

Consumption Segregation^{*}

Corina Boar[†]

Elisa Giannone[‡]

March 2025

Abstract

This paper introduces consumption segregation as a new dimension of residential segregation, examines its patterns and determinants, and discusses its implications. Using novel longitudinal and highly granular data, we measure consumption segregation in the United States and document that, while it remains high, it has declined over the past 15 years and exhibits substantial regional variation. We find a strong correlation between consumption segregation and income segregation. We interpret it through the lens of a standard consumption model and argue that the correlation arises primarily due to incomplete markets that limit consumption insurance across income groups.

^{*}We thank Federico Kochen for superb research assistance, our discussants Rebecca Diamond, Eduardo Morales, and David Cutler for insightful feedback, as well as seminar participants at the ASSA Meetings, SED Meeting, the Federal Reserve Bank of Richmond, Dartmouth Mini Conference, CREI, OSUS, Barcelona Summer Forum, NBER SI Income distribution and NBER SI Urban for their helpful comments. Elisa Giannone acknowledges nancial support from the Spanish Agencia Estatal de Investigación (AEI), through the Severo Ochoa Programme for Centres of Excellence in RD (Barcelona School of Economics CEX2019-000915-S)”

[†]New York University and NBER, corina.boar@nyu.edu.

[‡]Center de Recerca en Economia Internacional and Barcelona School of Economics, egiannone@crei.cat.

1 Introduction

Over the past few decades, the United States has experienced a steady increase in economic inequality, accompanied by an increase in residential segregation by income and education. These trends have created pressing socio-economic challenges, leading to extensive research aimed at understanding their causes and consequences.¹ Although there is ample work on income and other forms of segregation, less is known about the segregation of consumption. This is important because consumption reflects individuals’ realized access to goods and services. Measuring consumption segregation can thus offer deeper insight into the spatial distribution of well-being and enhance our understanding of how economic disparities manifest in everyday life.

Despite its relevance, analyzing differences in consumption over time and across space remains challenging due to significant data availability constraints. In this paper, we leverage new longitudinal and highly granular consumer-level data to characterize patterns of residential consumption segregation, examine its determinants, and assess its potential effects on aggregate outcomes. We proceed in three steps. First, we provide the first systematic account of consumption segregation in the United States after validating the new data. Second, we analyze the determinants of consumption segregation and document a strong relationship with income segregation. Third, we draw on insights from a standard consumption model to argue that this relationship emerges as a consequence of incomplete financial markets that prevent consumption insurance. Based on these findings, we discuss potential implications for wealth inequality and policy measures to reduce consumption segregation.

We measure consumption segregation using the entropy and rank-order index, which evaluate how representative the consumption distribution within a geographic unit (e.g. a neighborhood) is relative to a broader region (e.g. a city). The index ranges from 0 to 1, where 0 indicates perfect integration—every narrow region has a consumption distribution identical to the broader region—while 1 indicates perfect segregation, meaning narrow regions are entirely homogeneous in their consumption levels. This measure is well-established in the segregation literature and outperforms alternative indices on key methodological properties.²

We begin our analysis by providing a thorough characterization of consumption segregation. We utilize two sources of data. First, we leverage a newly built, large-scale dataset from

¹A large body of work has studied the relationship between the two, emphasizing the role of human capital formation as a driver (Fernández and Rogerson, 1998; Fogli and Guerrieri, 2019; Zheng and Graham, 2020; Eckert and Kleineberg, 2021).

²See Reardon and Firebaugh (2002) for an in-depth comparison of segregation indices.

Infutor, a private company. This dataset covers nearly 250 million individuals, 160 million vehicles, and over 150 million properties. We use the extensive information on car and home ownership in this dataset, connected to demographic information, to analyze the extent of durable consumption segregation. Second, we complement it with Nielsen Homescan data to measure the segregation of non-service retail consumption, which we refer to as non-durable consumption for brevity throughout the paper. One key limitation of the data we use is that it does not include services. Therefore, our focus is on durable and non-service retail consumption, which together account for 51% of total household spending.

Since a secondary goal of our paper is to construct a comprehensive dataset on consumption differences across space and time that others can utilize, before measuring consumption segregation, we perform an extensive validation exercise on the Infutor data, which, to the best of our knowledge, has not been used to measure durable consumption. We find that it accurately reflects the demographic composition, as evidenced by comparisons with the American Community Survey (ACS). It also aligns well with various statistics on housing values, homeownership rates, and the vehicle fleet in the United States.

The Infutor dataset has multiple advantages. It has detailed geographic information which, combined with the large sample size, allows for a precise analysis of consumption differences across space. It has a relatively extended time coverage which allows us to explore trends in consumption differences. Additionally, it includes both values and characteristics of properties, as well as types of vehicles, thus allowing us to make a distinction between spending and consumption and bypass the lack of fine regional price data over time.

We highlight four main aspects of consumption segregation. First, we examine both the level of and trends in consumption segregation. We find that consumption segregation is as high as racial segregation. However, it has declined over the past 15 years, with a more pronounced decline in the segregation of non-durables, while the segregation of durables has remained relatively stable. Among all consumption categories, unsurprisingly, housing consumption is the most segregated. Second, we analyze regional differences in the degree of consumption segregation. We find substantial heterogeneity: consumption is most segregated in New York, where nearly 30% of the state-level diversity is not reflected in the average PUMA.³ In contrast, Wyoming—the state with the lowest consumption segregation—has an average PUMA that is nearly fully representative of the state. Third, we explore how consumption segregation varies with the socio-economic characteristics of regions. We find

³A PUMA is the smallest geographic unit available across all the datasets we use.

that it is highest in richer, more highly educated, and larger regions, as well as in those where a smaller percentage of the population is white. Fourth, we examine how consumption segregation varies with individuals’ demographic characteristics. Our findings indicate that most of the observed segregation in consumption occurs within, rather than across, demographic groups.

To further understand consumption segregation, we connect to a long line of research that examines segregation in the United States across various socio-economic dimensions, such as race, education, income and ethnicity. Specifically, we analyze the extent to which residential segregation along these previously documented dimensions shapes consumption segregation, thereby establishing a link between these factors and the residential segregation of well-being, as captured by consumption segregation. We find an important role for income segregation as a determinant of consumption segregation, and a more muted role for racial segregation. Drawing on insights from a standard consumption model and variation in consumption segregation across narrowly defined demographic groups, we argue that the strong correlation between income and consumption segregation arises because incomplete markets constrain consumption insurance among individuals with different income levels.

Lastly, we discuss the implications of our findings for understanding inequality more broadly. First, consumption segregation provides a nuanced perspective on economic disparity by capturing not only financial resources but also access to goods that shape quality of life, offering deeper insight into the spatial distribution of well-being. Second, consumption segregation may itself contribute to wealth inequality through social comparison effects. In less segregated areas, lower-income individuals may increase borrowing to match their wealthier neighbors’ consumption, reducing their wealth and exacerbating inequality. Conversely, in highly segregated areas, weaker social comparison pressures may mitigate this effect, though the prominence of visible goods in net worth measures could counterbalance this dynamic. Finally, our finding that market incompleteness reinforces the link between income and consumption segregation suggests policy interventions aimed at improving financial access—such as digital banking solutions in underserved areas—could help mitigate consumption segregation even in the presence of income segregation.

Related Work. By studying the segregation of consumption, we contribute to the vast literature on residential segregation that studies trends in segregation by race ([Cutler and Glaeser, 1997](#); [Cutler et al., 1999, 2008](#)) or by income ([Reardon and Firebaugh, 2002](#)). We refer the reader to [Trounstine \(2018\)](#) for a comprehensive description of this work. We

contribute to this literature by moving beyond income and studying consumption segregation.

Our paper also relates to recent work studying consumption differences across space. [Agarwal et al. \(2017\)](#) study how consumers’ willingness to travel shapes the characteristics of industries that deliver final consumption. [Davis et al. \(2017\)](#) use Yelp reviews to study the role of spatial and social frictions in influencing restaurant visits in New York City. [Allcott et al. \(2019\)](#) study the extent to which the presence of supermarkets affects nutritional inequality. [Diamond and Moretti \(2022\)](#) provide evidence on consumption differences across individuals and commuting zones and connects income and consumption inequality through sorting across space. Lastly, [Couture et al. \(2025\)](#) find large disparities across groups in their exposure to high-income co-patrons at restaurants and estimate preferences for the racial and income composition at such establishments. We complement these studies in three main dimensions. First, we focus on the segregation of durable consumption categories, as well as that of non-service retail consumption for the entire country. Second, we explore the time dimension of spatial differences in consumption. Third, the large and detailed scale of our data allows to compare consumption segregation to other forms of segregation and understand its determinants.

Lastly, by emphasizing the role of income segregation as a driver of consumption segregation, our paper is also related to papers that study the relationship between income and consumption inequality over time: [Krueger and Perri \(2006\)](#), [Blundell et al. \(2008\)](#), [Fisher et al. \(2013\)](#), [Aguiar and Bils \(2015\)](#). Differently from these papers, we focus on segregation as the geographical element of inequality.

The remainder of the paper proceeds as follows. Section 2 describes the data sources used in the analysis. Section 3 documents the patterns of consumption segregation. Section 4 analyzes the drivers of consumption segregation. Finally, section 5 concludes.

2 Data and Validity

In this section, we introduce the data we use in our analysis. In addition to learning about consumption segregation, an objective of our paper is to construct a comprehensive dataset on consumption across space and time that others can utilize. As a result, a substantial portion of this section focuses on explaining and validating the data and the construction of our primary variables of interest.

Our empirical assessment of consumption segregation and its relationship with other dimensions of residential segregation is based on three datasets that enable us to observe the

income and consumption choices of individuals, their demographic characteristics, and track their locations precisely. Specifically, we analyze the consumption segregation of durables such as cars and houses using Infutor data, non-service retail consumption segregation using Nielsen Homescan data, and residential segregation by income and demographic characteristics using data from the American Community Survey (ACS). We provide a detailed discussion of these datasets below, with particular emphasis on Infutor data, which is less commonly used, especially in the study of durable consumption. For validation, we compare key summary statistics from Infutor with counterparts from the ACS and the Bureau of Transportation Statistics (BTS).⁴

2.1 Infutor Data

This dataset is compiled by Infutor, a private data vendor based in Chicago that aggregates publicly available information on individuals and makes it available for purchase to both researchers and business.⁵ Infutor uses multiple data sources: phone books, voter files, credit header files, public government records, property deeds, county property records, vehicle warranties, and data from vehicle repair and maintenance providers. Most of these sources are publicly available. The key innovation is that Infutor developed an algorithm to identify individuals across these different datasets and link them to each other.

The resulting data contains information on more than 250 million individuals, who are linked to 160 million vehicles and more than 150 million properties spanning from 2000 to 2018. This dataset includes demographic information for consumers (gender, age, marital status), vehicle information (make, model, year), household income, property information (value, construction quality), length of residence, type of dwelling, geographic delineations (time zone, county), and address information (street address, state, city, zip). Identifiable information within this dataset includes Social Security Numbers, Vehicle Identification Numbers, date of birth, and name. A noteworthy feature of this dataset is that it allows us to explore some dimensions of the quality of goods, which arguably offers insight into effective consumption beyond spending.

The information it contains and the size of the sample make for an unprecedented dataset. Additionally, it contains person identifiers that could facilitate merging with other datasets, thus increasing its value. While other studies have used the component on the history of addresses from Infutor, to the best of our knowledge, this is the first study that employs the

⁴In Appendix A, we provide further details about these datasets.

⁵Penn State University purchased these data in 2019.

vehicle, property, and demographic components together. As a result, we devote a significant part of our analysis to testing the validity and representativeness of the data. In the following sections, we describe the structure of the dataset and discuss its representativeness, comparing it with data from the Census, the American Community Survey (ACS), and the Bureau of Transportation Statistics (BTS). We also discuss its advantages and shortcomings. Appendix A contains a more detailed description of our treatment of the data.

2.1.1 Structure of the Infutor Data

The Infutor data are organized into four linkable files: Auto Profiles, Property Profiles, Demographic Profiles, and History of Addresses. In this project, we primarily use the first three but describe the fourth for completeness. We have access to different snapshots of these data. From 2012 to 2018, we observe the data at an annual frequency. After March 2019, we observe the data at a monthly frequency. For housing, the first available snapshot appears in 2015. Although the snapshots start in either 2012 or 2015, ownership records and the history of addresses go much further back in time. For example, we observe ownership records of vehicles starting in the 1980s, while property deed records go back to the 1950s. The longitudinal aspect of the data enables us to observe the history of car ownership and residences as well as study trends in consumption segregation across both time and regions. We next describe the content of the Infutor data files that are most relevant to our paper.⁶

Auto Profiles. This file provides information on individual car ownership, including characteristics such as the car’s manufacturer (make), model, manufacturing year, and the first 10 digits of the Vehicle Identification Number (VIN). It also includes the owner’s address, allowing us to identify the vehicle’s location and certain demographic and economic characteristics. This data file contains approximately 160 million observations during any given period (i.e., November 2012–2018 and every month of 2019), which we use to measure the segregation of car consumption.

In Section 3.1, we compute two measures of segregation: one along categorical dimensions (e.g. car model) and one along continuous dimensions (e.g. car value). To measure car values in the Infutor data, we merge these data with transaction prices for all dealer sales of new and used cars in Texas from 2012 to 2019, based on the VIN. When the VIN is not available, we infer prices using the average transaction price of a make-model-year combination.

⁶Figure 6 in the Appendix contains a visual summary of this section.

While we observe Auto Profile snapshots starting in 2012, information on vehicle ownership records dates back to the 1980s. With additional assumptions, this allows us to study a more extended time dimension of segregation. Specifically, for years prior to the first time we observe a vehicle, we set the location of the vehicle to the ZIP code of the first snapshot in which we observe it. We are confident that this does not introduce substantial noise in our measurement, as between 2012 and 2018, only 3.3% of the vehicles in the sample changed Public Use Microdata Areas (PUMAs), while only 1.1% changed states. For years prior to 2012, we impute the value of a vehicle using a VIN-level annual depreciation rate. Appendix A contains a detailed discussion of these procedures.

Property Profiles. This file contains information on residential properties, including each home’s address, value, and characteristics such as the year it was built, square footage, and building quality. The dataset includes approximately 150 million observations at any given time, and we focus only on properties used for residential purposes rather than commercial ones. We use this information to measure housing consumption segregation.

The primary dimension of housing segregation we study is home value. To measure the value of a home, we use property deeds data, which provide information on the date and price at which a property was acquired by its current owner. We must note that we track properties rather than owners, as our objective is to include renters in the measurement. While snapshots of Property Profiles are only available starting in 2015, property deed records date back as early as the 1950s, offering a long historical perspective. However, a home’s market value is only recorded when it is sold and a deed is issued. To better capture segregation trends, we use ZIP code-level price indices based on transacted homes to impute home values between two consecutive sales. Appendix A.2 shows that this imputation procedure produces house price levels and dynamics that closely align with those published by the Federal Housing Finance Agency.

