Benjamin Enke, David Huffman & Uwe Sunde** (2018) “Global evidence on economic preferences*”, *The Quarterly Journal of Economics*, 133 (4), p. 1645–1692.
- Feldman, Maryann P & David B Audretsch** (1999) “Innovation in cities : Science-based diversity , specialization and localized competition”, *European Economic Review*, 43 (2), pp. 409–429.
- Ferrara, Andreas, Patrick A. Testa & Liyang Zhou** (2021) “New area- and population-based geographic crosswalks for U.S. counties and congressional districts, 1790–2020”, Working Paper 588, Centre for Competitive Advantage in the Global Economy (CAGE), University of Warwick, Coventry, UK.
- Fiszbein, Martin** (2022) “Agricultural Diversity, Structural Change, and Long-Run Development: Evidence from the United States”, *American Economic Journal: Macroeconomics*, 14 (2), pp. 1–43.
- Fulford, Scott L., Ivan Petkov & Fabio Schiantarelli** (2020) “Does it matter where you came from? Ancestry composition and economic performance of US counties, 1850–2010”, *Journal of Economic Growth*, 25 (3), pp. 341–380.
- Galor, Oded, Marc Klemp & Daniel C Wainstock** (2024) “Roots of Cultural Diversity”, Discussion Paper 17481, IZA Institute of Labor Economics.
- Ghosh, Arkadev, Sam Il Myoung Hwang & Munir Squires** (2023) “Economic Consequences of Kinship: Evidence From U.S. Bans on Cousin Marriage”, *The Quarterly Journal of Economics*, 138 (4), pp. 2559–2606.

- Giuliano, Paola & Nathan Nunn** (2021) “Understanding cultural persistence and change”, *The Review of Economic Studies*, 88 (4), p. 1541–1581.
- Glaeser, Edward** (2011) “Engines of Innovation”, *Scientific American*, p. 7.
- Glaeser, Edward L., Hedi D. Kallal, José A. Scheinkman & Andrei Shleifer** (1992) “Growth in cities”, *Journal of Political Economy*, 100 (6), p. 1126–1152.
- Goldsmith-Pinkham, Paul, Isaac Sorkin & Henry Swift** (2020) “Bartik Instruments: What, When, Why, and How”, *American Economic Review*, 110 (8), pp. 2586–2624.
- Gomez-Lievano, Andres, Oscar Patterson-Lomba & Ricardo Hausmann** (2017) “Explaining the prevalence, scaling and variance of urban phenomena”, *Nature Human Behaviour*, 1 (1), p. No. 0012.
- Gorodnichenko, Yuriy & Gerard Roland** (2016) “Culture, institutions, and the wealth of nations”, *The Review of Economics and Statistics*, 99 (3), p. 402–416.
- Griliches, Zvi** (1990) “Patent Statistics as Economic Indicators: A Survey”, *Journal of Economic Literature*, 28 (4), pp. 1661–1707.
- Güell, Maia, José V. Rodríguez Mora & Christopher I. Telmer** (2015) “The Informational Content of Surnames, the Evolution of Intergenerational Mobility, and Assortative Mating”, *The Review of Economic Studies*, 82 (2), pp. 693–735.
- Henrich, J., J. Ensminger, R. McElreath, A. Barr, C. Barrett, A. Bolyanatz, J. C. Cardenas, M. Gurven, E. Gwako, N. Henrich, C. Lesorogol, F. Marlowe, D. P. Tracer & J. Ziker** (2010) “Market, religion, community size and the evolution of fairness and punishment”, *Science*, 327, p. 1480–1484.
- Henrich, Joseph** (2016) *The Secret of Our Success: How Culture Is Driving Human Evolution, Domesticating Our Species, and Making Us Smarter*, Princeton, Princeton University Press.
- Henrich, Joseph** (2020) *The WEIRDest People in the World: How the West Became Psychologically Peculiar and Particularly Prosperous*, New York, Farrar, Straus and Giroux.
- Henrich, Joseph & James Broesch** (2011) “On the nature of cultural transmission networks: evidence from fijian villages for adaptive learning biases”, *Philosophical Transactions of the Royal Society B: Biological Sciences*, 366 (1567), p. 1139–1148.
- Henrich, Joseph & Michael Muthukrishna** (2024) “What makes us smart?”, *Topics in Cognitive Science*, 16 (2), pp. 322–342.

- Hewlett, Barry S.** (2016) “Social learning and innovation in hunter-gatherers”, Hideaki Terashima & Barry S. Hewlett eds. *Social Learning and Innovation in Contemporary Hunter-Gatherers: Evolutionary and Ethnographic Perspectives*, Tokyo, Springer, pp. 1–15.
- Jacobs, J.** (1985) *Cities and the Wealth of Nations: Principles of Economic Life*, Vintage Books.
- Jacobs, Jane** (1969) *The Economy of Cities*, New York, Random House.
- Jaffe, Adam B., Manuel Trajtenberg & Rebecca Henderson** (1993) “Geographic Localization of Knowledge Spillovers as Evidenced by Patent Citations*”, *The Quarterly Journal of Economics*, 108 (3), pp. 577–598.
- Jones, Charles I** (2023) “Recipes and economic growth: A combinatorial march down an exponential tail”, *Journal of Political Economy*, 131 (8), pp. 1994–2031.
- Kelly, Bryan, Dimitris Papanikolaou, Amit Seru & Matt Taddy** (2021) “Measuring Technological Innovation over the Long Run”, *American Economic Review: Insights*, 3 (3), pp. 303–320.
- Kerr, William R.** (2008) “The ethnic composition of us inventors”, Working Paper 08-006, Harvard Business School, Cambridge, MA.
- Kerr, William R.** (2010) “Breakthrough inventions and migrating clusters of innovation”, *Journal of Urban Economics*, 67 (1), pp. 46–60.
- Lerner, Josh & Amit Seru** (2022) “The use and misuse of patent data: Issues for finance and beyond”, *The Review of Financial Studies*, 35 (6), p. 2667–2704.
- Lerner, Josh & Julie Wulf** (2007) “Innovation and Incentives: Evidence from Corporate R&D”, *the Review of Economics and Statistics*, 89 (4), pp. 634–644.
- Lindquist, Matthew J, Joeri Sol & Mirjam Van Praag** (2015) “Why do entrepreneurial parents have entrepreneurial children?”, *Journal of Labor Economics*, 33 (2), pp. 269–296.
- Lissoni, Francesco & Ernest Miguelez** (2024) “Migration and innovation: Learning from patent and inventor data”, *Journal of Economic Perspectives*, 38 (1), p. 27–54.
- Logan, Trevon D. & John M. Parman** (2017) “The National Rise in Residential Segregation”, *The Journal of Economic History*, 77 (1), pp. 127–170.
- Maddux, William W, Hajo Adam & Adam D Galinsky** (2010) “When in rome... learn why the romans do what they do: How multicultural learning experiences facilitate creativity”, *Personality and Social Psychology Bulletin*, 36 (6), p. 731–741.

- Mill, J.S.** (1871) *Principles of Political Economy: With Some of Their Applications to Social Philosophy*, Principles of Political Economy: With Some of Their Applications to Social Philosophy, Longmans, Green, Reader, and Dyer.
- Mokyr, Joel** (2002) *The Gifts of Athena: Historical Origins of the Knowledge Economy*, Princeton, NJ, Princeton University Press.
- Moser, Petra** (2013) “Patents and Innovation: Evidence from Economic History”, *The Journal of Economic Perspectives*, 27 (1), pp. 23–44.
- Moser, Petra & Shmuel San** (2020) “Immigration, science, and invention. lessons from the quota acts”, Working Paper 3558718, Social Science Research Network.
- Moser, Petra, Alessandra Voena & Fabian Waldinger** (2014) “German Jewish Émigrés and US Invention”, *American Economic Review*, 104 (10), pp. 3222–3255.
- Muthukrishna, Michael & Joseph Henrich** (2016) “Innovation in the collective brain”, *Philosophical Transactions of the Royal Society B: Biological Sciences*, 371 (1690), pp. 1–14.
- Nguyen, Kieu-Trang** (2025) “Trust and innovation within the firm: Evidence from matched CEO-Firm data”.
- Nisbett, Richard E** (2003) *The Geography of Thought: How Asians and Westerners Think Differently– and Why*, New York, Free Press.
- Nunn, Nathan** (2021) “History as evolution”, *The handbook of historical economics*, Elsevier, pp. 41–91.
- Ottaviano, Gianmarco IP & Giovanni Peri** (2006) “The economic value of cultural diversity: Evidence from US cities”, *Journal of Economic Geography*, 6 (1), pp. 9–44.
- Packalen, Mikko & Jay Bhattacharya** (2015) “Cities and ideas”, Working Paper 20921, National Bureau of Economic Research, Cambridge, MA.
- Philips, Lawrence** (1990) “Hanging on the metaphone”, *Computer Language*, 7 (12), pp. 39–43.
- Posch, Max, Jonathan Schulz & Joseph Henrich** (2025) “Replication Data for: How Cultural Diversity Drives Innovation: Surnames and Patents in U.S. History”.
- Raz, Itzhak T.** (2025) “Soil heterogeneity, social learning, and the formation of close-knit communities”, *Journal of Political Economy*, 133 (8).

- Ruggles, Steven, Catherine A. Fitch, Ronald Goeken, J. David Hacker, Matt A. Nelson, Evan Roberts, Megan Schouweiler & Matthew Sobek** (2021) *IPUMS Ancestry Full Count Data: Version 3.0 [Dataset]*, Minneapolis, MN, IPUMS.
- Rustagi, Devesh** (2024) “Market exposure, civic values, and rules”.
- Schimmelpfennig, Robin, Layla Razek, Eric Schnell & Michael Muthukrishna** (2022) “Paradox of diversity in the collective brain”, *Philosophical Transactions of the Royal Society B: Biological Sciences*, 377 (1843), p. 20200316.
- Schniter, Eric, Hillard S. Kaplan & Michael Gurven** (2023) “Cultural transmission vectors of essential knowledge and skills among tsimane forager-farmers”, *Evolution and Human Behavior*, 44 (6), p. 530–540.
- Schulz, Jonathan F** (2022) “Kin Networks and Institutional Development”, *The Economic Journal*, 132 (647), pp. 2578–2613.
- Schulz, Jonathan F, Duman Bahrami-Rad, Jonathan P. Beauchamp & Joseph Henrich** (2019) “The Church, Intensive Kinship, and Global Psychological Variation”, *Science*, 366 (6466).
- Sequeira, Sandra, Nathan Nunn & Nancy Qian** (2020) “Immigrants and the Making of America”, *Review of Economic Studies*, 87, pp. 382–419.
- Shannon, Claude Elwood** (1948) “A mathematical theory of communication”, *The Bell system technical journal*, 27 (3), pp. 379–423.
- Spolaore, Enrico & Romain Wacziarg** (2009) “The Diffusion of Development”, *The Quarterly Journal of Economics*, 124 (2), pp. 469–529.
- Squicciarini, Mara P & Nico Voigtländer** (2015) “Human capital and industrialization: Evidence from the age of the Enlightenment”, *The Quarterly Journal of Economics*, 130 (4), pp. 1825–1883.
- Stirling, Andy** (2007) “A general framework for analysing diversity in science, technology and society”, *Journal of The Royal Society Interface*, 4 (15), pp. 707–719.
- Strumsky, Deborah, Jose Lobo & Sander van der Leeuw** (2011) “Measuring the relative importance of reusing, recombining and creating technologies in the process of invention”, Working Paper 11-02-003, Santa Fe Institute, Santa Fe, NM.

Tabellini, Marco (2020) “Gifts of the Immigrants, Woes of the Natives: Lessons from the Age of Mass Migration”, *The Review of Economic Studies*, 87 (1), pp. 454–486.

Terry, Stephen J, Thomas Chaney, Konrad B Burchardi, Lisa Tarquinio & Tarek A Hassan (2024) “Immigration, Innovation, and Growth”.

Usher, A.P. (2013) *A History of Mechanical Inventions: Revised Edition*, Dover Publications.

Weitzman, Martin L. (1998) “Recombinant Growth”, *The Quarterly Journal of Economics*, 63 (2).

Youn, H J, D Strumsky, L M A Bettencourt & J Lobo (2015) “Invention as a combinatorial process: Evidence from US patents”, *Journal of the Royal Society Interface*, 12 (106).

Tables

Table 1: Surname-county level estimates of the effects of push-pull instruments on surname stocks

	No. of people in county i with surname k in year:															
	1900				1905				1910				1915			
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)	(11)	(12)	(13)	(14)	(15)	(16)
$I_{k,-r(t)}^{1880} \times \frac{I_{k,-r(t)}^{1880}}{I_{k,-r(t)}^{1885}}$	3.263** (0.104)	3.151** (0.106)	3.271** (0.096)	3.140** (0.099)	3.812** (0.109)	3.639** (0.113)	3.846** (0.083)	3.668** (0.087)	4.257** (0.089)	4.058** (0.094)	4.418** (0.091)	4.182** (0.096)	4.889** (0.102)	4.596** (0.099)	5.016** (0.098)	4.704** (0.095)
$I_{k,-r(t)}^{1895} \times \frac{I_{k,-r(t)}^{1895}}{I_{k,-r(t)}^{1895}}$	0.722** (0.295)	0.931** (0.289)	1.081 (0.700)	1.289* (0.769)	1.443* (0.802)	1.652* (0.881)	1.805** (0.567)	2.033** (0.627)	2.038** (0.613)	2.299** (0.677)	2.869** (0.607)	3.183** (0.667)	3.183** (0.248)	3.539** (0.225)	3.245** (0.314)	3.640** (0.293)
$I_{k,-r(t)}^{1900} \times \frac{I_{k,-r(t)}^{1900}}{I_{k,-r(t)}^{1900}}$			-0.516 (2.547)	-1.381 (2.845)	2.135 (2.993)	0.952 (3.355)	-0.160 (2.376)	-0.937 (2.609)	1.199 (2.639)	0.429 (2.890)	-4.712* (2.843)	-5.265* (3.093)	-3.889 (3.505)	-4.545 (3.309)	-2.471 (3.207)	-2.631 (2.980)
$I_{k,-r(t)}^{1905} \times \frac{I_{k,-r(t)}^{1905}}{I_{k,-r(t)}^{1905}}$					14.507** (1.208)	15.342** (1.437)	18.806** (1.112)	19.327** (1.275)	21.604** (1.243)	22.245** (1.419)	27.458** (1.315)	27.682** (1.465)	30.517** (1.171)	30.812** (1.120)	33.261** (1.094)	33.180** (1.058)
$I_{k,-r(t)}^{1910} \times \frac{I_{k,-r(t)}^{1910}}{I_{k,-r(t)}^{1910}}$							18.276** (2.214)	17.835** (2.401)	20.650** (2.369)	20.248** (2.574)	27.962** (2.373)	27.438** (2.559)	31.176** (3.365)	30.101** (3.365)	33.057** (2.942)	31.915** (2.967)
$I_{k,-r(t)}^{1915} \times \frac{I_{k,-r(t)}^{1915}}{I_{k,-r(t)}^{1915}}$									8.119** (0.840)	8.409** (0.888)	14.490** (1.087)	14.751** (1.145)	16.782** (1.432)	16.942** (1.315)	19.092** (1.611)	19.215** (1.493)
$I_{k,-r(t)}^{1920} \times \frac{I_{k,-r(t)}^{1920}}{I_{k,-r(t)}^{1920}}$											-1.490 (1.968)	-1.515 (2.047)	-1.018 (1.085)	-0.982 (1.004)	2.825* (1.490)	2.644* (1.377)
$I_{k,-r(t)}^{1925} \times \frac{I_{k,-r(t)}^{1925}}{I_{k,-r(t)}^{1925}}$													26.776** (1.214)	27.354** (1.286)	32.424** (1.499)	32.829** (1.564)
$I_{k,-r(t)}^{1930} \times \frac{I_{k,-r(t)}^{1930}}{I_{k,-r(t)}^{1930}}$															-34.025** (3.464)	-29.284** (3.320)
Constant	0.010** (0.000)		0.008** (0.000)		0.007** (0.000)		0.004** (0.000)		0.003** (0.000)		0.001** (0.000)		-0.002** (0.000)		0.000 (0.000)	
F-stat	824	942	1341	938	1161	821	1338	1002	1119	837	885	675	620	640	698	681
Observations	5,346,601	5,346,601	6,604,833	6,604,833	7,012,154	7,012,154	7,284,327	7,284,327	7,364,850	7,364,850	8,067,190	8,067,190	8,143,606	8,143,606	8,980,212	8,980,212
R ²	0.604	0.705	0.581	0.681	0.600	0.695	0.604	0.692	0.607	0.696	0.578	0.662	0.624	0.705	0.613	0.695
Origin country-County fixed effects	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓
Surname-Census division fixed effects	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓
I_i/I_i controls	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓

Notes: This table reports OLS estimates for the specification described in equation (3), corresponding to step 1 of the instrument construction. An observation is a surname-county in a period from 1900 to 1940. The table includes the F-statistic for the Wald test for the joint nullity of the push-pull instruments. Standard errors clustered at the surname level. **, *, and * indicate significance at the 1%, 5%, and 10% levels.

Table 2: County-level estimates of the effect of surname diversity on innovation

	Patents per 1,000 people (mean = 1.05, sd = 1.73)				Breakthrough patents per 1,000 people (mean = 0.13, sd = 0.28)			
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
<i>Panel A: Least-squares estimates</i>								
Surname diversity	1.322*** (0.169)	1.125*** (0.202)	1.081*** (0.263)	1.338*** (0.174)	0.148*** (0.016)	0.107*** (0.015)	0.099*** (0.017)	0.099*** (0.018)
Population		0.674*** (0.106)	0.645*** (0.106)	0.736*** (0.189)		0.139*** (0.020)	0.125*** (0.018)	0.099*** (0.029)
<i>Panel B: Reduced-form estimates</i>								
Predicted surname diversity	0.956*** (0.179)	0.698*** (0.235)	0.667** (0.283)	0.939*** (0.236)	0.111*** (0.014)	0.046** (0.021)	0.043* (0.024)	0.061*** (0.019)
Predicted population		0.546*** (0.147)	0.521*** (0.161)	0.205 (0.205)		0.138*** (0.026)	0.122*** (0.025)	0.051 (0.035)
<i>Panel C: Instrumental-variable estimates</i>								
Surname diversity	1.388*** (0.171)	1.161*** (0.205)	1.128*** (0.262)	1.392*** (0.191)	0.161*** (0.018)	0.104*** (0.018)	0.104*** (0.019)	0.101*** (0.019)
Population		0.664*** (0.137)	0.640*** (0.120)	0.731*** (0.236)		0.166*** (0.028)	0.137*** (0.021)	0.118** (0.054)
Sanderson-Windmeijer F-stat	129	162; 426	147; 227	149; 29	129	162; 426	147; 227	149; 29
<i>Panel D: First-stage estimates</i>								
	Surname diversity				Population			
Predicted surname diversity	0.689*** (0.061)	0.693*** (0.075)	0.723*** (0.073)	0.729*** (0.074)	-0.160*** (0.033)	-0.231*** (0.042)	-0.103** (0.041)	
Predicted population		-0.009 (0.041)	-0.076* (0.040)	-0.142** (0.063)	0.837*** (0.053)	0.948*** (0.073)	0.550*** (0.107)	
County fixed effects	✓	✓	✓	✓	✓	✓	✓	✓
Period fixed effects	✓	✓			✓	✓		
State-Period fixed effects			✓	✓			✓	✓
County-specific linear trends				✓				✓
Observations	21,984	21,984	21,984	21,984	21,984	21,984	21,984	21,984

Notes: The table reports OLS, reduced-form and IV estimates for the specifications described in equation (2) as well as IV first-stage estimates. An observation is a county in a period from 1900 to 1940. The endogenous variables are county-level surname diversity and population in t . In Columns 1 to 4, the dependent variable is number of patents filed in the county in the five-year period starting in t divided by county population size in 1900. In Columns 5 to 8, it is number of breakthrough patents filed in the county in the five-year period starting in t divided by county population size in 1900. Standard errors are clustered at the state level. All independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table 3: County-surname-level estimates of the effect of surname diversity on innovation

	Patents per 1,000 people (mean = 1.02, sd = 13.31)				Breakthrough patents per 1,000 people (mean = 0.1, sd = 2.23)			
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
<i>Panel A: Least-squares estimates</i>								
Surname diversity	1.173*** (0.194)	0.930*** (0.174)	0.948*** (0.178)	1.143*** (0.182)	0.128*** (0.031)	0.087*** (0.026)	0.088*** (0.027)	0.061*** (0.015)
Population		3.702*** (0.600)	3.720*** (0.613)	3.679*** (0.597)		0.617*** (0.107)	0.624*** (0.108)	0.728*** (0.100)
<i>Panel B: Reduced-form estimates</i>								
Predicted surname diversity	0.742*** (0.160)	0.521*** (0.177)	0.530*** (0.182)	0.675*** (0.161)	0.080*** (0.021)	0.046* (0.024)	0.046* (0.025)	0.044 (0.030)
Predicted population		3.072*** (0.946)	3.089*** (0.971)	0.958 (0.971)		0.472*** (0.155)	0.479*** (0.158)	0.112 (0.280)
<i>Panel C: Instrumental-variable estimates</i>								
Surname diversity	1.382*** (0.251)	1.207*** (0.245)	1.227*** (0.252)	1.358*** (0.248)	0.149*** (0.035)	0.122*** (0.035)	0.123*** (0.036)	0.091** (0.043)
Population		3.060*** (0.551)	3.071*** (0.568)	1.478 (1.509)		0.474*** (0.100)	0.480*** (0.102)	0.170 (0.424)
Sanderson-Windmeijer F-stat	178	201; 50	201; 49	215; 300	178	201; 50	201; 49	215; 300
<i>Panel D: First-stage estimates</i>								
	Surname diversity				Population			
Predicted surname diversity	0.537*** (0.040)	0.532*** (0.041)	0.532*** (0.041)	0.520*** (0.046)		-0.040** (0.016)	-0.040** (0.016)	-0.021*** (0.004)
Predicted population		0.058* (0.030)	0.060* (0.031)	-0.028 (0.059)		0.981*** (0.142)	0.982*** (0.144)	0.674*** (0.049)
County-Surname fixed effects	✓	✓	✓	✓	✓	✓	✓	✓
State-Period fixed effects	✓	✓	✓	✓	✓	✓	✓	✓
Surname-Period fixed effects			✓	✓			✓	✓
County-specific linear trends				✓				✓
Observations	19,482,024	19,482,024	19,482,024	19,482,024	19,482,024	19,482,024	19,482,024	19,482,024

Notes: The table reports OLS, reduced-form and IV estimates for the specifications described in equation (8) as well as IV first-stage estimates. An observation is a surname in a given county in a period from 1900 to 1940. Observations are weighted based on their population shares within counties. The endogenous variables are county-level surname entropy and population in t . In Columns 1 to 4, the dependent variable is number of patents filed by individuals with surname n residing in county in i in the five-year period starting in t divided by surname population size in 1900. In Columns 5 to 8, the dependent variable is number of breakthrough patents filed by individuals with surname n residing in county in i in the five-year period starting in t divided by surname population size in 1900. Standard errors are clustered at the state level. All independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table 4: Inventor-level estimates of the effect of surname diversity on innovation

	Patents (mean = 2.24, sd = 1.69)		Breakthrough patents (mean = 0.26, sd = 0.44)	
	(1)	(2)	(3)	(4)
<i>Panel A: Least-squares estimates</i>				
Surname diversity	0.220*** (0.062)	0.219*** (0.061)	0.043 (0.027)	0.043 (0.026)
Population		-0.016 (0.053)		-0.011 (0.015)
<i>Panel B: Reduced-form estimates</i>				
Predicted surname diversity	0.077*** (0.022)	0.078*** (0.022)	0.017** (0.008)	0.016** (0.008)
Predicted population		0.003 (0.040)		-0.008 (0.010)
<i>Panel C: Instrumental-variable estimates</i>				
Surname diversity	0.463** (0.188)	0.438** (0.182)	0.104 (0.063)	0.094 (0.059)
Population		-0.045 (0.044)		-0.019 (0.012)
Sanderson-Windmeijer <i>F</i> -stat	12	17; 504	12	17; 504
<i>Panel D: First-stage estimates</i>				
	Surname diversity		Population	
Predicted surname diversity	0.167*** (0.048)	0.179*** (0.049)		0.020 (0.048)
Predicted population		0.095*** (0.019)		0.867*** (0.101)
Inventor fixed effects	✓	✓	✓	✓
Period-State fixed effects	✓	✓	✓	✓
Observations	209,944	209,944	209,944	209,944

Notes: The table reports OLS, reduced-form and IV estimates for the specifications described in equation (9) as well as IV first-stage estimates. An observation is an inventor residing in a given county in a period from 1900 to 1940. The endogenous variables are county-level surname diversity and population in year t , and the dependent variables are number of patents filed by the inventor in the five-year period starting in t (Columns 1–2) and number of breakthrough patents filed by the inventor in the five-year period starting in t (Columns 3–4). Standard errors are clustered at the state level. The independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table 5: County-level horse race between surname diversity and other diversities

	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)	(11)	(12)
<i>Panel A:</i>												
	Patents per 1,000 people											
Surname diversity	1.081*** (0.263)		1.076*** (0.264)		1.134*** (0.262)		1.078*** (0.249)		0.978*** (0.228)	1.099*** (0.240)	1.065*** (0.230)	1.356*** (0.152)
Birth-country diversity		0.302*** (0.084)	-0.047 (0.121)								-0.042 (0.131)	0.032 (0.143)
Ancestral-country diversity				0.233*** (0.083)	-0.110** (0.054)						-0.080* (0.044)	-0.150** (0.056)
Racial diversity						0.007 (0.201)	-0.162 (0.119)				-0.150 (0.112)	-0.470*** (0.140)
Occupational diversity								0.507*** (0.157)	0.260*** (0.091)		0.259*** (0.092)	0.155 (0.100)
Population	0.645*** (0.106)	0.853*** (0.111)	0.647*** (0.107)	0.813*** (0.109)	0.652*** (0.107)	0.847*** (0.110)	0.651*** (0.107)	0.888*** (0.121)	0.687*** (0.110)	0.659*** (0.096)	0.697*** (0.100)	0.747*** (0.189)
R ²	0.810	0.800	0.812	0.798	0.810	0.799	0.812	0.803	0.813	0.814	0.815	0.877
<i>Panel B:</i>												
	Breakthrough patents per 1,000 people											
Surname diversity	0.099*** (0.017)		0.100*** (0.019)		0.105*** (0.019)		0.102*** (0.018)		0.088*** (0.019)	0.104*** (0.018)	0.095*** (0.021)	0.099*** (0.020)
Birth-country diversity		0.039*** (0.013)	0.007 (0.014)								0.004 (0.015)	0.038 (0.026)
Ancestral-country diversity				0.019*** (0.007)	-0.012 (0.009)						-0.008 (0.008)	-0.018** (0.008)
Racial diversity						-0.004 (0.030)	-0.020 (0.025)				-0.023 (0.025)	-0.052 (0.038)
Occupational diversity								0.060*** (0.011)	0.038*** (0.009)		0.037*** (0.008)	0.015 (0.009)
Population	0.125*** (0.018)	0.144*** (0.018)	0.125*** (0.018)	0.141*** (0.018)	0.126*** (0.019)	0.144*** (0.018)	0.125*** (0.018)	0.149*** (0.019)	0.130*** (0.019)	0.121*** (0.017)	0.127*** (0.017)	0.094*** (0.028)
R ²	0.678	0.675	0.679	0.674	0.678	0.675	0.679	0.677	0.680	0.681	0.682	0.755
County fixed effects	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓
State-Period fixed effects	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓
Immigrant shares by 59 birth countries										✓	✓	✓
County-specific linear trends												✓
Observations	21,984	21,873	21,873	21,984	21,984	21,873	21,873	21,871	21,871	21,873	21,871	21,871

Notes: The table reports OLS estimates of regressions of innovation outcomes on surname diversity and other diversities. An observation is a county in a period from 1900 to 1940. In Panel A (Panel B), the dependent variable is number of patents and breakthrough patents filed in the county in the five-year period starting in t divided by county population size in 1900. Standard errors are clustered on states and reported in parentheses. Independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table 6: Patent-level estimates of the effect of surname diversity on innovation

	Breakthrough patent indicator (mean = 0.17, sd = 0.38)			Technologies per patent (mean = 2.49, sd = 1.82)		
	(1)	(2)	(3)	(4)	(5)	(6)
<i>Panel A: Least-squares estimates</i>						
Surname diversity of patent	0.015*** (0.001)	0.010*** (0.001)	0.010*** (0.001)	0.032*** (0.005)	0.020*** (0.005)	0.019*** (0.005)
<i>Panel B: Reduced-form estimates</i>						
Predicted surname diversity	0.010** (0.004)	0.008** (0.004)		0.039*** (0.015)	0.032** (0.013)	
<i>Panel C: Instrumental-variable estimates</i>						
Surname diversity of patent	1.127* (0.583)	0.908* (0.520)		4.341** (2.178)	3.731* (2.022)	
Sanderson-Windmeijer F-stat	10	9		10	9	
<i>Panel D: First-stage estimates</i>						
	Surname diversity of patent					
Predicted surname diversity	0.009*** (0.003)	0.008*** (0.003)				
Team size fixed effects	✓	✓	✓	✓	✓	✓
County fixed effects	✓	✓		✓	✓	
State-Period fixed effects	✓	✓		✓	✓	
Patent technology field-Period fixed effects		✓	✓		✓	✓
County-Period fixed effects			✓			✓
Observations	1,451,459	1,451,459	1,451,459	1,451,459	1,451,459	1,451,459

Notes: The table reports least squares, reduced-form and IV estimates for the specifications described in equation (10). An observation is a patent with inventors residing in a given county in a five-year period from 1900 to 1940. If there are multiple inventors who reside in different counties, the patent is assigned to all counties and the observation is weighted by 1 / number of inventors. The endogenous variables is surname diversity of patent. The dependent variables are breakthrough patent indicator in Columns 1 to 3 and the number of technology codes per patent in Columns 4 to 6. Columns 2 and 5 add six broad patent technology fields (chemical, computers, drugs, electrical, mechanical and 'other') fixed effects interacted with period fixed effects. Columns 3 and 6 additionally include country-period fixed effects. Standard errors are clustered at the county level. All diversity variables are standardized to have a mean of zero and a unit variance. The sources and construction of all variables are detailed in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table 7: Surnames cluster in education, occupations, birthplaces, and patent technologies

	Years of schooling	Occupation				Country of birth		Region of birth	USPTO tech class	
Sample:		All		Immigrants				Germans	Inventors	
Year:	1940	1880	1940	1880	1940	1880	1940	1880	1880-9	1940-9
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)
Surname	0.089	0.167	0.056	0.121	0.073	0.415	0.246	0.270	0.092	0.068
Country of birth	0.086	0.173	0.052	0.117	0.064					
Country of ancestry	0.081	0.158	0.049	0.104	0.050				0.004	0.003
U.S. county of residence	0.102	0.231	0.088	0.183	0.084	0.273	0.123	0.248	0.014	0.017

Notes: This table reports normalized Herfindahl-Hirschman Indices (HHI), where larger values indicate greater concentration. The indices are calculated as the average HHI of the variable in the header computed for each value of the variables on the left. For example, we calculate the normalized Herfindahl index of occupations for each surname and then average over all surnames using the number of individuals with a given surname as weights. Column 2 (1 and 3) includes all male individuals aged 25-60 in the 1880 (1940) census. Columns 4 and 6 (5 and 7) include all male immigrants aged 25-60 in the 1880 (1940) census. Column 8 restricts the sample to German male immigrants aged 25-60 in 1880. We use the 32 subnational regional origins for German immigrants recorded by the Census (the variable `bp1d` with codes 45301 to 45362). Column 9 (10) includes all inventors of patents issued from 1880 to 1889 (1940 to 1949). We use the main technology class on the patent.

