

WORKING PAPER · NO. 2018-47

Coming apart? Cultural distances in the United States over time

Marianne Bertrand and Emir Kamenica

July 2018

Coming apart? Cultural distances in the United States over time

Marianne Bertrand and Emir Kamenica*

June 2018

Abstract

We analyze temporal trends in cultural distance between groups in the US defined by income, education, gender, race, and political ideology. We measure cultural distance between two groups as the ability to infer an individual's group based on his or her (i) media consumption, (ii) consumer behavior, (iii) time use, or (iv) social attitudes. Gender difference in time use decreased between 1965 and 1995 and has remained constant since. Differences in social attitudes by political ideology and income have increased over the last four decades. Whites and non-whites have converged somewhat on attitudes but have diverged in consumer behavior. For all other demographic divisions and cultural dimensions, cultural distance has been broadly constant over time.

* University of Chicago Booth School of Business. We thank the University of Chicago Booth School of Business for financial support. Jihoon Sung provided excellent research assistance. We thank Jann Spiess and Clara Marquardt for their advice and guidance with the implementation of the machine learning algorithms.

What's great about this country is that America started the tradition where the richest consumers buy essentially the same things as the poorest. You can be watching TV and see Coca-Cola, and you know that the President drinks Coke, Liz Taylor drinks Coke, and just think, you can drink Coke, too. A Coke is a Coke and no amount of money can get you a better Coke than the one the bum on the corner is drinking. All the Cokes are the same and all the Cokes are good. Liz Taylor knows it, the President knows it, the bum knows it, and you know it.

- Andy Warhol

1 Introduction

While rural America watches *Duck Dynasty* and goes fishing and hunting, urban America watches *Modern Family* and does yoga in the park.¹ The economically better-off travel the world and seek out ethnic restaurants in their neighborhoods, while the less well-off don't own a passport and eat at McDonald's.² Conservatives give their boys masculine names like Kurt, while liberals opt for the more feminine-sounding options such as Liam.³ While men play video games and watch pornography, women browse Pinterest and post pictures on Instagram.⁴ These are just a few examples of the cultural distances across groups within America today. The presence of such cultural divides is not new – in the early 2000s, scholars emphasized racial differences in music tastes, language use, media consumption, and consumer behavior⁵ – but there is a perception that cultural distances are growing,⁶ with a particular emphasis on increasing political polarization.⁷

These cultural distances may have important consequences. A large empirical literature in political economy documents that high levels of ethno-linguistic fragmentation hinder public good provision (Alesina et al. 1999; Alesina and La Ferrara 2005), decrease social capital (Alesina and La Ferrara, 2000), and increase the probability of conflict (Montalvo and Reynal-Querol 2005). Moreover, Desmet et al. (2017) suggest that these outcomes especially worsen when cultural differences

¹<https://www.nytimes.com/interactive/2016/12/26/upshot/duck-dynasty-vs-modern-family-television-maps.html>

²<https://www.theatlantic.com/national/archive/2011/03/americas-great-passport-divide/72399/>

³See Oliver, Wood, and Bass (2016).

⁴<http://www.pewinternet.org/2005/08/18/adult-content-online/>; https://www.washingtonpost.com/news/the-switch/wp/2013/10/10/25-percent-of-men-watch-online-porn-and-other-facts-about-americans-online-video-habits/?utm_term=.450a3dfccb89; <http://www.pewresearch.org/fact-tank/2015/08/28/men-catch-up-with-women-on-overall-social-media-use>.

⁵See Waldfogel (2003), Wolfram and Thomas (2002), and Fryer and Levitt (2004).

⁶Fryer and Levitt (2004) document an increase in prevalence of distinctively black names over time. Focusing on differences across socio-economic groups within the white population, Murray (2012) writes, “It is not the existence of classes that is new, but the emergence of classes that diverge on core behavior and values – classes that barely recognize their underlying American kinship.”

⁷See Kaufman (2002) on the increasing gender gap in party affiliation and Gentzkow (2016) on trends in polarization across party lines.

across ethnic groups are greater.