Demographic Profiles. This file contains information on the demographic and economic characteristics of approximately 250 million individuals at any given time, along with their addresses. The file includes data on individuals’ age, gender, ethnicity, marital status, educational attainment, number of children, and estimated income, wealth, and home value. However, Infutor data do not contain information on race. To study the relationship between race, income, and consumption, we probabilistically assign race to individuals in the data using an algorithm proposed by the Census, which assigns race based on surnames from the

2010 Census Surname Table.⁷ We verify the accuracy of this imputation by comparing the race distribution in Infutor data with that in the ACS.

History of Addresses. This dataset contains information about individuals’ past and current addresses in the United States, including the time periods in which they lived there from 1990 to 2019. Although we do not utilize these data for our project, we mention them for completeness. [Diamond et al. \(2019\)](#) uses the detailed information in this dataset to examine the effects of rent control on relocation in San Francisco, while [Qian and Tan \(2020\)](#) investigates how the entry of new firms into a location impacts individuals’ choices.

2.1.2 Representativeness and Validity of the Infutor Data

Naturally, there is a concern about the representativeness of the Infutor dataset relative to the US population and specific demographic groups. Although the large number of observations suggests the reliability of the data, researchers such as [Bernstein et al. \(2018\)](#) and [Phillips \(2019\)](#) have systematically compared these data with Census data and found that they cover 78% of the adult population estimated by the US Census and match the cross-sectional distribution of the population across counties. In this section, we aim to answer two questions: *(i)* How representative is Infutor data of the overall US population and the different demographic groups we are interested in studying? and *(ii)* Is Infutor data suitable for studying questions related to durable consumption segregation across space and over time?

Demographics. To answer the first question, Table 1 presents a comparison of summary statistics for various demographic groups between the Infutor data and the ACS. Column 1 reports the distributions of gender, race, educational attainment, age, and household income in the ACS, while Column 2 provides the corresponding distributions in the Infutor data. To assess the spatial representativeness of Infutor data for the different demographic groups of interest, Columns 3 and 4 show the cross-state and cross-PUMA correlations of each demographic group’s share in each state and PUMA between the two datasets.

In the top panel of Table 1, we present the results for gender. The share of women in the Infutor data is 52%, which is very close to the 51% reported in the ACS. The cross-state

⁷We use file B from https://www.census.gov/topics/population/genealogy/data/2010_surnames.html. The Census reports last names that occurred more than 100 times in its collected 2010 data and the proportion of people with each last name who belong to each race category (White, Black, Asian/Pacific Islander, American Indian/Alaskan Native, bi/multi-racial, and Hispanic). We assigned the likelihood of belonging to each race group based on individuals’ last names using the Census race data. We were able to assign race probabilistically to 90.25% of the sample.

and cross-PUMA correlations of gender shares between the two datasets are 0.8 and 0.59, respectively, indicating a high level of correspondence.

In the second panel, we report the results for race. In the ACS data, 75% of individuals identify as White, 12% as Black, and the remaining 13% as belonging to other racial groups. In the Infutor data, the corresponding shares are 62%, 12%, and 26%, respectively. Although we use an imputation procedure to assign race probabilistically, the high correlations between the race shares in the two datasets for U.S. states and PUMAs suggest that Infutor accurately captures the distribution of racial groups across space. For example, the state-level correlation of the share of the White population is 0.71, while those for Black and other racial groups are 0.8 and 0.86, respectively. The accurate representation of the racial distribution across space is crucial for our analysis of the role of racial segregation in driving consumption segregation.

In the third panel, we report the shares of the population with and without a college degree. The share of the college-educated population in the Infutor data is higher, at 51%, compared to 29% in the ACS. However, the high correlations in Columns 3 and 4 suggest that Infutor data effectively capture which U.S. states and PUMAs have the highest share of college-educated individuals.

The fourth panel of the table shows that the age distribution in the Infutor data closely resembles that in the ACS, except for the underrepresentation of the 20–29 age group in the Infutor data. This discrepancy may explain the higher college-educated share in the Infutor data, as the dataset may not cover very young individuals who are still in college.

Lastly, the fifth panel of Table 1 reports the income distribution across the two datasets. Although income is imputed in the Infutor data, the income distribution in both datasets is similar, with approximately 10% of individuals earning less than \$20,000 per year and the most common income range falling between \$50,000 and \$75,000 per year. However, we note that the Infutor data underrepresent individuals with an annual income between \$20,000 and \$50,000, as well as those earning over \$125,000 per year. The high correlations in Columns 3 and 4 indicate that Infutor data accurately capture the distribution of income both nationally and at the level of U.S. states and PUMAs.

Auto. To answer the second question, Table 2 reports summary statistics on car ownership and the age of cars in the Infutor data, comparing them with statistics from the Bureau of Transportation Statistics (BTS). The table reports the total number of vehicles and the average age of a vehicle. According to the BTS, the average age of a car is 11.5 years. Cars in the Infutor dataset are, on average, seven months older. According to the BTS,

Table 1: Summary Statistics of Demographics and Representativeness of the Infutor Data

		Population shares		Correlation	
		ACS	Infutor	State	PUMA
Gender					
	Female	0.51	0.52	0.80	0.59
	Male	0.49	0.48	0.80	0.59
Race					
	White	0.75	0.62	0.71	0.59
	Black	0.12	0.12	0.80	0.79
	Other	0.13	0.26	0.86	0.74
Education					
	Non-college	0.71	0.49	0.70	0.87
	College degree	0.29	0.51	0.70	0.87
Age					
	20-29	0.19	0.07	-0.13	0.38
	30-39	0.19	0.17	0.25	0.65
	40-49	0.18	0.20	0.83	0.67
	50-59	0.19	0.23	0.31	0.63
	60-69	0.16	0.20	0.37	0.78
	70-80	0.10	0.13	0.72	0.90
Income (\$)					
	<20,000	0.11	0.10	0.79	0.85
	20,000-29,999	0.08	0.06	0.75	0.82
	30,000-39,999	0.08	0.12	0.59	0.71
	40,000-49,999	0.08	0.12	0.83	0.68
	50,000-74,999	0.18	0.27	0.85	0.71
	75,000-99,999	0.14	0.16	0.52	0.55
	100,000-124,999	0.10	0.09	0.82	0.64
	$\geq 125,000$	0.22	0.08	0.95	0.89
# obs.		249,773,525	239,077,422	0.99	0.74

The table reports population shares by demographic group in the Infutor data and in the ACS, as well as the spatial correlation of these population shares across the two datasets, both at the state and at PUMA level.

there are 254,582,694 vehicles in the U.S., with Infutor covering 144,863,872, accounting for approximately 57% of the total. The state-level correlation between the number of vehicles reported by the BTS and those in the Infutor data is 0.54, suggesting that our data accurately capture the spatial distribution of vehicles. To provide further validation of the data, Table

10 in Appendix B reports the share of vehicles by make and model for 2009 and 2017 relative to the National Household Travel Survey (NHTS). Overall, we find that the two datasets exhibit similar distributions of cars by make and model at both the national and state levels.

Table 2: Summary Statistics of Vehicles and Representativeness of the Infutor Data

	National level		Correlation
	BTS	Infutor	State-level
Vehicle age	11.5	12.2	
# obs.	254,582,694	144,863,872	0.54

Note: The table reports the average age and number of vehicles in the BTS and in Infutor, as well as the state-level correlation between the number of vehicles in the two datasets.

Properties. To further examine the second question, Table 3 presents summary statistics on homeownership rates, house values, house sizes, and years of construction in the Infutor dataset. Columns 1 and 2 compare homeownership rates and house values between the ACS and the Infutor data. In the ACS, the homeownership rate is 67%, while in Infutor it is 62%. We also observe a high correlation between homeownership rates at the state and PUMA levels in the two datasets, indicating that Infutor data capture not only the average homeownership rate but also regional patterns.

Furthermore, we find that Infutor data accurately capture house values. The average house value is \$319,000 in the ACS and \$360,000 in the Infutor data, while the median home value is \$218,000 in the ACS and \$222,000 in the Infutor data. The regional correlation of house values is also high, as shown in Columns 3 and 4. Finally, the table reports that the average house size is 1,900 square feet, and the average year of construction is 1974, although this information is not available in the ACS.

2.1.3 Strengths and Limitations of the Infutor Data for Studying Consumption

A considerable body of literature employs the Panel Study of Income Dynamics (PSID) and the Consumer Expenditure Survey (CEX) as primary sources of consumption data (Krueger and Perri, 2006; Aguiar and Hurst, 2007; Aguiar and Bils, 2015; Andreski et al., 2014; Boar, 2021; Aguiar et al., 2024). One advantage of these datasets is that they cover a significant share of spending: PSID captures roughly 70% of National Income and Product Accounts (NIPA) consumption, and CEX covers an even larger share. However, the sample sizes of

Table 3: Summary Statistics of Properties and Representativeness of the Infutor Data

	National level		Correlation	
	ACS	Infutor	State	PUMA
Homeownership share	0.67	0.62	0.88	0.88
House value (\$)				
Mean	319,081	361,336	0.58	0.60
Median	218,000	221,667	0.96	0.98
Average house size (SqFt)		1900		
Year of construction		1974		
# obs.	137,389,926	109,691,219	0.98	

Note: The table reports homeownership shares and house values in the ACS and in the Infutor data, as well as the spatial correlation of these statistics across the two datasets.

both PSID and CEX are only around 5,000 households, which is too small to undertake our analysis at finer geographic levels (i.e., below the U.S. state level).

Recent studies of consumption patterns rely on data from credit and debit card transactions or personal finance management applications (Agarwal et al., 2017; Ganong and Noel, 2019; Dolfen et al., 2019; Olafsson and Pagel, 2018; Diamond and Moretti, 2022). Compared to these studies, Infutor data offer three distinct advantages. First, it enables us to accurately capture durable consumption categories such as automobiles and housing, which are less likely to be purchased using credit cards. Second, it provides broader temporal coverage, allowing us to track trends over two decades. Third, in addition to capturing the value of goods, Infutor data provide information on the types of goods purchased, distinguishing consumption from expenditure—a particularly useful feature given the challenges of observing prices at granular geographic levels.

However, a key limitation of using Infutor data to study consumption segregation is that it captures only durable consumption. To address this, we supplement our analysis with Nielsen Homescan data, which provides detailed information about household purchasing patterns at major retailers such as Walmart and Whole Foods. This allows us to study non-durable consumption segregation. Although the inclusion of Nielsen Homescan data expands the range of consumption categories we analyze, it does not allow us to measure spending on services such as restaurant meals, haircuts, taxi rides, or other leisure activities.

This dimension of consumption segregation has been explored in part by [Davis et al. \(2017\)](#) in their study of restaurant consumption patterns in New York City using Yelp data and theoretically by [Couture et al. \(2019\)](#). Thus, we view our work as complementary to these studies. Overall, the consumption categories covered in the Infutor and Nielsen datasets account for 51% of total household spending.

2.2 Nielsen Homescan Data

Our data source for analyzing non-durable consumption is the Nielsen Homescan dataset, provided by the Kilts Marketing Data Center at the University of Chicago Booth School of Business. This dataset contains longitudinal panel data from approximately 40,000 to 60,000 U.S. households covering the period from 2004 to 2018. It provides detailed information about the products households purchase, including when and where these purchases occur. The Nielsen Homescan data cover the entire nation and include a wide range of household locations and demographics. In total, the dataset contains records for nearly 250 million unique items. One key advantage of this dataset is its granular level of detail: it records barcodes at a highly specific level and tracks expenditures on each item, allowing us to construct multiple measures of non-durable consumption segregation.

2.3 American Community Survey

The American Community Survey (ACS) is a yearly survey conducted by the U.S. Census Bureau, which randomly samples individuals in all 50 states, the District of Columbia, and Puerto Rico. We use data on income, education, age, race, and geography to validate the representativeness of Infutor data and to measure residential segregation by these characteristics. These variables are important to our analysis but are imputed in Infutor, making ACS a valuable benchmark. We limit the ACS sample to individuals aged 22 to 80 and exclude observations with non-positive income. The income variable we focus on is total household income. Although car and home characteristics in Infutor data are recorded at the individual level, we consider these consumption categories to be shared among all household members.

3 Patterns of Consumption Segregation

In this section, we describe patterns of consumption segregation across space, time, consumption categories, and socio-economic groups. We highlight four main findings. First,

consumption segregation has declined over the past 15 years. Second, there is substantial regional variation in consumption segregation. Third, consumption is more segregated in wealthier, younger, larger, more educated, and less white regions. Fourth, consumption segregation is driven by differences within, rather than across, demographic groups.

3.1 Measuring Segregation

Throughout the paper, we consider two measures of segregation, depending on whether the outcome of interest is categorical or continuous: the *entropy index* for categorical outcomes and the *rank-order index* for continuous outcomes. Both indices quantify the extent to which the distribution of a given economic outcome within a geographic sub-unit deviates from its distribution in a broader geographic unit. They are commonly employed in the segregation literature and possess several desirable properties.⁸ We next describe them in more detail.

The Entropy Index. We use the entropy index to measure segregation along categorical dimensions, such as car models, housing characteristics, race, and grocery categories.

To construct the entropy index for a broad geographic unit, we begin by defining the entropy score, a measure of diversity. The entropy score of a geographic unit *ie.g.*, a street, Census tract, ZIP code, or city) is

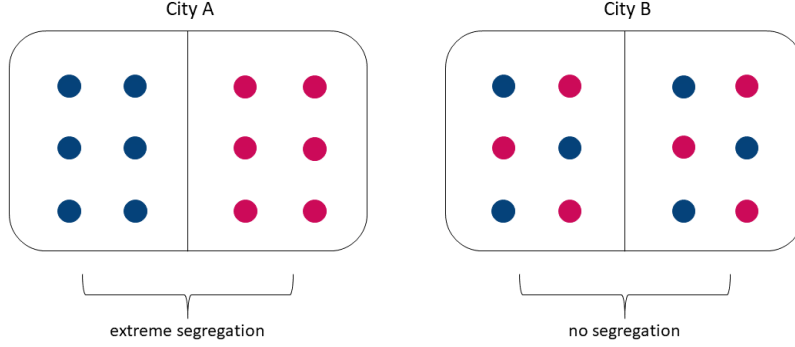
$$h_i = - \sum_{j=1}^J p_{ij} \ln(p_{ij}), \quad (1)$$

where J is the number of mutually exclusive groups (*e.g.*, car or house characteristics in the Infutor data) and p_{ij} is the share of the population of geographic unit i that belongs to group j . Specifically, $p_{ij} = \frac{n_{ij}}{n_i}$, where n_i is the total population of the geographic unit i and n_{ij} is the number of individuals in i that belong to group j . For example, if i is a zip code, then in our case n_i is the total number of cars in geographic unit i and n_{ij} is the total number of cars of model j in geographic unit i . The maximum value for h_i is $\ln(J)$, with a higher h_i indicating greater diversity. The extreme case $h_i = 0$ occurs when geographic unit i contains only one group.