Table 8: Effects on proxy measures of informational and social-behavioral mechanisms

	Occupational diversity		USPTO tech codes diversity		Residential segregation of surname groups		Residential segregation of ancestral-country groups		Strength of family ties	
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)
<i>Panel A: Least-squares estimates</i>										
Surname diversity	0.301*** (0.079)	0.328*** (0.082)	0.303*** (0.051)	0.362*** (0.066)	-0.387*** (0.080)	-0.432*** (0.086)	-1.079*** (0.162)	-1.320*** (0.135)	-0.202** (0.086)	-0.157* (0.087)
Ancestral-country diversity		-0.051 (0.044)		-0.096** (0.036)		0.078 (0.050)		0.417*** (0.028)		-0.080** (0.035)
Population	-0.050* (0.029)	-0.049* (0.027)	0.059 (0.036)	0.062* (0.037)	0.044 (0.031)	0.040 (0.032)	0.091** (0.043)	0.073* (0.037)	-0.088 (0.057)	-0.085 (0.057)
<i>Panel B: Reduced-form estimates</i>										
Predicted surname diversity	0.281*** (0.068)	0.278*** (0.065)	0.233*** (0.047)	0.247*** (0.045)	-0.314*** (0.080)	-0.352*** (0.074)	-0.824*** (0.217)	-0.878*** (0.215)	-0.207*** (0.062)	-0.237*** (0.066)
Predicted ancestral-country diversity		0.003 (0.012)		-0.015 (0.015)		0.037* (0.020)		0.054*** (0.013)		0.031 (0.027)
Predicted population	-0.118** (0.046)	-0.117** (0.046)	-0.029 (0.038)	-0.034 (0.036)	0.129** (0.057)	0.146** (0.055)	0.455*** (0.159)	0.479*** (0.161)	-0.089 (0.066)	-0.079 (0.062)
<i>Panel C: Instrumental-variable estimates</i>										
Surname diversity	0.372*** (0.086)	0.398*** (0.100)	0.350*** (0.060)	0.457*** (0.095)	-0.407*** (0.091)	-0.586*** (0.127)	-1.031*** (0.192)	-1.340*** (0.118)	-0.322*** (0.096)	-0.467*** (0.130)
Ancestral-country diversity		-0.057 (0.083)		-0.193 (0.130)		0.345** (0.165)		0.597*** (0.098)		0.299* (0.173)
Population	-0.122* (0.071)	-0.103 (0.066)	0.002 (0.064)	0.045 (0.090)	0.114* (0.059)	-0.006 (0.095)	0.506*** (0.161)	0.298** (0.137)	-0.240** (0.118)	-0.337* (0.174)
Sanderson-Windmeijer <i>F</i> -stat	159; 28	57; 90; 123	163; 18	38; 43; 60	145; 25	39; 63; 86	145; 25	39; 63; 86	158; 28	55; 86; 124
County fixed effects	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓
State-Period fixed effects	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓
County-specific linear trends	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓
Observations	21,871	21,871	10,353	10,353	13,727	13,727	13,727	13,727	21,968	21,968

Notes: The table reports OLS, reduced-form and IV estimates for the the specification described in equation (2) but with occupational diversity, USPTO technology codes diversity, residential segregation of surname groups, residential segregation of ancestral-country group and strength of family ties as the outcome variables. Even columns add ancestral-country diversity to the specification. An observation is a county in a period from 1900 to 1940. Standard errors are clustered at the state level. All variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table 9: Spillover analysis

	Patents per 1,000 people (mean = 1.05, sd = 1.73)			Breakthrough patents per 1,000 people (mean = 0.13, sd = 0.28)		
	(1)	(2)	(3)	(4)	(5)	(6)
<i>Panel A: Least-squares estimates</i>						
Surname diversity	1.288*** (0.181)	1.290*** (0.182)	1.288*** (0.181)	0.094*** (0.017)	0.095*** (0.017)	0.095*** (0.018)
Surname diversity (< 100 miles)	0.460* (0.261)	0.442* (0.252)	0.437* (0.252)	0.033 (0.041)	0.025 (0.043)	0.025 (0.043)
Surname diversity (100 < 200 miles)		0.122 (0.235)	0.131 (0.235)		0.058 (0.083)	0.059 (0.083)
Surname diversity (200 < 300 miles)			-0.082 (0.193)			0.008 (0.038)
Population	0.737*** (0.190)	0.737*** (0.190)	0.736*** (0.190)	0.098*** (0.029)	0.098*** (0.028)	0.098*** (0.028)
Population (< 100 miles)	-0.011 (0.163)	-0.022 (0.174)	-0.074 (0.181)	0.081* (0.045)	0.071* (0.041)	0.066 (0.040)
Population (100 < 200 miles)		-0.136 (0.204)	-0.206 (0.222)		-0.094** (0.039)	-0.102** (0.045)
Population (200 < 300 miles)			-0.295* (0.172)			-0.037 (0.043)
<i>Panel B: Reduced-form estimates</i>						
Predicted surname diversity	0.894*** (0.241)	0.891*** (0.242)	0.891*** (0.241)	0.059*** (0.019)	0.058*** (0.019)	0.058*** (0.018)
Predicted surname diversity (< 100 miles)	0.334* (0.171)	0.328* (0.165)	0.332* (0.167)	0.007 (0.028)	0.006 (0.027)	0.005 (0.027)
Predicted surname diversity (100 < 200 miles)		0.132 (0.151)	0.136 (0.153)		0.066** (0.030)	0.065** (0.031)
Predicted surname diversity (200 < 300 miles)			-0.026 (0.106)			-0.030 (0.027)
Predicted population	0.224 (0.206)	0.225 (0.205)	0.224 (0.205)	0.049 (0.035)	0.049 (0.035)	0.049 (0.034)
Predicted population (< 100 miles)	0.040 (0.135)	0.025 (0.141)	-0.039 (0.146)	0.083** (0.035)	0.067* (0.033)	0.061* (0.032)
Predicted population (100 < 200 miles)		-0.097 (0.180)	-0.180 (0.191)		-0.084** (0.035)	-0.093** (0.040)
Predicted population (200 < 300 miles)			-0.356** (0.135)			-0.032 (0.037)
<i>Panel C: Instrumental-variable estimates</i>						
Surname diversity	1.075*** (0.276)	1.069*** (0.275)	1.069*** (0.287)	0.105*** (0.022)	0.103*** (0.022)	0.099*** (0.023)
Surname diversity (< 100 miles)	0.508 (0.317)	0.561** (0.278)	0.559** (0.271)	-0.015 (0.056)	-0.018 (0.054)	-0.020 (0.054)
Surname diversity (100 < 200 miles)		-0.090 (0.775)	-0.079 (0.913)		0.100 (0.138)	0.139 (0.163)
Surname diversity (200 < 300 miles)			-0.049 (0.945)			-0.186 (0.175)
Population	0.641*** (0.112)	0.642*** (0.108)	0.643*** (0.107)	0.133*** (0.019)	0.132*** (0.018)	0.131*** (0.018)
Population (< 100 miles)	0.128 (0.127)	0.133 (0.128)	0.136 (0.126)	0.034 (0.033)	0.038 (0.035)	0.039 (0.033)
Population (100 < 200 miles)		-0.131 (0.139)	-0.126 (0.141)		-0.058** (0.027)	-0.057* (0.029)
Population (200 < 300 miles)			0.037 (0.161)			0.027 (0.030)
Sanderson-Windmeijer <i>F</i> -stat: Surname diversity	164; 125	145; 88; 28	215; 136; 25; 25	164; 125	145; 88; 28	215; 136; 25; 25
Sanderson-Windmeijer <i>F</i> -stat: Population	31; 485	39; 817; 161	42; 644; 234; 414	31; 485	39; 817; 161	42; 644; 234; 414
County fixed effects	✓	✓	✓	✓	✓	✓
State-Period fixed effects	✓	✓	✓	✓	✓	✓
County-specific linear trends	✓	✓	✓	✓	✓	✓
Observations	21,942	21,942	21,942	21,942	21,942	21,942

Notes: The table reports OLS, reduced-form and IV estimates of regressions of innovation outcomes on surname diversity and population. The unit of observation is a county-period from 1900 to 1940 (including the midyears). The table sequentially adds surname diversity and population in areas within 100 miles (excluding i), 100 miles to 200 miles, 200 miles to 300 miles of county i , and the associated instruments. For each specification we report the first-stage Sanderson-Windmeijer F -stats for all endogenous variables. Standard errors are clustered on states and reported in parentheses. All independent variables are standardized to mean zero and unit variance. ***, **, and * indicate significance at the 1%, 5%, and 10% levels. First-stage estimates are reported in Table B53.

Figures

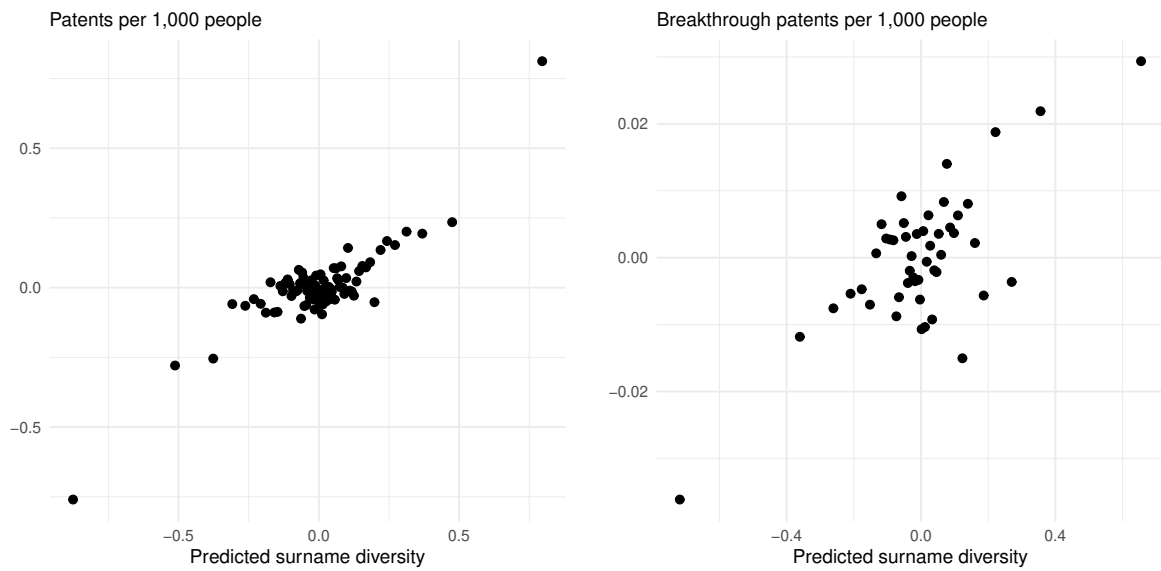


Figure 1: Reduced-from relationships between predicted surname diversity and innovation outcomes

Notes: County-level data from 1900 to 1940. Observations are residualized by predicted population, county fixed effects and state-period fixed effects.

Online Appendix

to “How Cultural Diversity Drives Innovation: Surnames and Patents in U.S. History”

Max Posch, Jonathan Schulz, and Joseph Henrich

A Data Sources and Construction

Surnames and baseline surname diversity measure

To construct county-level surname diversities up until the year 1940, we use the 1850, 1860, 1870, 1880, 1900, 1910, 1920, 1930, and 1940 waves of the full-count Integrated Public Use Microdata Series (IPUMS) compiled by [Ruggles et al. \(2021\)](#) and available on the NBER servers. For each wave, we obtain county identifiers and the variable `name1ast` of all individuals. We perform the following steps to clean the surname variable. First, we transform non-ASCII characters into ASCII characters—e.g., we convert characters with accents or umlauts to the closest letter in English. Second, we convert all characters to upper case. Third, we remove all non-alphabetic characters, including all spaces (e.g., ‘MAC ARTHUR’ becomes ‘MACARTHUR’). Fourth, we drop entries with one or fewer letters. Last, to address potential biases in surnames from misspellings and Anglicization, we apply the [Philips \(1990\)](#) phonetic algorithm *metaphone* to standardize name strings. This method reduces the risk that our analysis is biased by such variations. For instance, the algorithm treats phonetically similar names—like “Heinrich” and its Anglicized form “Henrich”—as identical (“HNRX”). This mitigates concerns that discrimination-induced name changes in less innovative regions could spuriously suggest a link between low surname entropy and lack of innovation. There are 10,633,834 non-standardized unique surnames in the data set, which are reduced to 1,084,159 unique surnames after standardization.

Following [Terry et al. \(2024\)](#), we also obtain individuals’ age and year of immigration, the variables `age` and `yrimmig`, to estimate surname diversities for the midyears 1895, 1905, 1915, and 1925 by removing all individuals who were born or immigrated after the midyear. Ideally, we would also remove all individuals who moved to the county after the midyear, but this information is unavailable.

We harmonize all historical Census data to the 1900 boundaries of US counties using the [Ferrara et al. \(2021\)](#) crosswalks. Specifically, we use the M4 weights, which account for urban and rural areas and topographic suitability. We use 1900 as the reference year

because this is the first year of our panel data set in our main analysis.

Counties harmonized with very few people in a given census year may exhibit very low surname diversity (notably counties in Texas in 1900), potentially due to small sample bias. To mitigate this, we winsorize all entropic surname diversity variables at the 1% level when we run the regressions. Surname fractionalization indices are winsorized at the 2% level. We also standardize all diversity measures to have a mean of zero and a standard deviation of one in the regression analyses, and we compute county-level population size by summing over all surname stocks within a country. We winsorize this variable at the 1% level and standardize it to have a mean of zero and a standard deviation of one when we run the regressions.

Alternative surname diversity measures

We also compute alternative measures of surname diversity using two approaches. First, we interact surnames with a male indicator (*sex*) or the main categories of race (*race*) or birthplace (*bp1*). For the latter, we recode US states and territories (*bp1* codes <10000) to a single code.

Second, we create a distance-weighted measure of surname diversity by integrating the genetic distances between a surname’s most common ancestral country to the US into the entropy measure. We identify the most frequent ancestral country for each phonetically grouped surname using the universe of immigrants to the US from 1880 to 1930—e.g. Germany for surnames like SXLS (Schulz), HNRX (Henrich) and PSX (Posch). This approach accurately identifies the ancestral country for approximately 47% of immigrants, for whom we know the country of birth.

We then link each surname to the genetic distance data from [Spolaore and Wacziarg \(2009\)](#). If no immigrants with a surname are found during this period (which occurs for about 30% of surnames), we categorize it a US surname and assign a distance of 0. We choose the genetic distance data because it serves as a common proxy for cultural distance, is well-defined at the country level, easy to link, and available for most countries. [Spolaore and Wacziarg \(2009\)](#) demonstrates its effectiveness in predicting diffusion of development.

Additionally, given the potentially moderate accuracy of this approach, we also compute a weighted average genetic distance for each surname instead of relying solely on the most frequent ancestral country. For example, 87% of immigrants named SXLS originated from Germany, 3% from Poland, 3% from Russia, 1.5% from Austria, etc. We weight their genetic distances to the US accordingly to get the distance measure for SXLS.

We then integrate the distances into our entropy measure by computing $-\sum dp \log p$, where d is the genetic distance. The two approaches (most frequent ancestral country and

weighted average) yield almost perfectly correlated measures ($\rho \approx 0.99$). The correlation with our baseline entropy measure is $\rho \approx 0.88$.

Ancestral-country diversity measures

To construct a measure of ancestral-country diversity, we use the most frequent ancestral country approach described above, but with phonetically ungrouped surnames instead of phonetically grouped ones. This improves the accuracy rate of identifying the ancestral country among immigrants to 68%. We then calculate the entropy of the distribution of the most frequent ancestral countries of surnames in each county. We winsorize this measure at the 1% level and standardize it to have a mean of zero and a standard deviation of one when we run the regressions. The correlation with our baseline surname entropy measure is $\rho \approx 0.41$ (see Table B6).

Construction of the instrumental variables

We build on the [Burchardi et al. \(2019\)](#) approach to construct an instrumental variable for surname diversity. We identify the number of individuals in a given US county i at the time of each census who immigrated to the US since the prior census and have the surname k . For the 1900 to 1930 census waves, we separate this immigration into five-year periods based on the year each migrant arrived in the US. We obtain immigration flows for the following bins: 1881-1895, 1896-1900, 1901-1905, 1906-1910, 1911-1915, 1916-1920, 1921-1925, and 1926-1930. From the 1880 census wave, we count all first-generation immigrants as well as second-generation immigrants who had an immigrant father, regardless of the date of arrival in the US.

When we predict the stock of people $N_{k,i}^t$ in equation (3), we obtain negative values for 6.3% of observations. The logarithmic transformation of a negative value is undefined. To obtain Shannon entropy for counties containing $N_{i,k}^t$ with negative values, we truncate those negative values at the smallest positive value we observe in the data in a given year. The resulting predicted stock variable is highly correlated with the original variable ($\rho = 0.962$).

We also construct an instrument for population size by summing up the predicted stocks in each county for each period.

Construction of other demographic measures

We construct county-level measures for other types of diversity (race, country of birth and occupation) for each census year from 1850 to 1940. All data are taken from the full-count

IPUMS available on the NBER servers. To calculate population, we sum over all surname stocks within a country. To compute race entropy, we obtain the variable `race` and use the nine main categories. To compute origin country entropy and immigrant shares, we draw on the variable `bp1` with 188 main categories. We recode US states and territories (codes <100) to a single category. To compute occupational entropy, we draw on the variable `occ1950`, dropping all observations with value greater or equal than 979. As before, we transform the data from each period to 1900 US counties using the M4 weights from the [Ferrara et al. \(2021\)](#) cross-walks and use the variables `age` and `yrimmig` to estimate these metrics for the midyears 1895, 1905, 1915, and 1925 by removing all individuals who were born or immigrated after the midyear. We winsorize all entropy variables at the 1% level and standardize them to have a mean of zero and a standard deviation when we run the regressions.

Following the methodology in [Raz \(2025\)](#), we construct the strength of family ties measure from the full-count census data for all census waves from 1860 to 1940. The strength of family ties is determined by the first principal component of four underlying variables: (i) the divorce-to-marriage ratio, (ii) the share of elderly people living without a relative, (iii) the share of people living with at least one non-relative, and (iv) the mean size of families. The variables `age` and `yrimmig` are also used to estimate the strength of family ties for the midyears by removing all individuals who arrived after the midyear. We winsorize and standardize this variable when running regressions as above.

We closely follow the methodology of [Logan and Parman \(2017\)](#) to construct a measure of residential segregation, focusing on surnames rather than race. We identify household heads using the variable `relate` and sort the dataset using `serial`. An indicator is set to one if a neighbor (above or below on the same page, as indicated by `pageno`) has a different surname. If the line above (or below) is missing, we only consider the available line. We then aggregate these indicators to the county-surname level to determine the number of households with at least one different-surname neighbor for each surname group in each county. We also tally households for which we observe both, one, or no neighbors at the county-surname level, as well as the county population with different surname neighbors. Data transformation to 1900 US counties uses the M4 weights from the [Ferrara et al. \(2021\)](#) cross-walks. Following [Logan and Parman \(2017\)](#), we compute segregation estimates under random assignment and complete segregation, calculating the final county-surname level segregation measure. County-level segregation estimates are then averaged across all surname groups, weighted by their population size.

We also apply this approach to origin countries to compute residential segregation of origin-country groups.

The segregation measures are constructed for all census years from 1880 to 1940, excluding midyears due to the impact of excluding post-midyear arrivals on household order.

Counties with very few people in a given census year may show anomalously high segregation values, likely due to small sample bias. To address this, we winsorize all segregation variables at the 1% level. We standardize these variables in the analyses as described above.

Finally, we compute the average years of schooling for each county in 1940 using the variable `higrade` and harmonize the data to 1900 county boundaries.

Construction of the innovation measures

We use the *Comprehensive Universe of US Patents* (CUSP) compiled by [Berkes \(2018\)](#). The data set contains US patents from 1836-2015 and is primarily constructed from Google Patents with supplementary information from other sources. For each patent, the data set provides inventor names and location of residence (geocoded and the 2000 county FIPS identifiers), filing and issuing years of patents, and the US Patent and Trademark Office technology classifications. We harmonize the data to 1900 US counties.

We also employ the breakthrough patent indicator created by [Kelly et al. \(2021\)](#). The authors use the text in patent documents to estimate patent novelty and impact. They assign a higher quality to patents that are novel in terms of cosine similarities. Patents are considered novel if they have low similarity with the existing stock of patents and are impactful in that they have high similarity with subsequent patents.

We construct innovation outcome variables at four levels: county, surname-county, inventor and patent. The county-level outcomes measure the number of patents and breakthrough patents filed by inventors residing in county i during period t , normalized by the county population in 1900. For patents filed by multiple inventors possibly from different counties, we divide the patent count by the number of inventors. We calculate county population sizes in 1900 using the full-count IPUMS and the [Ferrara et al. \(2021\)](#) border harmonization procedure.

In our least-squares analysis, depicted in Figure [B13](#), we use patent issuing years rather than filing years and normalize patents by the population size in 1850 because filing years are inconsistently recorded in the CUSP dataset before 1870.

The surname-county-level outcomes count the number of patents and breakthroughs filed by inventors with surname k residing in county i during period t , normalized by the 1900 surname population within county i . Constructing these variables requires inventor surnames. The CUSP data includes inventor names as a string variable containing the

surname, first name, and sometimes middle names or initials. Identifying surnames from this variable can be challenging due to inconsistent name order. We use punctuation marks such as semicolons, colons, or commas to identify surnames. When the string variable starts with initials followed by a token of two or more characters, or when it ends with a whitespace followed by “DE”, “DU”, “DE LA”, “DI”, “DEL”, “DELLA”, “VAN”, “VON”, “LE”, “LA”, or “ST”, we identify the surnames accordingly. For the remaining entries, we tokenize the string based on whitespace, keeping the first and last tokens, typically representing the first name and surname. To determine the surname, we compare the frequencies of all name combinations from pooled census years 1900, 1910, 1920, 1930, and 1940, identifying the surname based on the most common constellation. For example, for the tokens “JOHN” and “PETER”, we identify the surname based on whether there were more individuals named “JOHN PETER” or “PETER JOHN”. Finally, we clean the surname variable following the described steps.

The inventor-level outcomes count the number of patents and breakthroughs filed by inventor j during period t . The patent-level outcomes are the breakthrough indicator and the number of technology classes assigned to the patent.

For county-level observations, we winsorize innovation outcome variables at the 1% level to lessen the impact of outlier counties with a large number of patents. For surname-county observations, we winsorize the variables at 0.1% level, because few observations have a positive number of breakthrough patents filed by inventors with a specific surname in a given county during a specified period. For inventor-level observations, we winsorize the variables at 10% level, because the distribution is highly skewed, likely because of inaccuracies assigning unique inventor IDs (see below). We also present results using non-winsorized, log-transformed patent counts (see Tables B16, B18, B17, B19 and B20). We do not use the log-transformed variables in the main analysis to avoid the potential issues around logs with zero (Chen and Roth, 2023).

Construction of the inventor-level dataset

We assign unique inventor IDs to observations with the same first and last name and residing in the same county if they filed patents in either the preceding or following period (or both). These periods are grouped as follows: 1900-1904, 1905-1909, 1910-1914, 1915-1919, 1920-1924, 1925-1929, 1930-1934 and 1940-1944. For robustness, we also use the first letter of the first name instead of the full first name to assign inventor IDs. Using the full first name approach, we assign unique IDs for approximately 47% of the patents filed between 1900 and 1944, while the first-letter approach yields about 49%. Table B54 reports the estimates when we replicate the analysis using for the first-letter

approach. We acknowledge that this method is imperfect, as it lumps together individuals with common names like John Smith in big counties. To mitigate this skew, we winsorize the patent outcomes at the 10% level (see above).

B Additional Tables and Figures

Table B1: County-level summary statistics

	N	Mean	SD	Min	Max
Surname Diversity and Population					
Surname diversity	21984	9.552	0.902	7.13	11.767
Predicted surname diversity	21984	10.635	1.032	7.827	13.301
Population	21984	30.721	50.04	1.174	367.82
Predicted population	21984	36.098	41.52	3.415	283.595
Patents and Breakthroughs					
Patents per 1,000 people	21984	1.053	1.733	0	11.196
Breakthrough patents per 1,000 people	21984	0.126	0.282	0	1.918
Other diversities, Segregation and Family Ties					
Birth-country diversity	21873	0.487	0.512	0.003	1.954
Ancestral-country diversity	21984	2.864	0.253	2.189	3.552
Predicted ancestral-country diversity	21984	4.768	0.338	3.219	5.283
Racial diversity	21873	0.334	0.361	0	1.003
Occupational diversity	21871	3.833	0.946	1.797	5.711
USPTO tech codes diversity	10353	4.323	2.611	0	12.264
Segregation of surname groups	13727	0.026	0.018	0.004	0.096
Segregation of ancestral-country groups	13737	1.008	0.008	1	1.051
Strength of family ties	21968	0.023	1.567	-5.346	2.807
Spillovers					
Surname diversity (< 100 miles)	21942	11.205	0.641	9.922	12.565
Predicted surname diversity (< 100 miles)	21942	9.177	0.523	7.746	10.606
Surname diversity (100 < 200 miles)	21942	11.472	0.59	10.398	12.566
Predicted surname diversity (100 < 200 miles)	21942	9.492	0.461	8.554	10.639
Surname diversity (200 < 300 miles)	21942	11.602	0.528	10.513	12.542
Predicted surname diversity (200 < 300 miles)	21942	9.624	0.433	8.735	10.628
Population (< 100 miles)	21942	1,594.204	1,722.36	26.672	11,537.515
Predicted population (< 100 miles)	21942	1,706.251	1,650.026	65.373	11,214.863
Population (100 < 200 miles)	21942	4,320.107	3,496.394	173.435	17,387.047
Predicted population (100 < 200 miles)	21942	4,633.992	3,442.929	343.034	17,602.293
Population (200 < 300 miles)	21942	6,433.896	4,661.43	234.4	21,431.629
Predicted population (200 < 300 miles)	21942	6,894.486	4,525.089	518.537	21,003.459

Notes: The table reports the number of observations, mean, standard deviation, minimum and maximum for all variables considered in the main county-level tables. The first section of the table lists summary statistics for surname diversity and population (in 1,000s of people). The second section provides summary statistics for patents and breakthrough patents filed in the subsequent 5-year period per 1,000 people. The third section contains summary statistics for birth-country diversity, ancestral-country diversity, racial diversity, occupational diversity, USPTO tech codes diversity, segregation of surname and ancestral-country groups, and strength of family ties. The fourth section reports the summary statistics for the surname diversity and population variables used in the spillover analysis.

Table B2: Surname-county, inventor and patent-level summary statistics

	N	Mean	SD	Min	Max
<i>Panel A: Surname-county level</i>					
Surname Diversity and Population					
Surname diversity	19482024	10.334	1.012	8.118	12.512
Predicted surname diversity	19482024	11.516	1.111	9.139	14.017
Population	19482024	174.544	440.804	4.51	2,764.837
Predicted population	19482024	152.471	368.978	8.212	2,623.951
Patents and Breakthroughs					
Patents per 1,000 people	19482024	2.092	23.858	0	500
Breakthrough patents per 1,000 people	19482024	0.193	3.664	0	95.238
<i>Panel B: Inventor level</i>					
Surname Diversity and Population					
Surname diversity	209944	11.474	0.815	9.122	12.652
Predicted surname diversity	209944	12.643	0.859	10.189	14.564
Population	209944	701.607	916.742	10.722	3,984.478
Predicted population	209944	517.423	676.921	16.939	3,047.96
Patents and Breakthroughs					
Patents	209944	2.24	1.692	1	6
Breakthrough patents	209944	0.259	0.438	0	1
<i>Panel C: Patent level</i>					
Surname Diversity and Team Size					
Surname diversity of patent	1451459	0.102	0.302	0	1
Predicted surname diversity	1451459	12.672	0.991	9.911	14.604
Team size	1451459	1.122	0.358	1	7
Patents and Breakthroughs					
Technologies per patent	1451459	2.486	1.822	1	50
Breakthrough patent indicator	1451459	0.17	0.376	0	1

Notes: The table reports the number of observations, mean, standard deviation, minimum and maximum for all variables considered in the main surname-county (Panel A), inventor (Panel B) and patent-level tables (Panel C). The first section of each panel lists summary statistics for county or patent-level surname diversity and population (in 1,000s of people) or team size. The second section of each panel provides summary statistics for patents and breakthrough patents filed in the subsequent 5-year period per 1,000 people (Panel A), not normalized patents and breakthrough patents (Panel B), the breakthrough patent indicator and the number of technologies per patent (Panel C).

Table B3: Surname diversity over time of the counties with the highest and lowest surname diversity in 1940

		Population in year:		Surname entropy in year:					
		1900	1940	1900	1910	1920	1930	1940	
<i>Highest surname diverse counties in 1940</i>									
Cook County, IL	Chicago	1,842,480	4,066,566	12.21	12.52	12.65	12.59	12.61	
Kings County, NY	Brooklyn	1,167,062	2,697,634	11.91	12.28	12.43	12.43	12.56	
New York County, NY	Manhattan	2,056,956	3,283,844	12.16	12.44	12.49	12.44	12.55	
Milwaukee County, MN	Milwaukee	331,519	766,385	12.09	12.37	12.45	12.47	12.48	
Queens County, NY	Queens	153,329	1,298,365	11.64	12.05	12.19	12.34	12.46	
Wayne County, MI	Detroit	347,974	2,017,092	11.95	12.28	12.51	12.44	12.46	
Cuyahoga County, OH	Cleveland	437,737	1,217,402	12.01	12.35	12.45	12.44	12.44	
Hudson County, NJ	Jersey City	386,474	651,688	11.74	12.12	12.34	12.33	12.43	
Erie County, NY	Buffalo	432,269	797,945	11.98	12.23	12.31	12.28	12.39	
Passaic County, NJ	Paterson	155,265	309,293	11.59	12.02	12.18	12.25	12.35	
Sample mean		27,123	46,580	9.28	9.51	9.53	9.65	9.71	
<i>Lowest surname diverse counties in 1940</i>									
Loving County, TX		43	289	4.09	5.69	4.28	5.89	6.01	
Alpine County, CA		379	323	6.56	6.01	5.69	5.88	6.34	
Hinsdale County, CO		1,608	361	8.58	7.53	6.87	6.95	6.64	
Greene County, VA		6,210	5,220	7.03	6.94	6.68	7.03	6.79	
Wayne County, UT		1,899	2,390	7.08	7.03	6.89	7.42	6.86	
Charles City, VA		5,088	4,278	7.30	7.27	6.92	7.08	7.03	
Kane County, UT		1,744	2,563	6.68	6.65	6.65	6.87	7.03	
Rich County, UT		1,946	2,069	7.02	7.01	6.88	7.12	7.13	
Hyde County, NC		9,187	7,864	7.26	7.17	7.08	7.16	7.15	
Craig County, VA		4,286	3,769	7.20	7.44	7.14	7.09	7.18	

Notes: The table lists the surname diversity (entropy) from 1900 to 1940 and population in 1900 and 1940 of the counties with the highest and lowest surname diversity in 1940 as well as the sample means. Column 2 reports the largest city within the county.