Sociologists such as Pierre Bourdieu (1984 [1979]) provide some theoretical foundations for the findings in the political economy literature. Bourdieu was concerned with the concept of cultural capital, which he associates with the set of tastes, mannerisms, or material belongings that one holds. Sharing cultural capital with others, Bourdieu argues, creates a sense of having a common identity. When cultural differences between groups increase, these groups find it more difficult to interact, communicate, and trust each other. Bourdieu was particularly concerned about how cultural differences between rich and poor damage social mobility. For example, students from poorer backgrounds might better integrate into college life if they can connect with better-off peers (Zimmerman 2017) but having little in common with those peers (e.g., having a different favorite TV show, different hobbies, different food preferences, etc.) may result in lower chances of forming new friendships across income lines. The lack of a shared culture may thus reduce the accumulation of both social and human capital. Bourdieu’s logic also extends to groups not defined by income. African Americans and women may struggle to succeed in a predominantly white and male corporate America because of the greater difficulty of connecting with their majority-culture peers.⁸

Why might cultural divides be greater today than in the past? Technological progress could lead to cultural divergence: when there is only one channel to watch on television and only one brand of ketchup to buy, all groups are mechanically constrained to share the same culture on these dimensions. Thus, increased choice sets might have fueled cultural divergence. This, however, is not a foregone conclusion. First, it is possible that with only two TV shows, each show caters to one group or the other, but with thousands of shows, idiosyncratic preferences unrelated to group membership become the predominant driver of cultural choices. Second, universally-adopted new technologies might wipe out cultural differences; perhaps the rich and the poor used to spend their time differently from each other, but in the future everyone will just monitor their Facebook feed all day.

In this paper, we measure the extent of cultural distance across various groups in the US over time. In particular, we define groups of Americans based on their income, education, gender, race, and political ideology.⁹ We assemble multiple datasets that allow us to capture as many aspects of people’s cultural lives as possible, for as long as possible. This includes detailed information

⁸<http://fortune.com/2016/08/11/african-american-executives-diversity-racism/>

⁹In our Online Appendix, we also examine cultural distances by urbanicity (cf: Figure A.5) and age (cf: Figure A.6). Due to data constraints, we analyze cultural distances by urbanicity only in time use and social attitudes.