Using this measure, we can calculate segregation for a broader geographic unit, such as a

⁸For example, [Reardon and Firebaugh \(2002\)](#) evaluate multiple segregation indices, including the dissimilarity index and the Gini index, which are also widely used, and conclude that the entropy index is the most conceptually and mathematically satisfactory across a wide range of considerations.

Figure 1: Measuring Segregation: An Example



Note: The figure illustrates with an example our measure of segregation. Cities A and B are divided into two neighborhoods: left and right. Each neighborhood contains individuals that are either blue or pink. In city A there is extreme segregation and in city B there is no segregation.

state. Specifically, the entropy index of a broader geographic unit is defined as

$$H = \frac{\hat{H} - \bar{H}}{\hat{H}}, \quad (2)$$

where $\hat{H}$ is the entropy index calculated at the level of a broader geographic unit using equation (1), and $\bar{H}$ is the population-share-weighted average of h_i , for all geographic units i belonging to the broader unit. The index H is bounded between 0 and 1. A value of $H = 1$ indicates extreme segregation, where each geographic unit i contains only one group (i.e., $\bar{H} = 0$), as illustrated in the left panel of Figure 1. Conversely, a value of $H = 0$ indicates no segregation, meaning that each geographic unit i is representative of the broader geographic unit (i.e., $\hat{H} = \bar{H}$), as illustrated in the right panel of Figure 1.

The Rank-order Index. We use the rank-order information theory index to measure segregation along continuous dimensions such as income and spending. The rank-order index of a continuous variable y is given by

$$H^R(y) = 2 \int_0^1 \left\{ \sum_{i=1}^I \omega_i \left[\bar{F}_{i,p} \ln \left(\frac{\bar{F}_{i,p}}{p} \right) + (1 - \bar{F}_{i,p}) \ln \left(\frac{1 - \bar{F}_{i,p}}{1 - p} \right) \right] \right\} dp \quad (3)$$

where ω_i is the population share of geographic unit i , $\bar{F}_{i,p} = F_i(y(p))$ is the cumulative function in the geographic unit i evaluated at $y(p)$. Here, both y and the percentiles p correspond to the income distribution of a geographic unit that is broader than i . The rank-order index H^R is also bounded between 0 and 1, where $H^R = 1$ indicates extreme segregation and $H^R = 0$ indicates no segregation. We note that our results remain consistent

if, instead, we measure segregation in continuous variables by applying the entropy index to bins of the variable, such as income deciles.

Measuring Segregation in the Data. When constructing the segregation measures reported in this paper, the geographic unit i is a PUMA (Public Use Microdata Area), which is the smallest unit of geography available across all the datasets we use. The broader geographic unit is either a Core Based Statistical Area (CBSA) or a U.S. state.⁹ In reporting trends, the broader geographic unit is the entire country.¹⁰

We exclude CBSAs that contain fewer than two PUMAs, as the entropy index calculation requires the broader geographic unit to include at least two sub-units. Out of the 933 CBSAs initially considered, we drop 295 that do not include any PUMA (very small CBSAs) and 422 that correspond one-to-one with a PUMA. The remaining 216 CBSAs include, on average, nine PUMAs and cover approximately 80% of the total population.

Since our consumption data contain information on both the dollar value of goods consumed (e.g., dollars spent on groceries, the value of a car, the value of a home) and the actual goods consumed (e.g., types of groceries purchased, types of cars owned), we report statistics on both spending segregation (using the rank-order index) and consumption segregation (using the entropy index). We also analyze segregation in durable and non-durable consumption, which are known to respond differently to economic shocks and policy (Berger and Vavra, 2014, 2015).

We aggregate the segregation indices for different consumption categories using expenditure shares as weights. Using data from the Consumer Expenditure Survey (CEX), we calculate that the spending share on rent and rent equivalents for homeowners is 32%, the spending share on food at home, alcohol at home, tobacco, personal care products, and other personal goods is 15%, and the spending share on new and used motor vehicles is 4%. These

⁹Because CBSAs and states have different numbers of PUMAs, a natural concern is that a higher number of PUMAs could mechanically lead to a higher segregation index at the CBSA or state level. We performed Monte Carlo simulations to assess this possibility. In each simulation, we randomly partitioned states into geographic sub-units and computed the resulting segregation index. We conducted this analysis using the ACS for both categorical variables (race, measured using the entropy index) and continuous variables (income, measured using the rank-order index). We then regressed the entropy index on the number of sub-units and compared the distribution of slope coefficients across 500 simulations with the slope coefficient observed in the actual data. While we generally find a positive relationship between the number of geographic sub-units and the segregation index, the empirical slope coefficient is much larger than those from the simulated data. This suggests that the regional differences in segregation reported below are not driven by differences in the number of PUMAs across states.

¹⁰Table 13 in Appendix B shows that our results are robust to considering zip codes or counties as the geographic unit i .

are the three largest components of the household consumption basket (Theloudis, 2020).¹¹

3.2 Trends in Consumption Segregation

We begin by examining trends in consumption segregation in the United States. Figure 2 illustrates these trends from 2000 onward. Panel (a) presents the evolution of three broad measures of consumption—total consumption, durable consumption, and non-durable consumption—measured in terms of expenditure. Two key observations emerge from Panel (a). First, durable consumption is more segregated than non-durable consumption. The average entropy index for durable consumption is 25%, approximately 2.5 times higher than that of non-durable consumption. This means that, on average, 25% of the diversity in durable consumption observed at the national level is not reflected within individual PUMAs. Given the significant expenditure share allocated to durables, total consumption segregation closely mirrors that of durables, with an average total consumption entropy index of 20%. Second, the segregation of non-durable consumption has declined by approximately 25% since 2005, the start of our Nielsen Homescan sample, whereas durable and total consumption segregation exhibit no clear trend.

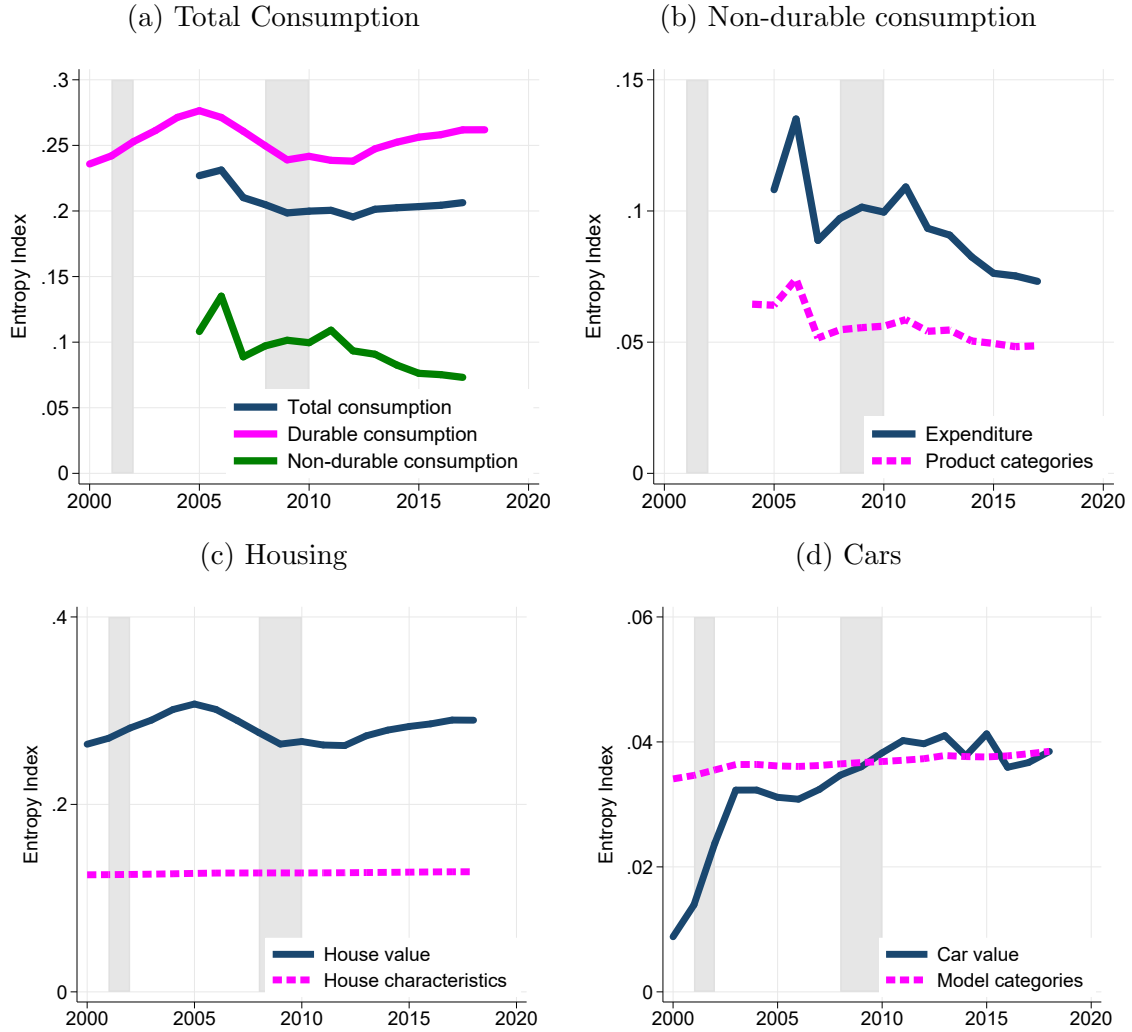
In Panels (b)-(d) of Figure 2, we focus on the segregation of specific categories. Panel (b) illustrates the evolution of non-durable consumption segregation, measured both using non-service retail expenditure and specific product categories purchased. On average, non-durable consumption segregation measured by expenditure is higher than that measured by product categories. However, both measures are declining over time and are lower than the average consumption segregation. These patterns may reflect both demand and supply forces. That segregation levels are both low and declining is consistent with Allcott et al. (2019), who find that only 10% of food inequality between rich and poor households stems from supply-side factors such as food deserts. The fact that segregation in non-durable spending is higher than in product categories suggests that factors beyond product availability shape non-durable consumption segregation. We explore these factors below.

Panels (c) and (d) show the evolution of housing and car consumption segregation between 2000 and 2018. The first notable feature is that housing consumption segregation is significantly higher than car consumption segregation. This result holds regardless of whether house values or housing characteristics are used to measure the entropy index over time. It

¹¹Our entropy indices are conditional on having positive expenditures on groceries and owning at least one vehicle or home.

also remains robust when using home values directly reported in the ACS and Census.¹² Car consumption, in fact, is the least segregated consumption category, regardless of whether it is measured by car values or car models. The second key feature is that housing consumption segregation has remained relatively stable over time. When measured by house values, it mirrors overall house price trends, consistent with differential house price appreciation across U.S. regions (Guren et al., 2020). In contrast, car consumption segregation exhibits a mild but steady upward trend.

Figure 2: Consumption Segregation Over Time

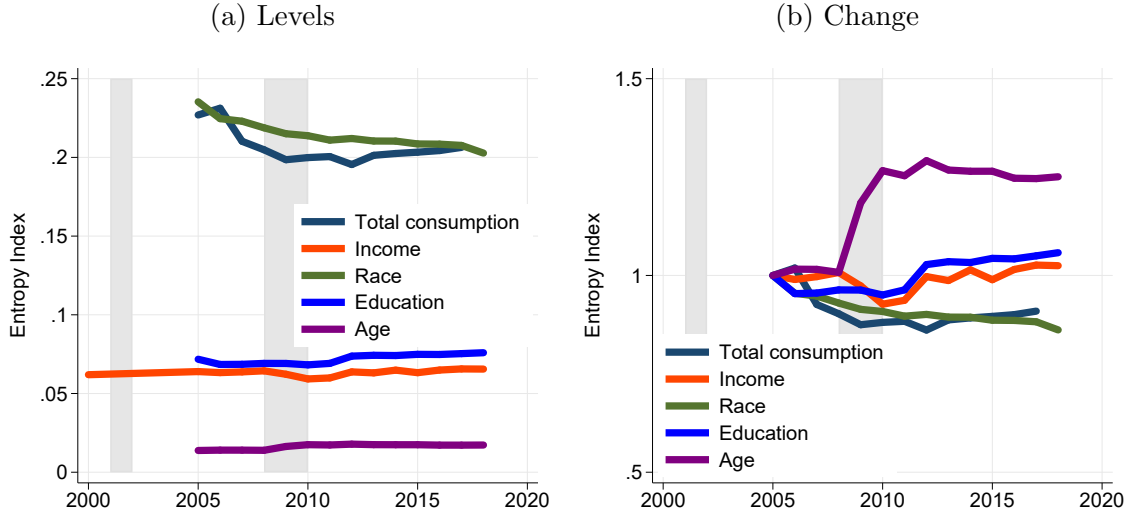


How does consumption segregation compare with other measures of segregation? We benchmark both the level and time evolution of consumption segregation against income,

¹²See Figure 8 in Appendix B.

racial, educational, and age segregation in Figure 3. Overall, our findings align with previous research (e.g., Massey, 2016 and Fogli and Guerrieri, 2019), confirming that income segregation has increased over time but moderately after 2000, whereas racial segregation has declined throughout the same time period. However, racial segregation remains at a higher level than other forms of segregation.¹³ Surprisingly, consumption segregation reaches levels comparable to racial segregation. However, like racial segregation, it has been steadily declining. Between 2005 and 2018, the average entropy index for racial segregation was 0.21, while consumption segregation averaged 0.25 over the same period. In contrast, income and education segregation exhibit similar levels and trends, with entropy index averages of 0.06 and 0.07, respectively. If we focus exclusively on non-durable consumption, the level of segregation would be closer to that of income and education rather than race, while still displaying a downward trend.

Figure 3: Consumption and Other Dimensions of Segregation



3.3 The Geography of Consumption Segregation

In this section, we examine consumption segregation across U.S. regions. Figure 4 provides a visual representation of its geographic distribution. First, we compute entropy indices for each consumption category, defining the narrow geographic unit as a PUMA and the broader

¹³These results hold even though our measures of segregation are based on the entropy and rank-order indices, whereas Massey (2016) primarily used the dissimilarity index.

geographic unit as a U.S. state. Next, we report each state’s average entropy index for the period 2016–2018 in the figure.

Panel (a) of Figure 4 reveals substantial variation in total consumption segregation across U.S. states, with the most segregated state exhibiting levels 11 times higher than the least segregated state. The entropy index ranges from 28% in New York to 2.5% in Wyoming, indicating that in Wyoming, the average PUMA is almost fully representative of the state whereas in New York nearly 30% of the state-level diversity is not reflected within individual PUMAs. Overall, total consumption segregation is highest on the West Coast, in the Northeast, and in Texas. Conversely, the Rocky Mountain Region exhibits the lowest levels of consumption segregation.

Panel (b) of Figure 4 illustrates the geographic distribution of durable consumption segregation. The patterns closely resemble those of total consumption segregation, reflecting the large share of durable goods in overall spending. Specifically, durable consumption segregation is highest in New York, while Wyoming, South Dakota, and Alaska exhibit the lowest levels. Figure 5 provides a more detailed breakdown of the sources of durable consumption segregation. Panel (a) shows that durable consumption segregation largely reflects segregation in housing consumption. Panel (b), which focuses on vehicle consumption segregation, reveals a slightly different geographic pattern, with higher segregation in the South and lower segregation in the Northeast. Nonetheless, segregation in vehicle consumption is generally much lower than in housing consumption and approaches near-perfect diversity.