Table B4: Surname diversity over time of the counties with the highest and lowest immigrant shares in 1900

	Population in year:	Immigrant share in year:	Surname entropy in year:				
	1900	1900	1900	1910	1920	1930	1940
<i>Counties with largest immigrant share in 1900</i>							
Logan, ND	1,577	52.9%	7.50	9.08	9.05	8.86	8.74
Stark, ND	7,614	51.4%	9.57	10.80	10.73	10.76	10.66
Webb, TX	21,798	51.3%	8.56	8.63	8.85	8.96	9.32
Pembina, ND	17,850	50.9%	9.98	9.85	9.84	9.86	9.93
Kittson, MN	7,864	49.9%	8.63	8.77	8.76	8.81	8.81
Gogebic, MI	16,723	49.5%	10.06	10.59	10.84	10.65	10.85
Chippewa, MI	21,070	49.5%	10.15	10.41	10.32	10.38	10.55
Cavalier, ND	12,271	49.3%	9.88	10.10	9.96	9.92	9.73
Luce, MI	2,982	49.0%	8.85	9.25	9.64	9.78	9.83
Lake, MN	4,639	48.6%	8.99	9.28	9.13	8.86	8.86
<i>Counties with lowest immigrant share in 1900</i>							
Harris, GA	18,014	0.006%	8.55	8.61	8.61	8.48	8.86
Jackson, TN	15,010	0.007%	8.37	8.29	8.15	8.24	8.05
Johnson, KY	13,625	0.007%	7.77	8.03	8.17	8.16	8.19
Heard, GA	11,134	0.009%	8.67	8.63	8.64	8.43	8.57
Lee, GA	10,310	0.010%	8.56	8.61	8.59	8.71	8.86
Franklin, GA	17,671	0.011%	8.90	9.20	9.09	9.22	9.05
Perry, KY	8,275	0.012%	6.79	6.83	8.53	8.93	8.73
Alleghany, NC	7,731	0.013%	7.64	7.68	7.47	7.69	7.90
Owsley, KY	6,845	0.015%	7.70	7.45	7.41	7.48	7.63
Baker, GA	6,743	0.015%	8.44	8.59	8.41	8.52	8.47

Notes: The table lists the surname diversity (entropy) from 1900 to 1940, population and immigrant shares in 1900 of the U.S. counties with the highest and lowest non-zero immigrant shares in 1900.

Table B5: Correlations between baseline surname diversity and alternative surname diversity measures

Surname fract. index	Surname, uncorrected	Surname, men	Surname, household heads	Surname, whites	Surname-race	Surname-county of birth	Genetic distance weighted
0.87	0.99	1.00	1.00	0.99	0.98	0.98	0.91

Notes: This table reports the correlations between our baseline county-level surname diversity measure (surname entropy) and (i) a surname fractionalization index, (ii) entropy of surnames that are not phonetically corrected, surname entropy among (iii) men, (iv) household heads, (v) white individuals, (vi) alternative entropy measures that interact surnames with race or country of birth, and (vii) genetic distance weighted surname entropy. An observation is a county from 1900 to 1940 (including the midyears). The sources and construction of all variables are explained in Appendix A.

Table B6: Correlations between surname diversity and other dimensions of cultural diversity

	Ancestral-country diversity	Birth-country diversity	Racial diversity	Occupational diversity
Partial Corr. (Period FE)	0.41	0.56	-0.27	0.71
Partial Corr. (State-Period FE)	0.33	0.40	0.23	0.64
Partial Corr. (County FE, State-Period FE)	0.40	0.24	0.11	0.27
Partial Corr. (Period FE, Population)	0.36	0.49	-0.33	0.61
Partial Corr. (State-Period FE, Population)	0.26	0.22	0.19	0.55
Partial Corr. (County FE, State-Period FE, Population)	0.38	0.25	0.10	0.29

Notes: This table reports conditional correlations between surname diversity and other dimensions of cultural diversity. An observation is a county from 1900 to 1940 (including the midyears). The sources and construction of all variables are explained in Appendix A.

Table B7: Controlling for immigration (county level)

	Patents per 1,000 people (mean = 1.09, sd = 1.68)				Breakthrough patents per 1,000 people (mean = 0.12, sd = 0.27)			
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
<i>Panel A: Least-squares estimates</i>								
Surname diversity	1.079*** (0.152)	1.067*** (0.154)	1.026*** (0.204)	1.216*** (0.196)	0.087*** (0.014)	0.085*** (0.014)	0.075*** (0.015)	0.121*** (0.018)
Immigration		0.099** (0.033)	0.112** (0.032)	0.025 (0.036)		0.019** (0.009)	0.025*** (0.007)	0.013 (0.011)
Population	0.779*** (0.099)	0.778*** (0.093)	0.743*** (0.094)	0.345* (0.189)	0.139*** (0.020)	0.139*** (0.019)	0.119*** (0.017)	0.043 (0.031)
<i>Panel B: Reduced-form estimates</i>								
Predicted surname diversity	0.641*** (0.194)	0.623*** (0.197)	0.621** (0.234)	0.862*** (0.227)	0.032** (0.016)	0.029* (0.016)	0.031* (0.017)	0.082*** (0.020)
Immigration		0.099** (0.047)	0.118** (0.045)	0.058 (0.043)		0.020* (0.010)	0.026*** (0.009)	0.017 (0.011)
Predicted population	0.628*** (0.136)	0.634*** (0.134)	0.584*** (0.151)	-0.097 (0.176)	0.135*** (0.026)	0.137*** (0.025)	0.109*** (0.024)	-0.008 (0.026)
<i>Panel C: Instrumental-variable estimates</i>								
Surname diversity	1.087*** (0.164)	1.063*** (0.168)	1.066*** (0.207)	1.239*** (0.207)	0.083*** (0.016)	0.079*** (0.016)	0.079*** (0.015)	0.118*** (0.019)
Immigration		0.099*** (0.033)	0.110*** (0.032)	0.030 (0.037)		0.019** (0.008)	0.025*** (0.006)	0.014 (0.011)
Population	0.800*** (0.120)	0.807*** (0.114)	0.736*** (0.104)	0.239 (0.249)	0.168*** (0.026)	0.169*** (0.025)	0.127*** (0.020)	0.026 (0.048)
Sanderson-Windmeijer F-stat	183; 287	180; 308	167; 158	124; 16	183; 287	180; 308	167; 158	124; 16
<i>Panel D: First-stage estimates</i>								
	Surname diversity				Population			
Predicted surname diversity	0.705*** (0.071)	0.707*** (0.071)	0.733*** (0.070)	0.714*** (0.078)	-0.157*** (0.033)	-0.159*** (0.033)	-0.216*** (0.041)	-0.093** (0.043)
Immigration		-0.009 (0.006)	0.002 (0.006)	0.018*** (0.006)		0.012 (0.021)	0.007 (0.023)	0.023 (0.017)
Predicted population	-0.025 (0.042)	-0.025 (0.042)	-0.082* (0.041)	-0.170** (0.070)	0.818*** (0.062)	0.819*** (0.061)	0.913*** (0.085)	0.477*** (0.122)
County fixed effects	✓	✓	✓	✓	✓	✓	✓	✓
Period fixed effects	✓	✓			✓	✓		
State-Period fixed effects			✓	✓			✓	✓
County-specific linear trends				✓				✓
Observations	19,236	19,236	19,236	19,236	19,236	19,236	19,236	19,236

Notes: The table reports OLS, reduced-form, and IV estimates for the specifications described in equation (2). An observation is a county in a period from 1900 to 1930. The endogenous variables are county-level surname diversity and population (both in t). Columns 2–4 and 6–8 control for the number of recent immigrants (between $t - 1$ and t). In columns 1 to 4, the dependent variable is the number of patents filed in the county in the five-year period starting in t divided by county population size in 1900. In columns 5 to 8, the dependent variable is the number of breakthrough patents filed in the county in the five-year period starting in t divided by county population size in 1900. Standard errors are clustered at the state level. Independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B8: Controlling for immigration (county-surname level)

	Patents per 1,000 people (mean = 1.05, sd = 13.43)				Breakthrough patents per 1,000 people (mean = 0.1, sd = 2.16)			
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
<i>Panel A: Least-squares estimates</i>								
Surname diversity	0.940*** (0.160)	0.932*** (0.160)	0.949*** (0.164)	0.945*** (0.238)	0.072*** (0.022)	0.070*** (0.022)	0.070*** (0.022)	0.055*** (0.016)
Immigration		0.565** (0.248)	0.585** (0.255)	0.165 (0.161)		0.115* (0.058)	0.118* (0.059)	0.088 (0.074)
Population	4.201*** (0.819)	4.104*** (0.742)	4.120*** (0.757)	2.881*** (0.370)	0.633*** (0.105)	0.613*** (0.100)	0.621*** (0.101)	0.445*** (0.081)
<i>Panel B: Reduced-form estimates</i>								
Predicted surname diversity	0.487*** (0.177)	0.471*** (0.174)	0.479*** (0.178)	0.500*** (0.140)	0.036 (0.021)	0.033 (0.021)	0.032 (0.021)	0.048 (0.033)
Immigration		0.859** (0.395)	0.887** (0.407)	0.260 (0.165)		0.159* (0.084)	0.164* (0.086)	0.093 (0.064)
Predicted population	3.542*** (1.135)	3.529*** (1.022)	3.543*** (1.043)	0.796 (0.586)	0.477*** (0.143)	0.474*** (0.134)	0.481*** (0.137)	-0.137 (0.345)
<i>Panel C: Instrumental-variable estimates</i>								
Surname diversity	1.169*** (0.222)	1.148*** (0.222)	1.166*** (0.227)	1.089*** (0.249)	0.102*** (0.030)	0.098*** (0.030)	0.098*** (0.030)	0.086* (0.043)
Immigration		0.597** (0.274)	0.619** (0.282)	0.198 (0.190)		0.124* (0.068)	0.128* (0.070)	0.102 (0.085)
Population	3.588*** (0.684)	3.580*** (0.618)	3.593*** (0.633)	1.235 (0.957)	0.486*** (0.091)	0.484*** (0.094)	0.491*** (0.096)	-0.189 (0.478)
Sanderson-Windmeijer <i>F</i> -stat	211; 40	209; 45	209; 44	170; 68	211; 40	209; 45	209; 44	170; 68
<i>Panel D: First-stage estimates</i>								
	Surname diversity				Population			
Predicted surname diversity	0.539*** (0.041)	0.539*** (0.041)	0.539*** (0.041)	0.493*** (0.044)	-0.040** (0.016)	-0.041** (0.016)	-0.042** (0.016)	-0.029*** (0.006)
Immigration		0.008 (0.005)	0.008 (0.006)	0.004 (0.008)		0.071 (0.046)	0.072 (0.047)	0.046** (0.021)
Predicted population	0.046 (0.032)	0.045 (0.031)	0.047 (0.031)	-0.059 (0.072)	0.972*** (0.154)	0.971*** (0.145)	0.971*** (0.147)	0.697*** (0.090)
County-Surname fixed effects	✓	✓	✓	✓	✓	✓	✓	✓
State-Period fixed effects	✓	✓	✓	✓	✓	✓	✓	✓
Surname-Period fixed effects			✓	✓			✓	✓
County-specific linear trends				✓				✓
Observations	17,046,771	17,046,771	17,046,771	17,046,771	17,046,771	17,046,771	17,046,771	17,046,771

Notes: The table reports OLS, reduced-form, and IV estimates for the specifications described in equation (8). An observation is a surname in a given county in a period from 1900 to 1930. Observations are weighted based on their population shares within counties. The endogenous variables are county-level surname diversity and population (both in t). Columns 2–4 and 6–8 control for the number of recent immigrants (between $t - 1$ and t). In columns 1 to 4, the dependent variable is number of patents filed by individuals with surname n residing in county i in the five-year period starting in t divided by surname population size in 1900. In columns 5 to 8, the dependent variable is number of breakthrough patents filed by individuals with surname n residing in county i in the five-year period starting in t divided by surname population size in 1900. Standard errors are clustered at the state level. Independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B9: Controlling for immigration (inventor level)

	Patents (mean = 2.22, sd = 1.68)		Breakthrough patents (mean = 0.26, sd = 0.44)	
	(1)	(2)	(3)	(4)
<i>Panel A: Least-squares estimates</i>				
Surname diversity	0.228*** (0.062)	0.229*** (0.062)	0.046* (0.026)	0.045* (0.027)
Immigration		0.007 (0.015)		-0.004 (0.004)
Population	-0.019 (0.061)	-0.021 (0.058)	-0.016 (0.012)	-0.015 (0.013)
<i>Panel B: Reduced-form estimates</i>				
Predicted surname diversity	0.080** (0.024)	0.080** (0.024)	0.019** (0.008)	0.019** (0.008)
Immigration		0.002 (0.015)		-0.006** (0.003)
Predicted population	-0.008 (0.040)	-0.007 (0.043)	-0.013* (0.008)	-0.016** (0.007)
<i>Panel C: Instrumental-variable estimates</i>				
Surname diversity	0.456** (0.187)	0.451** (0.182)	0.110* (0.064)	0.112* (0.064)
Immigration		0.009 (0.016)		-0.004 (0.004)
Population	-0.060 (0.043)	-0.056 (0.047)	-0.028*** (0.009)	-0.029*** (0.008)
Sanderson-Windmeijer F-stat	17; 504	16; 622	17; 504	16; 622
<i>Panel D: First-stage estimates</i>				
	Surname diversity		Population	
Predicted surname diversity	0.179*** (0.048)	0.180*** (0.047)	0.026 (0.052)	0.021 (0.044)
Immigration		-0.008 (0.011)		0.066*** (0.022)
Predicted population	0.098*** (0.020)	0.095*** (0.020)	0.869*** (0.096)	0.891*** (0.092)
Inventor fixed effects	✓	✓	✓	✓
Period-State fixed effects	✓	✓	✓	✓
Observations	198,805	198,805	198,805	198,805

Notes: The table reports OLS, reduced-form, and IV estimates for the specifications described in equation (9). An observation is an inventor residing in a given county in a period from 1900 to 1930. The endogenous variables are county-level surname diversity and population (both in t). Columns 2 and 4 control for the number of recent immigrants (between $t - 1$ and t). The dependent variables are number of patents filed by the inventor in the five-year period starting in t (column 1–2) and number of breakthrough patents filed by the inventor in the five-year period starting in t (column 3–4). Standard errors are clustered at the state level. The independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B10: Placebo and dynamic effects (county level)

	Patents per 1,000 people				Breakthrough patents per 1,000 people			
	$t - 5$ (1)	t (2)	$t + 5$ (3)	$t + 10$ (4)	$t - 5$ (5)	t (6)	$t + 5$ (7)	$t + 10$ (8)
<i>Panel A: Least-squares estimates</i>								
Surname diversity	0.119 (0.129)	1.338*** (0.174)	0.797*** (0.120)	0.112 (0.205)	-0.015 (0.023)	0.099*** (0.018)	0.142*** (0.023)	-0.009 (0.032)
Population	0.222 (0.186)	0.736*** (0.189)	-0.081 (0.169)	-0.164 (0.233)	0.049 (0.030)	0.099*** (0.029)	-0.110* (0.062)	0.011 (0.034)
<i>Panel B: Reduced-form estimates</i>								
Predicted surname diversity	0.077 (0.127)	0.939*** (0.236)	0.566*** (0.127)	0.049 (0.143)	-0.025 (0.022)	0.061*** (0.019)	0.107*** (0.025)	0.000 (0.017)
Predicted population	0.140 (0.131)	0.205 (0.205)	-0.125 (0.127)	0.006 (0.155)	0.019 (0.024)	0.051 (0.035)	-0.019 (0.041)	0.011 (0.027)
<i>Panel C: Instrumental-variable estimates</i>								
Surname diversity	0.154 (0.205)	1.392*** (0.191)	0.784*** (0.138)	0.073 (0.188)	-0.035 (0.032)	0.101*** (0.019)	0.149*** (0.024)	0.003 (0.021)
Population	0.358 (0.297)	0.731*** (0.236)	0.080 (0.248)	0.039 (0.294)	0.038 (0.058)	0.118** (0.054)	0.025 (0.081)	0.023 (0.056)
Sanderson-Windmeijer F -stat	91; 22	149; 29	132; 12	135; 15	91; 22	149; 29	132; 12	135; 15
<i>Panel D: First-stage estimates</i>								
	Surname diversity				Population			
	$t - 5$ (1)	t (2)	$t + 5$ (3)	$t + 10$ (4)	$t - 5$ (5)	t (6)	$t + 5$ (7)	$t + 10$ (8)
Predicted surname diversity	0.639*** (0.089)	0.729*** (0.074)	0.730*** (0.082)	0.719*** (0.077)	-0.061 (0.038)	-0.103** (0.041)	-0.074 (0.047)	-0.088* (0.045)
Predicted population	-0.079 (0.067)	-0.142** (0.063)	-0.205** (0.079)	-0.164** (0.072)	0.426*** (0.104)	0.550*** (0.107)	0.447*** (0.131)	0.476*** (0.127)
County fixed effects	✓	✓	✓	✓	✓	✓	✓	✓
State-Period fixed effects	✓	✓	✓	✓	✓	✓	✓	✓
County-specific linear trends	✓	✓	✓	✓	✓	✓	✓	✓
Observations	16,488	21,984	16,488	16,488	16,488	21,984	16,488	16,488

Notes: The table reports OLS, reduced-form, and IV estimates of the leads and lags of innovation outcomes on surname diversity and population for the specifications described in equation (2). Columns 1 and 5 use the five-year lag of the dependent variables. Columns 2 and 6 repeat the baseline specification (contemporaneous values of the dependent variables). Columns 3–4 and 7–8 use the five-year and ten-year lead of the dependent variables, respectively. An observation is a county in a period. Standard errors are clustered on states and reported in parentheses. Independent variables are standardized to mean zero and unit variance. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B11: Placebo and dynamic effects (county-surname level)

	Patents per 1,000 people				Breakthrough patents per 1,000 people			
	$t - 5$ (1)	t (2)	$t + 5$ (3)	$t + 10$ (4)	$t - 5$ (5)	t (6)	$t + 5$ (7)	$t + 10$ (8)
<i>Panel A: Least-squares estimates</i>								
Surname diversity	-0.016 (0.274)	1.143*** (0.182)	1.113*** (0.219)	0.629*** (0.130)	-0.069** (0.032)	0.061*** (0.015)	0.149*** (0.033)	0.053* (0.029)
Population	2.631*** (0.703)	3.679*** (0.597)	-2.454 (2.224)	-2.484 (1.622)	0.216* (0.128)	0.728*** (0.100)	-1.543** (0.688)	0.327*** (0.076)
<i>Panel B: Reduced-form estimates</i>								
Predicted surname diversity	-0.196 (0.185)	0.675*** (0.161)	0.538*** (0.158)	0.404*** (0.116)	-0.064* (0.034)	0.044 (0.030)	0.092** (0.042)	0.038* (0.020)
Predicted population	2.834* (1.480)	0.958 (0.971)	-0.543 (1.779)	-2.185 (1.590)	0.312 (0.333)	0.112 (0.280)	-0.428 (0.670)	0.106 (0.095)
<i>Panel C: Instrumental-variable estimates</i>								
Surname diversity	-0.238 (0.362)	1.358*** (0.248)	1.041*** (0.203)	0.627*** (0.141)	-0.118** (0.058)	0.091** (0.043)	0.148*** (0.036)	0.086** (0.040)
Population	4.038** (1.859)	1.478 (1.509)	-0.346 (2.048)	-2.894 (1.933)	0.450 (0.445)	0.170 (0.424)	-0.481 (0.736)	0.151 (0.134)
Sanderson-Windmeijer F -stat	250; 188	215; 300	156; 26	146; 70	250; 188	215; 300	156; 26	146; 70
<i>Panel D: First-stage estimates</i>								
	Surname diversity				Population			
	$t - 5$ (1)	t (2)	$t + 5$ (3)	$t + 10$ (4)	$t - 5$ (5)	t (6)	$t + 5$ (7)	$t + 10$ (8)
Predicted surname diversity	0.462*** (0.032)	0.520*** (0.046)	0.505*** (0.045)	0.499*** (0.046)	-0.021*** (0.005)	-0.021*** (0.004)	-0.036*** (0.010)	-0.031*** (0.007)
Predicted population	0.045 (0.052)	-0.028 (0.059)	-0.251* (0.135)	-0.072 (0.078)	0.705*** (0.079)	0.674*** (0.049)	0.812*** (0.162)	0.739*** (0.090)
County-Surname fixed effects	✓	✓	✓	✓	✓	✓	✓	✓
Surname-Period fixed effects	✓	✓	✓	✓	✓	✓	✓	✓
State-Period fixed effects	✓	✓	✓	✓	✓	✓	✓	✓
County-specific linear trends	✓	✓	✓	✓	✓	✓	✓	✓
Observations	14,611,518	19,482,024	14,611,518	14,611,518	14,611,518	19,482,024	14,611,518	14,611,518

Notes: The table reports OLS, reduced-form, and IV estimates of the leads and lags of innovation outcomes on surname diversity and population for the specifications described in equation (8). Columns 1 and 5 use the five-year lag of the dependent variables. Columns 2 and 6 repeat the baseline specification (contemporaneous values of the dependent variables). Columns 3–4 and 7–8 use the five-year and ten-year lead of the dependent variables, respectively. An observation is a surname in a given county in a period from 1900 to 1940. Observations are weighted based on their population shares within counties. Standard errors are clustered on states and reported in parentheses. Independent variables are standardized to mean zero and unit variance. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B12: Placebo and dynamic effects (inventor level)

	Patents				Breakthrough patents			
	$t - 5$ (1)	t (2)	$t + 5$ (3)	$t + 10$ (4)	$t - 5$ (5)	t (6)	$t + 5$ (7)	$t + 10$ (8)
<i>Panel A: Least-squares estimates</i>								
Surname diversity	0.118 (0.078)	0.219*** (0.061)	0.089 (0.113)	0.165** (0.074)	-0.001 (0.024)	0.043 (0.026)	0.030** (0.014)	0.030 (0.020)
Population	0.026* (0.014)	-0.016 (0.053)	0.018 (0.017)	0.017 (0.013)	0.006* (0.004)	-0.011 (0.015)	-0.007** (0.003)	-0.005 (0.004)
<i>Panel B: Reduced-form estimates</i>								
Predicted surname diversity	0.024 (0.054)	0.078*** (0.022)	-0.012 (0.045)	0.095* (0.047)	0.003 (0.013)	0.016** (0.008)	0.022** (0.009)	0.009 (0.017)
Predicted population	0.039** (0.016)	0.003 (0.040)	0.039** (0.019)	0.025* (0.014)	0.010 (0.007)	-0.008 (0.010)	-0.005 (0.006)	-0.003 (0.004)
<i>Panel C: Instrumental-variable estimates</i>								
Surname diversity	0.152 (0.294)	0.438** (0.182)	-0.017 (0.144)	0.298** (0.137)	0.023 (0.072)	0.094 (0.059)	0.066** (0.026)	0.025 (0.051)
Population	0.031 (0.040)	-0.045 (0.044)	0.044 (0.034)	-0.001 (0.023)	0.009 (0.014)	-0.019 (0.012)	-0.012* (0.007)	-0.006 (0.005)
Sanderson-Windmeijer F -stat	22; 19	17; 504	39; 28	37; 33	22; 19	17; 504	39; 28	37; 33
<i>Panel D: First-stage estimates</i>								
	Surname diversity				Population			
	$t - 5$ (1)	t (2)	$t + 5$ (3)	$t + 10$ (4)	$t - 5$ (5)	t (6)	$t + 5$ (7)	$t + 10$ (8)
Predicted surname diversity	0.174*** (0.035)	0.179*** (0.049)	0.301*** (0.046)	0.318*** (0.051)	-0.080 (0.072)	0.020 (0.048)	-0.164 (0.127)	-0.270** (0.126)
Predicted population	0.081*** (0.013)	0.095*** (0.019)	0.089*** (0.013)	0.087*** (0.013)	0.855*** (0.129)	0.867*** (0.101)	0.931*** (0.129)	0.919*** (0.125)
Inventor fixed effects	✓	✓	✓	✓	✓	✓	✓	✓
Period-State fixed effects	✓	✓	✓	✓	✓	✓	✓	✓
Observations	176,961	209,944	174,858	158,072	176,961	209,944	174,858	158,072

Notes: The table reports OLS, reduced-form, and IV estimates of the leads and lags of innovation outcomes on surname diversity and population for the specifications described in equation (9). An observation is an inventor residing in a given county in a period from 1900 to 1940. Columns 1 and 5 use the five-year lag of the dependent variables. Columns 2 and 6 repeat the baseline specification (contemporaneous values of the dependent variables). Columns 3–4 and 7–8 use the five-year and ten-year lead of the dependent variables, respectively. An observation is an inventor residing in a given county in a period from 1900 to 1940. Standard errors are clustered on states and reported in parentheses. The independent variables are standardized to mean zero and unit variance. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B13: Regional heterogeneity in the effect of surname diversity on innovation (county level)

	Patents per 1,000 people		Breakthrough patents per 1,000 people	
	(1)	(2)	(3)	(4)
<i>Panel A: Least-squares estimates</i>				
Surname diversity × Region = Midwest	0.821*** (0.148)	1.352*** (0.230)	0.058** (0.025)	0.087* (0.050)
Surname diversity × Region = Northeast	0.681 (0.496)	0.423*** (0.124)	0.173 (0.106)	0.180* (0.097)
Surname diversity × Region = South	1.158*** (0.282)	1.083*** (0.325)	0.099*** (0.017)	0.065*** (0.014)
Surname diversity × Region = West	1.481*** (0.243)	2.267*** (0.423)	0.173*** (0.034)	0.200*** (0.035)
Population × Region = Midwest	0.974*** (0.181)	1.299*** (0.201)	0.149*** (0.037)	0.139*** (0.033)
Population × Region = Northeast	0.968*** (0.202)	0.519*** (0.049)	0.173*** (0.034)	0.139*** (0.042)
Population × Region = South	0.349*** (0.126)	0.147 (0.279)	0.095*** (0.021)	0.049 (0.064)
Population × Region = West	0.548** (0.233)	1.780** (0.733)	0.136*** (0.040)	0.150** (0.059)
<i>Panel B: Reduced-form estimates</i>				
Predicted surname diversity × Region = Midwest	0.496*** (0.169)	1.119*** (0.270)	0.028 (0.023)	0.070 (0.048)
Predicted surname diversity × Region = Northeast	0.045 (0.337)	0.160 (0.144)	-0.013 (0.069)	-0.051* (0.029)
Predicted surname diversity × Region = South	0.852** (0.331)	0.803* (0.411)	0.063*** (0.024)	0.049** (0.021)
Predicted surname diversity × Region = West	1.106*** (0.255)	1.586*** (0.257)	0.093*** (0.030)	0.135*** (0.025)
Predicted population × Region = Midwest	0.927*** (0.200)	0.675* (0.354)	0.147*** (0.036)	0.113* (0.064)
Predicted population × Region = Northeast	0.905*** (0.208)	0.352** (0.133)	0.191*** (0.039)	0.147*** (0.050)
Predicted population × Region = South	0.077 (0.108)	-0.396 (0.292)	0.076*** (0.023)	-0.006 (0.036)
Predicted population × Region = West	0.176 (0.334)	0.705** (0.345)	0.108** (0.042)	0.013 (0.061)
<i>Panel C: Instrumental-variable estimates</i>				
Surname diversity × Region = Midwest	0.831*** (0.172)	1.464*** (0.267)	0.069** (0.028)	0.101* (0.054)
Surname diversity × Region = Northeast	0.791 (0.512)	0.549*** (0.146)	0.136 (0.120)	0.000 (0.055)
Surname diversity × Region = South	1.155*** (0.280)	1.135*** (0.352)	0.099*** (0.021)	0.073*** (0.022)
Surname diversity × Region = West	1.738*** (0.300)	2.439*** (0.454)	0.190*** (0.034)	0.195*** (0.029)
Population × Region = Midwest	0.989*** (0.225)	0.901** (0.420)	0.156*** (0.042)	0.165* (0.089)
Population × Region = Northeast	0.930*** (0.189)	0.510** (0.208)	0.198*** (0.028)	0.238** (0.097)
Population × Region = South	0.334*** (0.108)	-0.388 (0.342)	0.113*** (0.019)	0.039 (0.123)
Population × Region = West	0.217 (0.389)	1.692** (0.736)	0.133*** (0.045)	0.079 (0.062)
Sanderson-Windmeijer <i>F</i> -stat: Surname diversity	547; 448; 569; 213	390; 101; 210; 237	547; 448; 569; 213	390; 101; 210; 237
Sanderson-Windmeijer <i>F</i> -stat: Population	616; 109; 133; 668	183; 122; 23; 47	616; 109; 133; 668	183; 122; 23; 47
County fixed effects	✓	✓	✓	✓
Period fixed effects	✓	✓	✓	✓
County-specific linear trends		✓		✓
Observations	21,984	21,984	21,984	21,984

Notes: The table reports regional heterogeneity in the estimates for the specifications described in equation (2). An observation is a county in a period from 1900 to 1940. Standard errors are clustered at the state level. Independent variables are standardized to mean zero and unit variance. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

**Table B14: Regional heterogeneity in the effect of surname diversity on innovation
(county-surname level)**

	Patents per 1,000 people	Breakthrough patents per 1,000 people
	(1)	(2)
<i>Panel A: Least-squares estimates</i>		
Surname diversity × Region = Midwest	1.140*** (0.239)	0.125*** (0.042)
Surname diversity × Region = Northeast	2.021 (1.247)	0.392 (0.281)
Surname diversity × Region = South	0.521** (0.200)	0.025 (0.016)
Surname diversity × Region = West	1.655*** (0.513)	0.144*** (0.032)
Population × Region = Midwest	3.197*** (0.439)	0.375*** (0.088)
Population × Region = Northeast	2.530*** (0.485)	0.586*** (0.125)
Population × Region = South	3.721 (2.283)	0.632*** (0.233)
Population × Region = West	5.092*** (0.127)	0.843*** (0.020)
<i>Panel B: Reduced-form estimates</i>		
Predicted surname diversity × Region = Midwest	1.023*** (0.249)	0.115*** (0.041)
Predicted surname diversity × Region = Northeast	1.549* (0.805)	0.343 (0.212)
Predicted surname diversity × Region = South	0.174 (0.159)	-0.013 (0.015)
Predicted surname diversity × Region = West	0.801*** (0.258)	0.043* (0.025)
Predicted population × Region = Midwest	2.335** (0.892)	0.256** (0.095)
Predicted population × Region = Northeast	1.863*** (0.308)	0.407*** (0.101)
Predicted population × Region = South	3.698 (3.025)	0.670* (0.362)
Predicted population × Region = West	7.058*** (0.471)	1.100*** (0.061)
<i>Panel C: Instrumental-variable estimates</i>		
Surname diversity × Region = Midwest	1.433*** (0.310)	0.161*** (0.054)
Surname diversity × Region = Northeast	2.882* (1.434)	0.638 (0.381)
Surname diversity × Region = South	0.564** (0.267)	0.017 (0.019)
Surname diversity × Region = West	2.383*** (0.748)	0.207** (0.085)
Population × Region = Midwest	2.719*** (0.585)	0.298*** (0.090)
Population × Region = Northeast	1.863*** (0.295)	0.407*** (0.100)
Population × Region = South	4.782* (2.558)	0.799*** (0.245)
Population × Region = West	4.496*** (0.190)	0.714*** (0.022)
Sanderson-Windmeijer <i>F</i> -stat: Surname diversity	473; 144; 159; 220	473; 144; 159; 220
Sanderson-Windmeijer <i>F</i> -stat: Population	311; 726; 94; 10035	311; 726; 94; 10035
County-Surname fixed effects	✓	✓
Surname-Period fixed effects	✓	✓
State-Period fixed effects	✓	✓
Observations	19,482,024	19,482,024

Notes: The table reports regional heterogeneity in the estimates for the specifications described in equation (8). An observation is a surname in a given county in a period from 1900 to 1940. Observations are weighted based on their population shares within counties. Standard errors are clustered at the state level. Independent variables are standardized to mean zero and unit variance. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B15: Regional heterogeneity in the effect of surname diversity on innovation (inventor level)