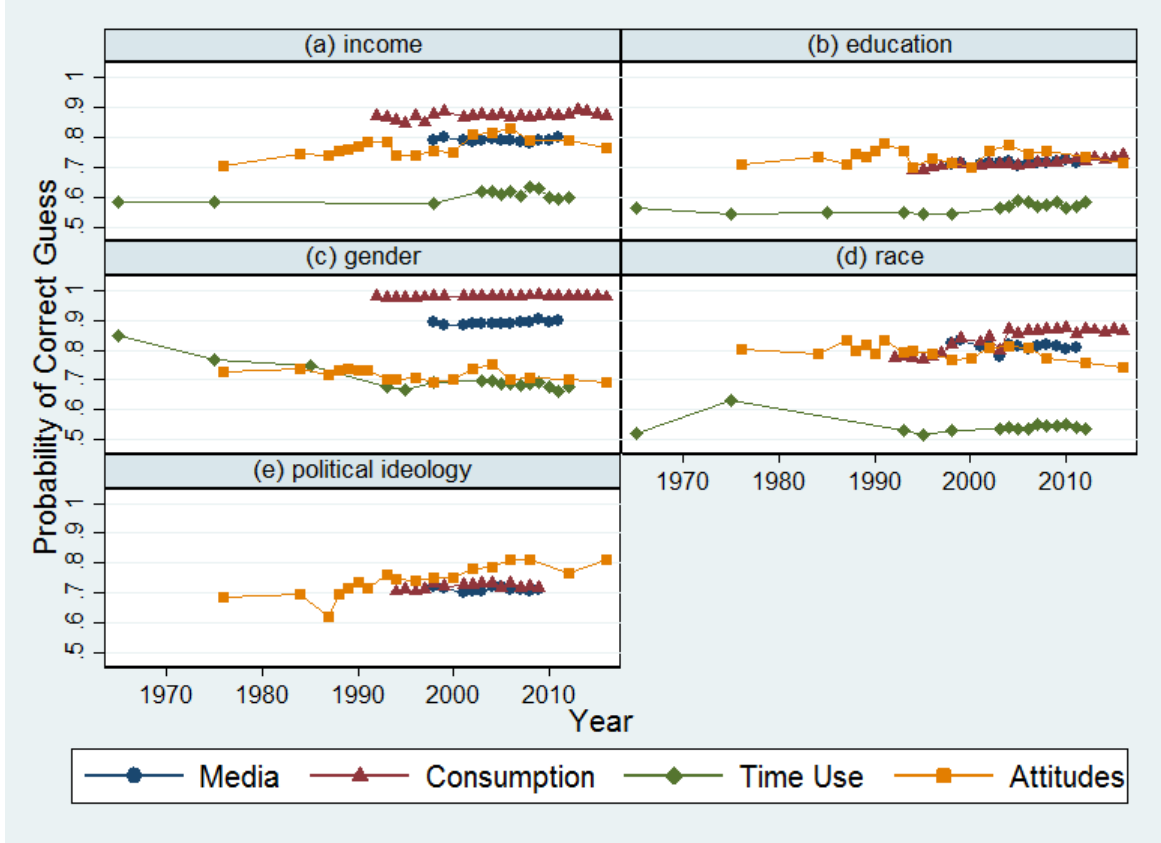


Figure 1: Cultural distances over time

Note: Figure shows the likelihood, in each year, of correctly guessing an individual's group membership based on his/her media diet, consumer behavior, time use, or social attitudes.

on media consumption and consumer behavior (from 1992 onward), attitudes (from 1976 onward), and time use (from 1965 onward).¹⁰ We define cultural distance in media consumption between the rich and the poor in a given year by our ability to predict whether an individual is rich or poor based on her media consumption that year.¹¹ We use an analogous definition for the other three dimensions of culture (consumer behavior, attitudes, and time use) and other group memberships. We use a machine learning approach to determine how predictable group membership is from a set of variables in a given year. In particular, we use an ensemble method that combines predictions from three distinct approaches, namely elastic net, regression tree, and random forest (Mullainathan and Spiess 2017).

Figure 1 summarizes our findings. The results overall refute the hypothesis of growing cultural

¹⁰As we discuss at greater length in the next section, we use Mediamark Research Intelligence data for media consumption and consumer behavior, General Social Survey for attitudes, and American Heritage Time Use Study for time use.

¹¹This is the approach taken by Gentzkow et al. (2017) to measure differences between Democrats and Republicans in their Congressional speech.

divides. With few exceptions, the extent of cultural distance has been broadly constant over time. One (unsurprising) exception is that men and women’s time use became more similar from 1965 to 1995; perhaps more surprisingly, there has been no subsequent change in the gender differences in time use over the last 20 years.¹² We also find that differences in social attitudes by political ideology and income have increased since the 1970s. Finally, whites and non-whites have converged somewhat on social attitudes but have diverged in consumer behavior. Nevertheless, our headline result is that for all other demographic divisions and cultural dimensions, cultural distance has been broadly constant over time.¹³

Two papers closest to ours are Alesina et al. (2017) and contemporaneous work by Desmet and Wacziarg (2018). Alesina et al. (2017) employ the European Value Survey and the General Social Survey (GSS) and find that, from the early 1980s to 2010, cultural differences across countries in the EU and across nine large American states have somewhat increased. Desmet and Wacziarg (2018) define cultural distance between two groups as the share of total heterogeneity in responses to questions in the GSS that is not attributable to within-group heterogeneity. They examine eleven demographic divisions, including our five, and also report that cultural distances have been “remarkably” stable over time. Their results do somewhat contrast with ours as we find steady cultural divergences in social attitudes based on political ideology and income.¹⁴ In contrast to both Alesina et al. (2017) and Desmet and Wacziarg (2018), we examine other dimensions of culture besides social attitudes, namely media consumption, consumer behavior, and time use.