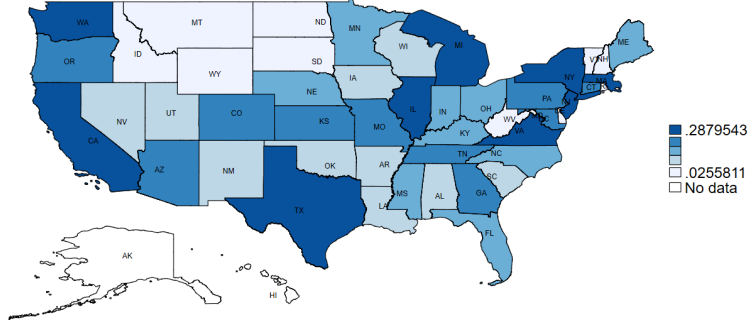
Turning to non-durable consumption, Panel (c) of Figure 4 highlights that its geographic distribution differs from that of durable consumption. Specifically, while New York and California remain among the most segregated states, the figure also reveals particularly high levels of non-durable consumption segregation in the South. Texas exhibits the highest level, with 10% of the state-level diversity not reflected in the average PUMA. At the other extreme, in Vermont, individual PUMAs are nearly representative of the state.

We next investigate the locus of consumption segregation. To do so, we leverage the additive organizational decomposability of the entropy index, which allows us to break down overall segregation into components that reflect differences both within and between clusters, helping to pinpoint where segregation is most pronounced. Specifically, as shown in [Reardon and Firebaugh \(2002\)](#), assuming that J geographic units grouped into K clusters ($K < J$), the entropy index H can be expressed as the sum of $K + 1$ additive components

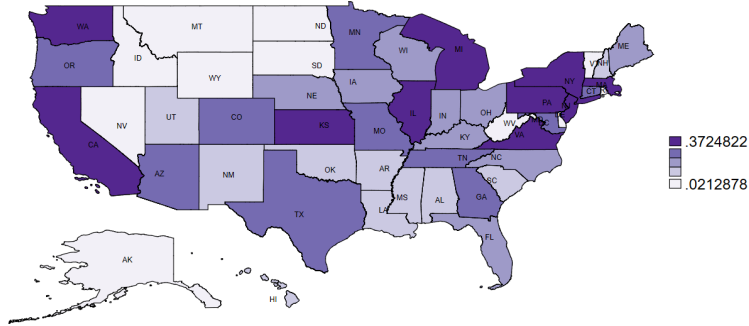
$$H = H_K + \sum_k \omega_k \frac{h_k}{h} H_k, \quad (4)$$

Figure 4: State-level Consumption Segregation

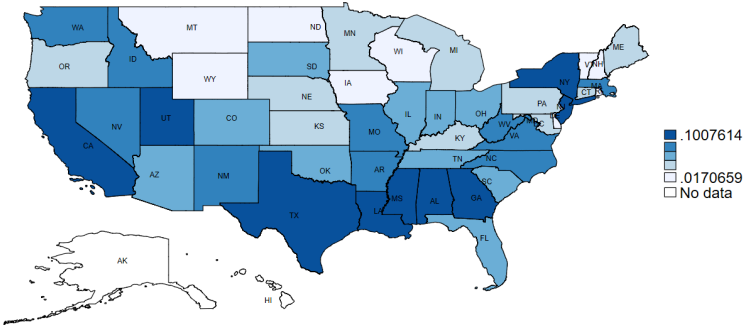
(a) Total consumption



(b) Durable consumption



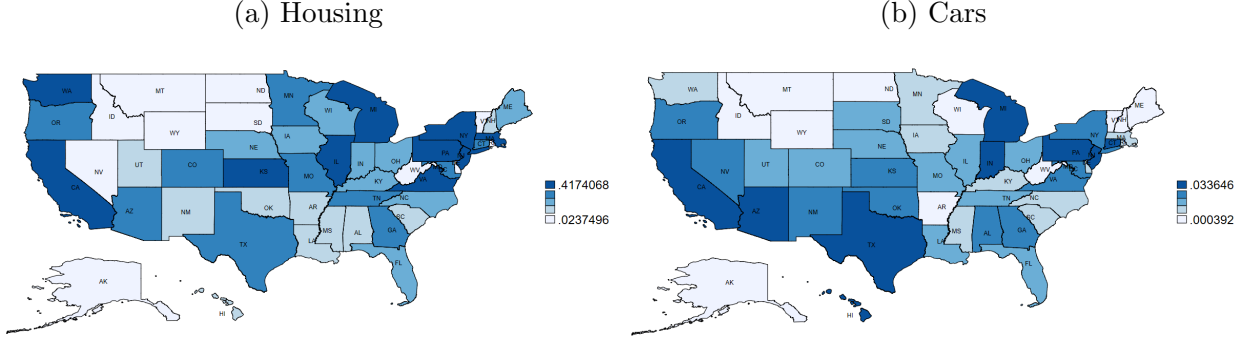
(c) Non-durable consumption



Notes: The figure plots state-level consumption entropy indices averaged over the period 2016-2018.

where ω_k is the population share of cluster k , h_k is the entropy score of cluster k , h is the national-level entropy score, and H_k is the information theory index of cluster k , calculated according to equation (2). The term H_K , determined residually, represents the *between*-cluster component, while the remaining K terms account for the portion of total segregation due to *within*-cluster segregation. We derive a similar decomposition for the rank-order

Figure 5: State-level Durable Consumption Segregation



Notes: The figure plots state-level consumption entropy indices averaged over the period 2016-2018.

index H^R . Specifically,

$$H^R = H_K^R + \sum_k \omega_k H_k^R, \quad (5)$$

where, relative to equation (4), the term $\frac{h_k}{h}$ equals 1. The share of total segregation attributable to within-cluster segregation represents the extent to which total segregation would decrease if segregation within cluster k were eliminated by redistributing individuals among its geographic units while leaving all other units unchanged.

We report the decomposition results in Table 4. The top panel of the table shows that two-thirds of consumption segregation occurs within states. In other words, total consumption segregation would decrease by approximately 70% if within-state segregation were eliminated by appropriately redistributing individuals across PUMAs within each state. The second and third columns of the table break down the segregation of durable and non-durable consumption into within-state and between-state components. While most consumption segregation is driven by within-state differences across both broad consumption categories, we note that in the case of non-durable consumption, nearly all segregation is attributable to the within-state component. Finally, the last two columns of the table provide further insight into durable consumption segregation. The results show that the strong within-state component of durable consumption segregation is largely driven by housing consumption, which is more segregated within states than across states. By contrast, vehicle consumption segregation is primarily explained by differences between states.

The bottom panel of Table 4 performs the same decomposition, this time grouping geographic units into CBSAs instead of states. Given that CBSAs are smaller than U.S. states, and that the state-level decomposition revealed a significant role for within-state segregation, this analysis provides additional insight into the locus of within-state segregation. We find

that within- and between-CBSA components contribute almost equally to total consumption segregation, as well as to the segregation of durable and non-durable consumption.

Table 4: Segregation Within and Between Regions

	Total consumption	Durable consumption	Non-durable consumption	Housing
<i>State-level segregation</i>				
Between states	0.285	0.392	0.036	0.350
Within state	0.715	0.608	0.964	0.650
<i>CBSA-level segregation</i>				
Between CBSAs	0.444	0.569	0.557	0.670
Within CBSA	0.556	0.431	0.443	0.330

Notes: The table reports averages of the between and within region components of the consumption segregation indices over the period 2016-2018.

Table 5 examines the relationship between consumption segregation and the socioeconomic characteristics of CBSAs, building on the observed heterogeneity in segregation across the U.S. We calculate the average entropy index for total, durable, and non-durable consumption in CBSAs that fall within the top 25% or bottom 25% of the national distribution in terms of income, education, race, age, and population size. The results indicate that consumption segregation is higher in CBSAs that are richer, younger, more populous, have a higher share of college-educated individuals, and a lower share of white residents. For instance, total consumption is 2.5 times more segregated in CBSAs in the top 25% of the income distribution than in those in the bottom 25%. This income-based segregation pattern holds across both durable and non-durable consumption, with durable consumption being 2.7 times more segregated and non-durable consumption 1.9 times more segregated in richer CBSAs. The findings for education closely mirror those for income. In terms of race, total consumption is twice as segregated in CBSAs in the bottom 25% of the white population share. This pattern is more pronounced in durable consumption than in non-durable consumption, while segregation differences by age are minimal. Finally, consumption segregation is 4.3 times higher in CBSAs in the top 25% of the population distribution. These results suggest that consumption segregation varies significantly across key demographic characteristics.

We next analyze how consumption segregation varies with individuals' demographic characteristics. In Table 6, we present consumption segregation levels for different demographic groups. The top panel of the table categorizes individuals into three racial groups and cal-

Table 5: Segregation and Regional Characteristics

	Total consumption	Durable consumption	Non-durable consumption
<i>By income</i>			
Bottom 25%	0.058	0.063	0.040
Top 25%	0.144	0.169	0.075
<i>By education (college share)</i>			
Bottom 25%	0.063	0.068	0.044
Top 25%	0.134	0.154	0.075
<i>By race (share white)</i>			
Bottom 25%	0.136	0.152	0.087
Top 25%	0.060	0.069	0.036
<i>By age (share ≤ 35)</i>			
Bottom 25%	0.076	0.088	0.046
Top 25%	0.103	0.110	0.072
<i>By population</i>			
Bottom 25%	0.028	0.032	0.018
Top 25%	0.121	0.136	0.076

Notes: The table reports population-weighted averages of CBSA-level consumption entropy indices over the period 2016-2018.

culates the entropy index, H^R , as defined in equation (3), for each group. We observe little variation in total consumption segregation across racial groups, primarily due to similar levels of durable consumption segregation. However, notable differences emerge for non-durable consumption. Specifically, white individuals exhibit the highest segregation in non-durable consumption, at a level comparable to that of Black individuals but 37% higher than that of other racial groups. The middle panel of the table displays the average consumption segregation for college-educated and non-college-educated individuals. Total consumption is 13% more segregated among the college-educated, with an 11% higher durable consumption segregation and a 19% higher non-durable consumption segregation within this group.

These findings suggest that differences across demographic groups play a minor role in consumption segregation. To investigate this further, we decompose consumption segrega-

Table 6: Segregation by Demographic Group

	Total consumption	Durable consumption	Non-durable consumption
<i>By race</i>			
White	0.102	0.111	0.071
Black	0.108	0.116	0.068
Other	0.115	0.124	0.052
<i>By education</i>			
No college degree	0.102	0.107	0.079
College degree	0.115	0.119	0.094
<i>By age</i>			
Younger than 35	0.104	0.118	0.058
Older than 35	0.104	0.117	0.064

Notes: The table reports population-weighted averages of CBSA-level consumption entropy indices by demographic groups over the period 2016-2018.

tion into within- and between-demographic group components. We examine both the broad demographic groups defined in Table 6 and narrower demographic groups, which are defined as (race, education, age) tuples.¹⁴ We employ a decomposition that is analogous to that described in equation (5), leveraging the additive group decomposability property of the entropy and rank-order indices (Reardon and Firebaugh, 2002).¹⁵ Table 7 summarizes the decomposition results and shows that the majority of consumption segregation stems from within-group differences. Consequently, differences in segregation across demographic groups have minimal influence on overall consumption segregation patterns.

4 Understanding Consumption Segregation

In this section, we examine potential determinants of consumption segregation. We document that differences in consumption segregation across space and time are strongly associated with differences in income segregation, and that this association likely reflects households' inability

¹⁴We consider all combinations of three racial groups (White, Black, Other), two education groups (college degree, no college degree), and two age groups (younger than 35, older than 35).

¹⁵For this decomposition, the geographic unit i in equation (1) corresponds to a combination of a PUMA and a demographic group.

Table 7: Segregation Within and Between Demographic Groups

	Total consumption	Durable consumption	Non-durable consumption
<i>Race</i>			
Between group	0.033	0.009	0.091
Within group	0.967	0.991	0.909
<i>Education</i>			
Between group	0.028	0.034	0.012
Within group	0.972	0.966	0.988
<i>Age</i>			
Between group	0.014	0.007	0.029
Within group	0.986	0.993	0.971
<i>Race $\times$ Education $\times$ Age</i>			
Between group	0.033	0.042	0.007
Within group	0.967	0.958	0.993

Notes: The table reports averages of the between and within demographic group components of the CBSA-level consumption entropy indices by demographic groups over the period 2016-2018.

to arbitrage away income differences through financial markets.

4.1 Determinants of Consumption Segregation

A long line of research examines segregation in the United States over the past century across various socio-economic dimensions—such as race, education, income, and ethnicity—with a particular focus on racial and income segregation.¹⁶ In this section, we analyze the extent to which residential segregation along these previously documented dimensions shapes consumption segregation, thereby connecting earlier studies of residential segregation to the residential segregation of well-being, as captured by consumption segregation.

To that end, we project consumption segregation onto other dimensions of residential segregation and estimate the following regression specification

$$\ln H_{C,rt} = \alpha_0 + \alpha_1 \ln H_{Y,rt} + \alpha_2 \ln H_{R,rt} + \alpha_3 \ln H_{E,rt} + \alpha_4 \ln H_{A,rt} + \gamma \mathbf{X}_{rt} + \delta_t + \varepsilon_{rt}, \quad (6)$$

¹⁶See [Trounstein \(2018\)](#) for a comprehensive summary of this literature.

where $H_{C,rt}$, $H_{Y,rt}$, $H_{R,rt}$, $H_{E,rt}$, and $H_{A,rt}$ denote the entropy indices of consumption, income, race, education and age in CBSA r and year t , respectively, $\mathbf{X}_{rt}$ is a vector of time-varying controls for CBSA r , and δ_t represents year fixed effects. The measures of consumption and income segregation are calculated as in equation (3), while the measures of racial, education and age segregation are calculated as in equation (2), considering the same racial, education and age groups as in Section 3.3. Motivated by the analysis in Section 3.3, the vector of controls includes average income in CBSA r in year t , along with the population share of white, Black, college-educated individuals and of those younger than 35.

Table 8 reports the estimation results, sequentially adding segregation controls.¹⁷ First, the specification in column (1) only controls for segregation by income, $H_{Y,rt}$, in addition to the set of controls $\mathbf{X}_{rt}$ and the time fixed effects. We then progressively add controls for racial segregation, $H_{R,rt}$, in column (2) and for segregation by education and age, $H_{E,rt}$, and $H_{A,rt}$, in column (3).