	Patents		Breakthrough patents	
	(1)	(2)	(3)	(4)
<i>Panel A: Least-squares estimates</i>				
Surname diversity × Region = Midwest	0.420*** (0.081)	0.420*** (0.086)	0.048** (0.020)	0.058*** (0.019)
Surname diversity × Region = Northeast	-0.034 (0.041)	-0.028 (0.059)	0.009 (0.022)	-0.003 (0.022)
Surname diversity × Region = South	-0.105 (0.098)	-0.158 (0.113)	0.035 (0.028)	0.017 (0.027)
Surname diversity × Region = West	0.417*** (0.055)	0.360*** (0.049)	0.090*** (0.021)	0.076*** (0.022)
Population × Region = Midwest		0.006 (0.049)		-0.021*** (0.004)
Population × Region = Northeast		0.003 (0.131)		0.010 (0.011)
Population × Region = South		0.595** (0.256)		0.094 (0.077)
Population × Region = West		0.061*** (0.013)		0.005 (0.006)
<i>Panel B: Reduced-form estimates</i>				
Predicted surname diversity × Region = Midwest	0.010 (0.053)	0.032 (0.054)	0.016 (0.013)	0.012 (0.013)
Predicted surname diversity × Region = Northeast	-0.032* (0.016)	-0.015 (0.020)	0.002 (0.012)	0.001 (0.013)
Predicted surname diversity × Region = South	-0.058 (0.057)	-0.016 (0.099)	0.017 (0.015)	0.012 (0.022)
Predicted surname diversity × Region = West	0.047*** (0.013)	0.072*** (0.016)	0.005 (0.004)	0.011*** (0.004)
Predicted population × Region = Midwest		0.063 (0.039)		-0.010** (0.004)
Predicted population × Region = Northeast		-0.028 (0.105)		-0.009 (0.019)
Predicted population × Region = South		-0.331 (1.067)		0.051 (0.160)
Predicted population × Region = West		0.174*** (0.022)		0.036*** (0.007)
<i>Panel C: Instrumental-variable estimates</i>				
Surname diversity × Region = Midwest	0.069 (0.246)	0.191 (0.184)	0.083 (0.057)	0.044 (0.043)
Surname diversity × Region = Northeast	-0.043 (0.077)	0.019 (0.089)	0.032 (0.046)	0.016 (0.054)
Surname diversity × Region = South	-0.106 (0.165)	0.001 (0.224)	0.074 (0.046)	0.045 (0.050)
Surname diversity × Region = West	0.574* (0.330)	0.608* (0.339)	0.088 (0.057)	0.078 (0.060)
Population × Region = Midwest		0.057 (0.066)		-0.018** (0.007)
Population × Region = Northeast		0.000 (0.157)		-0.014 (0.046)
Population × Region = South		-0.289 (1.387)		0.055 (0.230)
Population × Region = West		-0.004 (0.053)		0.009 (0.014)
Sanderson-Windmeijer F-stat: Surname diversity	55; 51; 101; 17	149; 205; 187; 3935	55; 51; 101; 17	149; 205; 187; 3935
Sanderson-Windmeijer F-stat: Population		840; 3069; 174; 15344		840; 3069; 174; 15344
Inventor fixed effects	✓	✓	✓	✓
Period fixed effects	✓	✓	✓	✓
Observations	209,944	209,944	209,944	209,944

Notes: The table reports regional heterogeneity in the estimates for the specifications described in equation (9). An observation is an inventor residing in a given county in a period from 1900 to 1940. Standard errors are clustered at the state level. Independent variables are standardized to mean zero and unit variance. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B16: Log-transformed population-normalized innovation outcomes (county level)

	Log Patents per 1,000 people (mean = 0.54, sd = 0.56)			Log Breakthrough patents per 1,000 people (mean = 0.1, sd = 0.22)		
	(1)	(2)	(3)	(4)	(5)	(6)
<i>Panel A: Least-squares estimates</i>						
Surname diversity	0.411*** (0.064)	0.395*** (0.087)	0.409*** (0.058)	0.103*** (0.023)	0.098*** (0.029)	0.069*** (0.013)
Population	0.151*** (0.024)	0.141*** (0.023)	0.115*** (0.031)	0.105*** (0.019)	0.096*** (0.017)	0.081*** (0.018)
<i>Panel B: Reduced-form estimates</i>						
Predicted surname diversity	0.279*** (0.075)	0.270*** (0.093)	0.307*** (0.075)	0.051* (0.029)	0.049 (0.034)	0.045*** (0.015)
Predicted population	0.106*** (0.039)	0.091** (0.043)	-0.030 (0.049)	0.099*** (0.025)	0.088*** (0.025)	0.040* (0.024)
<i>Panel C: Instrumental-variable estimates</i>						
Surname diversity	0.433*** (0.060)	0.415*** (0.084)	0.429*** (0.062)	0.102*** (0.026)	0.100*** (0.031)	0.074*** (0.014)
Population	0.131*** (0.033)	0.129*** (0.027)	0.057 (0.048)	0.119*** (0.027)	0.100*** (0.020)	0.092** (0.034)
Sanderson-Windmeijer <i>F</i> -stat	162; 426	147; 227	149; 29	162; 426	147; 227	149; 29
<i>Panel D: First-stage estimates</i>						
	Surname diversity			Population		
Predicted surname diversity	0.693*** (0.075)	0.723*** (0.073)	0.729*** (0.074)	-0.160*** (0.033)	-0.231*** (0.042)	-0.103** (0.041)
Predicted population	-0.009 (0.041)	-0.076* (0.040)	-0.142** (0.063)	0.837*** (0.053)	0.948*** (0.073)	0.550*** (0.107)
County fixed effects	✓	✓	✓	✓	✓	✓
Period fixed effects	✓			✓		
State-Period fixed effects		✓	✓		✓	✓
County-specific linear trends			✓			✓
Observations	21,984	21,984	21,984	21,984	21,984	21,984

Notes: The table reports OLS, reduced-form, and IV estimates for the specifications described in equation (2). An observation is a county in a period from 1900 to 1940. The endogenous variables are county-level surname diversity and population in t . In columns 1 to 3, the dependent variable is log number of patents filed in the county in the five-year period starting in t divided by county population size in 1900 (plus one, to avoid dropping observations with zero patents). In columns 4 to 6, the dependent variable is log number of breakthrough patents filed in the county in the five-year period starting in t divided by county population size in 1900 (plus one). Standard errors are clustered at the state level. Independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

**Table B17: Log-transformed population-normalized innovation outcomes
(county-surname level)**

	Log Patents per 1,000 people (mean = 0.09, sd = 0.52)			Log Breakthrough patents per 1,000 people (mean = 0.01, sd = 0.2)		
	(1)	(2)	(3)	(4)	(5)	(6)
<i>Panel A: Least-squares estimates</i>						
Surname diversity	0.049*** (0.009)	0.051*** (0.009)	0.058*** (0.008)	0.009*** (0.003)	0.009*** (0.003)	0.007*** (0.002)
Population	0.209*** (0.022)	0.208*** (0.023)	0.251*** (0.074)	0.090*** (0.012)	0.091*** (0.012)	0.112*** (0.030)
<i>Panel B: Reduced-form estimates</i>						
Predicted surname diversity	0.026*** (0.009)	0.027*** (0.009)	0.036*** (0.009)	0.003 (0.003)	0.003 (0.003)	0.004 (0.005)
Predicted population	0.180*** (0.042)	0.179*** (0.043)	0.106 (0.077)	0.072*** (0.020)	0.073*** (0.020)	0.039 (0.056)
<i>Panel C: Instrumental-variable estimates</i>						
Surname diversity	0.062*** (0.012)	0.064*** (0.012)	0.076*** (0.013)	0.011*** (0.004)	0.012*** (0.004)	0.010 (0.007)
Population	0.180*** (0.020)	0.178*** (0.020)	0.160 (0.122)	0.072*** (0.011)	0.073*** (0.011)	0.058 (0.086)
Sanderson-Windmeijer F-stat	201; 50	201; 49	215; 300	201; 50	201; 49	215; 300
<i>Panel D: First-stage estimates</i>						
	Surname diversity			Population		
Predicted surname diversity	0.532*** (0.041)	0.532*** (0.041)	0.520*** (0.046)	-0.040** (0.016)	-0.040** (0.016)	-0.021*** (0.004)
Predicted population	0.058* (0.030)	0.060* (0.031)	-0.028 (0.059)	0.981*** (0.142)	0.982*** (0.144)	0.674*** (0.049)
County-Surname fixed effects	✓	✓	✓	✓	✓	✓
State-Period fixed effects	✓	✓	✓	✓	✓	✓
Surname-Period fixed effects		✓	✓		✓	✓
County-specific linear trends			✓			✓
Observations	19,482,024	19,482,024	19,482,024	19,482,024	19,482,024	19,482,024

Notes: The table reports OLS, reduced-form and IV estimates for the specifications described in equation (8). An observation is a surname in a given county in a period from 1900 to 1940. Observations are weighted based on their population shares within counties. The endogenous variables are county-level surname entropy and population in t . In columns 1 to 3, the dependent variable is log number of patents filed by individuals with surname n residing in county i in the five-year period starting in t divided by surname population size in 1900 (plus one, to avoid dropping observations with zero patents). In columns 4 to 6, the dependent variable is log number of breakthrough patents filed by individuals with surname n residing in county i in the five-year period starting in t divided by surname population size in 1900 (plus one). Standard errors are clustered at the state level. Independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B18: Log-transformed, not population normalized innovation outcomes (county level)

	Log Patents (mean = 2.13, sd = 1.56)			Log Breakthrough patents (mean = 0.77, sd = 1.14)		
	(1)	(2)	(3)	(4)	(5)	(6)
<i>Panel A: Least-squares estimates</i>						
Surname diversity	0.640*** (0.034)	0.596*** (0.030)	0.422*** (0.041)	0.254*** (0.037)	0.210*** (0.026)	0.154*** (0.029)
Population	0.340*** (0.039)	0.302*** (0.036)	0.215*** (0.060)	0.244*** (0.041)	0.218*** (0.038)	0.153*** (0.044)
<i>Panel B: Reduced-form estimates</i>						
Predicted surname diversity	0.437*** (0.051)	0.431*** (0.063)	0.338*** (0.044)	0.163*** (0.026)	0.119*** (0.027)	0.110*** (0.025)
Predicted population	0.279*** (0.050)	0.210*** (0.049)	0.015 (0.046)	0.211*** (0.049)	0.192*** (0.049)	0.045 (0.053)
<i>Panel C: Instrumental-variable estimates</i>						
Surname diversity	0.710*** (0.041)	0.686*** (0.031)	0.485*** (0.043)	0.295*** (0.047)	0.235*** (0.032)	0.169*** (0.034)
Population	0.341*** (0.050)	0.277*** (0.039)	0.152** (0.073)	0.255*** (0.055)	0.222*** (0.042)	0.125 (0.082)
Sanderson-Windmeijer F-stat	162; 426	147; 227	149; 29	162; 426	147; 227	149; 29
<i>Panel D: First-stage estimates</i>						
	Surname diversity			Population		
Predicted surname diversity	0.693*** (0.075)	0.723*** (0.073)	0.729*** (0.074)	-0.160*** (0.033)	-0.231*** (0.042)	-0.103** (0.041)
Predicted population	-0.009 (0.041)	-0.076* (0.040)	-0.142** (0.063)	0.837*** (0.053)	0.948*** (0.073)	0.550*** (0.107)
County fixed effects	✓	✓	✓	✓	✓	✓
Period fixed effects	✓			✓		
State-Period fixed effects		✓	✓		✓	✓
County-specific linear trends			✓			✓
Observations	21,984	21,984	21,984	21,984	21,984	21,984

Notes: The table reports OLS, reduced-form and IV estimates for the specifications described in equation (2). An observation is a county in a period from 1900 to 1940. The endogenous variables are county-level surname diversity and population in year t . In columns 1 to 3, the dependent variable is log number of patents filed in county i in the five-year period starting in t (plus one, to avoid dropping observations with zero patents). In columns 4 to 6, the dependent variable is log number of breakthrough patents filed in county i in the five-year period starting in t (plus one). Standard errors are clustered at the state level. Independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B19: Log-transformed, not population normalized innovation outcomes (county-surname level)

	Log Patents per 1,000 people (mean = 0.03, sd = 0.21)			Log Breakthrough patents per 1,000 people (mean = 0.01, sd = 0.09)		
	(1)	(2)	(3)	(4)	(5)	(6)
<i>Panel A: Least-squares estimates</i>						
Surname diversity	0.013*** (0.003)	0.013*** (0.003)	0.013*** (0.003)	0.002 (0.001)	0.002 (0.001)	0.001 (0.001)
Population	0.105*** (0.015)	0.105*** (0.015)	0.175*** (0.050)	0.049*** (0.006)	0.050*** (0.006)	0.100*** (0.032)
<i>Panel B: Reduced-form estimates</i>						
Predicted surname diversity	0.004 (0.003)	0.004 (0.003)	0.006 (0.005)	0.000 (0.001)	-0.001 (0.001)	-0.001 (0.003)
Predicted population	0.093*** (0.021)	0.093*** (0.022)	0.091* (0.050)	0.041*** (0.009)	0.041*** (0.009)	0.057 (0.042)
<i>Panel C: Instrumental-variable estimates</i>						
Surname diversity	0.015*** (0.004)	0.015*** (0.004)	0.018** (0.007)	0.002 (0.002)	0.002 (0.002)	0.001 (0.004)
Population	0.094*** (0.013)	0.094*** (0.013)	0.136* (0.081)	0.041*** (0.005)	0.042*** (0.005)	0.085 (0.067)
Sanderson-Windmeijer F-stat	201; 50	201; 49	215; 300	201; 50	201; 49	215; 300
<i>Panel D: First-stage estimates</i>						
	Surname diversity			Population		
Predicted surname diversity	0.532*** (0.041)	0.532*** (0.041)	0.520*** (0.046)	-0.040** (0.016)	-0.040** (0.016)	-0.021*** (0.004)
Predicted population	0.058* (0.030)	0.060* (0.031)	-0.028 (0.059)	0.981*** (0.142)	0.982*** (0.144)	0.674*** (0.049)
County-Surname fixed effects	✓	✓	✓	✓	✓	✓
State-Period fixed effects	✓	✓	✓	✓	✓	✓
Surname-Period fixed effects		✓	✓		✓	✓
County-specific linear trends			✓			✓
Observations	19,482,024	19,482,024	19,482,024	19,482,024	19,482,024	19,482,024

Notes: The table reports OLS, reduced-form and IV estimates for the specifications described in equation (8). An observation is a surname in a given county in a period from 1900 to 1940. Observations are weighted based on their population shares within counties. The endogenous variables are county-level surname entropy and population in t . In columns 1 to 3, the dependent variable is log number of patents filed by individuals with surname n residing in county i in the five-year period starting in t (plus one, to avoid dropping observations with zero patents). In columns 4 to 6, the dependent variable is log number of breakthrough patents filed by individuals with surname n residing in county i in the five-year period starting in t (plus one). Standard errors are clustered at the state level. Independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B20: Log-transformed, not population normalized innovation outcomes (inventor level)

	Patents (mean = 1.11, sd = 0.58)		Breakthrough patents (mean = 0.25, sd = 0.49)	
	(1)	(2)	(3)	(4)
<i>Panel A: Least-squares estimates</i>				
Surname diversity	0.081*** (0.022)	0.081*** (0.021)	0.028 (0.040)	0.026 (0.040)
Population		-0.005 (0.017)		-0.026 (0.016)
<i>Panel B: Reduced-form estimates</i>				
Predicted surname diversity	0.028*** (0.007)	0.028*** (0.007)	0.020* (0.011)	0.017 (0.011)
Predicted population		0.000 (0.015)		-0.022* (0.012)
<i>Panel C: Instrumental-variable estimates</i>				
Surname diversity	0.170** (0.065)	0.160** (0.062)	0.118 (0.088)	0.098 (0.081)
Population		-0.017 (0.018)		-0.036** (0.015)
Sanderson-Windmeijer F-stat	12	17; 504	12	17; 504
<i>Panel D: First-stage estimates</i>				
	Surname diversity		Population	
Predicted surname diversity	0.167*** (0.048)	0.179*** (0.049)		0.020 (0.048)
Predicted population		0.095*** (0.019)		0.867*** (0.101)
Inventor fixed effects	✓	✓	✓	✓
Period-State fixed effects	✓	✓	✓	✓
Observations	209,944	209,944	209,944	209,944

Notes: The table reports OLS, reduced-form and IV estimates for the specifications described in equation (9). An observation is an inventor residing in a given county in a period from 1900 to 1940. The endogenous variables are county-level surname diversity and population in year t , and the dependent variables are log one plus number of patents filed by the inventor in the five-year period starting in t (column 1–2) and log one plus number of breakthrough patents filed by the inventor in the five-year period starting in t (column 3–4). Standard errors are clustered at the state level. The independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B21: Baseline vs. distance-weighted surname diversity (county level)

	Patents per 1,000 people (mean = 1.05, sd = 1.73)					Breakthrough patents per 1,000 people (mean = 0.13, sd = 0.28)				
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)
<i>Panel A: Least-squares estimates</i>										
Genetic distance weighted surname diversity	1.123*** (0.109)	0.917*** (0.134)	0.866*** (0.175)	1.097*** (0.126)	-0.205 (0.155)	0.139*** (0.022)	0.099*** (0.014)	0.090*** (0.013)	0.093*** (0.018)	0.034 (0.021)
Surname diversity					1.522*** (0.150)					0.068*** (0.024)
Population		0.706*** (0.100)	0.681*** (0.106)	0.802*** (0.201)	0.743*** (0.190)		0.140*** (0.020)	0.127*** (0.018)	0.101*** (0.030)	0.098*** (0.029)
<i>Panel B: Reduced-form estimates</i>										
Predicted genetic distance weighted surname diversity	0.867*** (0.161)	0.609*** (0.206)	0.529** (0.244)	0.659*** (0.193)	-0.146 (0.135)	0.104*** (0.013)	0.044** (0.018)	0.037* (0.021)	0.043** (0.016)	-0.009 (0.020)
Predicted surname diversity					1.081*** (0.214)					0.070*** (0.023)
Predicted population		0.602*** (0.132)	0.614*** (0.145)	0.441** (0.197)	0.211 (0.205)		0.139*** (0.024)	0.126*** (0.024)	0.066* (0.034)	0.051 (0.035)
<i>Panel C: Instrumental-variable estimates</i>										
Genetic distance weighted surname diversity	1.324*** (0.147)	1.108*** (0.172)	1.029*** (0.243)	1.190*** (0.173)	-0.167 (0.281)	0.159*** (0.018)	0.106*** (0.016)	0.100*** (0.019)	0.088*** (0.019)	-0.005 (0.038)
Surname diversity					1.543*** (0.244)					0.106*** (0.036)
Population		0.677*** (0.119)	0.697*** (0.110)	0.985*** (0.239)	0.743*** (0.241)		0.164*** (0.026)	0.140*** (0.021)	0.135** (0.053)	0.119** (0.054)
Sanderson-Windmeijer F-stat	106	176; 450	167; 342	145; 36	246; 58; 54	106	176; 450	167; 342	145; 36	246; 58; 54
<i>Panel D: First-stage estimates</i>										
	Genetic distance weighted surname diversity					Population				
Predicted genetic distance weighted surname diversity	0.655*** (0.064)	0.639*** (0.087)	0.647*** (0.075)	0.623*** (0.068)	0.494*** (0.091)	-0.146*** (0.029)	-0.197*** (0.033)	-0.084** (0.031)	-0.025* (0.015)	
Predicted surname diversity					0.174*** (0.058)					-0.078** (0.037)
Predicted population		0.038 (0.058)	-0.029 (0.050)	-0.072 (0.062)	-0.109* (0.062)		0.827*** (0.053)	0.923*** (0.069)	0.535*** (0.098)	0.551*** (0.106)
	Surname diversity									
Predicted genetic distance weighted surname diversity					-0.029 (0.019)					
Predicted surname diversity					0.757*** (0.087)					
Predicted population					-0.141** (0.063)					
County fixed effects	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓
Period fixed effects	✓	✓				✓	✓			
State-Period fixed effects			✓	✓	✓			✓	✓	✓
County-specific linear trends				✓	✓				✓	✓
Observations	21,984	21,984	21,984	21,984	21,984	21,984	21,984	21,984	21,984	21,984

Notes: The table reports OLS, reduced-form, and IV estimates for the specifications described in equation (2). An observation is a county in a period from 1900 to 1940. The endogenous variables are county-level genetic distance-weighted surname diversity, baseline surname diversity and population in t . Standard errors are clustered at the state level. Independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B22: Baseline vs. distance-weighted surname diversity (county-surname level)

	Patents per 1,000 people (mean = 1.02, sd = 13.31)					Breakthrough patents per 1,000 people (mean = 0.1, sd = 2.23)				
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)
<i>Panel A: Least-squares estimates</i>										
Genetic distance weighted surname diversity	0.952*** (0.174)	0.707*** (0.136)	0.718*** (0.139)	0.891*** (0.145)	-0.112 (0.152)	0.113*** (0.030)	0.073*** (0.023)	0.073*** (0.023)	0.054*** (0.015)	0.020 (0.027)
Population		3.723*** (0.597)	3.742*** (0.609)	3.749*** (0.626)	3.685*** (0.597)		0.617*** (0.107)	0.624*** (0.108)	0.729*** (0.099)	0.727*** (0.099)
Surname diversity					1.246*** (0.254)					0.043 (0.027)
<i>Panel B: Reduced-form estimates</i>										
Predicted genetic distance weighted surname diversity	0.615*** (0.132)	0.429*** (0.149)	0.435*** (0.152)	0.517*** (0.115)	-0.026 (0.129)	0.065*** (0.018)	0.037 (0.022)	0.036 (0.022)	0.033* (0.018)	-0.004 (0.032)
Predicted surname diversity					0.700*** (0.215)					0.048 (0.058)
Predicted population		3.140*** (0.947)	3.157*** (0.971)	1.162 (0.927)	0.960 (0.971)		0.479*** (0.156)	0.485*** (0.159)	0.126 (0.263)	0.112 (0.278)
<i>Panel C: Instrumental-variable estimates</i>										
Genetic distance weighted surname diversity	1.177*** (0.210)	1.052*** (0.200)	1.066*** (0.204)	1.126*** (0.188)	-0.119 (0.302)	0.124*** (0.033)	0.105*** (0.030)	0.105*** (0.030)	0.075*** (0.025)	-0.013 (0.076)
Surname diversity					1.471*** (0.420)					0.103 (0.110)
Population		3.077*** (0.564)	3.090*** (0.580)	1.728 (1.465)	1.488 (1.509)		0.476*** (0.102)	0.483*** (0.104)	0.188 (0.403)	0.171 (0.420)
Sanderson-Windmeijer F-stat	215	284; 43	275; 40	245; 358	262; 64; 439	215	284; 43	275; 40	245; 358	262; 64; 439
<i>Panel D: First-stage estimates</i>										
	Genetic distance weighted surname diversity					Population				
Predicted genetic distance weighted surname diversity	0.522*** (0.040)	0.515*** (0.041)	0.515*** (0.041)	0.485*** (0.038)	0.463*** (0.063)	-0.037** (0.014)	-0.037** (0.014)	-0.017*** (0.004)		-0.003 (0.005)
Predicted surname diversity					0.029 (0.059)					-0.018*** (0.005)
Predicted population		0.124*** (0.036)	0.124*** (0.037)	0.006 (0.064)	-0.003 (0.058)		0.978*** (0.139)	0.979*** (0.142)	0.669*** (0.050)	0.674*** (0.050)
	Surname diversity									
Predicted genetic distance weighted surname diversity					0.023 (0.020)					
Predicted surname diversity					0.497*** (0.059)					
Predicted population					-0.030 (0.059)					
County-Surname fixed effects	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓
State-Period fixed effects	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓
Surname-Period fixed effects			✓	✓	✓			✓	✓	✓
County-specific linear trends				✓	✓				✓	✓
Observations	19,482,024	19,482,024	19,482,024	19,482,024	19,482,024	19,482,024	19,482,024	19,482,024	19,482,024	19,482,024

Notes: The table reports OLS, reduced-form, and IV estimates for the specifications described in equation (8). An observation is a surname in a given county in a period from 1900 to 1940. Observations are weighted based on their population shares within counties. The endogenous variables are county-level genetic distance-weighted surname diversity, baseline surname diversity and population in t . Standard errors are clustered at the state level. Independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B23: Baseline vs. distance-weighted surname diversity (inventor level)

	Patents (mean = 2.24, sd = 1.69)			Breakthrough patents (mean = 0.26, sd = 0.44)		
	(1)	(2)	(3)	(4)	(5)	(6)
<i>Panel A: Least-squares estimates</i>						
Genetic distance weighted surname diversity	0.128** (0.050)	0.131** (0.052)	-0.225** (0.087)	0.027 (0.026)	0.028 (0.026)	-0.032 (0.036)
Surname diversity			0.449** (0.104)			0.076** (0.030)
Population		-0.025 (0.055)	-0.003 (0.055)		-0.013 (0.016)	-0.009 (0.016)
<i>Panel B: Reduced-form estimates</i>						
Predicted genetic distance weighted surname diversity	0.056** (0.023)	0.056** (0.023)	-0.098** (0.047)	0.015* (0.009)	0.014* (0.008)	-0.007 (0.013)
Predicted surname diversity			0.170*** (0.048)			0.023* (0.012)
Predicted population		-0.008 (0.038)	0.006 (0.041)		-0.009 (0.010)	-0.007 (0.011)
<i>Panel C: Instrumental-variable estimates</i>						
Genetic distance weighted surname diversity	0.290** (0.141)	0.265** (0.132)	-0.674** (0.259)	0.076 (0.052)	0.066 (0.046)	-0.084 (0.066)
Surname diversity			1.116*** (0.368)			0.178* (0.099)
Population		-0.054 (0.039)	-0.006 (0.053)		-0.022* (0.013)	-0.014 (0.014)
Sanderson-Windmeijer <i>F</i> -stat	18	21; 481	176; 166; 508	18	21; 481	176; 166; 508
<i>Panel D: First-stage estimates</i>						
	Genetic distance weighted surname diversity			Population		
Predicted genetic distance weighted surname diversity	0.194*** (0.046)	0.210*** (0.045)	0.302*** (0.041)	0.003 (0.037)	-0.097* (0.051)	
Predicted surname diversity			-0.102 (0.063)		0.110 (0.094)	
Predicted population		0.145*** (0.030)	0.137*** (0.028)	0.861*** (0.094)	0.870*** (0.100)	
	Surname diversity					
Predicted genetic distance weighted surname diversity			0.094*** (0.031)			
Predicted surname diversity			0.091 (0.055)			
Predicted population			0.092*** (0.020)			
Inventor fixed effects	✓	✓	✓	✓	✓	✓
Period-State fixed effects	✓	✓	✓	✓	✓	✓
Observations	209,944	209,944	209,944	209,944	209,944	209,944

Notes: The table reports OLS, reduced-form and IV estimates for the specifications described in equation (9). An observation is an inventor residing in a given county in a period from 1900 to 1940. The endogenous variables are county-level genetic distance-weighted surname diversity, baseline surname diversity and population in *t*. Standard errors are clustered at the state level. Independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B24: Baseline surname diversity and its components (county level)

	Patents per 1,000 people (mean = 1.05, sd = 1.73)			Breakthrough patents per 1,000 people (mean = 0.13, sd = 0.28)		
	(1)	(2)	(3)	(4)	(5)	(6)
Surname diversity			1.025*** (0.222)			0.161*** (0.031)
Log distinct surnames	0.995*** (0.122)	1.218*** (0.119)	0.354*** (0.114)	0.051** (0.022)	0.065* (0.036)	-0.071** (0.030)
Surname dispersion		0.335*** (0.091)	0.052 (0.063)		0.020 (0.024)	-0.025 (0.016)
Population	0.747*** (0.187)	0.765*** (0.195)	0.708*** (0.186)	0.109*** (0.031)	0.110*** (0.031)	0.101*** (0.029)
County fixed effects	✓	✓	✓	✓	✓	✓
State-Period fixed effects	✓	✓	✓	✓	✓	✓
County-specific linear trends	✓	✓	✓	✓	✓	✓
Observations	21,984	21,983	21,983	21,984	21,983	21,983

Notes: The table reports OLS estimates for the specifications described in equation (2). An observation is a county in a period from 1900 to 1940. The independent variables are county-level entropic surname diversity (variety and dispersion), log number of distinct surnames (variety), surname dispersion and population in t . Standard errors are clustered at the state level. Independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B25: Baseline surname diversity and its components (county-surname level)

	Patents per 1,000 people (mean = 1.02, sd = 13.31)			Breakthrough patents per 1,000 people (mean = 0.1, sd = 2.23)		
	(1)	(2)	(3)	(4)	(5)	(6)
Surname diversity			0.524 (0.372)			0.071*** (0.019)
Log distinct surnames	1.021*** (0.204)	1.310*** (0.234)	0.768 (0.490)	0.054*** (0.014)	0.055*** (0.017)	-0.019 (0.017)
Surname dispersion		0.276*** (0.088)	0.102 (0.140)		0.001 (0.010)	-0.023** (0.009)
Population	3.624*** (0.572)	3.617*** (0.581)	3.593*** (0.577)	0.725*** (0.099)	0.725*** (0.099)	0.722*** (0.098)
County-Surname fixed effects	✓	✓	✓	✓	✓	✓
Surname-Period fixed effects	✓	✓	✓	✓	✓	✓
State-Period fixed effects	✓	✓	✓	✓	✓	✓
County-specific linear trends	✓	✓	✓	✓	✓	✓
Observations	19,482,024	19,482,024	19,482,024	19,482,024	19,482,024	19,482,024

Notes: The table reports OLS estimates for the specifications described in equation (8). An observation is a surname in a given county in a period from 1900 to 1940. Observations are weighted based on their population shares within counties. The independent variables are county-level entropic surname diversity (variety and dispersion), log number of distinct surnames (variety), surname dispersion and population in t . Standard errors are clustered at the state level. Independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B26: Baseline surname diversity and its components (inventor level)

	Patents (mean = 2.24, sd = 1.69)			Breakthrough patents (mean = 0.26, sd = 0.44)		
	(1)	(2)	(3)	(4)	(5)	(6)
Surname diversity			0.290* (0.161)			0.012 (0.063)
Log distinct surnames	0.229*** (0.058)	0.247** (0.096)	-0.166 (0.217)	0.032 (0.022)	0.071** (0.033)	0.053 (0.076)
Surname dispersion		0.017 (0.090)	-0.155 (0.104)		0.036** (0.017)	0.029 (0.037)
Population	-0.053 (0.044)	-0.051 (0.041)	-0.030 (0.047)	-0.017 (0.016)	-0.013 (0.016)	-0.012 (0.015)
Inventor fixed effects	✓	✓	✓	✓	✓	✓
Period-State fixed effects	✓	✓	✓	✓	✓	✓
Observations	209,944	209,944	209,944	209,944	209,944	209,944

Notes: The table reports OLS estimates for the specifications described in equation (9). An observation is an inventor residing in a given county in a period from 1900 to 1940. The independent variables are county-level entropic surname diversity (variety and dispersion), log number of distinct surnames (variety), surname dispersion and population in t . Standard errors are clustered at the state level. Independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B27: Estimates of the effect of surname fractionalization on innovation (county level)