The rest of the paper is organized as follows. We briefly describe our datasets in Section 2. Section 3 lays out our empirical strategy and provides a discussion of our definition of cultural distance. The main results are reported in Section 4. Section 5 concludes.

¹²Both of the aforementioned patterns also hold if we only examine how men and women spend their time when they are not at work.

¹³The results are broadly the same if instead of our machine learning approach, we measure cultural distance simply as the Euclidean distance between average responses across groups (cf: Figure A.7).

¹⁴The difference in our findings is presumably due to our different definitions of cultural distance. The most important way in which our approaches differ is that Desmet and Wacziarg (2018) ignore correlation in answers across different questions in the GSS, but there are other differences as well. Suppose there are four equally sized groups, A, B, A’, and B’, who are asked a single binary question. Suppose the share of the individuals giving a specific response to this question is 0, 5%, 10%, and 20% in the four groups, respectively. Our notion of cultural distance would indicate that groups A and B are closer together than groups A’ and B’, whereas Desmet and Wacziarg (2018) would say that A and B exhibit a greater cultural distance than A’ and B’. More generally, their definition allows for substantial cultural distance to be driven by arbitrarily rare behaviors, whereas our inference approach does not.

2 Short data description

In this section, we provide a short description of the data that we utilize. The Data Appendix provides a more detailed description of our variables and sample construction. All of our datasets study individuals in the United States.

Mediamark Research Intelligence (MRI) data contains two questionnaires conducted each year between 1992 and 2016.¹⁵ Demographic information (including household income) and some data pertaining to media exposure is obtained in a personal, face-to-face interview. The second questionnaire is left to be completed by the “principal shopper” of the household. This questionnaire asks whether the household has purchased, used, or owns a number of brands and products and services. It also solicits data on which magazines the respondent reads and which TV shows and recently released movies the respondent has seen. For our media consumption analysis, we use 871 to 1,186 binary (yes/no) answers about consumption of magazines, TV shows, and movies.¹⁶ An example of a variable about media consumption is “Have you seen the movie *Birdman* in the last 6 months?” For our consumer behavior analysis, we use 7,130 to 9,385 variables on brands and products or services used, purchased or owned. Variables include questions such as “Has your household used Grey Poupon Dijon mustard in the last six months?”, “Have you personally used a lipstick in the last six months?”, “Do you own a dishwasher?”, and “Have you personally used dry cleaning services in the last six months?”¹⁷ For ease of exposition, from here on we will use the word ‘product’ to refer to ‘products and services.’ We restrict attention to respondents who are between 20 and 64 years old.¹⁸ The sample size of the the MRI annual sample ranges from 15,352 to 22,033.¹⁹

¹⁵Not all of the variables are available in every year. For example, we use data on magazines from 1992 to 2011, data on TV shows from 1992 to 2016, and data on movies from 1998 to 2016. That means that when we report trends in cultural distance based on overall media consumption, we restrict our attention to years when all three of these subcomponents are available, namely 1998 to 2011.

¹⁶The MRI also asks questions about listening to radio and reading newspapers, but its coverage is too sparse to be useful for our purposes.

¹⁷As we discuss in greater length in next section, differences in how the rich and poor answer some of these questions (such as, “Do you own a dishwasher?”) surely reflect the way that income affects budget sets rather than some notion of “cultural distance”. We acknowledge the important distinction between income-constrained variables (such as the dishwasher one) and income-unconstrained ones (such as did you watch this movie or that one).

¹⁸The MRI only captures age by 5 year-brackets.