Focusing on total consumption, we find an elasticity with respect to income segregation of 0.67, meaning that 67% of income segregation translates into consumption segregation. The R^2 from estimating equation (6) without any residential segregation controls is 0.41. Comparing with the R^2 reported in column (1) suggests that differences in income segregation across CBSAs explain 31% of the differences in consumption segregation (0.72-0.41). When we include a control for racial segregation, we find that the elasticity of consumption segregation with respect to racial segregation is 0.14—positive but four times smaller than that for income segregation. Additionally, the R^2 increases by only two percentage points. Lastly, when we add controls for segregation by education and by age, we continue to find that income segregation is the main socio-economic dimension driving consumption segregation. We also find a sizable role for segregation by education but note that segregation by income and by education are strongly correlated, with a correlation coefficient of 0.77. We interpret both as broadly representing economic status. Segregation by education and age explains an additional 4% of the variation in consumption segregation across CBSAs, further reinforcing the role of income segregation as the primary factor shaping consumption segregation.

That income segregation emerges as the main contributor to consumption segregation is consistent with findings from the existing empirical literature on residential segregation and standard consumption theory. Specifically, Massey et al. (2009) document that, over the last third of the twentieth century, the United States transitioned from a regime of segregation

¹⁷See Table 17 in Appendix B for the corresponding results that exploit variation across U.S. states.

Table 8: Understanding Consumption Segregation

	Total consumption			Durable consumption			Non-durable consumption		
	(1)	(2)	(3)	(1)	(2)	(3)	(1)	(2)	(3)
$\ln H_Y$	0.67*** (0.03)	0.58*** (0.04)	0.39*** (0.06)	0.87*** (0.05)	0.73*** (0.05)	0.43*** (0.07)	0.44*** (0.03)	0.35*** (0.04)	0.24*** (0.05)
$\ln H_R$		0.14*** (0.05)	0.10** (0.05)		0.21*** (0.07)	0.16*** (0.06)		0.14*** (0.04)	0.12*** (0.03)
$\ln H_E$			0.21*** (0.05)			0.33*** (0.06)			0.11*** (0.03)
$\ln H_A$			0.02 (0.02)			0.01 (0.03)			0.04 (0.03)
R^2	0.72	0.74	0.78	0.65	0.66	0.73	0.57	0.59	0.61

Notes: The table reports the estimates of α_1 , α_2 , α_3 and α_4 in equation (6). The estimate of α_0 , the time fixed effects and the estimate of the vector γ are omitted. Robust standard errors are reported in parentheses. The regression is estimated with population weights based on the 2010 Census. ***, **, and *, represent statistical significance at 1%, 5% and 10%, respectively.

based on race and ethnicity to one based on economic status, defined by income and education. Standard consumption theory predicts that differences in income across individuals translate into differences in consumption if income differences are persistent and/or if credit market frictions or other forms of market incompleteness prevent individuals from insuring against income fluctuations. When extended to a spatial dimension, as in [Giannone et al. \(2019\)](#), this argument suggests that, under the same conditions, residential segregation by income will translate into residential segregation by consumption.

Turning to the determinants of durable and non-durable consumption segregation, we draw a similar conclusion. Specifically, we find that in the most conservative specification, 43% of income segregation translates into durable consumption segregation. This elasticity is approximately half as large for non-durable consumption, with only 24% of income segregation translating into non-durable consumption segregation. Relative to the R^2 from a regression that abstracts from any segregation controls, we find that differences in income segregation across CBSAs account for 33% of the variation in durable consumption segregation, while segregation by race, education, and age accounts for an additional 8%. Similarly, differences in non-durable consumption segregation across CBSAs are accounted for in proportion of 22% by differences in income segregation and by an additional 4% by differences

in segregation by race, education and age.¹⁸

4.2 Consumption and Income Segregation

Motivated by the results above, in this section, we examine the mechanism underlying the strong positive relationship between income and consumption segregation. Our starting point is standard consumption theory, which posits two potential explanations: market incompleteness and persistent income differences between individuals.¹⁹ Market incompleteness—due to limited access to financial products, either because they are unavailable or due to underlying informational frictions—and persistent income differences both have the potential to generate this positive relationship, as they restrict consumption insurance among individuals with different incomes.

4.2.1 An Analytical Framework

To guide the empirics, we outline a simple consumption model with heterogeneous households and incomplete markets. We consider the utility maximization problem of a household i who lives in region r . Each period t the household earns stochastic income y_{rt}^i and chooses the optimal path of consumption c_{rt}^i and savings a_{rt+1}^i that maximizes its expected life-time utility

$$\mathbf{E}_0 \sum_{t=0}^{\infty} \beta^t u(c_{rt}^i),$$

subject to the budget constraint

$$c_{rt}^i + \frac{a_{rt+1}^i}{R} = y_{rt}^i + a_{rt}^i$$

and a borrowing constraint $a_{rt+1}^i \geq \underline{a}$. Here, β is the discount factor of the household, R is the gross real interest rate on savings and $\underline{a}$ is the borrowing limit. The optimal consumption-saving choice is characterized by the Euler equation

$$u'(c_{rt}^i) = \beta R \mathbf{E}_t u'(c_{rt+1}^i) + \mu_t,$$

where μ_t is the multiplier on the borrowing constraint. That is, when deciding how much to save, the household compares the cost of saving an additional unit of income, expressed in

¹⁸The R^2 from regressing durable consumption segregation on the controls in equation (6), excluding the segregation controls, is 0.32. The corresponding R^2 for non-durable consumption segregation is 0.35.

¹⁹We also conjecture that our results may reflect persistent preference differences between individuals.

terms of the marginal utility of lost consumption, with the discounted benefit of having an additional unit of resources available for consumption in the following period.

In order to derive an analytical characterization of optimal consumption choices, we specialize this model to what is commonly referred to as the strict permanent income hypothesis. Specifically, we assume that (i) households have quadratic utility $u(c) = b_1c + \frac{1}{2}b_2c^2$, where the parameters b_1 and b_2 are such that the utility function is strictly increasing and strictly concave, (ii) the interest on savings equals the inverse of the discount rate $\beta R = 1$, and (iii) there is no borrowing constraint. Under these assumptions, it can be shown that consumption is a martingale

$$c_{rt}^i = \mathbf{E}_t c_{rt+1}^i$$

and that, more generally,

$$c_{rt}^i = \mathbb{E}_t c_{rt+j}^i, \forall j \geq 0.$$

Iterating forward on the budget constraint J times, taking the limit as $J \rightarrow \infty$ and using a No-Ponzi scheme condition gives

$$c_{rt}^i = \frac{R-1}{R} \left[a_{rt}^i + \sum_{j=0}^{\infty} \left(\frac{1}{R} \right)^j \mathbb{E}_t y_{rt+j}^i \right].$$

That is, consumption is the annuity value of human and financial wealth. It then follows that the change in consumption $\Delta c_{rt}^i \equiv c_{rt}^i - c_{rt-1}^i$ is equal to

$$\Delta c_{rt}^i = \frac{R-1}{R} \sum_{j=0}^{\infty} \left(\frac{1}{R} \right)^j (\mathbb{E}_t - \mathbb{E}_{t-1}) y_{rt+j}^i$$

and is proportional to the revision in expected earnings due to the new information accruing between $t-1$ and t .

To make further analytical progress on the predictions of the model regarding the relationship between income and consumption, we specialize the process for income by assuming the commonly used specification

$$y_{rt}^i = \rho y_{rt-1}^i + \varepsilon_{rt}^i,$$

where ε_{rt} is a mean-zero shock with variance σ_ε^2 that is iid across households, regions and over time, and ρ is the persistence parameter. Under this specialized income process, the change in consumption is

$$\Delta c_{rt}^i = \frac{R-1}{R-\rho} \varepsilon_{rt}^i.$$

This expression encompasses two special cases that are of interest for our empirical test, as they both predict a strong relationship between income and consumption. First, when $\rho = 0$ shocks to income are transitory and consumption evolves according to $\Delta c_{rt}^i = \frac{R-1}{R} \varepsilon_{rt}^i$, implying that the households consume only the annuity value of the shock while most of the income change is saved. However, in the absence of financial markets—an extreme case of market incompleteness— $c_{rt}^i = y_{rt}^i$ and, consequently, $\Delta c_{rt}^i = \varepsilon_{rt}^i$, so income shocks translate fully into consumption. Second, when $\rho = 1$, shocks to income are permanent, and consumption evolves according to $\Delta c_{rt}^i = \varepsilon_{rt}^i$, implying that the shocks again translate fully into consumption.

Zooming out to the distributions of income and consumption within a given region, the discussion above makes it clear that if income shocks are permanent, the two distributions will mirror each other. If we consider regions r as part of a broader geographic unit, such as a state or a CBSA, then income segregation within the broader geographic unit should be fully reflected in consumption segregation. More generally, the extent to which the distribution of consumption follows that of income depends on the persistence of income shocks: the more persistent the shocks, the closer the two distributions become.

If, instead, income shocks are transitory, the distribution of consumption becomes less reflective of the distribution of income (e.g. the term $\frac{R-1}{R} \approx 0.02$ for values of R typically used in the literature), as agents can use saving and borrowing to smooth consumption in response to these shocks. Once more, considering regions r as part of a broader geographic unit, this implies that income segregation translates into consumption segregation only to a limited extent. However, if asset markets that facilitate saving and borrowing are unavailable (i.e. if markets are incomplete), then the two distributions become more tightly linked, as in the case of permanent or persistent income shocks.

4.2.2 Empirical Evidence

We propose an empirical test to assess whether the correlation between income and consumption segregation observed in the data reflects permanent income differences between socio-economic groups or incomplete asset markets that limit consumption insurance against potentially insurable income shocks.

To that end, we decompose the relationship between income and consumption segregation into two components, capturing how income segregation translates into within- and between-group consumption segregation. In doing this decomposition, we follow [Krueger and Perri](#)

(2006) and adopt a specific perspective on the sources of income differences between and within socio-economic groups. Specifically, we assume that income differences between socio-economic groups stem from fixed household characteristics. Since it is likely difficult to insure against these differences, an increase in income segregation should translate into an increase in between-group consumption segregation. Within a socio-economic group, we assume that income differences arise from idiosyncratic income shocks. If financial markets provide partial insurance against these shocks, then income segregation should translate into within-group consumption segregation only to a limited extent.²⁰ A potential limitation of our approach is that if some persistent within-group income differences exist, our decomposition may inadvertently attribute them to market incompleteness.

Recall that in Section 3.3 we decomposed consumption segregation into a within demographic group and a between demographic group component, where a demographic group is defined as a $(race, education, age)$ tuple. Letting $H_{C,rt}$ denote the level of consumption segregation in region r in year t , the decomposition allows us to write

$$H_{C,rt} = H_{C,rt}^B + H_{C,rt}^W,$$

where $H_{C,rt}^B$ denotes the between-group component of consumption segregation and $H_{C,rt}^W$ denotes the within-group component. We then investigate the relative extent to which income segregation affects these two components of consumption segregation by estimating the following regression specification

$$H_{C,rt}^i = \alpha_0^i + \alpha_1^i H_{Y,rt} + \gamma^i \mathbf{X}_{rt} + \delta_t^i + \varepsilon_{rt}^i,$$

where $i \in \{B, W\}$ and $\mathbf{X}_{rt}$ is the same vector of controls as in equation (6). The coefficients α_1^B and α_1^W can be shown to sum up to the coefficient α_1 from estimating

$$H_{C,rt} = \alpha_0 + \alpha_1 H_{Y,rt} + \gamma \mathbf{X}_{rt} + \delta_t + \varepsilon_{rt}.$$

Therefore, the share $\frac{\alpha_1^B}{\alpha_1}$ is informative of the relative importance of permanent income differences in explaining consumption segregation, while $\frac{\alpha_1^W}{\alpha_1}$ is indicative of the importance of market incompleteness.

Table 9 presents estimates of α_1 , α_1^B and α_1^W for total consumption and consumption sub-categories. The between-group component accounts for 14% of the relationship between

²⁰ An extreme example is hand-to-mouth consumers, who cannot insure against transitory income shocks at all. For such consumers, income segregation translates one-to-one into consumption segregation. At the other extreme, when markets are complete, income segregation does not translate into consumption segregation.

income and total consumption segregation at the CBSA-level, while the within-group component accounts for the majority of this relationship at 86%. This pattern holds true for broad consumption categories, durable and non-durable consumption, as well as for specific durable consumption categories. The within-group component is particularly important in explaining the relationship between income segregation and the segregation of durable consumption, especially housing, which is consistent with market frictions that limit households' ability to smooth consumption fluctuations (Boar et al., 2021). Specifically, 90% of the mapping between income segregation and housing consumption segregation is accounted for by the within-group component, while in the case of non-durable consumption, this component explains 80% of the relationship.

Table 9: Consumption and Income Segregation

	Total	Between	Within
Total consumption	1.52	0.21	1.32
α_1^i/α_1		13.8%	86.2%
Non-durable consumption	0.64	0.13	0.51
α_1^i/α_1		20.3%	79.7%
Durable Consumption	1.94	0.16	1.76
α_1^i/α_1		9.7%	90.3%
Cars	0.17	0.03	0.14
α_1^i/α_1		17.6%	83.4%
Housing	2.16	0.21	1.94
α_1^i/α_1		9.7%	90.3%

Notes: The table reports the estimates of α_1 , α_1^B and α_1^W . All estimates are significant at 1% significance. Robust standard errors, time fixed effects and the estimates of γ^i and γ are omitted but are included in the estimation. The regression is estimated with population weights based on the 2010 Census, using data for the 2016-2018 period.

Overall, we conclude that while there is a role for permanent income differences in explaining why consumption is not equalized across space, market incompleteness plays a major role in explaining why income segregation translates into consumption segregation.

4.3 Implications of Consumption Segregation

Our findings have implications for our understanding of inequality more broadly. We briefly discuss a few key aspects.

The spatial distribution of well-being. First, consumption segregation provides a more nuanced perspective on the lived experience of economic disparity, capturing not just financial resources but also access to goods that contribute directly to quality of life. Thus, it provides insight into the spatial distribution of well-being, above and beyond what is captured by previously studied dimensions of residential segregation. This is important for understanding how economic disparities manifest in everyday life. Analyzing the interplay between these dimensions of segregation provides a more comprehensive understanding of the complex drivers of inequality and their impact on communities.

Wealth inequality. Second, consumption segregation itself may contribute to inequality. If income segregation translates into consumption segregation, as we have documented, then in the presence of social comparisons (e.g., keeping up with the Joneses consumption motives), whether individuals live in segregated areas or not can affect their wealth accumulation by shaping their spending and borrowing behavior.

For example, in less segregated locations, where the poor have rich neighbors, they may borrow to finance conspicuous consumption, attempting to match their neighbors' spending. This reduces their wealth, thereby increasing wealth inequality. In contrast, in more segregated locations, where the poor do not have rich neighbors, social comparison pressures are weaker and so is the effect on inequality. This mechanism suggests that, interestingly, consumption segregation could, in fact, reduce wealth inequality.

However, if the goods used for social comparison are highly visible (e.g., houses, cars, or jewelry), these possessions are typically included in net worth measures—beyond liquid wealth—introducing a counteracting force through which consumption segregation could, instead, exacerbate wealth inequality. Determining the dominant effect is an important quantitative question, which we aim to explore in future research.