	Patents per 1,000 people (mean = 1.05, sd = 1.73)				Breakthrough patents per 1,000 people (mean = 0.13, sd = 0.28)			
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
<i>Panel A: Least-squares estimates</i>								
Surname fractionalization index	0.798*** (0.145)	0.745*** (0.163)	0.702*** (0.203)	0.764*** (0.114)	0.065*** (0.010)	0.055*** (0.011)	0.050*** (0.012)	0.039** (0.015)
Population		0.878*** (0.106)	0.811*** (0.111)	1.028*** (0.226)		0.160*** (0.020)	0.141*** (0.018)	0.124*** (0.029)
<i>Panel B: Reduced-form estimates</i>								
Predicted surname fractionalization index	0.563*** (0.135)	0.464*** (0.160)	0.481** (0.184)	0.546*** (0.138)	0.049*** (0.008)	0.030** (0.012)	0.033*** (0.012)	0.035*** (0.011)
Predicted population		0.820*** (0.116)	0.757*** (0.120)	0.767*** (0.174)		0.154*** (0.022)	0.134*** (0.021)	0.088*** (0.029)
<i>Panel C: Instrumental-variable estimates</i>								
Surname fractionalization index	1.050*** (0.108)	0.984*** (0.133)	0.989*** (0.173)	1.025*** (0.104)	0.091*** (0.015)	0.078*** (0.012)	0.082*** (0.011)	0.070*** (0.023)
Population		1.039*** (0.127)	0.957*** (0.122)	1.708*** (0.309)		0.201*** (0.024)	0.170*** (0.019)	0.199*** (0.052)
Sanderson-Windmeijer F-stat	33	32; 268	38; 200	37; 49	33	32; 268	38; 200	37; 49
<i>Panel D: First-stage estimates</i>								
	Surname fractionalization index				Population			
Predicted surname fractionalization index	0.536*** (0.094)	0.531*** (0.096)	0.556*** (0.094)	0.570*** (0.098)	-0.056*** (0.012)	-0.071*** (0.013)	-0.023** (0.009)	
Predicted population		0.043** (0.021)	0.004 (0.023)	0.020 (0.032)	0.749*** (0.058)	0.787*** (0.067)	0.437*** (0.072)	
County fixed effects	✓	✓	✓	✓	✓	✓	✓	✓
Period fixed effects	✓	✓			✓	✓		
State-Period fixed effects			✓	✓			✓	✓
County-specific linear trends				✓				✓
Observations	21,984	21,984	21,984	21,984	21,984	21,984	21,984	21,984

Notes: The table reports OLS, reduced-form and IV estimates for the specifications described in equation (2). An observation is a county in a period from 1900 to 1940. The endogenous variables are county-level surname fractionalization index, defined as $1 - \sum p$, and population in t . In columns 1 to 4, the dependent variable is number of patents filed in the county in the five-year period starting in t divided by county population size in 1900. In columns 5 to 8, it is number of breakthrough patents filed in the county in the five-year period starting in t divided by county population size in 1900. Standard errors are clustered at the state level. Independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B28: Estimates of the effect of surname fractionalization on innovation
(county-surname level)

	Patents per 1,000 people (mean = 1.02, sd = 13.31)				Breakthrough patents per 1,000 people (mean = 0.1, sd = 2.23)			
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
<i>Panel A: Least-squares estimates</i>								
Surname fractionalization index	0.296*** (0.058)	0.279*** (0.062)	0.285*** (0.064)	0.450*** (0.103)	0.017** (0.008)	0.014* (0.008)	0.014* (0.008)	0.010* (0.005)
Population		3.882*** (0.564)	3.902*** (0.575)	4.012*** (0.654)		0.634*** (0.102)	0.641*** (0.103)	0.749*** (0.104)
<i>Panel B: Reduced-form estimates</i>								
Predicted surname fractionalization index	0.220*** (0.058)	0.110 (0.067)	0.113 (0.069)	0.299*** (0.091)	0.016 (0.010)	-0.002 (0.011)	-0.001 (0.012)	0.016 (0.011)
Predicted population		3.799*** (0.746)	3.833*** (0.764)	2.811*** (0.596)		0.613*** (0.130)	0.621*** (0.133)	0.425*** (0.091)
<i>Panel C: Instrumental-variable estimates</i>								
Surname fractionalization index	0.534*** (0.125)	0.712*** (0.153)	0.722*** (0.159)	0.860*** (0.184)	0.039 (0.026)	0.068* (0.036)	0.070* (0.038)	0.064** (0.031)
Population		4.681*** (0.496)	4.719*** (0.503)	6.067*** (1.498)		0.765*** (0.104)	0.775*** (0.106)	0.955*** (0.233)
Sanderson-Windmeijer F-stat	64	65; 49	62; 47	61; 43	64	65; 49	62; 47	61; 43
<i>Panel D: First-stage estimates</i>								
	Surname fractionalization index				Population			
Predicted surname fractionalization index	0.411*** (0.052)	0.407*** (0.052)	0.409*** (0.053)	0.434*** (0.060)		-0.038*** (0.012)	-0.039*** (0.013)	-0.012*** (0.004)
Predicted population		0.162** (0.063)	0.167** (0.063)	0.241*** (0.083)		0.787*** (0.113)	0.787*** (0.115)	0.429*** (0.102)
County-Surname fixed effects	✓	✓	✓	✓	✓	✓	✓	✓
State-Period fixed effects	✓	✓	✓	✓	✓	✓	✓	✓
Surname-Period fixed effects			✓	✓			✓	✓
County-specific linear trends				✓				✓
Observations	19,482,024	19,482,024	19,482,024	19,482,024	19,482,024	19,482,024	19,482,024	19,482,024

Notes: The table reports OLS, reduced-form and IV estimates for the specifications described in equation (8). An observation is a surname in a given county in a period from 1900 to 1940. Observations are weighted based on their population shares within counties. The endogenous variables are county-level surname fractionalization index, defined as $1 - \sum p$, and population in t . In columns 1 to 3, the dependent variable is number of patents filed by individuals with surname n residing in county i in the five-year period starting in t divided by surname population size in 1900. In columns 4 to 6, the dependent variable is number of breakthrough patents filed by individuals with surname n residing in county i in the five-year period starting in t divided by surname population size in 1900. Standard errors are clustered at the state level. Independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B29: Estimates of the effect of surname fractionalization on innovation (inventor level)

	Patents (mean = 2.24, sd = 1.69)		Breakthrough patents (mean = 0.26, sd = 0.44)	
	(1)	(2)	(3)	(4)
<i>Panel A: Least-squares estimates</i>				
Surname fractionalization index	0.124** (0.057)	0.123** (0.053)	0.041*** (0.014)	0.040*** (0.013)
Population		-0.011 (0.056)		-0.009 (0.016)
<i>Panel B: Reduced-form estimates</i>				
Predicted surname fractionalization index	0.030 (0.035)	0.034 (0.038)	0.019** (0.008)	0.020** (0.008)
Predicted population		0.042 (0.053)		0.010 (0.011)
<i>Panel C: Instrumental-variable estimates</i>				
Surname fractionalization index	0.416 (0.499)	0.491 (0.453)	0.261** (0.098)	0.249*** (0.078)
Population		0.019 (0.061)		-0.003 (0.014)
Sanderson-Windmeijer F-stat	17	32; 84	17	32; 84
<i>Panel D: First-stage estimates</i>				
	Surname fractionalization index		Population	
Predicted surname fractionalization index	0.072*** (0.017)	0.077*** (0.018)		-0.190*** (0.070)
Predicted population		0.052* (0.029)		0.870*** (0.106)
Inventor fixed effects	✓	✓	✓	✓
Period-State fixed effects	✓	✓	✓	✓
Observations	209,944	209,944	209,944	209,944

Notes: The table reports OLS, reduced-form and IV estimates for the specifications described in equation (9). An observation is an inventor residing in a given county in a period from 1900 to 1940. The endogenous variables are county-level surname fractionalization index, defined as $1 - \sum p$, and population in t . In columns 1 to 4, the dependent variable is number of patents filed in the county in the five-year period starting in t divided by county population size in 1900. In columns 5 to 8, it is number of breakthrough patents filed in the county in the five-year period starting in t divided by county population size in 1900. Standard errors are clustered at the state level. Independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B30: Excluding controls from zero-stage specification (county level)

	Patents per 1,000 people (mean = 1.05, sd = 1.73)				Breakthrough patents per 1,000 people (mean = 0.13, sd = 0.28)			
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
<i>Panel A: Reduced-form estimates</i>								
Predicted surname diversity	0.945*** (0.178)	0.732*** (0.214)	0.689*** (0.249)	0.877*** (0.204)	0.108*** (0.014)	0.059*** (0.018)	0.057*** (0.020)	0.055*** (0.015)
Predicted population		0.714*** (0.153)	0.698*** (0.164)	0.472*** (0.161)		0.162*** (0.027)	0.142*** (0.027)	0.096*** (0.030)
<i>Panel B: Instrumental-variable estimates</i>								
Surname diversity	1.393*** (0.171)	1.146*** (0.211)	1.098*** (0.269)	1.347*** (0.189)	0.159*** (0.017)	0.101*** (0.018)	0.099*** (0.021)	0.087*** (0.019)
Population		0.741*** (0.145)	0.732*** (0.150)	0.938*** (0.248)		0.172*** (0.026)	0.147*** (0.025)	0.180*** (0.053)
Sanderson-Windmeijer <i>F</i> -stat	130	127; 355	118; 323	148; 103	130	127; 355	118; 323	148; 103
<i>Panel C: First-stage estimates</i>								
	Surname diversity				Population			
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
Predicted surname diversity	0.678*** (0.059)	0.672*** (0.067)	0.669*** (0.066)	0.661*** (0.068)		-0.052** (0.021)	-0.062*** (0.021)	-0.015 (0.020)
Predicted population		0.020 (0.032)	-0.009 (0.029)	-0.029 (0.045)		0.932*** (0.053)	0.967*** (0.059)	0.545*** (0.055)
County fixed effects	✓	✓	✓	✓	✓	✓	✓	✓
Period fixed effects	✓	✓			✓	✓		
State-Period fixed effects			✓	✓			✓	✓
County-specific linear trends				✓				✓
Observations	21,984	21,984	21,984	21,984	21,984	21,984	21,984	21,984

Notes: The table reports reduced-form and IV estimates for the specifications described in equation (2). An observation is a county in a period from 1900 to 1940. The endogenous variables are county-level surname diversity and population in t . The instruments are computed from estimates of a version of equation (3) that includes only the push-pull instruments, not the fixed effects and controls. In columns 1 to 4, the dependent variable is number of patents filed in the county in the five-year period starting in t divided by county population size in 1900. In columns 5 to 8, it is number of breakthrough patents filed in the county in the five-year period starting in t divided by county population size in 1900. Standard errors are clustered at the state level. Independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B31: Excluding controls from zero-stage specification (county-surname level)

	Patents per 1,000 people (mean = 1.02, sd = 13.31)				Breakthrough patents per 1,000 people (mean = 0.1, sd = 2.23)			
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
<i>Panel A: Reduced-form estimates</i>								
Predicted surname diversity	0.714*** (0.157)	0.641*** (0.162)	0.651*** (0.166)	0.695*** (0.135)	0.075*** (0.019)	0.064*** (0.020)	0.065*** (0.020)	0.050*** (0.019)
Predicted population		2.594*** (0.721)	2.602*** (0.742)	0.223 (1.522)		0.385*** (0.114)	0.390*** (0.116)	-0.020 (0.381)
<i>Panel B: Instrumental-variable estimates</i>								
Surname diversity	1.367*** (0.249)	1.221*** (0.249)	1.243*** (0.255)	1.384*** (0.244)	0.143*** (0.032)	0.121*** (0.033)	0.122*** (0.034)	0.100** (0.042)
Population		2.761*** (0.523)	2.765*** (0.543)	0.252 (2.837)		0.416*** (0.085)	0.421*** (0.087)	-0.048 (0.715)
Sanderson-Windmeijer <i>F</i> -stat	170	286; 76	288; 74	162; 216	170	286; 76	288; 74	162; 216
<i>Panel C: First-stage estimates</i>								
	Surname diversity				Population			
Predicted surname diversity	0.522*** (0.040)	0.520*** (0.040)	0.519*** (0.040)	0.501*** (0.046)		0.002 (0.008)	0.002 (0.008)	0.003 (0.003)
Predicted population		0.086*** (0.015)	0.087*** (0.015)	0.063* (0.036)		0.902*** (0.105)	0.902*** (0.106)	0.539*** (0.041)
County-Surname fixed effects	✓	✓	✓	✓	✓	✓	✓	✓
State-Period fixed effects	✓	✓	✓	✓	✓	✓	✓	✓
Surname-Period fixed effects			✓	✓			✓	✓
County-specific linear trends				✓				✓
Observations	19,482,024	19,482,024	19,482,024	19,482,024	19,482,024	19,482,024	19,482,024	19,482,024

Notes: The table reports reduced-form and IV estimates for the specifications described in equation (8). An observation is a surname in a given county in a period from 1900 to 1940. Observations are weighted based on their population shares within counties. The endogenous variables are county-level surname diversity and population in t . The instruments are computed from estimates of a version of equation (3) that includes only the push-pull instruments, not the fixed effects and controls. In columns 1 to 4, the dependent variable is number of patents filed by individuals with surname n residing in county i in the five-year period starting in t divided by surname population size in 1900. In columns 5 to 8, the dependent variable is number of breakthrough patents filed by individuals with surname n residing in county i in the five-year period starting in t divided by surname population size in 1900. Standard errors are clustered at the state level. Independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B32: Excluding controls from zero-stage specification (inventor level)

	Patents (mean = 2.24, sd = 1.69)		Breakthrough patents (mean = 0.26, sd = 0.44)	
	(1)	(2)	(3)	(4)
<i>Panel A: Reduced-form estimates</i>				
Predicted surname diversity	0.086*** (0.027)	0.080*** (0.027)	0.023** (0.011)	0.021** (0.010)
Predicted population		-0.023 (0.038)		-0.009 (0.008)
<i>Panel B: Instrumental-variable estimates</i>				
Surname diversity	0.313** (0.124)	0.285** (0.114)	0.085* (0.047)	0.076* (0.043)
Population		-0.067 (0.045)		-0.022* (0.011)
Sanderson-Windmeijer F-stat	26	40; 2224	26	40; 2224
<i>Panel C: First-stage estimates</i>				
	Surname diversity		Population	
Predicted surname diversity	0.276*** (0.054)	0.302*** (0.052)		0.080 (0.055)
Predicted population		0.101*** (0.030)		0.768*** (0.043)
Inventor fixed effects	✓	✓	✓	✓
Period-State fixed effects	✓	✓	✓	✓
Observations	209,944	209,944	209,944	209,944

Notes: The table reports reduced-form and IV estimates for the specifications described in equation (9). An observation is an inventor residing in a given county in a period from 1900 to 1940. The endogenous variables are county-level surname diversity and population in year t . The instruments are computed from estimates of a version of equation (3) that includes only the push-pull instruments, not the fixed effects and controls. The dependent variables are number of patents filed by the inventor in the five-year period starting in t (column 1–2) and number of breakthrough patents filed by the inventor in the five-year period starting in t (column 3–4). Standard errors are clustered at the state level. The independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B33: Dropping counties with highest surname diversity in 1900 (county level)

	Patents per 1,000 people			Breakthrough patents per 1,000 people		
	(1)	(2)	(3)	(4)	(5)	(6)
<i>Panel A: Least-squares estimates</i>						
Surname diversity	1.360*** (0.173)	1.367*** (0.181)	1.370*** (0.189)	0.105*** (0.021)	0.102*** (0.022)	0.100*** (0.024)
Population	0.493 (0.341)	0.322 (0.453)	0.239 (0.604)	0.007 (0.067)	0.007 (0.089)	-0.003 (0.126)
<i>Panel B: Reduced-form estimates</i>						
Predicted surname diversity	1.163*** (0.239)	1.245*** (0.242)	1.234*** (0.256)	0.088*** (0.020)	0.096*** (0.018)	0.090*** (0.020)
Predicted population	-0.410* (0.224)	-0.700*** (0.255)	-0.717** (0.292)	-0.028 (0.039)	-0.060 (0.037)	-0.056 (0.041)
<i>Panel C: Instrumental-variable estimates</i>						
Surname diversity	1.454*** (0.199)	1.489*** (0.196)	1.466*** (0.208)	0.110*** (0.023)	0.115*** (0.024)	0.107*** (0.025)
Population	0.170 (0.423)	-0.313 (0.535)	-0.241 (0.719)	0.024 (0.109)	-0.049 (0.133)	-0.035 (0.163)
Sanderson-Windmeijer <i>F</i> -stat	252; 13	291; 8	243; 5	252; 13	291; 8	243; 5
<i>Panel D: First-stage estimates</i>						
	Surname diversity			Population		
Predicted surname diversity	0.803*** (0.076)	0.836*** (0.082)	0.844*** (0.087)	-0.023 (0.031)	0.001 (0.032)	0.010 (0.033)
Predicted population	-0.320*** (0.103)	-0.417*** (0.145)	-0.453** (0.173)	0.327*** (0.097)	0.254** (0.103)	0.217* (0.109)
County fixed effects	✓	✓	✓	✓	✓	✓
State-Period fixed effects	✓	✓	✓	✓	✓	✓
County-specific linear trends	✓	✓	✓	✓	✓	✓
Observations	20,880	19,784	17,584	20,880	19,784	17,584
Dropping top % of counties	5%	10%	20%	5%	10%	20%

Notes: The table reports OLS, reduced-form and IV estimates for the specifications described in equation (2). An observation is a county in a period from 1900 to 1940. The endogenous variables are county-level surname diversity and population in t . In columns 1 to 4, the dependent variable is number of patents filed in the county in the five-year period starting in t divided by county population size in 1900. In columns 5 to 8, it is number of breakthrough patents filed in the county in the five-year period starting in t divided by county population size in 1900. Column 1 and 5 drop the 5% of counties with the highest surname diversity in 1900; Column 2 and 6 drop the top 10%; and Column 3 and 7 drop the top 20%. Standard errors are clustered at the state level. All independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B34: Dropping counties with highest surname diversity in 1900 (county-surname level)

	Patents per 1,000 people			Breakthrough patents per 1,000 people		
	(1)	(2)	(3)	(4)	(5)	(6)
<i>Panel A: Least-squares estimates</i>						
Surname diversity	1.044*** (0.207)	1.143*** (0.212)	1.127*** (0.233)	0.041*** (0.014)	0.054*** (0.012)	0.046*** (0.013)
Population	7.630* (4.434)	3.197 (5.958)	2.889 (7.967)	0.932** (0.347)	0.304 (0.342)	0.568 (0.449)
<i>Panel B: Reduced-form estimates</i>						
Predicted surname diversity	0.695*** (0.176)	0.884*** (0.188)	0.853*** (0.197)	0.029** (0.012)	0.046*** (0.012)	0.036*** (0.011)
Predicted population	0.279 (2.186)	-3.201 (2.580)	-2.994 (2.831)	0.190 (0.155)	-0.147 (0.158)	-0.040 (0.169)
<i>Panel C: Instrumental-variable estimates</i>						
Surname diversity	1.170*** (0.260)	1.369*** (0.297)	1.312*** (0.344)	0.051*** (0.017)	0.071*** (0.019)	0.054*** (0.019)
Population	6.159 (6.060)	-1.808 (9.297)	-0.081 (12.278)	0.868** (0.419)	-0.013 (0.511)	0.427 (0.631)
Sanderson-Windmeijer <i>F</i> -stat	186; 10	164; 6	108; 4	186; 10	164; 6	108; 4
<i>Panel D: First-stage estimates</i>						
	Surname diversity			Population		
Predicted surname diversity	0.609*** (0.049)	0.646*** (0.053)	0.650*** (0.058)	-0.003 (0.005)	0.000 (0.005)	0.002 (0.005)
Predicted population	-1.323*** (0.400)	-2.022*** (0.652)	-2.270*** (0.820)	0.296*** (0.102)	0.239** (0.112)	0.193* (0.113)
County-Surname fixed effects	✓	✓	✓	✓	✓	✓
Surname-Period fixed effects	✓	✓	✓	✓	✓	✓
State-Period fixed effects	✓	✓	✓	✓	✓	✓
County-specific linear trends	✓	✓	✓	✓	✓	✓
Observations	14,590,920	12,536,792	9,720,304	14,590,920	12,536,792	9,720,304
Dropping top % of counties	5%	10%	20%	5%	10%	20%

Notes: The table reports OLS, reduced-form and IV estimates for the specifications described in equation (8). An observation is a surname in a given county in a period from 1900 to 1940. Observations are weighted based on their population shares within counties. The endogenous variables are county-level surname diversity and population in t . In columns 1 to 4, the dependent variable is number of patents filed by individuals with surname n residing in county i in the five-year period starting in t divided by surname population size in 1900. In columns 5 to 8, it is number of breakthrough patents filed by individuals with surname n residing in county i in the five-year period starting in t divided by surname population size in 1900. Column 1 and 5 drop the 5% of counties with the highest surname diversity in 1900; Column 2 and 6 drop the top 10%; and Column 3 and 7 drop the top 20%. Standard errors are clustered at the state level. All independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B35: Dropping counties with highest surname diversity in 1900 (inventor level)

	Patents			Breakthrough patents		
	(1)	(2)	(3)	(4)	(5)	(6)
<i>Panel A: Least-squares estimates</i>						
Surname diversity	0.268*** (0.079)	0.207** (0.097)	0.071 (0.104)	0.082*** (0.021)	0.085*** (0.028)	0.041 (0.026)
Population	0.036 (0.827)	0.047 (1.415)	1.061 (2.204)	-0.361 (0.227)	-0.019 (0.324)	0.253 (0.423)
<i>Panel B: Reduced-form estimates</i>						
Predicted surname diversity	0.098 (0.074)	0.176* (0.093)	0.204** (0.095)	0.053** (0.020)	0.058** (0.022)	0.033 (0.022)
Predicted population	0.049 (0.814)	-1.018 (1.613)	-3.724*** (1.013)	-0.219 (0.232)	-0.208 (0.372)	-0.202 (0.316)
<i>Panel C: Instrumental-variable estimates</i>						
Surname diversity	0.199 (0.130)	0.301* (0.159)	0.352* (0.192)	0.101*** (0.033)	0.103*** (0.036)	0.060 (0.041)
Population	0.399 (1.012)	-1.256 (2.153)	-6.946 (4.526)	-0.226 (0.392)	-0.245 (0.463)	-0.333 (0.674)
Sanderson-Windmeijer F-stat	142; 84	202; 43	136; 17	142; 84	202; 43	136; 17
<i>Panel D: First-stage estimates</i>						
	Surname diversity			Population		
Predicted surname diversity	0.507*** (0.046)	0.543*** (0.029)	0.541*** (0.036)	-0.008 (0.005)	-0.011* (0.006)	-0.002 (0.007)
Predicted population	-0.898 (0.666)	-0.216 (0.752)	-0.541 (0.641)	0.570*** (0.208)	0.759*** (0.203)	0.509 (0.313)
Inventor fixed effects	✓	✓	✓	✓	✓	✓
Period-State fixed effects	✓	✓	✓	✓	✓	✓
Observations	55,450	35,192	19,910	55,450	35,192	19,910
Dropping top % of counties	5%	10%	20%	5%	10%	20%

Notes: The table reports OLS, reduced-form and IV estimates for the specifications described in equation (9). An observation is an inventor residing in a given county in a period from 1900 to 1940. The endogenous variables are county-level surname diversity and population in year t . The instruments are computed from estimates of a version of equation (1) that includes only the push-pull instruments, not the fixed effects and controls. The dependent variables are number of patents filed by the inventor in the five-year period starting in t (column 1–3) and number of breakthrough patents filed by the inventor in the five-year period starting in t (column 4–6). Column 1 and 5 drop the 5% of counties with the highest surname diversity in 1900; Column 2 and 6 drop the top 10%; and Column 3 and 7 drop the top 20%. Standard errors are clustered at the state level. The independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B36: Excluding top and bottom 5% of surnames within each county (county level)

	Patents per 1,000 people (mean = 1.05, sd = 1.73)				Breakthrough patents per 1,000 people (mean = 0.13, sd = 0.28)			
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
<i>Panel A: Least-squares estimates</i>								
Surname diversity	1.279*** (0.169)	1.040*** (0.195)	1.023*** (0.245)	1.308*** (0.177)	0.153*** (0.016)	0.099*** (0.014)	0.095*** (0.017)	0.100*** (0.015)
Population		0.585*** (0.107)	0.557*** (0.105)	0.723*** (0.177)		0.133*** (0.022)	0.118*** (0.019)	0.097*** (0.025)
<i>Panel B: Reduced-form estimates</i>								
Predicted surname diversity	0.975*** (0.171)	0.795*** (0.224)	0.817*** (0.269)	1.104*** (0.231)	0.115*** (0.014)	0.055*** (0.021)	0.057*** (0.023)	0.074*** (0.018)
Predicted population		0.288*** (0.128)	0.229* (0.128)	-0.115 (0.179)		0.097*** (0.023)	0.080*** (0.019)	0.016 (0.029)
<i>Panel C: Instrumental-variable estimates</i>								
Surname diversity	1.243*** (0.171)	1.076*** (0.204)	1.084*** (0.252)	1.345*** (0.210)	0.147*** (0.015)	0.091*** (0.017)	0.095*** (0.019)	0.092*** (0.017)
Population		0.446*** (0.163)	0.388*** (0.136)	0.177 (0.444)		0.150*** (0.032)	0.116*** (0.021)	0.089 (0.084)
Sanderson-Windmeijer <i>F</i> -stat	375	692; 413	751; 247	897; 17	375	692; 413	751; 247	897; 17
<i>Panel D: First-stage estimates</i>								
	Surname diversity				Population			
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
Predicted surname diversity	0.784*** (0.041)	0.786*** (0.050)	0.817*** (0.049)	0.824*** (0.052)		-0.113*** (0.038)	-0.175*** (0.045)	-0.028 (0.038)
Predicted population		-0.002 (0.024)	-0.051** (0.024)	-0.126** (0.051)		0.652*** (0.050)	0.733*** (0.065)	0.309*** (0.089)
County fixed effects	✓	✓	✓	✓	✓	✓	✓	✓
Period fixed effects	✓	✓			✓	✓		
State-Period fixed effects			✓	✓			✓	✓
County-specific linear trends				✓				✓
Observations	21,981	21,981	21,981	21,981	21,981	21,981	21,981	21,981

Notes: The table reports OLS, reduced-form and IV estimates for the specifications described in equation (2). An observation is a county in a period from 1900 to 1940. The endogenous variables are county-level surname diversity and population in t . When constructing these variables, we dropped the most and least frequent 5% surnames within each county. In columns 1 to 4, the dependent variable is number of patents filed in the county in the five-year period starting in t divided by county population size in 1900. In columns 5 to 8, it is number of breakthrough patents filed in the county in the five-year period starting in t divided by county population size in 1900. Standard errors are clustered at the state level. All independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

**Table B37: Excluding top and bottom 5% of surnames within each county
(county-surname level)**

	Patents per 1,000 people (mean = 1.02, sd = 13.31)				Breakthrough patents per 1,000 people (mean = 0.1, sd = 2.23)			
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
<i>Panel A: Least-squares estimates</i>								
Surname diversity	1.703*** (0.322)	1.036*** (0.221)	1.056*** (0.227)	1.293*** (0.182)	0.208*** (0.051)	0.092*** (0.033)	0.093*** (0.034)	0.072*** (0.020)
Population		3.376*** (0.703)	3.397*** (0.720)	4.302*** (0.753)		0.586*** (0.122)	0.595*** (0.124)	0.761*** (0.132)
<i>Panel B: Reduced-form estimates</i>								
Predicted surname diversity	1.036*** (0.217)	0.355 (0.239)	0.363 (0.246)	0.555*** (0.200)	0.124*** (0.032)	0.000 (0.038)	0.000 (0.039)	-0.004 (0.034)
Predicted population		2.592*** (0.751)	2.608*** (0.771)	1.880* (0.950)		0.470*** (0.138)	0.477*** (0.142)	0.399* (0.210)
<i>Panel C: Instrumental-variable estimates</i>								
Surname diversity	1.589*** (0.297)	1.019*** (0.228)	1.037*** (0.235)	1.124*** (0.185)	0.310*** (0.088)	0.232*** (0.083)	0.216*** (0.080)	0.082* (0.044)
Population		3.274*** (0.699)	3.300*** (0.716)	5.828** (2.172)		0.192*** (0.063)	0.275*** (0.074)	0.757*** (0.216)
Sanderson-Windmeijer <i>F</i> -stat	324	420; 96	420; 96	384; 42	324	420; 96	420; 96	384; 42
<i>Panel D: First-stage estimates</i>								
	Surname diversity				Population			
Predicted surname diversity	0.652*** (0.036)	0.650*** (0.039)	0.649*** (0.039)	0.638*** (0.043)		-0.094*** (0.025)	-0.094*** (0.025)	-0.028** (0.010)
Predicted population		0.006 (0.025)	0.007 (0.025)	-0.120** (0.051)		0.790*** (0.094)	0.788*** (0.096)	0.346*** (0.070)
County-Surname fixed effects	✓	✓	✓	✓	✓	✓	✓	✓
State-Period fixed effects	✓	✓	✓	✓	✓	✓	✓	✓
Surname-Period fixed effects			✓	✓			✓	✓
County-specific linear trends				✓				✓
Observations	19,482,016	19,482,016	19,482,016	19,482,016	19,482,016	19,482,016	19,482,016	19,482,016

Notes: The table reports OLS, reduced-form and IV estimates for the specifications described in equation (8). An observation is a surname in a given county in a period from 1900 to 1940. Observations are weighted based on their population shares within counties. The endogenous variables are county-level surname entropy and population in t . When constructing these variables, we dropped the most and least frequent 5% surnames within each county. In columns 1 to 4, the dependent variable is number of patents filed by individuals with surname n residing in county in i in the five-year period starting in t divided by surname population size in 1900. In columns 5 to 8, the dependent variable is number of breakthrough patents filed by individuals with surname n residing in county in i in the five-year period starting in t divided by surname population size in 1900. Standard errors are clustered at the state level. All independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B38: Excluding top and bottom 5% of surnames within each county (inventor level)

	Patents (mean = 2.24, sd = 1.69)		Breakthrough patents (mean = 0.26, sd = 0.44)	
	(1)	(2)	(3)	(4)
<i>Panel A: Least-squares estimates</i>				
Surname diversity	0.267*** (0.048)	0.282*** (0.057)	0.035 (0.024)	0.043 (0.026)
Population		-0.040 (0.048)		-0.021* (0.012)
<i>Panel B: Reduced-form estimates</i>				
Predicted surname diversity	0.170*** (0.036)	0.166*** (0.040)	0.027 (0.017)	0.034* (0.019)
Predicted population		0.007 (0.025)		-0.010 (0.009)
<i>Panel C: Instrumental-variable estimates</i>				
Surname diversity	0.333*** (0.082)	0.333*** (0.084)	0.054 (0.039)	0.059 (0.040)
Population		-0.001 (0.044)		-0.017 (0.015)
Sanderson-Windmeijer <i>F</i> -stat	88	103; 215	88	103; 215
<i>Panel D: First-stage estimates</i>				
	Surname diversity		Population	
Predicted surname diversity	0.511*** (0.055)	0.496*** (0.059)		-0.270*** (0.057)
Predicted population		0.023 (0.025)		0.671*** (0.061)
Inventor fixed effects	✓	✓	✓	✓
Period-State fixed effects	✓	✓	✓	✓
Observations	209,943	209,943	209,943	209,943

Notes: The table reports OLS, reduced-form and IV estimates for the specifications described in equation (9) as well as IV first-stage estimates. An observation is an inventor residing in a given county in a period from 1900 to 1940. The endogenous variables are county-level surname entropy and population in year t . When constructing these variables, we dropped the most and least frequent 5% surnames within each county. The dependent variables are number of patents filed by the inventor in the five-year period starting in t (column 1–2) and number of breakthrough patents filed by the inventor in the five-year period starting in t (column 3–4). Standard errors are clustered at the state level. The independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B39: Surname diversity square (county level)