¹⁹An alternative dataset to MRI to predict group membership based on consumer behavior is the Kilts-Nielsen Consumer Panel data (see <https://research.chicagobooth.edu/nielsen>). The Nielsen data tracks households’ shopping behavior by asking participating households to scan the barcode of each purchased good after a shopping trip (using a scanning device provided by Nielsen). Nielsen differs from MRI in that it only covers products bought, not those used or owned; MRI also includes a broader set of products and services without barcode. The main disadvantage of Nielsen over MRI for our purpose is a shorter time series: the Nielsen data only starts in 2004. Also, the Nielsen data has income brackets that are too broad to be able to identify top and bottom quartile of the income distribution. Furthermore, the Nielsen data does not include information on political ideology. Figures A.8, A.9, A.10 compare

To measure time use, we use the American Heritage Time Use Study (AHTUS). The AHTUS is a harmonized collection of diary data on time use in the US from 1965 to 2012. For the early years, the AHTUS covers roughly one year per decade, but since 2003 it includes annual surveys conducted by BLS.²⁰ Our harmonized data consists of 78 variables that indicate time spent on a specific activity (e.g., gardening), and we also include 8 aggregate activities (e.g., non-market work) from Aguiar and Hurst (2009). We restrict our sample to individuals who are between 18 and 64 years old and are employed full time. The sample size ranges across years from 669 to 10,210.

We use the General Social Survey (GSS) as our source of data on social attitudes. The GSS has been conducted annually since 1972. It collects stated attitudes on topics such as civil liberties, government policies, and morality. An example question is, “Are we spending too much, too little, or about the right amount on foreign aid?” The GSS also asks some questions about behavior, such as whether the respondent voted in the most recent Presidential election, which we also include in our analysis. We exclude questions about the respondent’s assessment of his or her own financial situation.²¹ There are many questions that the GSS asks only intermittently, so we drop some years in order to have a larger number of questions which are asked in every year we consider. This leaves us with 84 questions and 18 interspersed years between 1976 and 2016. The GSS often presents a given question to only two thirds of the survey participants. We impute the missing values for the remaining third based on the marginal distribution of responses in a given group.²² We restrict attention to respondents who are between 18 and 64 years old. The sample size of the GSS annual samples ranges from 1,093 to 3,735.

predictability of respondents’ education, gender and race based on consumer behavior in the MRI and Nielsen data. Note that we restrict the analysis in these figures to single individuals, given Nielsen’s focus on the household and MRI’s focus on the “principal shopper.” The results for education and race are comparable across both datasets over the overlapping years. As we discuss in Section 4.3.2, our empirical approach to measuring trends in cultural distance is not suitable for comparing consumer behavior of men and women in the MRI since we can perfectly predict group membership in this case. We do not reach this upper bound in the Nielsen data, where we observe convergence between men and women in consumer behavior. See also footnote 50.

²⁰The AHTUS lacks information on household income in 1985, 1993, and 1995, so we drop those years from our analysis of cultural distance by income.

²¹When we examine cultural distance in social attitudes by ideology, we also exclude questions that directly pertain to ideology, namely political party affiliation and how the respondent voted in a presidential election.

²²So when we measure cultural distance by income, we impute the missing data based on the distribution of responses by the rich and the poor, whereas when we measure cultural distance by education, we impute the missing data based on the distribution of responses by the more and less educated, etc.

3 Empirical approach

3.1 Compositional changes

Our definition of each group allows for the demographic composition of the groups to change over time. For example, in the early 1970s, less than 10% of either the rich or the poor were Hispanic, but these days Hispanic individuals constitute 10% of the rich and 30% of the poor. Consequently, if Hispanic individuals are culturally distinct, this compositional change could lead to an increase in cultural distance between the rich and the poor. The same issue applies to other groups and other demographic characteristics. The share of people who are in our “more educated” group grows steadily over time and includes an ever-rising share of women. Figure A.11 in the Online Appendix reports compositional changes in each of our groups. In general, trends in cultural distance that are due to such compositional changes are something that we wish to capture rather than control for.²³

3.2 Predictability as a measure of cultural distance

In any given year, we say that two groups are further apart in their media consumption (or consumer behavior or time use or social attitudes) if we can predict more accurately which of the two groups a given individual belongs to based on his or her media consumption (or consumer behavior or time use or social attitudes). This approach follows Gentzkow et al. (2017), who measure partisanship of congressional speech by the ease with which one can infer a congressperson’s party from his or her speech.