Policies to reduce segregation. Third, our empirical finding that market incompleteness significantly contributes to the link between income segregation and consumption segregation has specific policy implications. In particular, it suggests that policies enhancing

risk diversification can mitigate consumption segregation, even in the presence of income segregation. Such policies could range from direct redistributive transfers to interventions that expand financial access in banking deserts (Morgan et al., 2016), such as the promotion of online banking and digital financial services to offset the lack of physical bank branches in areas where the poor cluster.

5 Conclusions

We leverage a new data source in conjunction with existing data to examine consumption segregation in the United States and investigate its underlying factors and consequences. We analyze how consumption segregation varies across time periods, regions, consumption categories, and demographic groups. Our findings reveal that consumption segregation is as large as racial segregation and has declined over the past two decades. However, there is significant spatial heterogeneity, with New York, the most segregated state, exhibiting 11 times the level of consumption segregation found in Wyoming, the least segregated state. Additionally, consumption segregation is higher in wealthier, more highly educated, and larger regions, as well as in those where a smaller percentage of the population is white. Interestingly, consumption segregation is similar across demographic groups, and most of it reflects differences in the spatial distribution of spending within demographic groups.

We examine the determinants of consumption segregation and find that income segregation plays a major explanatory role, even after accounting for previously documented forms of segregation, such as race, education, and age. We also find a role, albeit more limited, for racial segregation. We leverage insights from a standard consumption model and exploit variation in segregation within and across narrowly defined demographic groups to argue that the observed strong correlation between income and consumption segregation arises as a consequence of incomplete financial markets that prevent consumption insurance.

Our findings underscore the importance of consumption segregation in understanding inequality more broadly. By measuring consumption segregation, we gain deeper insight into the spatial distribution of well-being. Moreover, our results on the link between income and consumption segregation suggest that the latter influences wealth accumulation through social comparisons, thereby shaping wealth inequality. These findings also indicate that policies enhancing financial access, such as expanding online banking in banking deserts, could help reduce consumption segregation.

References

- Agarwal, Sumit, J. Bradford Jensen, and Ferdinando Monte**, “Consumer Mobility and the Local Structure of Consumption Industries,” Working Paper 23616, National Bureau of Economic Research July 2017.
- Aguiar, Mark and Erik Hurst**, “Measuring trends in leisure: The allocation of time over five decades,” *The Quarterly Journal of Economics*, 2007, *122* (3), 969–1006.
- **and Mark Bils**, “Has Consumption Inequality Mirrored Income Inequality?,” *American Economic Review*, September 2015, *105* (9), 2725–56.
- , — , **and Corina Boar**, “Who Are the Hand-to-Mouth?,” *The Review of Economic Studies*, 05 2024.
- Allcott, Hunt, Rebecca Diamond, Jean-Pierre Dubé, Jessie Handbury, Ilya Rahkovsky, and Molly Schnell**, “Food deserts and the causes of nutritional inequality,” *The Quarterly Journal of Economics*, 2019, *134* (4), 1793–1844.
- Andreski, Patricia, Geng Li, Mehmet Zahid Samancioglu, and Robert Schoeni**, “Estimates of annual consumption expenditures and its major components in the PSID in comparison to the CE,” *American Economic Review*, 2014, *104* (5), 132–35.
- Berger, David and Joseph Vavra**, “Measuring How Fiscal Shocks Affect Durable Spending in Recessions and Expansions,” *The American Economic Review*, 2014, *104* (5), 112–115.
- **and —** , “Consumption Dynamics During Recessions,” *Econometrica*, 2015, *83* (1), 101–154.
- Bernstein, Shai, Rebecca Diamond, Timothy J. McQuade, and Beatriz Pousada**, “The Contribution of High-Skilled Immigrants to Innovation in the United States,” Working Paper 2018.
- Blundell, Richard, Luigi Pistaferri, and Ian Preston**, “Consumption Inequality and Partial Insurance,” *American Economic Review*, December 2008, *98* (5).
- Boar, Corina**, “Dynastic Precautionary Savings,” *The Review of Economic Studies*, 03 2021, *88* (6), 2735–2765.

- , **Denis Gorea**, and **Virgiliu Midrigan**, “Liquidity Constraints in the U.S. Housing Market,” *The Review of Economic Studies*, 09 2021, *89* (3), 1120–1154.
- Couture, Victor**, **Cecile Gaubert**, **Jessie Handbury**, and **Erik Hurst**, “Income Growth and the Distributional Effects of Urban Spatial Sorting,” Working Paper 26142, National Bureau of Economic Research August 2019.
- , **Jonathan I Dingel**, **Allison E Green**, and **Jessie Handbury**, “Demographic Preferences and Income Segregation,” Working Paper 33386, National Bureau of Economic Research January 2025.
- Cutler, David M** and **Edward L Glaeser**, “Are ghettos good or bad?,” *The Quarterly Journal of Economics*, 1997, *112* (3), 827–872.
- , —, and **Jacob L Vigdor**, “The rise and decline of the American ghetto,” *Journal of political economy*, 1999, *107* (3), 455–506.
- , —, and —, “When are ghettos bad? Lessons from immigrant segregation in the United States,” *Journal of Urban Economics*, 2008, *63* (3), 759–774.
- Davis, Donald R**, **Jonathan I Dingel**, **Joan Monras**, and **Eduardo Morales**, “How Segregated is Urban Consumption?,” Technical Report, National Bureau of Economic Research 2017.
- Diamond, Rebecca** and **Enrico Moretti**, “Where is Standard of Living the Highest? Local Prices and the Geography of Consumption,” Technical Report 2022.
- , **Tim McQuade**, and **Franklin Qian**, “The effects of rent control expansion on tenants, landlords, and inequality: Evidence from San Francisco,” *American Economic Review*, 2019, *109* (9), 3365–94.
- Dolfen, Paul**, **Liran Einav**, **Peter J Klenow**, **Benjamin Klopach**, **Jonathan D Levin**, **Laurence Levin**, and **Wayne Best**, “Assessing the Gains from E-Commerce,” Working Paper 25610, National Bureau of Economic Research February 2019.
- Eckert, Fabian** and **Tatjana Kleineberg**, “Saving the American Dream? Education Policies in Spatial General Equilibrium,” *Education Policies in Spatial General Equilibrium (March 9, 2021)*, 2021.

- Fernández, Raquel and Richard Rogerson**, “Public Education and Income Distribution: A Dynamic Quantitative Evaluation of Education-Finance Reform,” *The American Economic Review*, 1998, 88 (4), 813–833.
- Fisher, Jonathan D., David S. Johnson, and Timothy M. Smeeding**, “Measuring the Trends in Inequality of Individuals and Families: Income and Consumption,” *American Economic Review*, May 2013, 103 (3), 184–88.
- Fogli, Alessandra and Veronica Guerrieri**, “The End of the American Dream? Inequality and Segregation in US Cities,” Working Paper 26143, National Bureau of Economic Research August 2019.
- Ganong, Peter and Pascal Noel**, “Consumer spending during unemployment: Positive and normative implications,” *American economic review*, 2019, 109 (7), 2383–2424.
- Giannone, Elisa, Qi Li, Nuno Paixao, and Xinle Pang**, “Unpacking Moving,” Working Paper 2019.
- Guren, Adam M, Alisdair McKay, Emi Nakamura, and Jón Steinsson**, “Housing Wealth Effects: The Long View,” *The Review of Economic Studies*, 04 2020.
- Krueger, Dirk and Fabrizio Perri**, “Does income inequality lead to consumption inequality? Evidence and theory,” *The Review of Economic Studies*, 2006, 73 (1), 163–193.
- Massey, Douglas S**, “Residential segregation is the linchpin of racial stratification,” *City & community*, 2016, 15 (1), 4.
- Massey, Douglas S., Jonathan Rothwell, and Thurston Domina**, “The Changing Bases of Segregation in the United States,” *The ANNALS of the American Academy of Political and Social Science*, 2009, 626 (1), 74–90.
- Morgan, Donald P., Maxim L. Pinkovskiy, and Bryan Yang**, “Banking Deserts, Branch Closings, and Soft Information,” Liberty Street Economics 20160307, Federal Reserve Bank of New York March 2016.
- Olafsson, Arna and Michaela Pagel**, “The Liquid Hand-to-Mouth: Evidence from Personal Finance Management Software,” *The Review of Financial Studies*, 04 2018, 31 (11), 4398–4446.

Phillips, David, “Measuring Housing Stability with Consumer Reference Data,” 2019.

Qian, Franklin and Rose Tan, “The Effects of High-skilled Firm Entry on Incumbent Residents,” 2020.

Reardon, Sean F and Glenn Firebaugh, “Measures of multigroup segregation,” *Sociological methodology*, 2002, 32 (1), 33–67.

Theloudis, Alexandros, “Consumption inequality across heterogeneous families,” Working Paper 2020.

Trounstein, Jessica, *Segregation by Design: Local Politics and Inequality in American Cities*, Cambridge University Press, 2018.

Zheng, Angela and James Graham, “Public Education Inequality and Intergenerational Mobility,” *Available at SSRN 3092714*, 2020.

Appendices

A Data Appendix

A.1 Data Sources

In this section, we describe in further detail the rest of the data sources we use. We first give a broad overview of the ACS and the Nielsen Homescan datasets. We also provide more detail on the structure of the Infutor data and the variables we use to measure the housing and car consumption segregation. Finally, we explain our imputation procedure for car and house prices and provide additional detail on the geographic crosswalk.

A.1.1 American Community Survey

The ACS is an annual survey conducted by the Census Bureau since 2003 where individuals are randomly sampled in each state, the District of Columbia and Puerto Rico. We use the information on income, education, age, race and geography (PUMA and state) from the ACS for two reasons. First, we use all the information above to validate the representativeness of the Infutor data. Second, we use the information on race and income to measure residential segregation by race, income, age and education, all of which are important variables in our analysis but are imputed in the Infutor data. We restrict the ACS sample to individuals between 22 and 80 and exclude observations with non-positive income. The income variable we focus on is total household income. Although the car and home characteristics data from Infutor is at the level of a car-individual owner or home-individual owner level, we view these two consumption categories as likely to be used jointly by all members of a household.

Although the ACS is available since the year 2000 we exclude the years before 2005 as the publicly available data does not report individuals' locations besides the state of residence. In all our analysis we use the finest unit of geography available in the ACS, which is the Public Use Microdata Area (PUMA). There are 2,351 different PUMAs covering all the U.S. territory and they are state-specific, which means that PUMAs do not cross state borders.

The main ACS variables we use in the analysis are total household income (`HHINCOME`) and house value (`VALUEH`). Our results for income segregation are robust to considering alternative income variables such as total personal earned income (`INCEARN`) or total personal income (`INCTOT`).

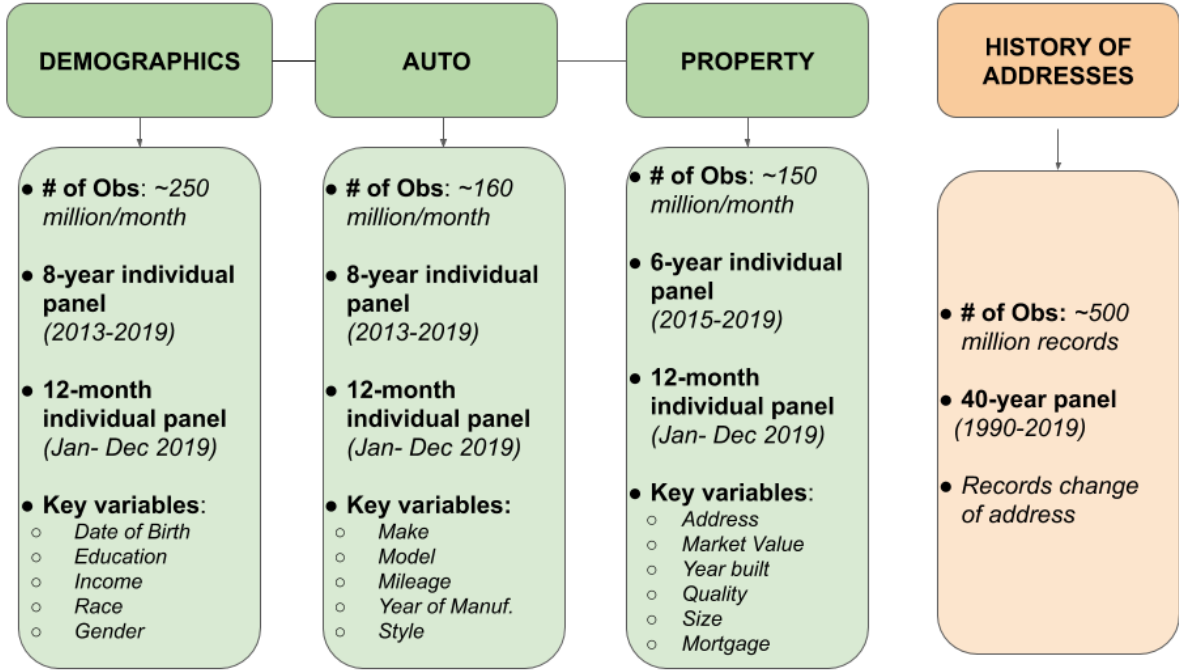
A.1.2 Census

To complement the ACS we use the 5% microdata sample from the 1990 and 2000 Census available in IPUMS. This sample contains information on approximately 10 million individuals and 5 million households. As with the ACS, from the Census we use in the analysis information on total household income and house value.

A.1.3 Infutor

The Infutor data is organized as described in Figure 6. Here we supplement the description of the dataset from the main text with additional details.

Figure 6: Structure of Infutor Data



Notes: Figure 6 describes the structure of the Infutor data divided in four main files. The first three files (Demographic Profiles, Auto Profiles and Property Profiles) have been linked through individual identifiers.

For both housing and vehicles we observe in the dataset the ZIP code where the house or the vehicle is located. ZIP codes are more than 10 times more disaggregated than PUMAs. However, given that our analysis for income is done at the PUMA level we aggregate ZIP codes to PUMAs, which we use as the finest geographic aggregation level for all the analysis, i.e. the geographic unit i in equation (1). Our segregation results for housing and vehicles are qualitatively robust to considering alternative definitions of the geographic unit i , such

as ZIP codes or counties, instead of PUMAs. In Section A.4 we describe the ZIP-PUMA crosswalk we use to assign the ZIP codes in the Infutor data to PUMAs.

Housing. For each property we observe a unique identifier (given by the property address), an owner unique identifier, its location, its characteristics, such as the year of construction (`PROP_YRBLD`) or the number of rooms, information about the property deed, and its history of mortgages. Infutor also reports whether the property is used as a business or as a residence. For all our analysis we restrict to properties that are used for residential purposes only.

For houses we mostly rely on the deed data from which we observe information of the date (`PROP_SALEDATE`) and the price (`PROP_SALEAMT`) at which the property was acquired by the current owner in a given snapshot. Furthermore, for approximately 68% of the deed records we also observe these two variables for the previous deed of the property (`PROP_SALEDATE_PRIOR` and `PROP_SALEAMT_PRIOR`). Combined with the multiple snapshots of the data, this implies that for several properties we observe more than one transaction. On average we observe 1.86 transactions per property with a maximum of 7 transactions for some of the properties.