	Patents per 1,000 people (mean = 1.05, sd = 1.73)				Breakthrough patents per 1,000 people (mean = 0.13, sd = 0.28)			
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
<i>Panel A: Least-squares estimates</i>								
Surname diversity	1.487*** (0.182)	1.132*** (0.161)	1.068*** (0.187)	1.513*** (0.171)	0.202*** (0.023)	0.140*** (0.020)	0.125*** (0.020)	0.147*** (0.019)
Surname diversity square	0.152** (0.068)	0.005 (0.068)	-0.009 (0.076)	0.115 (0.083)	0.049*** (0.011)	0.024*** (0.009)	0.019*** (0.007)	0.032*** (0.008)
Population		0.668*** (0.120)	0.653*** (0.127)	0.635*** (0.195)		0.116*** (0.019)	0.110*** (0.019)	0.072** (0.030)
<i>Panel B: Reduced-form estimates</i>								
Predicted surname diversity	0.961*** (0.169)	0.461*** (0.146)	0.396** (0.179)	0.749*** (0.171)	0.121*** (0.019)	0.028 (0.020)	0.025 (0.022)	0.059** (0.026)
Predicted surname diversity square	0.011 (0.062)	-0.178*** (0.059)	-0.192*** (0.063)	-0.120* (0.064)	0.022** (0.008)	-0.013** (0.006)	-0.013** (0.007)	-0.002 (0.008)
Predicted population		0.869*** (0.155)	0.894*** (0.170)	0.477** (0.198)		0.162*** (0.029)	0.147*** (0.029)	0.054 (0.044)
<i>Panel C: Instrumental-variable estimates</i>								
Surname diversity	1.560*** (0.201)	1.118*** (0.231)	1.036*** (0.250)	1.536*** (0.224)	0.215*** (0.025)	0.126*** (0.030)	0.121*** (0.027)	0.138*** (0.034)
Surname diversity square	0.175** (0.072)	-0.028 (0.086)	-0.065 (0.080)	0.094 (0.103)	0.055*** (0.012)	0.014 (0.012)	0.012 (0.009)	0.024 (0.015)
Population		0.711*** (0.179)	0.725*** (0.139)	0.525* (0.277)		0.142*** (0.033)	0.121*** (0.025)	0.065 (0.070)
Sanderson-Windmeijer F-stat	200; 312	105; 117; 105	135; 132; 165	99; 77; 84	200; 312	105; 117; 105	135; 132; 165	99; 77; 84
<i>Panel D: First-stage estimates</i>								
	Surname diversity				Surname diversity square			
Predicted surname diversity	0.657*** (0.035)	0.547*** (0.034)	0.581*** (0.038)	0.562*** (0.039)	-0.361*** (0.071)	-0.557*** (0.089)	-0.458*** (0.088)	-0.362*** (0.103)
Predicted surname diversity square	-0.068*** (0.014)	-0.110*** (0.013)	-0.100*** (0.012)	-0.106*** (0.011)	0.669*** (0.043)	0.595*** (0.067)	0.613*** (0.067)	0.619*** (0.072)
Predicted population		0.191*** (0.015)	0.119*** (0.016)	0.097*** (0.016)		0.340*** (0.123)	0.157 (0.129)	0.022 (0.167)
	Population							
Predicted surname diversity		-0.233*** (0.036)	-0.325*** (0.046)	-0.152*** (0.053)				
Predicted surname diversity square		-0.054*** (0.012)	-0.067*** (0.013)	-0.031*** (0.012)				
Predicted population		0.936*** (0.062)	1.077*** (0.088)	0.621*** (0.126)				
County fixed effects	✓	✓	✓	✓	✓	✓	✓	✓
Period fixed effects	✓	✓			✓	✓		
State-Period fixed effects			✓	✓			✓	✓
County-specific linear trends				✓				✓
Observations	21,984	21,984	21,984	21,984	21,984	21,984	21,984	21,984

Notes: The table reports OLS, reduced-form, and IV estimates for the specifications described in equation (2). An observation is a county in a period from 1900 to 1940. The endogenous variables are county-level surname diversity, surname diversity square and population in t . In columns 1 to 4, the dependent variable is number of patents filed in the county in the five-year period starting in t divided by county population size in 1900. In columns 5 to 8, it is number of breakthrough patents filed in the county in the five-year period starting in t divided by county population size in 1900. Standard errors are clustered at the state level. Independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B40: Surname diversity square (county-surname level)

	Patents per 1,000 people (mean = 1.02, sd = 13.31)			Breakthrough patents per 1,000 people (mean = 0.1, sd = 2.23)		
	(1)	(2)	(3)	(4)	(5)	(6)
<i>Panel A: Least-squares estimates</i>						
Surname diversity	1.721*** (0.334)	1.749*** (0.340)	1.840*** (0.339)	0.231*** (0.060)	0.235*** (0.061)	0.180*** (0.045)
Surname diversity square	0.474*** (0.124)	0.482*** (0.126)	0.364*** (0.119)	0.086*** (0.026)	0.088*** (0.027)	0.062*** (0.019)
Population	3.241*** (0.693)	3.256*** (0.709)	3.324*** (0.562)	0.532*** (0.123)	0.539*** (0.124)	0.668*** (0.088)
<i>Panel B: Reduced-form estimates</i>						
Predicted surname diversity	0.690** (0.258)	0.699** (0.262)	0.753*** (0.275)	0.083** (0.038)	0.084** (0.039)	0.064 (0.061)
Predicted surname diversity square	0.174* (0.096)	0.176* (0.098)	0.061 (0.114)	0.038** (0.018)	0.039** (0.018)	0.015 (0.026)
Predicted population	2.911*** (0.983)	2.927*** (1.009)	0.851 (1.121)	0.437*** (0.156)	0.443*** (0.160)	0.085 (0.316)
<i>Panel C: Instrumental-variable estimates</i>						
Surname diversity	2.158*** (0.501)	2.190*** (0.512)	2.255*** (0.824)	0.295*** (0.082)	0.298*** (0.084)	0.210 (0.183)
Surname diversity square	0.640*** (0.209)	0.651*** (0.214)	0.517 (0.363)	0.117*** (0.041)	0.118*** (0.042)	0.068 (0.082)
Population	2.366*** (0.643)	2.373*** (0.663)	0.685 (1.895)	0.348*** (0.105)	0.353*** (0.107)	0.065 (0.522)
Sanderson-Windmeijer F-stat	136; 314; 110	136; 314; 110	100; 143; 573	136; 314; 110	136; 314; 110	100; 143; 573
<i>Panel D: First-stage estimates</i>						
	Surname diversity			Surname diversity square		
Predicted surname diversity	0.470*** (0.035)	0.470*** (0.036)	0.416*** (0.040)	-0.345*** (0.058)	-0.346*** (0.058)	-0.325*** (0.053)
Predicted surname diversity square	-0.064*** (0.013)	-0.065*** (0.013)	-0.082*** (0.012)	0.505*** (0.034)	0.504*** (0.034)	0.479*** (0.037)
Predicted population	0.118*** (0.026)	0.119*** (0.026)	0.114** (0.052)	0.508*** (0.078)	0.501*** (0.078)	0.248* (0.134)
County-Surname fixed effects	✓	✓	✓	✓	✓	✓
State-Period fixed effects	✓	✓	✓	✓	✓	✓
Surname-Period fixed effects		✓	✓		✓	✓
	Population					
Predicted surname diversity	-0.044* (0.024)	-0.044* (0.024)	-0.025*** (0.007)			
Predicted surname diversity square	-0.004 (0.009)	-0.004 (0.009)	-0.003 (0.003)			
Predicted population	0.985*** (0.148)	0.986*** (0.150)	0.679*** (0.045)			
County fixed effects			✓			✓
County-specific linear trends			✓			✓
Observations	19,482,024	19,482,024	19,482,024	19,482,024	19,482,024	19,482,024

Notes: The table reports OLS, reduced-form, and IV estimates for the specifications described in equation (8). An observation is a surname in a given county in a period from 1900 to 1940. Observations are weighted based on their population shares within counties. The endogenous variables are county-level surname diversity, surname diversity square and population in t . In columns 1 to 4, the dependent variable is number of patents filed in the county in the five-year period starting in t divided by county population size in 1900. In columns 5 to 8, it is number of breakthrough patents filed in the county in the five-year period starting in t divided by county population size in 1900. Standard errors are clustered at the state level. Independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B41: Surname diversity square (inventor level)

	Patents (mean = 2.24, sd = 1.69)		Breakthrough patents (mean = 0.26, sd = 0.44)	
	(1)	(2)	(3)	(4)
<i>Panel A: Least-squares estimates</i>				
Surname diversity	0.217*** (0.058)	0.221*** (0.060)	0.042* (0.024)	0.044* (0.025)
Surname diversity square	-0.002 (0.020)	0.002 (0.020)	-0.002 (0.006)	0.001 (0.007)
Population		-0.017 (0.055)		-0.012 (0.016)
<i>Panel B: Reduced-form estimates</i>				
Predicted surname diversity	0.082*** (0.026)	0.083*** (0.027)	0.021** (0.009)	0.020** (0.009)
Predicted surname diversity square	-0.005 (0.009)	-0.005 (0.008)	-0.004 (0.003)	-0.003 (0.003)
Predicted population		0.005 (0.039)		-0.006 (0.011)
<i>Panel C: Instrumental-variable estimates</i>				
Surname diversity	0.584* (0.298)	0.542* (0.291)	0.116 (0.090)	0.106 (0.084)
Surname diversity square	0.099 (0.098)	0.123 (0.131)	0.009 (0.027)	0.015 (0.034)
Population		-0.127 (0.086)		-0.029 (0.023)
Sanderson-Windmeijer F-stat	15; 13	21; 20; 107	15; 13	21; 20; 107
<i>Panel D: First-stage estimates</i>				
	Surname diversity		Surname diversity square	
Predicted surname diversity	0.222*** (0.038)	0.242*** (0.040)	-0.479*** (0.077)	-0.412*** (0.072)
Predicted surname diversity square	-0.052*** (0.012)	-0.056*** (0.013)	0.259*** (0.055)	0.245*** (0.056)
Predicted population		0.118*** (0.023)		0.399*** (0.084)
	Population			
Predicted surname diversity		-0.021 (0.027)		
Predicted surname diversity square		0.036 (0.024)		
Predicted population		0.852*** (0.088)		
Inventor fixed effects	✓	✓	✓	✓
Period-State fixed effects	✓	✓	✓	✓
Observations	209,944	209,944	209,944	209,944

Notes: The table reports OLS, reduced-form and IV estimates for the specifications described in equation (9). An observation is an inventor residing in a given county in a period from 1900 to 1940. The endogenous variables are county-level surname diversity, surname diversity square and population in year t , and the dependent variables are number of patents filed by the inventor in the five-year period starting in t (column 1–2) and number of breakthrough patents filed by the inventor in the five-year period starting in t (column 3–4). Standard errors are clustered at the state level. The independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B42: Surname diversity vs. educational attainment diversity

	Patents per 1,000 people (mean = 0.33, sd = 0.6)			Breakthrough patents per 1,000 people (mean = 0.06, sd = 0.12)		
	(1)	(2)	(3)	(4)	(5)	(6)
<i>Panel A: Least-squares estimates</i>						
Surname diversity	0.110*** (0.028)		0.098*** (0.029)	0.020*** (0.006)		0.017*** (0.006)
Years of schooling diversity		0.059*** (0.015)	0.030** (0.013)		0.011*** (0.003)	0.006** (0.003)
Population	0.167*** (0.018)	0.209*** (0.022)	0.167*** (0.018)	0.030*** (0.003)	0.038*** (0.004)	0.030*** (0.003)
<i>Panel B: Reduced-form estimates</i>						
Predicted surname diversity	0.036 (0.025)		0.023 (0.026)	0.007 (0.005)		0.004 (0.005)
Years of schooling diversity		0.046*** (0.013)	0.041*** (0.012)		0.009*** (0.003)	0.008*** (0.003)
Predicted population	0.236*** (0.035)	0.253*** (0.031)	0.237*** (0.035)	0.042*** (0.006)	0.046*** (0.005)	0.043*** (0.006)
<i>Panel C: Instrumental-variable estimates</i>						
Surname diversity	0.093*** (0.032)		0.083** (0.034)	0.017*** (0.006)		0.015** (0.007)
Years of schooling diversity		0.050*** (0.014)	0.028** (0.013)		0.009*** (0.003)	0.005* (0.003)
Population	0.199*** (0.026)	0.245*** (0.028)	0.199*** (0.026)	0.036*** (0.005)	0.044*** (0.005)	0.036*** (0.005)
Sanderson-Windmeijer <i>F</i> -stat	1146; 703	937	947; 691	1146; 703	937	947; 691
<i>Panel D: First-stage estimates</i>						
	Surname diversity		Population			
Predicted surname diversity	0.763*** (0.025)		0.736*** (0.028)	-0.180*** (0.030)		-0.190*** (0.032)
Years of schooling diversity			0.082*** (0.020)		-0.015 (0.022)	0.029 (0.021)
Predicted population	0.033 (0.020)		0.035* (0.020)	1.172*** (0.044)	1.036*** (0.034)	1.173*** (0.044)
State fixed effects	✓	✓	✓	✓	✓	✓
Observations	2,748	2,737	2,737	2,748	2,737	2,737

Notes: The table reports OLS, reduced-form and IV estimates for estimates of regressions of number of patents and breakthrough patents filed between 1940 and 1944 per 1940 population (in 1,000) on surname diversity, population and individuals' years of schooling diversity in 1940. The endogenous variables are county-level surname diversity and population. The unit of observation is a county in 1940 harmonized to 1900 borders. Standard errors are clustered on states and reported in parentheses. Independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B43: Surname diversity vs. average educational attainment

	Patents per 1,000 people (mean = 0.33, sd = 0.6)			Breakthrough patents per 1,000 people (mean = 0.06, sd = 0.12)		
	(1)	(2)	(3)	(4)	(5)	(6)
<i>Panel A: Least-squares estimates</i>						
Surname diversity		0.081*** (0.025)	0.080*** (0.026)		0.015*** (0.005)	0.015*** (0.005)
Average years of schooling	0.126*** (0.032)	0.099*** (0.026)	0.148*** (0.033)	0.021*** (0.005)	0.016*** (0.004)	0.025*** (0.006)
Surname diversity × Average years of schooling			0.120*** (0.031)			0.023*** (0.006)
Population	0.204*** (0.021)	0.168*** (0.018)	0.130*** (0.019)	0.037*** (0.004)	0.030*** (0.003)	0.023*** (0.004)
<i>Panel B: Reduced-form estimates</i>						
Predicted surname diversity		0.017 (0.023)	0.023 (0.024)		0.004 (0.004)	0.005 (0.005)
Average years of schooling	0.113*** (0.028)	0.110*** (0.026)	0.133*** (0.027)	0.019*** (0.005)	0.018*** (0.004)	0.022*** (0.005)
Predicted surname diversity × Average years of schooling			0.089*** (0.024)			0.018*** (0.004)
Predicted population	0.246*** (0.030)	0.233*** (0.035)	0.192*** (0.033)	0.045*** (0.005)	0.042*** (0.006)	0.034*** (0.006)
<i>Panel C: Instrumental-variable estimates</i>						
Surname diversity		0.073** (0.028)	0.071** (0.028)		0.014** (0.006)	0.014** (0.005)
Average years of schooling	0.116*** (0.030)	0.095*** (0.025)	0.143*** (0.032)	0.019*** (0.005)	0.015*** (0.004)	0.025*** (0.005)
Population	0.238*** (0.027)	0.197*** (0.026)	0.154*** (0.027)	0.043*** (0.005)	0.035*** (0.005)	0.027*** (0.005)
Surname diversity × Average years of schooling			0.112*** (0.030)			0.023*** (0.006)
Sanderson-Windmeijer <i>F</i> -stat	969	943; 693	1154; 531; 260	969	943; 693	1154; 531; 260
<i>Panel D: First-stage estimates</i>						
	Surname diversity			Population		
Predicted surname diversity		0.738*** (0.029)	0.740*** (0.027)		-0.185*** (0.031)	-0.184*** (0.032)
Average years of schooling		0.145*** (0.030)	0.152*** (0.030)	-0.013 (0.027)	0.024 (0.024)	0.027 (0.028)
Predicted surname diversity × Average years of schooling			0.026 (0.025)			0.014 (0.028)
Predicted population		0.030 (0.022)	0.017 (0.027)	1.034*** (0.033)	1.172*** (0.045)	1.165*** (0.050)
	Surname diversity × Average years of schooling					
Predicted surname diversity			-0.015 (0.032)			
Average years of schooling			-0.222** (0.098)			
Predicted surname diversity × Average years of schooling			0.762*** (0.052)			
Predicted population			0.094** (0.040)			
State fixed effects	✓	✓	✓	✓	✓	✓
Observations	2,737	2,737	2,737	2,737	2,737	2,737

Notes: The table reports OLS, reduced-form and IV estimates for estimates of regressions of number of patents and breakthrough patents filed between 1940 and 1944 per 1940 population (in 1,000) on surname diversity, population and individuals' average years of schooling in 1940. The endogenous variables are county-level surname diversity, population and surname diversity interacted with years of schooling. The unit of observation is a county in 1940 harmonized to 1900 borders. Standard errors are clustered on states and reported in parentheses. Independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B44: Ancestral country-county level estimates of the effects of push-pull instruments on ancestry stocks

	No. of people in county i with ancestral country o in year:															
	1900				1905				1910				1915			
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)	(11)	(12)	(13)	(14)	(15)	(16)
$I_{0,-r(t)}^{1880} \times \frac{I_{0,-r(t)}^{1880}}{I_{-(-6)}}$	3.205*** (0.121)	3.124*** (0.120)	3.267*** (0.229)	3.252*** (0.273)	3.887*** (0.258)	3.739*** (0.287)	3.916*** (0.272)	3.648*** (0.299)	4.364*** (0.287)	4.053*** (0.306)	4.484*** (0.366)	4.143*** (0.406)	4.981*** (0.273)	4.708*** (0.355)	5.125*** (0.286)	4.863*** (0.364)
$I_{0,-r(t)}^{1895} \times \frac{I_{0,-r(t)}^{1895}}{I_{-(-6)}}$	2.796*** (0.434)	3.776*** (0.481)	3.541*** (1.173)	4.838*** (1.360)	4.786*** (1.033)	5.765*** (1.234)	5.721*** (1.272)	5.948*** (1.375)	6.536*** (1.446)	6.723*** (1.353)	8.153*** (1.589)	8.580*** (1.517)	9.097*** (0.287)	9.975*** (0.284)	9.533*** (0.366)	10.803*** (0.288)
$I_{0,-r(t)}^{1900} \times \frac{I_{0,-r(t)}^{1900}}{I_{-(-6)}}$			-2.993 (6.031)	-1.065 (8.458)	-0.465 (5.932)	-2.120 (7.932)	-4.675 (5.776)	-9.689 (8.382)	-3.842 (5.503)	-9.645 (8.549)	-14.394* (7.349)	-21.973** (10.152)	-14.400 (17.418)	-17.513 (22.636)	-13.957 (17.165)	-15.498 (21.836)
$I_{0,-r(t)}^{1905} \times \frac{I_{0,-r(t)}^{1905}}{I_{-(-6)}}$					18.867*** (3.301)	20.976*** (3.471)	22.565*** (4.732)	23.549*** (4.940)	25.879*** (5.542)	26.700*** (5.342)	30.794*** (5.641)	33.231*** (5.510)	34.257*** (4.125)	35.494*** (6.147)	37.721*** (4.871)	38.884*** (7.368)
$I_{0,-r(t)}^{1910} \times \frac{I_{0,-r(t)}^{1910}}{I_{-(-6)}}$							22.109*** (5.313)	18.560** (7.328)	25.430*** (6.226)	21.504*** (6.802)	37.181*** (6.292)	33.876*** (7.914)	41.072*** (11.306)	35.021** (14.001)	43.634*** (11.168)	38.027*** (13.019)
$I_{0,-r(t)}^{1915} \times \frac{I_{0,-r(t)}^{1915}}{I_{-(-6)}}$									12.814*** (3.450)	12.628*** (4.061)	22.541*** (5.257)	22.884*** (5.950)	25.652*** (3.907)	25.606*** (3.325)	29.307*** (4.657)	29.755*** (3.954)
$I_{0,-r(t)}^{1920} \times \frac{I_{0,-r(t)}^{1920}}{I_{-(-6)}}$											9.107 (10.879)	8.782 (11.473)	10.159* (5.689)	7.603* (4.355)	17.836** (7.107)	15.182*** (5.344)
$I_{0,-r(t)}^{1925} \times \frac{I_{0,-r(t)}^{1925}}{I_{-(-6)}}$													28.357*** (2.221)	27.532*** (3.398)	34.635*** (3.272)	34.147*** (4.152)
$I_{0,-r(t)}^{1930} \times \frac{I_{0,-r(t)}^{1930}}{I_{-(-6)}}$															-32.971*** (9.814)	-28.977*** (7.604)
Constant	0.392* (0.212)		0.393*** (0.071)		-0.019 (0.113)		-0.693*** (0.214)		-1.128*** (0.201)		-1.763*** (0.314)		-2.414*** (0.244)		-2.383*** (0.276)	
F-stat	360	337	76	48	68	44	54	41	63	55	58	45	15059	1246	59658	1492
Observations	72,905	72,905	76,156	76,156	77,160	77,160	77,711	77,711	77,978	77,978	79,740	79,740	79,934	79,934	81,602	81,602
R ²	0.775	0.881	0.757	0.867	0.795	0.893	0.778	0.878	0.790	0.887	0.750	0.860	0.793	0.886	0.789	0.885
Neighboring ancestral countries-County fixed effects	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓
Ancestral country-Census division fixed effects	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓
I_t^*/I^* controls	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓

Notes: This table reports OLS estimates for a version of the specification described in equation (3), corresponding to step 1 of the instrument construction for ancestral-country diversity. The dependent variable is the number of individuals from a given ancestral country residing in a given U.S. county in a period from 1900 to 1940. The table includes the F -statistic for the Wald test for the joint nullity of the push-pull instruments. Standard errors clustered at the surname level. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B45: Surname diversity vs. ancestral-country diversity (county level)

	Patents per 1,000 people (mean = 1.05, sd = 1.73)			Breakthrough patents per 1,000 people (mean = 0.13, sd = 0.28)		
	(1)	(2)	(3)	(4)	(5)	(6)
<i>Panel A: Least-squares estimates</i>						
Surname diversity	1.154*** (0.214)	1.134*** (0.262)	1.445*** (0.183)	0.105*** (0.016)	0.105*** (0.019)	0.110*** (0.019)
Ancestral-country diversity	-0.057 (0.059)	-0.110** (0.054)	-0.199*** (0.055)	0.005 (0.011)	-0.012 (0.009)	-0.020* (0.012)
Population	0.680*** (0.105)	0.652*** (0.107)	0.742*** (0.190)	0.139*** (0.020)	0.126*** (0.019)	0.100*** (0.029)
<i>Panel B: Reduced-form estimates</i>						
Predicted surname diversity	0.679*** (0.164)	0.644*** (0.205)	0.918*** (0.231)	0.056*** (0.019)	0.054** (0.021)	0.069*** (0.022)
Predicted ancestral-country diversity	0.019 (0.090)	0.023 (0.093)	0.022 (0.035)	-0.010 (0.008)	-0.010 (0.008)	-0.007 (0.008)
Predicted population	0.554*** (0.132)	0.531*** (0.139)	0.212 (0.202)	0.133*** (0.026)	0.117*** (0.024)	0.048 (0.036)
<i>Panel C: Instrumental-variable estimates</i>						
Surname diversity	1.193*** (0.144)	1.142*** (0.165)	1.457*** (0.229)	0.129*** (0.022)	0.139*** (0.027)	0.135*** (0.031)
Ancestral-country diversity	-0.089 (0.383)	-0.036 (0.518)	-0.140 (0.229)	-0.069 (0.043)	-0.088 (0.061)	-0.073 (0.052)
Population	0.690*** (0.202)	0.644*** (0.157)	0.778*** (0.257)	0.186*** (0.029)	0.148*** (0.023)	0.143** (0.060)
Sanderson-Windmeijer F-stat	88; 72; 129	122; 72; 307	75; 40; 120	88; 72; 129	122; 72; 307	75; 40; 120
<i>Panel D: First-stage estimates</i>						
	Surname diversity			Ancestral-country diversity		
	(1)	(2)	(3)	(4)	(5)	(6)
Predicted surname diversity	0.658*** (0.061)	0.694*** (0.062)	0.699*** (0.069)	0.011 (0.041)	0.107* (0.057)	0.153** (0.075)
Predicted ancestral-country diversity	0.035** (0.016)	0.028** (0.014)	0.031*** (0.009)	0.187*** (0.029)	0.151*** (0.028)	0.157*** (0.032)
Predicted population	0.007 (0.034)	-0.063* (0.034)	-0.132** (0.060)	0.328*** (0.052)	0.154*** (0.046)	0.170** (0.080)
	Population					
	(1)	(2)	(3)			
Predicted surname diversity	-0.151*** (0.039)	-0.225*** (0.047)	-0.101** (0.045)			
Predicted ancestral-country diversity	-0.009 (0.012)	-0.006 (0.012)	-0.002 (0.007)			
Predicted population	0.833*** (0.055)	0.945*** (0.075)	0.550*** (0.108)			
County fixed effects	✓	✓	✓	✓	✓	✓
Period fixed effects	✓			✓		
State-Period fixed effects		✓	✓		✓	✓
County-specific linear trends			✓			✓
Observations	21,984	21,984	21,984	21,984	21,984	21,984

Notes: The table reports OLS estimates for the specifications described in equation (2). An observation is a county in a period from 1900 to 1940. The endogenous variables are county-level ancestral-country diversity, surname diversity and population in t . In columns 1 to 3, the dependent variable is number of patents filed in the county in the five-year period starting in t divided by county population size in 1900. In columns 4 to 6, it is number of breakthrough patents filed in the county in the five-year period starting in t divided by county population size in 1900. Standard errors are clustered at the state level. All independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance 5%, and 10% levels.

Table B46: Share of patents with different inventor surname characteristics

Patents with	Number	Share (%)
1 inventor	1,289,035	88.81
> 1 inventors and 1 distinct surname	14,963	1.03
> 1 inventors and > 1 distinct surnames	147,461	10.16

Notes: The table reports the number of patents with (1) a single patentee, (2) multiple patentees who share a single surname and (3) multiple patentees who do not share a single surname.

Table B47: Alternative measures of surname diversity of patents

	Surname diversity of patent		Log distinct surnames of patent		Surname dispersion of patent		Genetic distance weighted surname diversity of patent	
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
<i>Panel A: Least-squares estimates</i>								
Surname diversity	0.020 (0.019)	0.011* (0.006)	0.021 (0.019)	0.012* (0.006)	0.020 (0.019)	0.012* (0.006)	0.020 (0.019)	0.010 (0.006)
<i>Panel B: Reduced-form estimates</i>								
Predicted surname diversity	0.014* (0.009)	0.009*** (0.003)	0.014 (0.009)	0.009*** (0.003)	0.014* (0.009)	0.009*** (0.003)	0.014* (0.008)	0.008*** (0.003)
<i>Panel C: Instrumental-variable estimates</i>								
Surname diversity	0.067* (0.039)	0.040*** (0.013)	0.066* (0.040)	0.041*** (0.013)	0.067* (0.039)	0.040*** (0.013)	0.067* (0.039)	0.040*** (0.013)
Sanderson-Windmeijer F-stat	40	40	40	40	40	40	40	40
<i>Panel D: First-stage estimates</i>								
Predicted surname diversity	Surname diversity of patent							
	0.212*** (0.033)	0.212*** (0.033)						
County fixed effects	✓	✓	✓	✓	✓	✓	✓	✓
State-Period fixed effects	✓	✓	✓	✓	✓	✓	✓	✓
Patent technology field-Period fixed effects	✓	✓	✓	✓	✓	✓	✓	✓
Team size fixed effects		✓		✓		✓		✓
Observations	1,451,459	1,451,459	1,451,459	1,451,459	1,451,459	1,451,459	1,451,459	1,451,459

Notes: The table reports least squares, reduced-form and IV estimates for the specifications described in equation (10) but with county-level surname diversity as the main independent variable instead of patent-level surname diversity. An observation is a patent with inventors residing in a given county in a five-year period from 1900 to 1940. Observations are weighted by the inverse of the number of authors multiplied by the weight used to assign patents to counties within their 1900 borders. Standard errors are clustered on counties and reported in parentheses. All variables are standardized to have a mean of zero and a unit variance. The sources and construction of all variables are detailed in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B48: Alternative measures of surname diversity of patents with 2+ inventors

	Surname diversity of patent		Log distinct surnames of patent		Surname dispersion of patent		Genetic distance weighted surname diversity of patent	
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
<i>Panel A: Least-squares estimates</i>								
Surname diversity	0.111** (0.056)	0.106* (0.056)	0.101* (0.054)	0.101* (0.054)	0.118* (0.061)	0.117* (0.061)	0.078 (0.053)	0.077 (0.053)
<i>Panel B: Reduced-form estimates</i>								
Predicted surname diversity	0.069** (0.028)	0.070*** (0.026)	0.062** (0.027)	0.066*** (0.025)	0.077*** (0.029)	0.077*** (0.029)	0.058** (0.026)	0.060** (0.025)
<i>Panel C: Instrumental-variable estimates</i>								
Surname diversity	0.332** (0.134)	0.334*** (0.124)	0.298** (0.128)	0.316*** (0.117)	0.368*** (0.138)	0.367*** (0.138)	0.277** (0.120)	0.286** (0.116)
Sanderson-Windmeijer <i>F</i> -stat	35	35	35	35	35	35	35	35
<i>Panel D: First-stage estimates</i>								
	Surname diversity of patent							
Predicted surname diversity	0.209** (0.035)	0.209** (0.035)						
County fixed effects	✓	✓	✓	✓	✓	✓	✓	✓
State-Period fixed effects	✓	✓	✓	✓	✓	✓	✓	✓
Patent technology field-Period fixed effects	✓	✓	✓	✓	✓	✓	✓	✓
Team size fixed effects		✓		✓		✓		✓
Observations	162,424	162,424	162,424	162,424	162,424	162,424	162,424	162,424

Notes: The table reports least squares, reduced-form and IV estimates for the specifications described in equation (10) but with county-level surname diversity as the main independent variable instead of patent-level surname diversity. An observation is a patent with inventors residing in a given county in a five-year period from 1900 to 1940. Observations are weighted by the inverse of the number of authors multiplied by the weight used to assign patents to counties within their 1900 borders. Standard errors are clustered on counties and reported in parentheses. All variables are standardized to have a mean of zero and a unit variance. The sources and construction of all variables are detailed in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B49: Patent-level estimates of the effect of genetic distance-weighted surname diversity on innovation

	Breakthrough patent indicator (mean = 0.17, sd = 0.38)			Technologies per patent (mean = 2.49, sd = 1.82)		
	(1)	(2)	(3)	(4)	(5)	(6)
<i>Panel A: Least-squares estimates</i>						
Genetic distance weighted surname diversity of patent	0.014*** (0.001)	0.009*** (0.001)	0.009*** (0.001)	0.031*** (0.005)	0.018*** (0.005)	0.017*** (0.005)
<i>Panel B: Reduced-form estimates</i>						
Predicted genetic distance weighted surname diversity	0.010** (0.004)	0.007* (0.004)		0.044*** (0.013)	0.034*** (0.011)	
<i>Panel C: Instrumental-variable estimates</i>						
Genetic distance weighted surname diversity of patent	1.279* (0.674)	0.966 (0.594)		5.562** (2.544)	4.730** (2.407)	
Sanderson-Windmeijer F-stat	8	7		8	7	
<i>Panel D: First-stage estimates</i>						
	Genetic distance weighted surname diversity of patent					
Predicted genetic distance weighted surname diversity	0.008*** (0.003)	0.007*** (0.003)				
Team size fixed effects	✓	✓	✓	✓	✓	✓
County fixed effects	✓	✓		✓	✓	
State-Period fixed effects	✓	✓		✓	✓	
Patent technology field-Period fixed effects		✓			✓	
County-Period fixed effects			✓			✓
Observations	1,451,459	1,451,459	1,451,459	1,451,459	1,451,459	1,451,459