Given some outcome space X , one could define the distance between two disjoint groups A and B based on any metric d on $\Delta(X)$ by letting the distance between the groups be equal to $d(\mu_A, \mu_B)$, where μ_A is the distribution of X in group A and μ_B is the distribution of X in group B .²⁴ Our predictability-based measure of distance implicitly sets d to be the total variation metric.²⁵

²³One exception may be the change in the age distribution of the rich and the poor. To the extent that trends in cultural distance are driven or hidden by the changes in the relative age between the rich and the poor, we may want to take those changes out, especially if we think that lifetime rather than contemporaneous income is a more meaningful way of defining who is rich and who is poor. In Figure A.12 in the Online Appendix, we define an individual as rich (poor) if he or she is in a household that is in the top (bottom) quartile, in terms of household income, among individuals in the same 5-year age bracket. We do not use information on household type (cf: discussion of household types in Section 4.1) since we do not have sufficient sample sizes to construct our groups based on both the age bracket and the household type. As seen in the figure, the results are mainly unaffected.

²⁴In our setting, X would be the set of all possible vectors of answers to questions about media consumption (or consumer behavior or time use or social attitudes).

²⁵If A and B are equally sized (as they are by construction in our approach), the ability to predict whether a person belongs to A or B is equal to $\frac{1}{2} + \frac{1}{2}d_{TV}(\mu_A, \mu_B)$, where d_{TV} denotes the total variation metric (cf: proof of

This measure has several features that are worth noting. First, the measure takes no stance on which elements of X are close to another.²⁶ Suppose X consists of four elements: vodka, Sprite, 7 Up, and water. Suppose there are three equally-sized groups: in group A , 80% of people drink vodka and 20% drink water; in group B , 80% drink Sprite and 20% drink water; in group C , 80% drink 7 Up and 20% drink water. Our approach would say that the cultural distance between A and B is the same as the cultural distance between B and C , despite the fact that one might argue that B and C are closer since Sprite and 7 Up are more similar to each other than either is to vodka. Second, our measure of cultural distance has an upper bound that is achieved when one can perfectly predict group membership. Consequently, if some subset of variables is always sufficient to reach the upper bound, we would not be able to detect any changes in how similar the groups are on variables outside of that set. This turns out not to be an issue, however, since we are always far from the upper bound, except in the case of predicting gender using consumer behavior.²⁷ Third, a nice feature of our measure is that its units are easily interpretable. Contrast this with a measure that uses normalized Euclidean distance between μ_A and μ_B as the notion of cultural distance. Formally, letting μ_G^x denote the share of individuals in group G with outcome x , we could measure the distance between A and B by $\frac{\sqrt{\sum_{x \in X} (\mu_A^x - \mu_B^x)^2}}{|X|}$. In Figure A.7 in the Online Appendix, we replicate our results using this measure and find qualitatively similar patterns, but with units of cultural distances that are harder to interpret.²⁸

Finally, note that our predictability-based approach does not allow us to aggregate across cultural dimensions that are not measured in the same dataset. For example, since we do not know the joint distribution of attitudes, time use, and income, we do not know how well one could predict income with both attitudes and time use. We do have data on media use and consumer behavior in the same dataset, so in Figure A.13 in the Online Appendix we report cultural distance over time for these two aggregated dimensions. Again, we find no trends over time.²⁹

Claim 3.30 in Mossel et al. 2014).

²⁶Formally, this is related to the fact that total variation (unlike say the Prokhorov metric) does not require X to be a metric space itself.

²⁷Thus, despite panel (c) in Figure 1, it is possible that men and women have become more or less similar over time in some aspects of their consumption patterns.

²⁸The memetic fractionalization approach by Desmet and Wacziarg (2018) also has interpretable units. In footnote 14 we discuss some differences between our approaches.

²⁹If distance is measured based on the normalized Euclidean distance between μ_A and μ_B , it is possible to aggregate across datasets. The overall trend in cultural distance is then just the (weighted) average of the trends depicted in Figure A.7.