Sample selection. Using the variables described above, we include a house in the year t sample if:

$$\min \left\{ \text{PROP_YRBLD}, \min_{\text{snapshot}} \{ \text{PROP_SALEDATE_PRIOR}, \text{PROP_SALEDATE} \} \right\} \leq t, \quad (7)$$

where all the date variables are in years and the minimum is computed allowing for possible missing values. This definition implies that we include in our sample all the properties which we can verify that have been built or sold at any previous, or current, date. In Section A.2 we explain how we interpolate selling prices to compute the value of all in-sample houses across time.

Vehicles. For vehicles we observe unique owner-vehicle identifier together with the owner's address which allows us to identify vehicles' location. Unlike houses, we cannot track the same vehicle across multiple owners. We also observe relevant characteristics of the vehicle such as the manufacturer (`MAKE`), the model (`MODEL`), the year (`YEAR`), and the first 10-digits of the Vehicle Identification Number (`VIN`). For example, an entry in our data is a BMW-5 Series-2015 with VIN code WBA5M6C5xF.

As with housing, we combine all the information available at the different snapshots. This allows us to increase the number of vehicles across time, particularly as the coverage

significantly increased in the most recent snapshots. Regarding the vehicle location, for the years between 2012 to 2018 we compute the segregation index using the current snapshot ZIP code, when available. For the rest of the years prior to the first time we observe a vehicle we impute the ZIP code of the first snapshot in which we observe it. This assumption is likely to only have a limited impact as between 2012 and 2018 on average only 3.3% of vehicles changed PUMAs and only 1.1% changed state.

Sample selection. A vehicle is included in the year t sample if

$$\text{YEAR} \leq t \leq \max \{\text{snapshot}\}, \quad (8)$$

where the max operator is computed across, potentially different snapshot years for which we observe the same vehicle.

Two comments about the vehicle sample selection are in order. First, the backward-looking nature of equation (8), going back to the vehicle’s manufacturing year, allows us to run our segregation analysis even for years before Infutor’s first snapshot, which corresponds to the year 2012 for the case of vehicles. Thus, for example, if we observe a Toyota-Corolla-2007 in the 2014 snapshot we include this vehicle in our segregation analysis starting from the year 2007 even though we observe this vehicle for the first time in 2014. In other words, what definitions (7) and (8) are doing is to include houses and vehicles since the first year we verify that they exist. Second, definition (8) allows for vehicle exit given by the last snapshot in which the vehicle was observed. We allow for that as for the cases in which Infutor confirmed that the vehicle was sold the owner-vehicle entry will stop appearing in the data in all subsequent snapshots. In Section A.3 we describe the detail of how we obtain the value of the vehicles in our sample and across time.

To test the validity of our vehicles’ dataset, we exploit information on age and manufacturers. Specifically, the top panel of Table 10 reports the average age of vehicles in Infutor and BTS between 2012 and 2018. The average age in Infutor is slightly higher but both can be approximated to be at 12 years. The table also reports the correlation at state level of the average age. We find this correlation to be 0.54. The central (bottom) panel reports the age of vehicles in Infutor and NHTS in 2009 (2017). We find that the average age is 2 years higher (3 years lower) in NHTS than in Infutor and the state-level correlation is 0.37 (0.33). We also analyze the distribution of make in both datasets. We find that overall the distributions are very similar both at national- and state-level. Table 11 reports the average age of vehicles at national level for each year and the BTS datasets and in both versions of the Infutor datasets that we built. Overall we find that comparing the snapshot Infutor

dataset to the BTS, the latter has on average slightly younger vehicles. Instead, when we compare the baseline dataset to the BTS, we find that the latter has slightly older cars for the years before 2016 and it reverts afterward. Overall, the differences are not very large.

A.1.4 Nielsen Homescan

For non-durable consumption we use the Nielsen Homescan Data from the Kilts Marketing Data Center at the University of Chicago Booth School of Business. These data consist on a longitudinal panel of approximately 40,000-60,000 U.S. households and contain information about the products they buy, as well as when and where they were purchased between 2004 and 2017. It tracks consumers' grocery purchases by asking them to scan the bar codes for each product they purchase after each shopping trip. The Nielsen Homescan data has national coverage and provides wide variation in household location and demographics. Overall, the data include purchases of almost 250 million different items. One of the advantages of this dataset is that it records the bar codes at a very fine level, as well as the expenditure on each of them. Moreover, it covers longer period of time than Infutor. A disadvantage of the Nielsen Homescan dataset relative to the Infutor data is that the sample is orders of magnitude smaller, which means that we would not be able to use these data at very fine geographical detail.

The Nielsen-Kilts data reports detailed location of households. In particular it reports households' ZIP code of residence. We use the same crosswalk used for the Infutor data to assign households to the different PUMAs.

Household purchases are reported at a very granular level. Specifically, at the Universal Product Code (UPC) level. This data reports both the quantity and the total expenses made at the UPC level. Besides UPC there are other two additional aggregation levels which are product modules, such as soap or beer, and ten different department codes such as dairy or health and beauty consumption.

A.2 House Prices

This section explains how we compute prices for in-sample houses in the Infutor data, defined in equation (7), at any time t . We use observed selling prices to compute ZIP code level price indices with which we interpolate property prices at different time periods.

In our analysis, we restrict attention to properties for which we observe the date and the price of at least one transaction. As explained above, for each property we keep all the selling

Table 10: Representativeness in Infutor and BTS datasets

		National-level		State-level	
		BTS	NHTS	Infutor	FHA NHTS
2012-2018					
Vehicle Age					
	Mean	11.5		12.2	0.54
2009					
Vehicle Age					
	Mean	10.3	9.4	7.1	0.37
Vehicle Make*					
	Chrysler		0.12	0.15	0.74
	Ford		0.19	0.19	0.59
	GM		0.22	0.22	0.62
	Honda		0.09	0.08	0.40
	Toyota		0.12	0.10	0.22
	Other		0.26	0.26	0.36
N		254,582,694	211,501,318	98,996,193	0.32 0.44
2017					
Vehicle Age					
	Mean	11.7	10.4	13.3	0.33
Vehicle Make*					
	Chrysler		0.11	0.13	0.71
	Ford		0.15	0.18	0.65
	GM		0.19	0.22	0.73
	Honda		0.11	0.08	0.57
	Toyota		0.15	0.11	0.40
	Other		0.30	0.28	0.55
N		254,582,694	222,578,947	158,898,357	0.56 0.54

Notes: The table reports the average age of a car across datasets, as well as the distribution of car makes.

Table 11: Age distribution of cars in Infutor and BTS

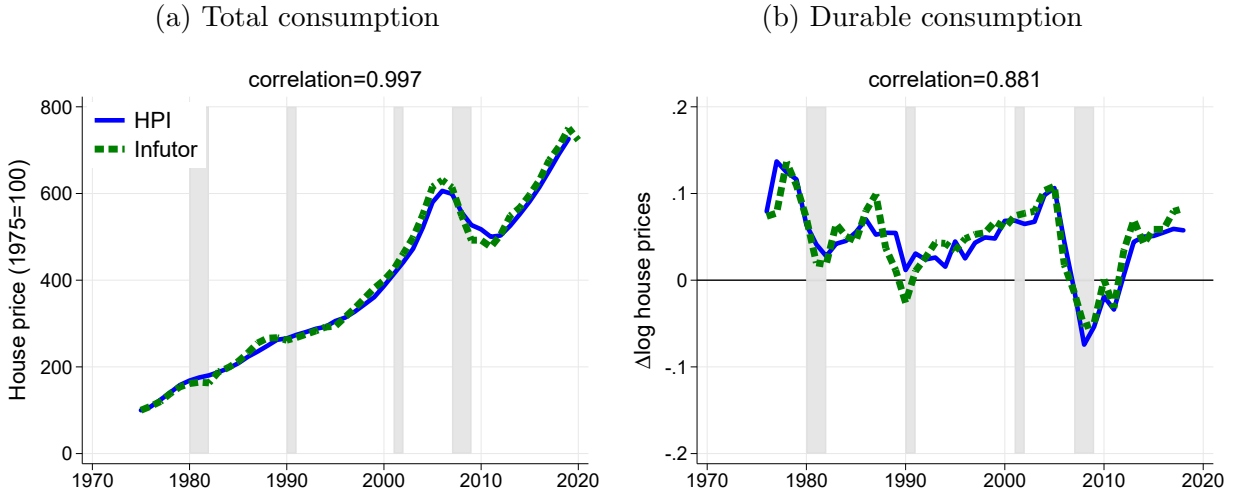
Vehicle Age	National-level		
	BTS	Snapshot	Baseline
2000	8.9		5.7
2001	8.9		5.9
2002	9.6		6.0
2003	9.7		6.1
2004	9.8		6.3
2005	9.8		6.6
2006	9.9		6.9
2007	10		7.2
2008	10.1		7.4
2009	10.3		8.0
2010	10.6		8.4
2011	10.9		8.8
2012	11.2	9.7	9.1
2013	11.4	12.0	9.5
2014	11.4	11.6	10.2
2015	11.5	12.1	10.9
2016	11.6	12.4	11.7
2017	11.7	13.3	12.6
2018	11.7	14.4	13.5
Corr with BTS			
2012-2018		0.950	0.952
2000-2018			0.935

Notes: The table reports the age distribution of cars in BTS and Infutor. For Infutor, we report both the “Snapshot” and the “Baseline” version of the data. The bottom of the table reports correlation coefficients of the average age by year between datasets.

dates and prices available in the data. The total number of transactions we observe varies by year. For example, for the years before 1990, we observe around 600,000 annual transactions. For the time period between 1990 and 2018, we observe more than 3 million transactions per year.

Using the selling prices we compute a time series for the median transacted house price at different geographies: at the national, state, PUMA and ZIP code level. To assess the quality of our house price data, Figure 7 displays the level and the annual log changes of the national-level median price against the House Price Index (HPI) from the Federal Housing Finance Agency. This figure shows that the national-level house prices obtained from Infutor are consistent with aggregate-level house price dynamics and properly capture the different boom and bust episodes observed in the last four decades.

Figure 7: House Prices



Notes: The figure compares the evolution over time of house prices and house price growth in Infutor and HPI. HPI denotes the All-Transactions House Price Index from the Federal Housing Finance Agency. Infutor denotes the national-level median selling price.

ZIP Code House Price Indices. We use the data to construct ZIP code-level price indices from 1950 to 2018. To that end, we first compute annual price changes across all four geographies for which we observe at least 50 transactions. If this constraint is not met for a given ZIP-year pair we try to use the price change of the immediate broader unit of geography level, which is the PUMA. If the corresponding PUMA year does not have the price change available we continue to the state or even the national level. These imputations are more common as we go back in time, particularly, before 1990.

Price Interpolation. With the corresponding ZIP code level price index in which the

property is located we interpolate prices to other years different from the property observed transactions. This strategy takes into account different boom and bust house price cycles observed at different geographies. To exemplify our interpolation strategy, suppose that we observe that house h , located in ZIP code z , was sold at price $p_{h,z,\tau}$ in period τ . We compute the price of house h at any other year t by assuming that its price between period t and τ followed the zip code level price dynamics. This assumption can be written as

$$\Delta_{t-\tau} \log p_{h,z,\tau} = \Delta_{t-\tau} \log p_{z,\tau},$$

which implies that house h price at year t is equal to

$$p_{h,z,t} = \exp(\log p_{h,z,\tau} + \Delta_{t-\tau} \log \log p_{z,\tau}),$$

where $p_{z,\tau}$ is ZIP code z price index in period τ .

Note that in the expression above we could have $t < \tau$ as long as equation (7) is satisfied. For the properties for which we observe more than one transaction we interpolate the price at period t using the closet (in terms of minimum time distance) selling price.

A.3 Vehicle Prices

One limitation of the Infutor vehicle data is that we do not observe prices or any other estimate for the value of vehicles. To circumvent this limitation we merged the Infutor data with transaction prices for all the dealer sales of new and used cars in Texas from 2012 to 2019. Dealer sales refer to all the sales done by licensed dealers.

We perform this merge in two steps, trying to use the most disaggregated data available. To that end, we first compute time series for 10-digit VIN-level mean selling prices for all the VINs for which we observe at least 50 transactions. Additionally, we compute mean selling prices, also restricting to the minimum 50 observations constraint, for all the make-model-year combinations. VIN codes are considerably narrower specifications compared to make-model-year tuples as they also distinguish other technical characteristics such as the vehicle engine or the type of fuel it uses. To put this in perspective, in Infutor, we observe more than 130,000 different 10-digit VINs and close to 15,000 make-model-year combinations. In our data around 50% of the make-model-year tuples have at most 4 different 10-digit VINs.

Using both sets of time series we compute VIN-level prices for 2012 to 2018. As we did for houses, we impute missing VIN-level prices using the make-model-year price when the VIN-level price is not available. For these years we drop observations below the 10^{th} and

above the 90th percentile annual price change distribution. This restriction rules out, for example, price increases in used cars, which probably reflect changes in the composition of sold vehicles over time.

Using prices from 2012 to 2018, we then interpolate prices for missing observations and for the years before 2012, for the older vehicles in our sample. For that, we use a VIN-level annual depreciation rate. Specifically, for each VIN v we compute the average depreciation rate between 2012 and 2018 as

$$\delta_v = -\frac{1}{\#\mathcal{T}_v} \sum_{t \in \mathcal{T}_v} \Delta \log p_{v,t},$$

where $\mathcal{T}_v \subset \{2012, \dots, 2018\}$ is the set of years for which we observe VIN v 's prices. The average depreciation rate (or average annual price decrease) in our VIN-level prices is 0.108.

With this VIN-level depreciation rate we interpolate missing prices for our in-sample vehicles in year t that satisfy equation (8). For example, for a 2005 VIN (YEAR= 2005) for which we observe selling prices starting in 2012 we compute its time $2005 \leq t < 2012$ price as

$$p_{v,t} = \exp [\log p_{v,2012} + (2012 - t)\delta_v],$$

where $p_{v,2012}$ is the VIN-level selling price observed in 2012.

A.4 Geographic Crosswalks

In addition to the country and U.S. states, the economic geography and segregation literature usually consider five other levels of geography: (1) ZIP-code, (2) county, (3) Public Use Microdata Area (PUMA), (4) Core Based Statistical Area (CBSA), and (5) Commuting Zone (CZ). These categories are presented in Table 12, ordered from left to right in terms of their level of disaggregation.