Notes: The table reports least squares, reduced-form and IV estimates for the specifications described in equation (10). An observation is a patent with inventors residing in a given county in a period from 1900 to 1940. If there are multiple inventors who reside in different counties, the patent is assigned to all counties and the observation is weighted by 1 / number of inventors. The endogenous variables is genetic distance-weighted surname diversity of patent. The dependent variables are a breakthrough patent indicator in columns 1 to 3 and the number of technology classes per patent in columns 4 to 6. Columns 2–3 and 5–6 add one of six broad patent technology fields (chemical, computers, drugs, electrical, mechanical, other) fixed effects interacted with period fixed effects. Standard errors are clustered clustered at the county level. Columns 3 and 6 additionally include country-period fixed effects. All diversity variables are standardized to have a mean of zero and a unit variance. The sources and construction of all variables are detailed in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B50: Patent-level estimates of the relationship between ancestral-country diversity and innovation

	Breakthrough patent indicator (mean = 0.17, sd = 0.38)			Technologies per patent (mean = 2.49, sd = 1.82)		
	(1)	(2)	(3)	(4)	(5)	(6)
<i>Panel A</i>						
Ancestral-country diversity of patent	0.005*** (0.001)	0.003*** (0.001)	0.003*** (0.001)	0.019*** (0.004)	0.014*** (0.004)	0.013*** (0.004)
<i>Panel B</i>						
Ancestral-country diversity of patent	0.003* (0.001)	0.001 (0.001)	0.002 (0.001)	0.019** (0.010)	0.018** (0.009)	0.015 (0.010)
Surname diversity of patent	0.014*** (0.001)	0.008*** (0.001)	0.008*** (0.001)	0.015** (0.007)	0.007 (0.007)	0.011 (0.007)
<i>Panel C</i>						
Ancestral-country diversity	-0.004 (0.005)	-0.004 (0.004)	-0.004 (0.004)	0.025 (0.018)	0.022 (0.016)	0.022 (0.016)
<i>Panel D</i>						
	Ancestral-country diversity of patent			Surname diversity of patent		
Ancestral-country diversity	-0.015 (0.026)	-0.017 (0.025)		-0.055* (0.030)	-0.056* (0.030)	
Team size fixed effects	✓	✓	✓	✓	✓	✓
County fixed effects	✓	✓		✓	✓	
State-Period fixed effects	✓	✓		✓	✓	
Patent technology field-Period fixed effects		✓	✓		✓	✓
County-Period fixed effects			✓			✓
Observations	1,451,459	1,451,459	1,451,459	1,451,459	1,451,459	1,451,459

Notes: The table reports least squares estimates for versions of the specifications described in equation (10), replacing surname diversity of patent with ancestral-country diversity of patent in Panel A and with county-level ancestral-country diversity in Panels C and D. An observation is a patent with inventors residing in a given county in a five-year period from 1900 to 1940. If there are multiple inventors who reside in different counties, the patent is assigned to all counties and the observation is weighted by 1 / number of inventors. The dependent variables are breakthrough patent indicator in Columns 1 to 3 and the number of technology classes per patent in Columns 4 to 6. In Panel D, the dependent variables are ancestral-country diversity of patent in Columns 1 and 2 and surname diversity of patent in Columns 4 and 5. Columns 2 and 5 add six broad patent technology fields (chemical, computers, drugs, electrical, mechanical and 'other') fixed effects interacted with period fixed effects. Columns 3 and 6 additionally include country-period fixed effects. Standard errors are clustered at the county level. All diversity variables are standardized to have a mean of zero and a unit variance. The sources and construction of all variables are detailed in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B51: Surnames clustering (full sample)

	Years of schooling	Occupation				Country of birth		Region of birth	USPTO tech class	
Sample:		All		Immigrants				Germans	Inventors	
Year:	1940	1880	1940	1880	1940	1880	1940	1880	1880-9	1940-9
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)
Surname	0.054	0.117	0.045	0.097	0.065	0.393	0.227	0.268	0.092	0.068
Country of birth	0.055	0.120	0.043	0.096	0.060					
Country of ancestry	0.051	0.112	0.041	0.087	0.048				0.004	0.003
U.S. county of residence	0.065	0.171	0.071	0.153	0.080	0.288	0.130	0.269	0.014	0.017

Notes: This table reports normalized Herfindahl indices, where larger values indicate greater concentration. The indices are calculated as the average Herfindahl indices of the variable in the header computed for each value of the variables on the left. For example, we calculate the normalized Herfindahl index of occupations for each surname and then average over all surnames using the number of individuals with a given surname as weights. Column 2 (1 and 3) includes all individuals in the 1880 (1940) census. Columns 4 and 6 (5 and 7) include all immigrants in the 1880 (1940) census. Column 8 restricts the sample to German immigrants in 1880. We use the 32 subnational regional origins for German immigrants recorded by the Census (the variable `bp1d` with codes 45301 to 45362). Column 9 (10) includes all inventors of patents issued from 1880 to 1889 (1940 to 1949). We use the main technology class on the patent.

Table B52: Effects on proxy measures of informational and social-behavioral mechanisms (birth-country diversity)

	Occupational diversity		USPTO tech codes diversity		Residential segregation of surname groups		Residential segregation of ancestral-country groups		Strength of family ties	
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)
Surname diversity	0.301*** (0.079)	0.292*** (0.086)	0.303*** (0.051)	0.312*** (0.057)	-0.387*** (0.080)	-0.361*** (0.086)	-1.079*** (0.162)	-1.148*** (0.145)	-0.202** (0.086)	-0.146** (0.069)
Birth-country diversity		0.050 (0.068)		-0.047 (0.043)		-0.133 (0.121)		0.373*** (0.058)		-0.392*** (0.059)
Population	-0.050* (0.029)	-0.053* (0.028)	0.059 (0.036)	0.061 (0.037)	0.044 (0.031)	0.049 (0.030)	0.091** (0.043)	0.071* (0.037)	-0.088 (0.057)	-0.068 (0.050)
County fixed effects	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓
State-Period fixed effects	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓
County-specific linear trends	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓
Observations	21,871	21,871	10,353	10,334	13,727	13,667	13,727	13,667	21,968	21,864

Notes: The table reports OLS estimates for the the specification described in equation (2) but with occupational diversity, USPTO tech codes diversity, residential segregation of surname groups, residential segregation of ancestral-country group and strength of family ties as the outcome variables. Even columns add birth-country diversity to the specification. An observation is a county in a period from 1900 to 1940. Standard errors are clustered at the state level. All variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B53: First-stage estimates of spillover analysis

	(1)	(2)	(3)	(4)	(5)	(6)
	Surname diversity			Surname diversity (< 100 miles)		
Predicted surname diversity	0.719*** (0.072)	0.718*** (0.072)	0.718*** (0.072)	0.018** (0.008)	0.017** (0.008)	0.017** (0.008)
Predicted surname diversity (< 100 miles)	0.077*** (0.021)	0.074*** (0.020)	0.073*** (0.020)	0.332*** (0.039)	0.327*** (0.038)	0.328*** (0.038)
Predicted surname diversity (100 < 200 miles)		0.037* (0.019)	0.036* (0.019)		0.052** (0.023)	0.054** (0.024)
Predicted surname diversity (200 < 300 miles)			-0.017 (0.029)			0.017 (0.028)
Predicted population	-0.138** (0.062)	-0.138** (0.062)	-0.138** (0.062)	0.003 (0.007)	0.004 (0.007)	0.004 (0.007)
Predicted population (< 100 miles)	0.047 (0.031)	0.052 (0.035)	0.053 (0.033)	0.102*** (0.034)	0.107*** (0.035)	0.105*** (0.036)
Predicted population (100 < 200 miles)		0.013 (0.031)	0.015 (0.032)		0.012 (0.035)	0.010 (0.037)
Predicted population (200 < 300 miles)			0.012 (0.030)			-0.013 (0.022)
	Surname diversity (100 < 200 miles)			Surname diversity (200 < 300 miles)		
Predicted surname diversity		-0.004 (0.007)	-0.004 (0.006)			-0.017 (0.011)
Predicted surname diversity (< 100 miles)		0.037* (0.019)	0.039** (0.019)			0.002 (0.017)
Predicted surname diversity (100 < 200 miles)		0.199*** (0.038)	0.201*** (0.038)			0.036 (0.030)
Predicted surname diversity (200 < 300 miles)			0.038* (0.019)			0.203*** (0.048)
Predicted population		-0.003 (0.006)	-0.003 (0.006)			0.007 (0.008)
Predicted population (< 100 miles)		-0.035** (0.017)	-0.032* (0.019)			-0.039** (0.016)
Predicted population (100 < 200 miles)		0.112** (0.044)	0.116** (0.045)			-0.011 (0.030)
Predicted population (200 < 300 miles)			0.016 (0.028)			0.151*** (0.049)
	Population			Population (< 100 miles)		
Predicted surname diversity	-0.106** (0.040)	-0.107** (0.040)	-0.107** (0.040)	0.004 (0.003)	0.004 (0.003)	0.004 (0.003)
Predicted surname diversity (< 100 miles)	0.021** (0.010)	0.018* (0.009)	0.018* (0.010)	-0.046** (0.019)	-0.048** (0.019)	-0.048** (0.019)
Predicted surname diversity (100 < 200 miles)		0.034*** (0.011)	0.034*** (0.011)		0.000 (0.011)	0.000 (0.011)
Predicted surname diversity (200 < 300 miles)			-0.001 (0.013)			0.003 (0.009)
Predicted population	0.552*** (0.107)	0.553*** (0.107)	0.553*** (0.107)	0.002 (0.004)	0.002 (0.004)	0.002 (0.004)
Predicted population (< 100 miles)	0.029 (0.028)	0.032 (0.026)	0.033 (0.027)	0.741*** (0.026)	0.749*** (0.025)	0.746*** (0.025)
Predicted population (100 < 200 miles)		0.007 (0.023)	0.009 (0.023)		0.036* (0.019)	0.033 (0.022)
Predicted population (200 < 300 miles)			0.008 (0.013)			-0.015 (0.018)
	Population (100 < 200 miles)			Population (200 < 300 miles)		
Predicted surname diversity		0.001 (0.002)	0.001 (0.002)			0.000 (0.001)
Predicted surname diversity (< 100 miles)		-0.004 (0.007)	-0.004 (0.007)			-0.002 (0.006)
Predicted surname diversity (100 < 200 miles)		-0.054** (0.025)	-0.055** (0.025)			-0.003 (0.007)
Predicted surname diversity (200 < 300 miles)			0.000 (0.009)			-0.043** (0.021)
Predicted population		0.003 (0.002)	0.003 (0.002)			0.000 (0.002)
Predicted population (< 100 miles)		0.006 (0.011)	0.009 (0.010)			-0.001 (0.009)
Predicted population (100 < 200 miles)		0.732*** (0.022)	0.736*** (0.023)			-0.014 (0.017)
Predicted population (200 < 300 miles)			0.017 (0.026)			0.793*** (0.024)
County fixed effects	✓	✓	✓	✓	✓	✓
State-Period fixed effects	✓	✓	✓	✓	✓	✓
County-specific linear trends	✓	✓	✓	✓	✓	✓
Observations	21,942	21,942	21,942	21,942	21,942	21,942

Notes: The table reports first-stage estimates of the IV regressions reported in Table 9. Standard errors are clustered on states and reported in parentheses. All variables are standardized to mean zero and unit variance. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

Table B54: Alternative disambiguation of inventors

	Patents (mean = 2.24, sd = 1.69)		Breakthrough patents (mean = 0.26, sd = 0.44)	
	(1)	(2)	(3)	(4)
<i>Panel A: Least-squares estimates</i>				
Surname diversity	0.245*** (0.053)	0.248*** (0.053)	0.047* (0.027)	0.048* (0.027)
Population		0.041 (0.061)		0.008 (0.016)
<i>Panel B: Reduced-form estimates</i>				
Predicted surname diversity	0.065*** (0.022)	0.071*** (0.023)	0.010 (0.007)	0.012* (0.006)
Predicted population		0.049 (0.044)		0.011 (0.014)
<i>Panel C: Instrumental-variable estimates</i>				
Surname diversity	0.395** (0.167)	0.403** (0.166)	0.062 (0.046)	0.066 (0.043)
Population		0.013 (0.045)		0.005 (0.016)
Sanderson-Windmeijer <i>F</i> -stat	11	18; 453	11	18; 453
<i>Panel D: First-stage estimates</i>				
	Surname diversity		Population	
Predicted surname diversity	0.163*** (0.049)	0.176*** (0.049)		0.018 (0.047)
Predicted population		0.094*** (0.019)		0.878*** (0.110)
Inventor fixed effects	✓	✓	✓	✓
Period-State fixed effects	✓	✓	✓	✓
Observations	218,618	218,618	218,618	218,618

Notes: The table reports OLS, reduced-form and IV estimates for the specifications described in equation (9). Inventors are assigned based on the first letter of the first name and county of residence. An observation is an inventor residing in a given county in a period from 1900 to 1940. The endogenous variables are county-level surname diversity and population in year t , and the dependent variables are number of patents filed by the inventor in the five-year period starting in t (column 1–2) and number of breakthrough patents filed by the inventor in the five-year period starting in t (column 3–4). Standard errors are clustered at the county level. The independent variables are standardized to mean zero and unit variance. The sources and construction of all variables are explained in Appendix A. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

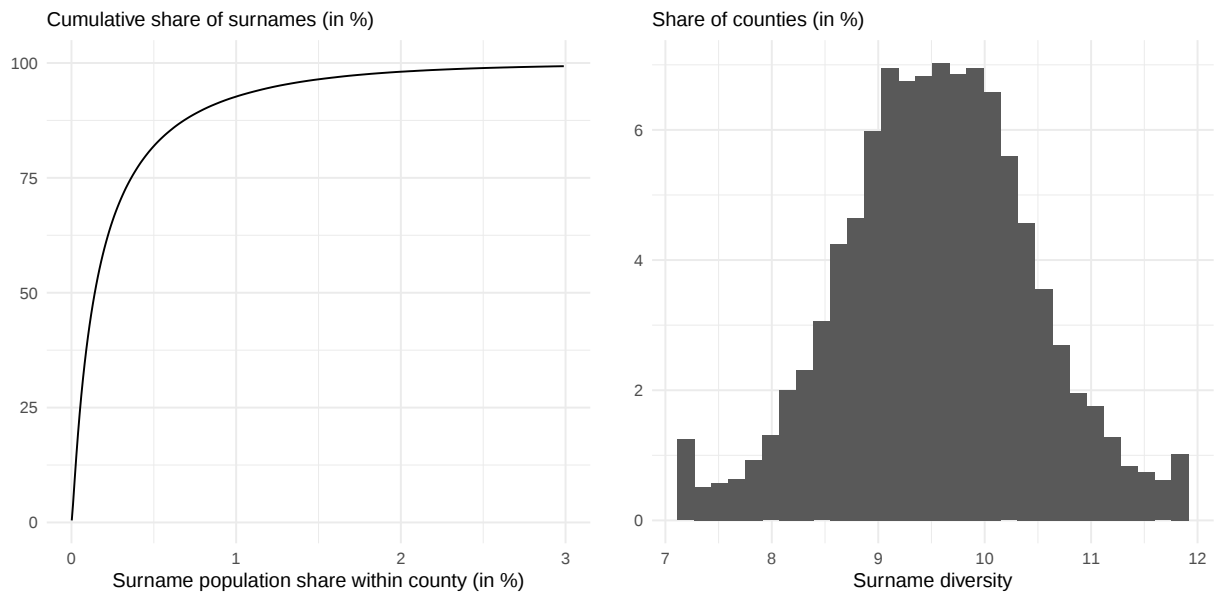


Figure B1: Distribution of surname population shares within counties and surname diversity

Notes: The figures show the distribution of surname population shares within counties (1850–1940, left plot) and surname diversity (1900–1940, right plot).

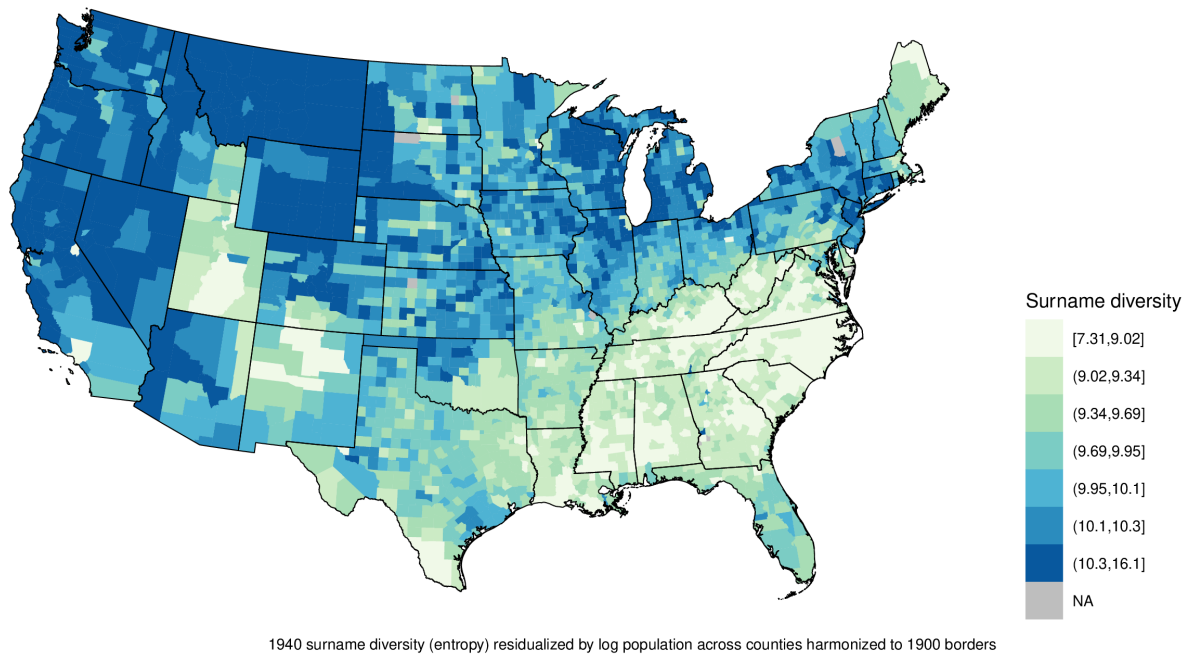
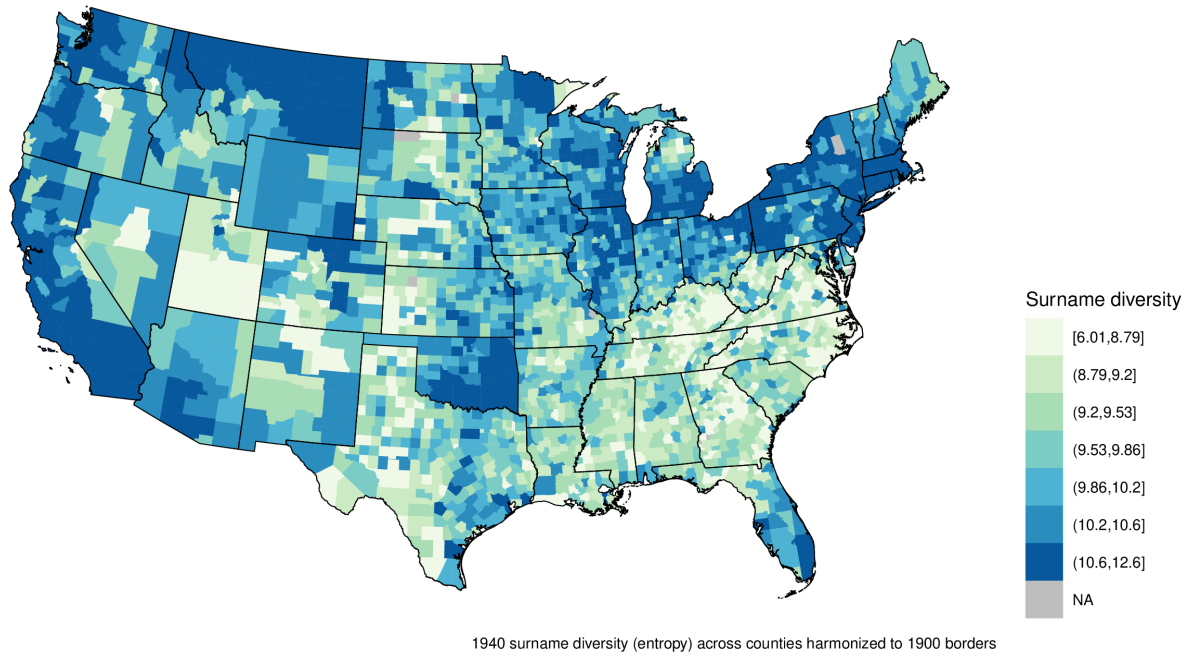


Figure B2: Surname diversity in 1940

Notes: The figures show the geographic variation in surname entropy in 1940 across counties harmonized to 1900 borders. Top: Raw data; bottom: residualized by log county population.

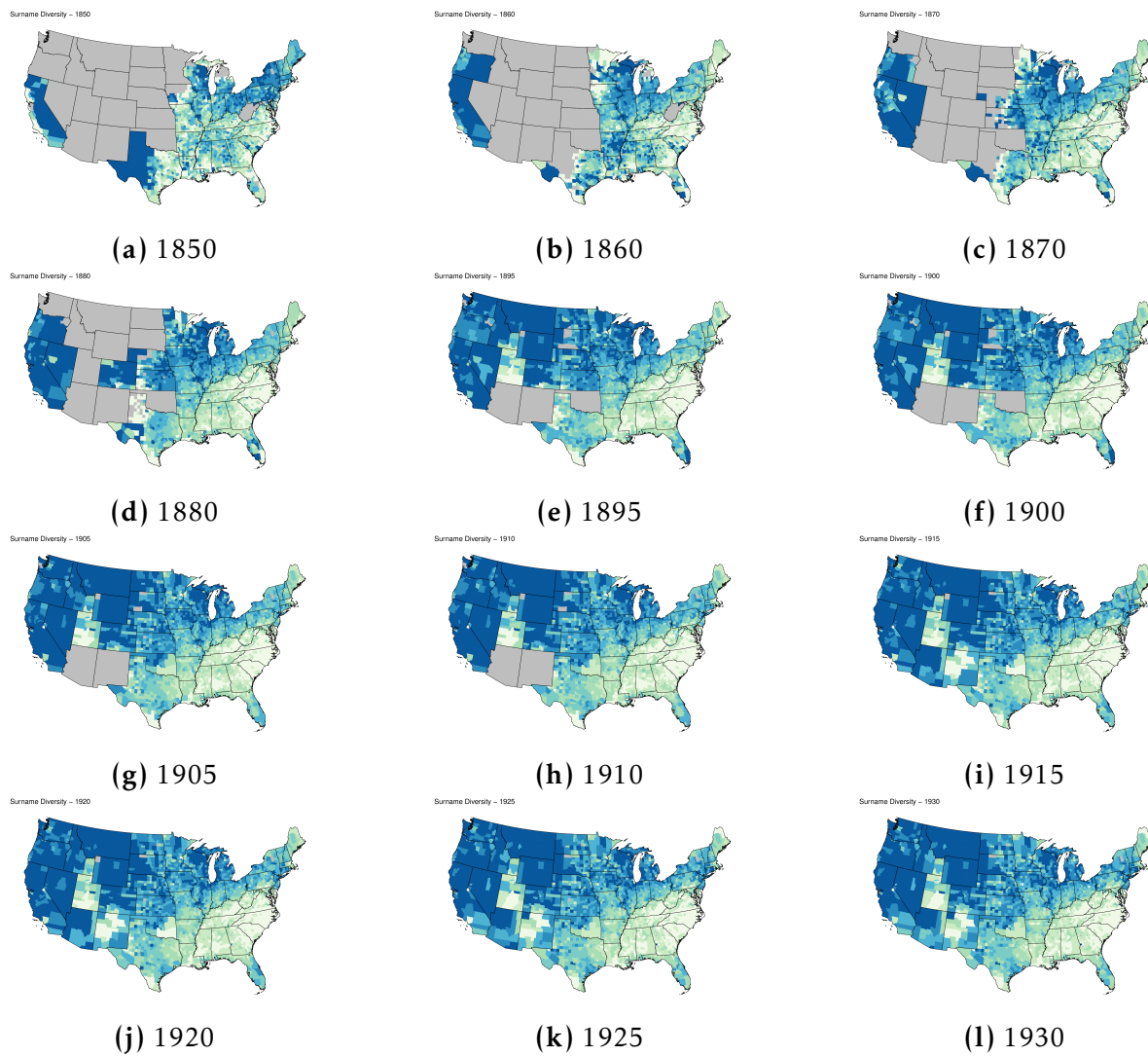


Figure B3: Surname diversity from 1850 to 1930

Notes: The figure shows surname entropy residualized by log county population in the respective year across counties harmonized to 1900 borders.

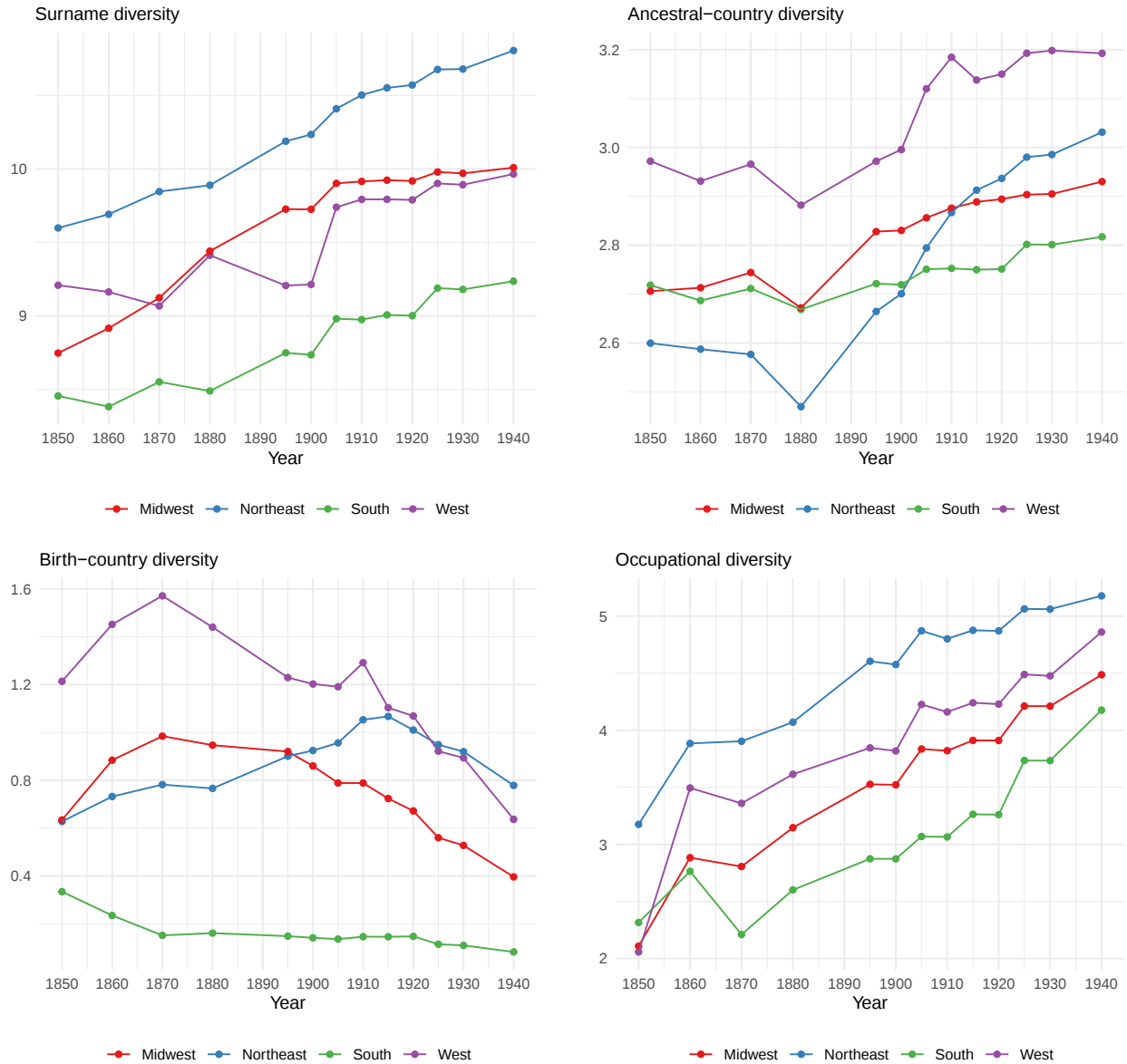


Figure B4: Surname diversity and other diversities from 1850 to 1940 across US census regions

Notes: The figure shows mean surname, ancestral-country, birth-country and occupational entropy in the respective year across census regions.

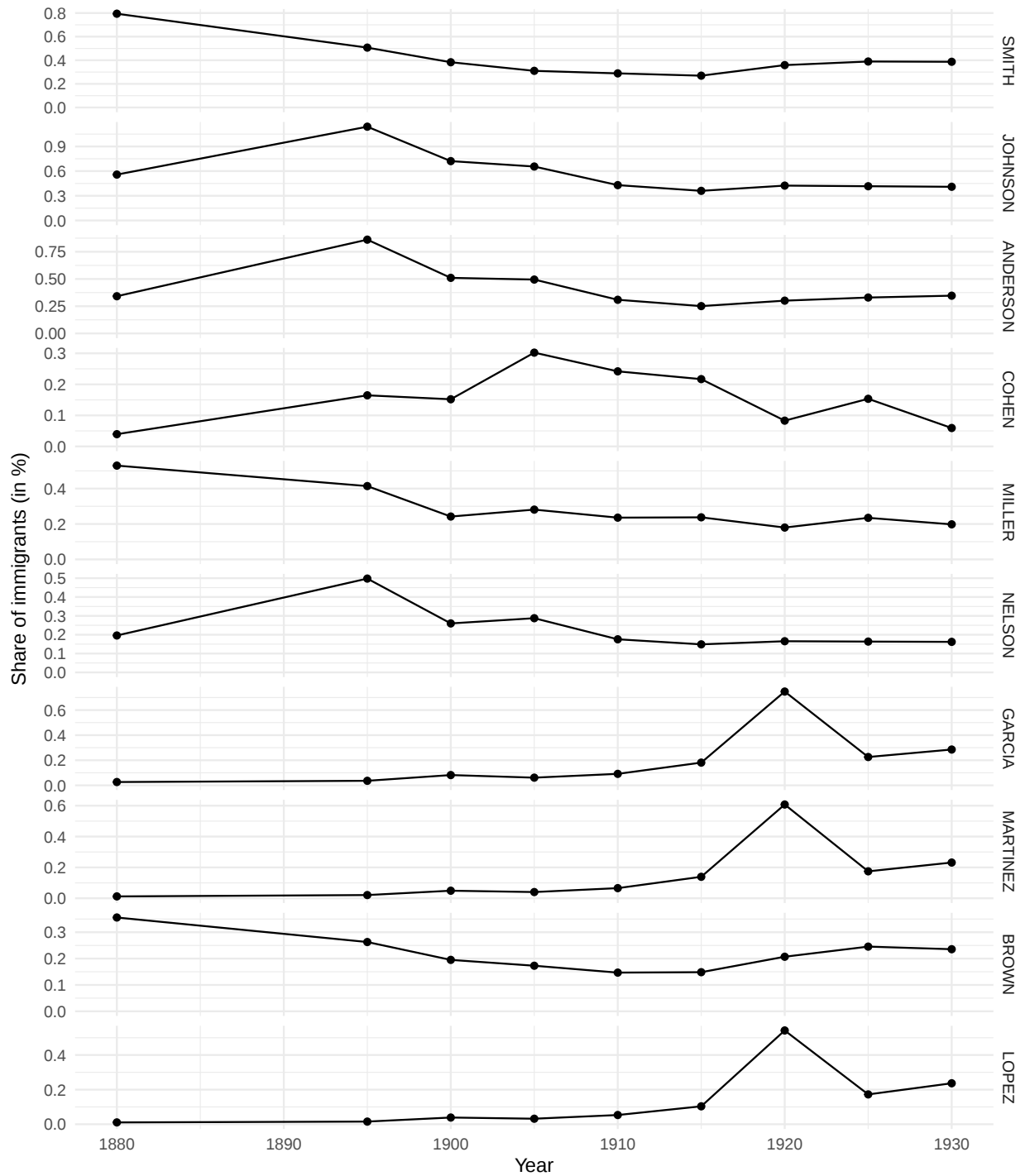


Figure B5: Share immigrants to the US by Surname

Notes: This figure illustrates the share of immigration into the US from the 10 surname groups with the highest immigration rates during specified periods: Smith, Johnson, Anderson, Cohen, Miller, Nelson, Garcia, Martinez, Brown, and Lopez. For each period, the most frequent surname was selected. If a top surname had already been selected for a previous period, the next most frequent surname was chosen, and so forth. The figure highlights variation in the push factors, illustrating how the number of migrants with each given surname changes over time.