3.3 Machine learning

We use a machine-learning ensemble method to determine how predictable group membership is from the variables in each dataset (i.e., time use, social attitudes, media consumption, and consumer behavior) in each year. The ensemble method consists of running separate prediction algorithms (we employ elastic net, regression tree, and random forest) and then combining the predictions of these algorithms with weights chosen by OLS (Mullainathan and Spiess 2017). For each dataset, year, and group division (e.g., time use data by gender in 2010), we first split the dataset into a training sample (70% of the data) and a hold-out sample (30% of the data). We empirically tune each algorithm on the training sample by cross-validation. In particular, we partition the training data into five folds. For a given fold, we fit the algorithm on the other folds for every value of the tuning parameter. Through this process, we obtain a prediction (e.g., probability that the respondent is a woman) for every observation in the training sample for every value of the tuning parameter. We then average the squared-error loss function for each tuning parameter over the full training sample and choose the tuning parameter that minimizes the loss. This gives us a prediction for every observation in the training sample for each of the three algorithms. We regress (using simple OLS) group membership on the three predictions (from the three algorithms) in the full training sample. We use the coefficients from this regression to combine the three algorithms into the ensemble prediction in the next step.

We then turn to our hold-out sample. For each observation in the hold-out sample, we derive the prediction of each algorithm using the model estimated in the training sample under the optimal tuning parameter. We then compute the ensemble prediction for that observation using the aforementioned OLS coefficients. We then guess a respondent’s group affiliation based on the ensemble prediction: if the probability that a respondent is in a group is above $\frac{1}{2}$, we guess that she is in that group; otherwise, we guess that she is in the other group. We define cultural distance (for each dataset, year, and demographic category) as the predictability of the group membership, i.e., the share of the guesses in the hold-out sample that are correct.³⁰

3.4 Data over time

We need to ensure that the “quality” of our datasets – in terms of number of observations and the availability of relevant variables – is constant over time. Otherwise, our ability to predict group

³⁰All of our results about trends over time are the same if we use any one of the algorithms (elastic net, regression tree, or random forest) rather than combining them into the ensemble prediction.

membership might change over time for reasons unrelated to any changes in cultural distance. The solution to time-varying sample sizes is straightforward. For each dataset and demographic group, we equalize the number of observations in each year and demographic group as follows. Denoting by n the minimum sample size across years and groups (e.g., when computing the cultural distance in time use by education, the smallest year-group are the less educated in 1965), we randomly select n observations for every year-group. This yields a “balanced” dataset with the same number of observations in each year and with half of the observations in each of the two groups. We then compute the predictability of group membership – as described in the previous subsection – in this balanced dataset. We repeat this procedure a number of times,³¹ drawing a new random sample each time, and then we take the average predictability of group membership (averaged across the draws) as our measure of cultural distance. Note that this means that the sample sizes reported in Section 2 are larger than the balanced sample sizes that we use to make each prediction of group membership.

Another important consideration is related to the changes over time in the particular questions asked to survey participants. When it comes to the GSS and AHTUS data, we insist on having the same set of variables in each year. When it comes to the time use data, we think the set of activities that people can spend their time on has not changed that much over time, with the exception of spending time on a computer. Therefore, if the set of variables in the time use data expanded over time, this would likely be a reflection of improvement in data collection rather than a reflection of actual changes in the ways people are spending their time. Therefore, we use the crosswalk provided by the University of Oxford Center for Time Use research³² to harmonize time use variables across years.³³ With regard to social attitudes, the GSS often asks a particular question only intermittently, and we do not believe that this is a reflection of the fact that this question was only relevant in the years the question was asked. Consequently, we limit the set of GSS variables and years we use in a way that ensures that each variable is available in each year.³⁴

³¹In the GSS and the AHTUS, we take 500 draws. In the MRI, which has much larger sample sizes, we take only 25 draws for media consumption and only 5 draws for consumer behavior.

³²See <https://www.timeuse.org/ahtus/documentation>.

³³The AHTUS asks about computer use only after 1985. We impute zero computer use for all respondents prior to 1985. The AHTUS does not ask about smartphone usage. All activities related to the use of computer and internet for leisure are aggregated under computer use.