We make some remarks about these units of geography. First, ZIP-codes do not cover the entire territory, approximately 0.002% of population is not assigned to a ZIP-code. Additionally, a ZIP-code can cover more than one state. Second, counties cover the entire territory and are state-specific. Third, PUMAs also cover the entire territory and are state-specific. As mentioned above, this is the finest aggregation level available in the ACS and Census public use microdata. Fourth, CBSAs do not cover the entire territory, approximately 6% of population is not assigned to a CBSA. CBSAs can cover more than one state. For example, the CBSA New York-Newark-Jersey City, spans three states: New York, New Jersey and

Table 12: Number of Locations by Geographic Aggregation Level

	ZIP	County-10	PUMA	CBSA-15	CZ
MCDC	32,845	3,143	2,351	933	741
ACS [2005-2018]		376-430	2,351		
Census 1990		385	1,726		
Census 2000		377	2,351		
Infutor [2012-2019]	39,131				

Notes: The table reports the number of geographic units from smaller to larger. MCDC denotes the Missouri Census Data Center. County-10 and CBSA-15 denote the list of counties and CBSAs in 2010 and 2015 respectively. The Missouri Census Data Center Geographic Correspondence Engine is a comprehensive source of definitions and geographical cross-walks for the United States. We compare the zipcodes available in Infutor with the ones in the MCDC data.

Pennsylvania. Fifth, MSAs (not reported in the table), are, broadly, a subgroup of CBSAs that only restricts to metropolitan areas. There are 384 MSAs (vs. 933 CBSAs). Lastly, CZs are clusters of U.S. counties that are characterized by strong within-cluster and weak between-cluster commuting ties.

For our analysis we constructed a crosswalk that uniquely assigns ZIPs to PUMAs (an injection correspondence), and consequently, to states. We started from the ZIP-PUMA crosswalk available in the MCDC. This crosswalk, besides assigning ZIPs to PUMAs, reports the ZIP-code population share that is assigned to the potentially different PUMAs. Indeed, for approximately 28% of the ZIP-codes there is more than one PUMA assigned. In these cases, we assigned each ZIP to a single PUMA using the one with the largest population share. For example, if a ZIP-code is assigned to two PUMAs with 70-30% population shares we assigned this ZIP-code to the former PUMA with a 70% population share. With this procedure we constructed an injective crosswalk between ZIP codes to PUMAs. This crosswalk is available upon request.

B Additional Figures and Tables

This section provides several robustness tests for the entropy measure. First, we change the subunit of geography from PUMA to zipcodes and report the correlation coefficients between entropy indices at the national or state level calculated using zipcodes, counties and PUMAs as the underlying geographic unit. Table 13 suggests that the correlations between the entropy indexes for car and housing consumption when the smallest unit of geography is the zipcode, county or PUMA are very highly correlated both at the national level and at

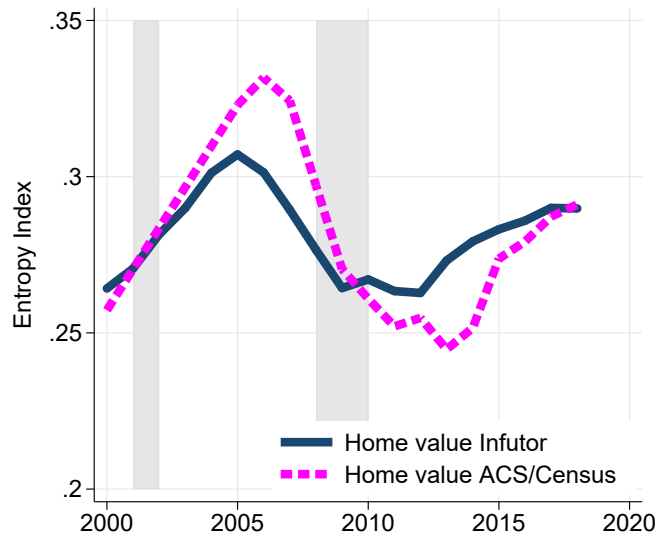
the state level. The correlations range between 0.639 and 1, suggesting that the analysis at PUMA level is robust of the analysis at the zipcode level. Second, Figure 8 compares the results for entropy measure in house values both in Infutor and ACS. Third, we change the larger unit of geography from states to CBSAs and report the main results for the entropy index at CBSA level rather than at the state level.

Table 13: Correlations between Entropy Indexes at Different Levels of Geography

	State-level		National-level	
	Housing - PUMA	Car - PUMA	Housing - PUMA	Car - PUMA
Housing - County	0.717***		0.971***	
Housing - Zipcode	0.971***		0.976***	
Car - County		0.812***		1.000***
Car - Zipcode		0.639***		0.992***
Observations	2499	1224	49	24

Notes: The table reports the correlations between H-indexes for housing and car consumption measured using the Infutor data for different levels of underlying geographies. On the left panel, we report correlations between H-indexes aggregated at the national level. On the right, we report correlations for H-indexes aggregated at state-level.

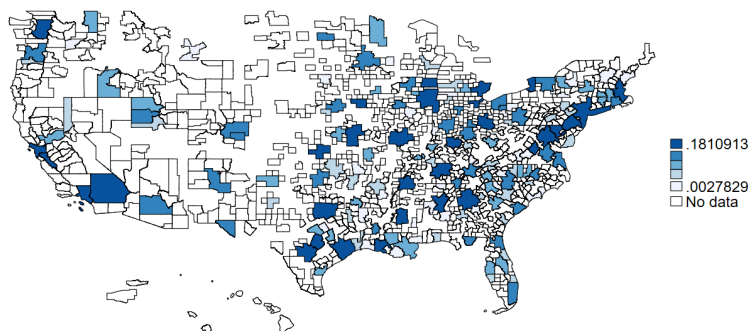
Figure 8: Comparing Trends in Entropy Index for Infutor and ACS Housing Data



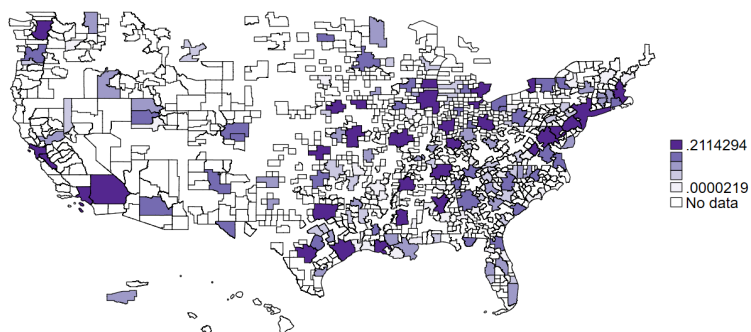
Note: The figure plots consumption entropy indices for house values in Infutor and ACS.

Figure 9: CBSA-level Consumption Segregation

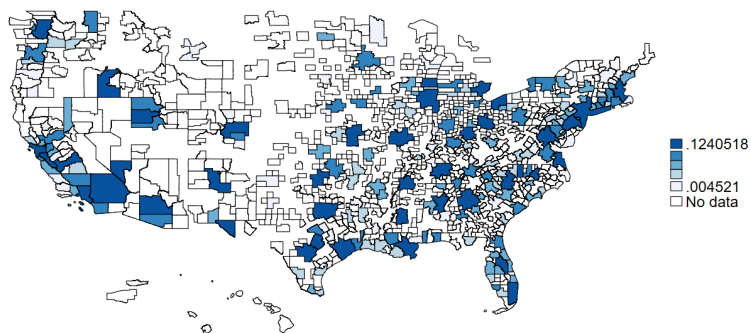
(a) Total consumption



(b) Durable consumption

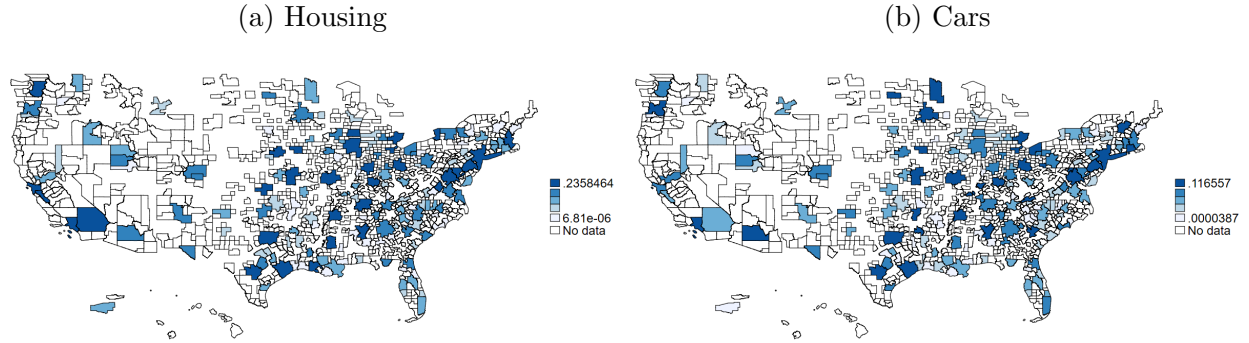


(c) Non-durable consumption



Notes: The figure plots CBSA-level consumption entropy indices for durables, housing and cars, averaged over the period 2016-2018.

Figure 10: CBSA-level Durable Consumption Segregation



Notes: The figure plots CBSA-level consumption entropy indices averaged over the period 2016-2018.

Table 14: Segregation and Regional Characteristics

	Total consumption	Durable consumption	Non-durable consumption	Income
<i>By income</i>				
Bottom 25%	0.096	0.107	0.067	0.037
Top 25%	0.200	0.251	0.085	0.064
<i>By education (college share)</i>				
Bottom 25%	0.101	0.113	0.068	0.039
Top 25%	0.188	0.236	0.072	0.063
<i>By race (share white)</i>				
Bottom 25%	0.200	0.246	0.085	0.062
Top 25%	0.098	0.115	0.056	0.039
<i>By age (share ≤ 35)</i>				
Bottom 25%	0.118	0.141	0.060	0.044
Top 25%	0.172	0.207	0.089	0.061
<i>By population</i>				
Bottom 25%	0.061	0.068	0.044	0.021
Top 25%	0.169	0.206	0.078	0.060

Notes: The table reports population weighted averages of state-level consumption entropy indices over the period 2016-2018.

Table 15: Segregation by Demographic Group

	Total consumption	Durable consumption	Non-durable consumption
<i>By race</i>			
White	0.143	0.166	0.083
Black	0.156	0.170	0.118
Other	0.152	0.183	0.070
<i>By education</i>			
No college degree	0.148	0.168	0.095
College degree	0.166	0.178	0.135
<i>By age</i>			
Younger than 35	0.167	0.184	0.123
Older than 35	0.148	0.176	0.076

Notes: The table reports population weighted averages of state-level consumption entropy indices over the period 2016-2018.

Table 16: Segregation by Demographic Group

	Total consumption	Durable consumption	Non-durable consumption	Income
<i>By race</i>				
White	0.143	0.166	0.083	0.050
Black	0.156	0.170	0.118	0.076
Other	0.152	0.183	0.070	0.093
<i>By education</i>				
No college degree	0.148	0.168	0.095	0.039
College degree	0.166	0.178	0.135	0.046
<i>By age</i>				
Younger than 35	0.167	0.184	0.123	0.062
Older than 35	0.148	0.176	0.076	0.053

Notes: The table reports population weighted averages of state-level consumption entropy indices over the period 2016-2018.

Table 17: Drivers of Consumption Segregation at State Level

	Total consumption			Durable consumption			Non-durable consumption		
	(1)	(2)	(3)	(1)	(2)	(3)	(1)	(2)	(3)
$\ln H_Y$	0.60*** (0.05)	0.49*** (0.06)	0.47*** (0.08)	0.77*** (0.05)	0.65*** (0.06)	0.49*** (0.09)	0.20*** (0.06)	0.25*** (0.07)	0.21** (0.09)
$\ln H_R$		0.14*** (0.07)	0.14* (0.08)		0.16*** (0.06)	0.14** (0.07)		-0.06 (0.05)	-0.07 (0.06)
$\ln H_E$			0.04 (0.12)			0.22** (0.11)			0.07 (0.09)
$\ln H_A$			-0.05 (0.05)			-0.08 (0.05)			-0.02 (0.07)
R^2	0.82	0.84	0.84	0.79	0.81	0.82	0.53	0.53	0.53

Note: The table reports the estimates of α_1 , α_2 , α_3 and α_4 in equation (6). The estimate of α_0 , the time fixed effects and the estimate of the vector γ are omitted. Robust standard errors are reported in parentheses. The regression is estimated with population weights based on the 2010 Census. ***, **, and *, represent statistical significance at 1%, 5% and 10%, respectively.

Table 18: Determinants of Consumption Segregation, Heterogeneity

	Total consumption	Durable consumption	Non-durable consumption
<i>By income</i>			
$\ln H_Y$	0.60*** (0.05)	0.83*** (0.08)	0.34*** (0.05)
$\ln H_Y \times \mathbf{1}_{\text{Top 25\%}}$	0.23 (0.16)	0.09 (0.18)	0.15* (0.08)
<i>By education (college share)</i>			
$\ln H_Y$	0.61*** (0.05)	0.82*** (0.08)	0.41*** (0.05)
$\ln H_Y \times \mathbf{1}_{\text{Top 25\%}}$	0.10 (0.08)	0.11 (0.13)	0.008 (0.07)
<i>By race (share white)</i>			
$\ln H_Y$	0.69*** (0.06)	0.87*** (0.05)	0.48*** (0.07)
$\ln H_Y \times \mathbf{1}_{\text{Top 25\%}}$	-0.06 (0.09)	-0.12 (0.09)	-0.11 (0.08)
<i>By age (share ≤ 35)</i>			
$\ln H_Y$	0.67*** (0.05)	0.77*** (0.06)	0.52*** (0.04)
$\ln H_Y \times \mathbf{1}_{\text{Top 25\%}}$	-0.03 (0.10)	0.06 (0.12)	-0.02 (0.07)
<i>By population</i>			
$\ln H_Y$	0.14 (0.09)	0.23** (0.09)	-0.02 (0.07)
$\ln H_Y \times \mathbf{1}_{\text{Top 25\%}}$	0.60*** (0.15)	0.71*** (0.16)	0.26** (0.11)

Notes: The table reports slope coefficients from regressing consumption segregation on income segregation and income segregation interacted with a dummy variable that equals 1 if a CBSA is in the top 25% of the national distribution of income, the share of the population that is college educated, white, younger than 35, and population, respectively. All regressions include separate controls for the dummy variable previously described and, as well as the set of controls in equation (6). All other coefficients are omitted in the interest of space. With the exception of the last panel, all regressions are estimated with population weights based on the 2010 Census. ***, **, and *, represent statistical significance at 1%, 5% and 10%, respectively.

Table 19: Consumption and Income Segregation, State Level

	Total	Between	Within
Total consumption	2.05	0.04	2.01
α_1^i/α_1		2%	98%
Non-durable consumption	0.71	0.25	0.47
α_1^i/α_1		35.2%	64.8%
Durable Consumption	2.99	0.15	2.84
α_1^i/α_1		5%	95%
Cars	0.14	0.02	0.12
α_1^i/α_1		14.3%	85.7%
Housing	3.34	0.13	3.21
α_1^i/α_1		3.9%	96.1%

Notes: The table reports the estimates of α_1 , α_1^B and α_1^W . All estimates are significant at 1% significance. Robust standard errors, time fixed effects and the estimates of γ^i and γ are omitted but are included in the estimation. The regression is estimated with population weights based on the 2010 Census, using data for the 2016-2018 period.