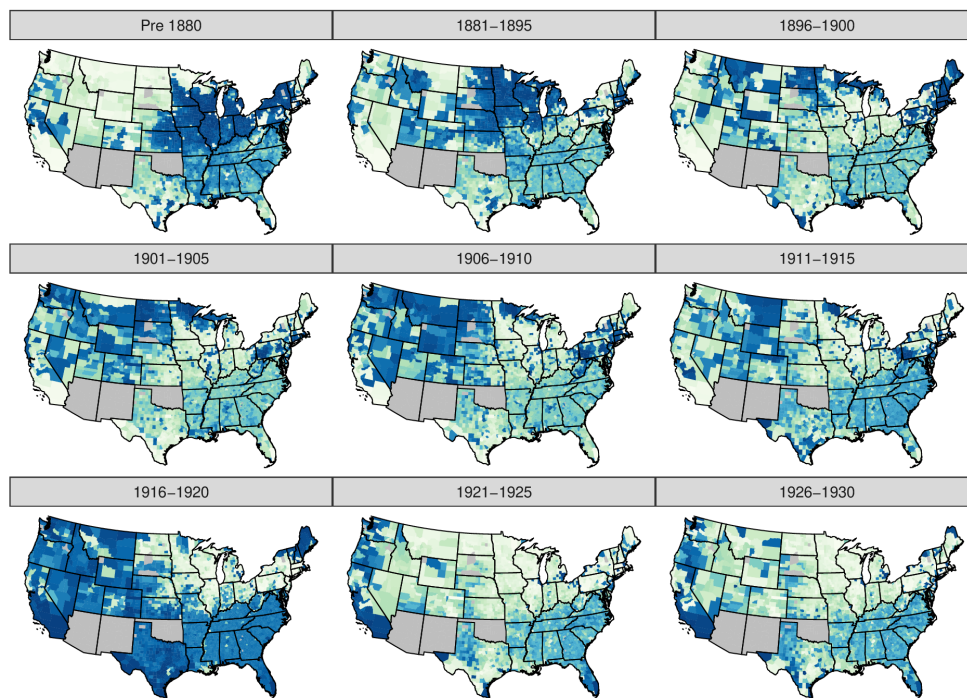
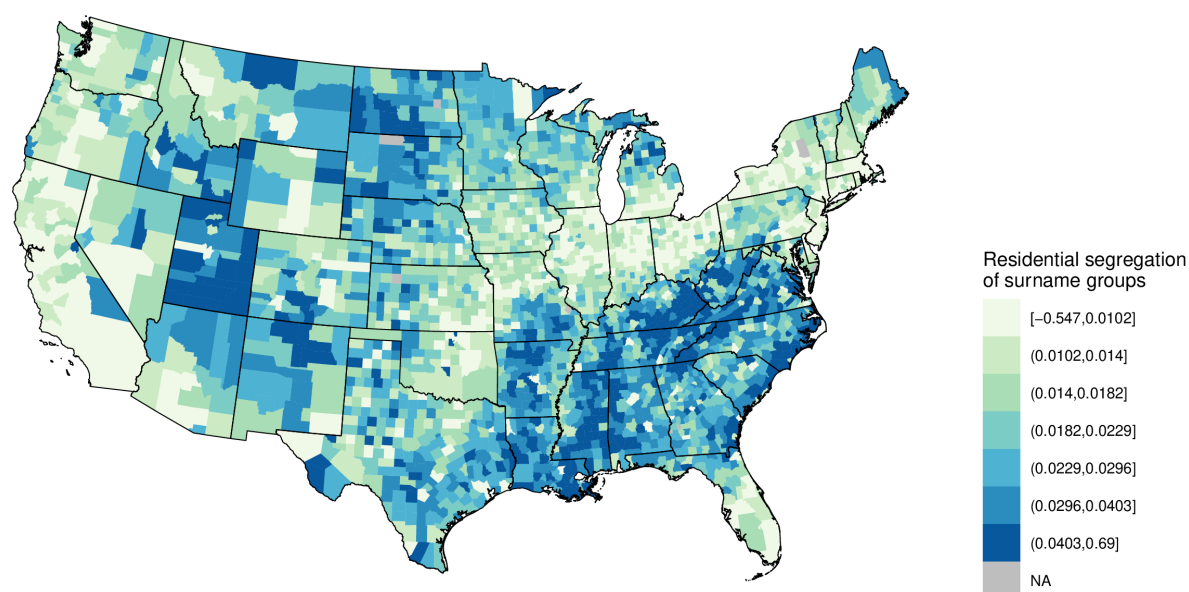
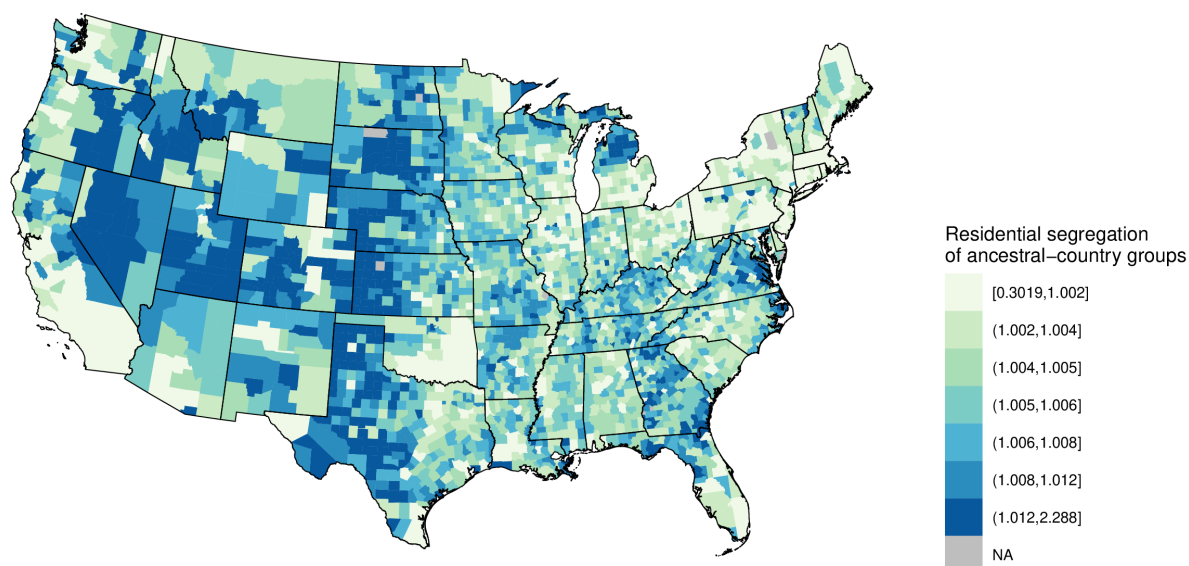


Figure B6: Destinations of immigrants to the US

Notes: This figure maps immigration flows into US counties by 5-year periods (except between 1880 and 1895). We regress the share of immigrants settling in US county i at time t , I_i^t , on destination county i and year t fixed effects, and calculate the residuals. The color coding depicts the 200 quantiles of the residuals across counties and within census periods. Darker colors indicate a higher quantile.



1940 residential segregation of surname group across counties harmonized to 1900 borders



1940 residential segregation of ancestral-country group across counties harmonized to 1900 borders

Figure B7: Residential segregation of surname and ancestral-country groups in 1940

Notes: The figures show the geographic variation in residential segregation of surname (top) and ancestral-country (bottom) groups across counties in 1940, harmonized to 1900 borders. The segregation measures are constructed adapting the [Logan and Parman \(2017\)](#) procedure to surnames and country of birth. The sources and construction of all variables are explained in Appendix A.

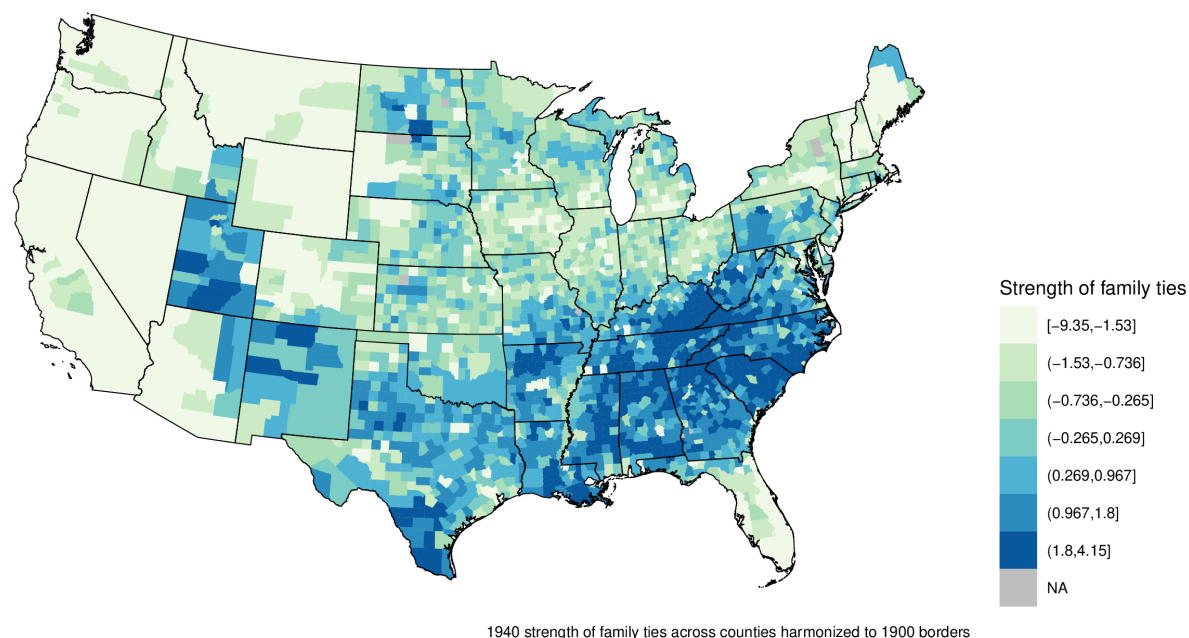


Figure B8: Strength of family ties in 1940

Notes: The figure shows the geographic variation in strength of family ties in 1940 across counties harmonized to 1900 borders. The measure is constructed following [Raz \(2025\)](#). The sources and construction of all variables are explained in Appendix A.

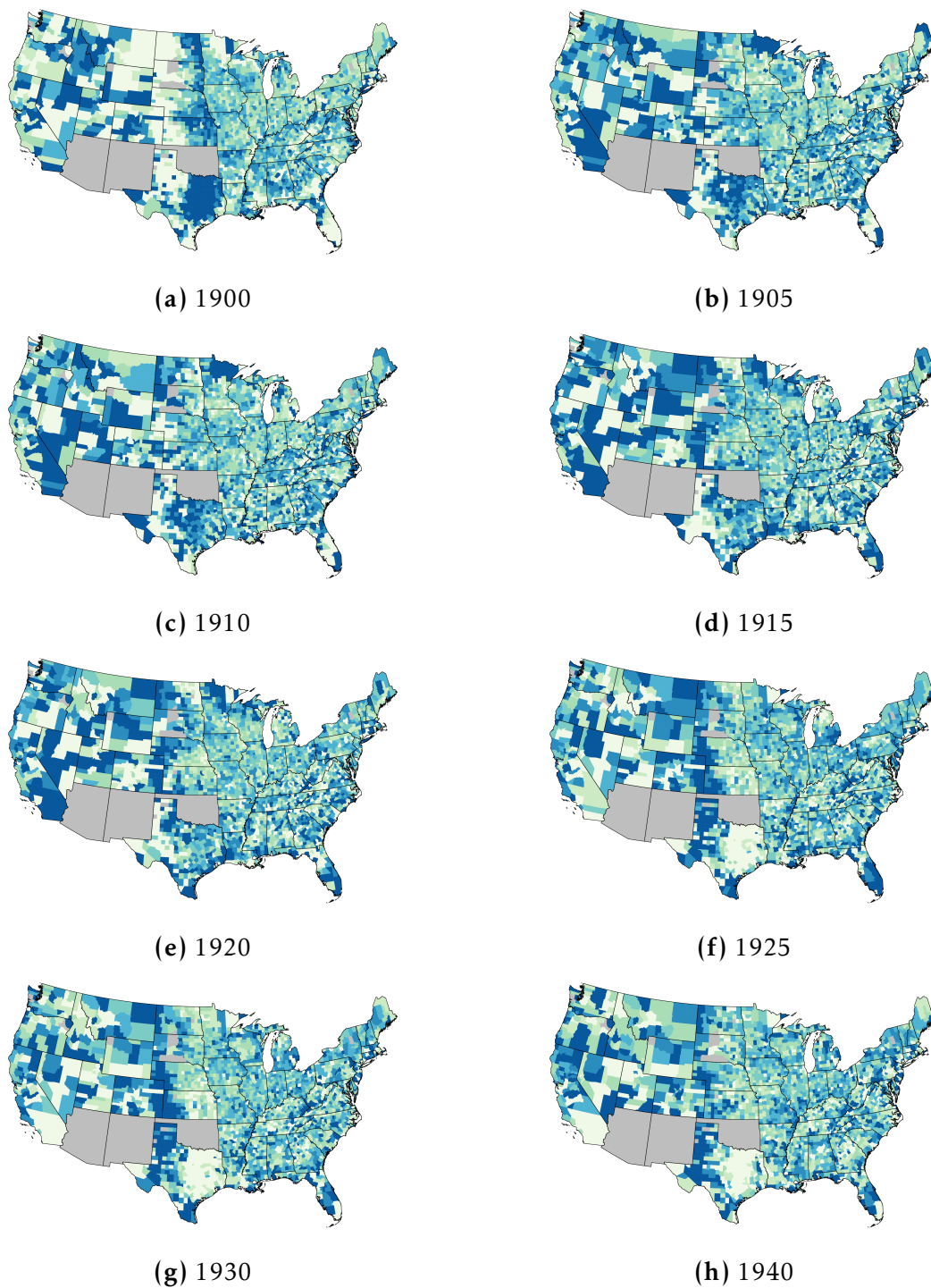


Figure B9: Predicted surname diversity (residuals)

Notes: This figure maps residualized instrumented surname diversity for each of the eight periods. We regress the instrument for surname diversity on the instrument for population and county and state-year fixed effects, and calculate the residuals. The color coding depicts 7 intervals across counties and within census periods, with darker colors indicating higher values. Grey indicates a lack of population data in 1900.

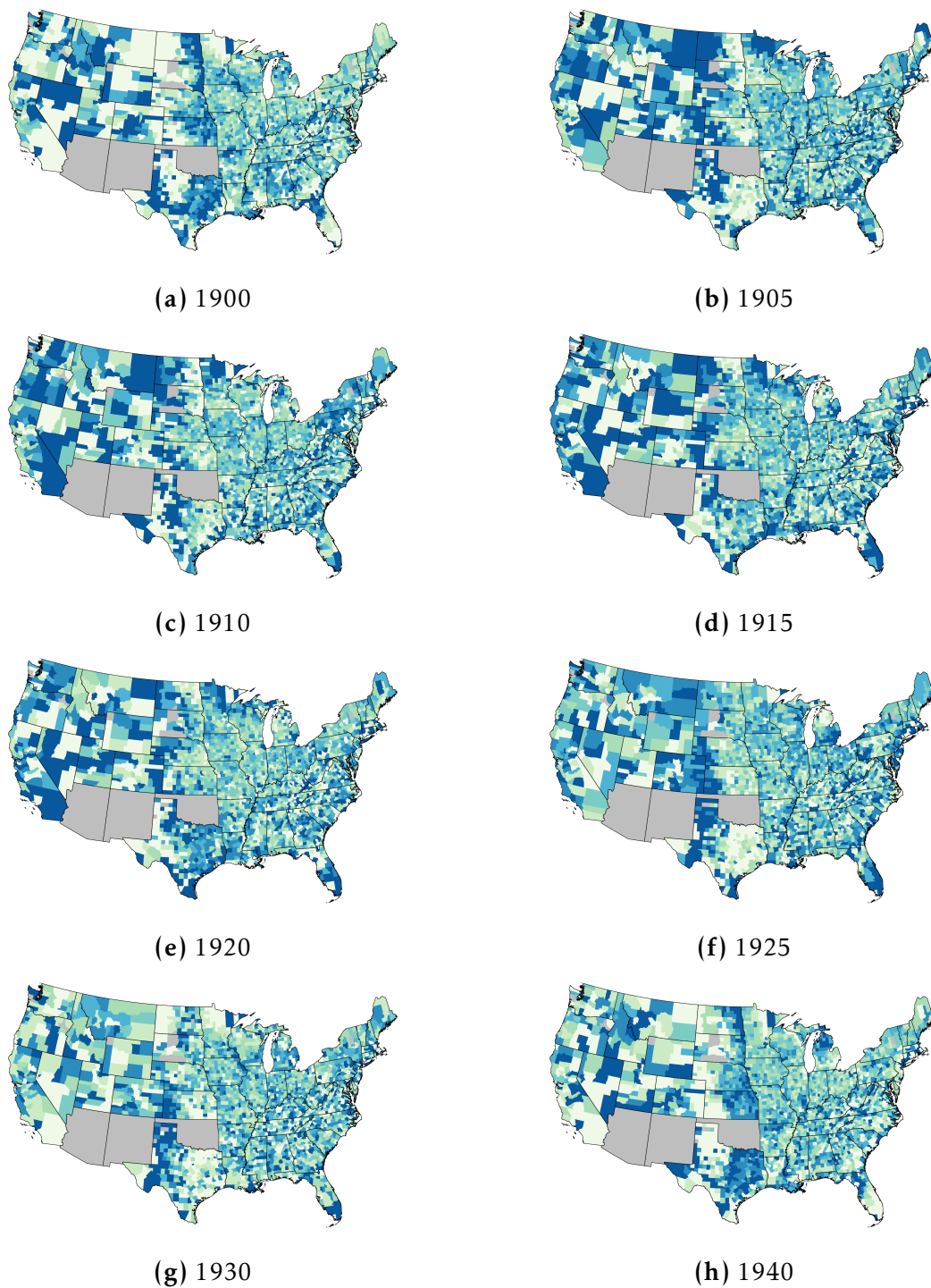


Figure B10: Predicted surname entropy (residuals)

Notes: This figure maps residualized instrumented surname entropy for each of the eight periods. We regress the instrument for surname entropy on the instrument for population, county and state-year fixed effects, and county specific linear time trends, and calculate the residuals. The color coding depicts 7 intervals across counties and within census periods, with darker colors indicating higher values. Grey indicates a lack of population data in 1900.

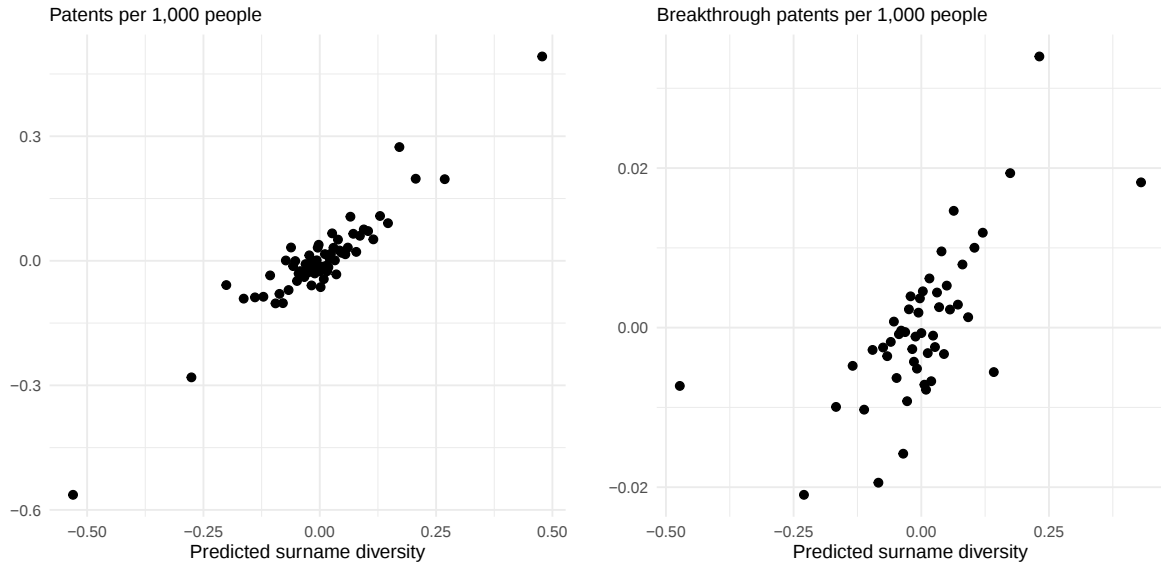


Figure B11: Reduced-form relationships of predicted surname diversity and innovation outcomes, additionally controlling for county-specific trends

Notes: County-level data from 1900 to 1940. Observations are residualized by predicted population, county fixed effects, state-period fixed effects and county-specific time trends.

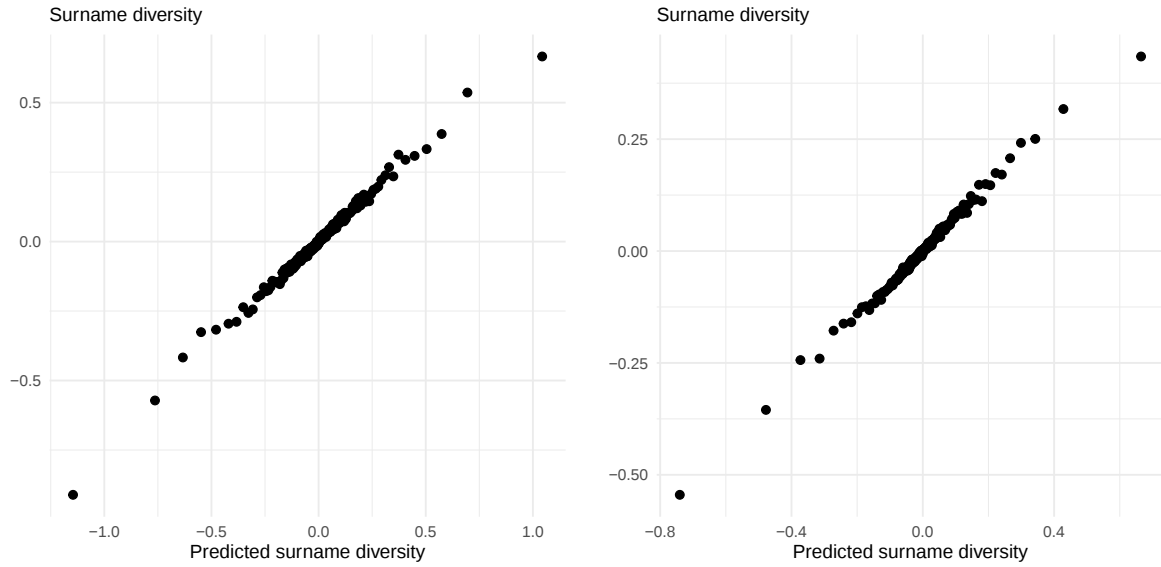


Figure B12: First Stage: Binned scatter plots of predicted surname diversity and actual surname diversity from 1900 to 1940

Notes: County-level data from 1900 to 1940 (including midyears). Observations are residualized by predicted population, county fixed effects and state-period fixed effects (left plot) and additionally county-specific trends (right plot).

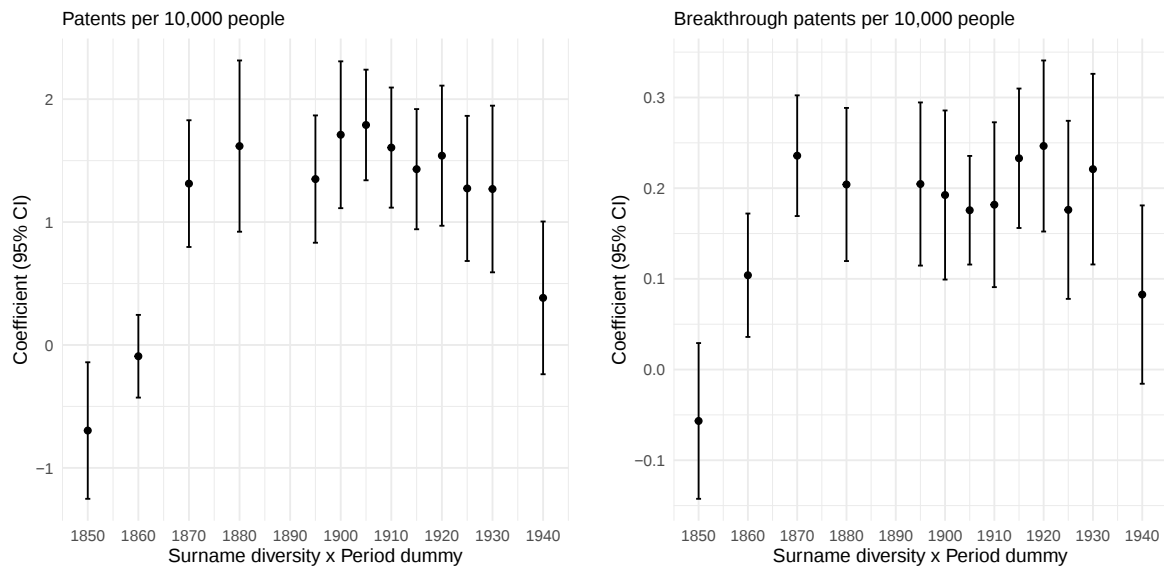


Figure B13: OLS estimates in years 1850-1940

Notes: The figures display coefficients from regressions of the number of patents and breakthrough patents normalized by county population in period t on surname diversity, interacted with period dummies, conditional on population size, county fixed effects and state-period fixed effects. Standard errors are clustered at the state level. Few patents were issued in 1850 compared to the latter years. Surname diversity is standardized to have a mean of zero and unit variance. The sources and construction of all variables are explained in Appendix A. The dependent variables are winsorized at the 5 per cent level. ***, **, and * indicate significance at the 1%, 5%, and 10% levels.

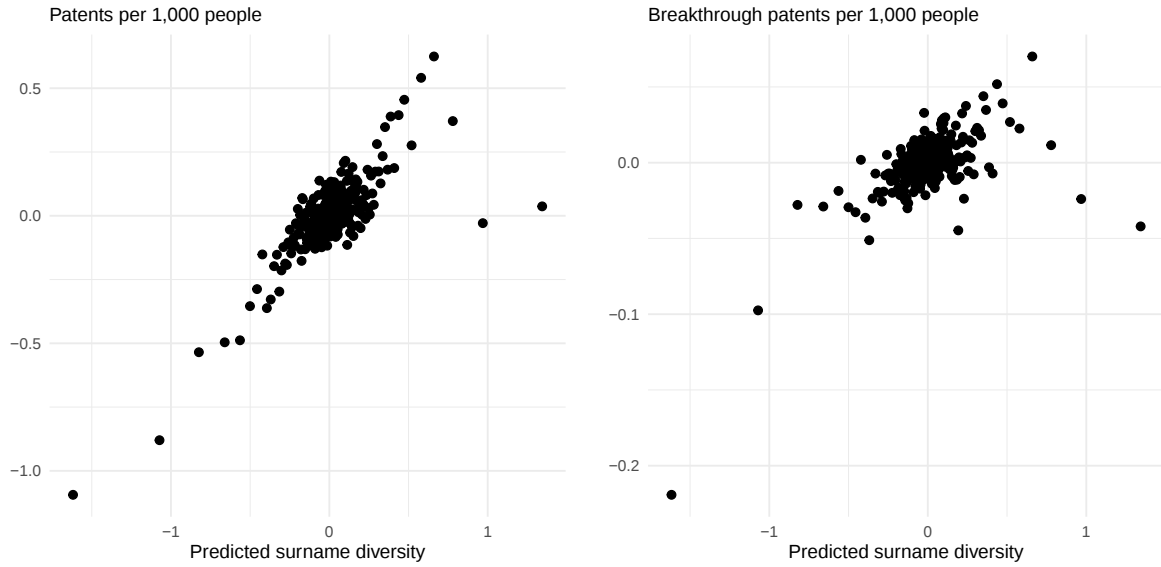


Figure B14: Reduced-from relationships of predicted surname diversity and innovation outcomes at the county-surname level

Notes: County-surname-level data from 1900 to 1940. Observations are residualized by predicted population, county-surname fixed effects and state-period fixed effects.

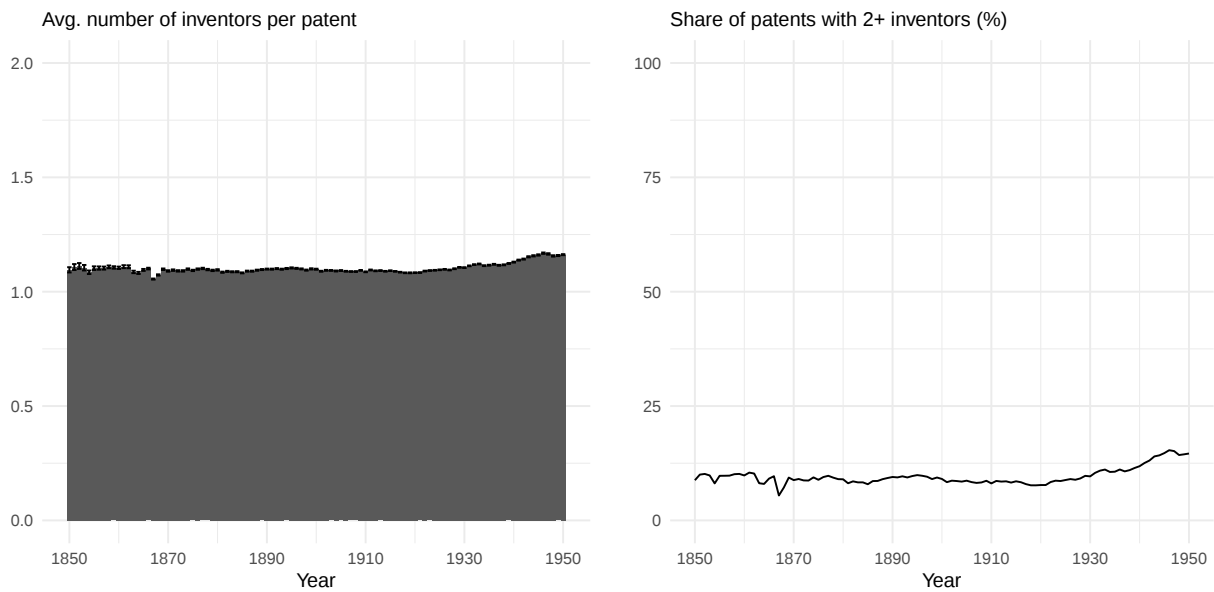


Figure B15: Patentees per patent over time

Notes: The left plot shows the average number of patentees per patent per year; The plot displays the share of patents with two or more patentees per year.