³⁴In contrast to the time use data, we are less confident in our decision to harmonize the set of GSS questions over time. It might very well be that the GSS changes questions it asks from one year to the next because the set of most important societal issues is changing, in which case there might be some argument for embracing the change in variables. Without a more specific model of how the GSS drafts their survey instrument each year, it is difficult to sign the potential bias induced by our harmonization choice. For example, if the GSS drops questions once everyone agrees on the answer and keeps only those questions where disagreement remains, our approach might

When it comes to the MRI data, we embrace the variations in the set of questions asked over time, both for media consumption and consumer behavior. Our understanding is that the MRI seeks to include questions about all media items (magazines, TV shows, movies) and consumer products that are relevant at the time. For example, each year the MRI asks respondents about whether they had seen a number of newly released movies. While the number of movies that the MRI asks about is reasonably constant, ranging from 83 to 97 across years, the set of movies they ask about of course changes completely from year to year, reflecting the new releases. We assume that the changes in the variables about TV shows, magazines, products, and brands similarly reflect real changes in consumers’ choice sets. While this assumption surely does not hold perfectly – for example, there is a big jump in the number of TV shows in the data in 2009 when the MRI added cable shows to the survey – it provides the most natural approach for measuring cultural distance when cultural elements are rapidly changing over time.

3.5 Confidence intervals

Throughout, we report our estimates of cultural distances without confidence intervals. One way to approach inference in our setting would be via subsampling (e.g., Politis et al. 1999), but our sample sizes are too small for the ensemble algorithm to perform well on partitioned data. That said, the fact that our measure of cultural distance tends to be pretty similar across years informally suggests that it is estimated reasonably precisely; otherwise, it would be highly unlikely for the estimates to fall so close to one another. We have also confirmed that if we add to the data a synthetically constructed variable whose correlation with group membership increases over time, we indeed observe a growing cultural distance using our method.

4 Results

We organize the results by group divisions: income, education, gender, race, and political ideology. For each group division, after a discussion of the overall patterns, we dive in greater detail into the four broad cultural components. For the media, we investigate the separate cultural influences due to TV watching, movie watching, and magazine readership. For consumption, we investigate the separate roles of products vs. brands. For social attitudes, we consider the separate influence of thematic sub-categories, such as views related to the role of government in society or views related

underestimate cultural convergence over time. Alternatively, if the GSS systematically adds questions that have become more controversial, our approach might underestimate the increase in cultural distance over time.

to civil liberties. Throughout, we try to enrich the results with a discussion of specific cultural traits that are most distinctive across groups at a given point in time. Rather than report every possible result for every group, we highlight the data we find most informative in the text and report the additional results in the Online Appendix.

4.1 Income

A vast literature in labor economics has documented the rise in income inequality in the US since the late 1970s. While a large share of this literature in recent years has focused on “top income inequality” (e.g., the share of total income going to the top 1 percent, or top 0.1 percent), it is also well understood that technological change and global competitive pressures have contributed to broader changes in income inequality across individuals and households (e.g., Autor et al. 2008, Meyer and Sullivan 2017). The causes of the rise in income inequality are now reasonably well understood, but the consequences are less clear. We are particularly interested in whether greater income inequality has led to a greater cultural gap between the rich and the poor. Technological change and a growing supply of goods and services also may have exacerbated or attenuated any changes in the cultural gap between the rich and the poor.³⁵ As discussed previously, increased cultural distance between rich and poor could be particularly damaging to social mobility. A high-income manager may promote the subordinate with whom she has the friendliest interactions around the water cooler, and that favorite subordinate will likely come from a high-income background if tastes, views, and experiences are sharply different between income classes.

We define an individual as rich (poor) if he or she is in a household that is in the the top (bottom) quartile of household income among households of the same type. We put households into four types: (i) a single adult with no dependents, (ii) two adults with no dependents, (iii) a single adult with dependent(s), and (iv) two adults with dependent(s).^{36,37} We use the Current

³⁵Jaravel (2017) documents that newly developed products in the US tend to target high-income households; this force could create a new set of goods around which a “culture of being rich” could coalesce. At the same time, other technological developments, such as certain forms of social media, could lead to cultural convergence between income groups by providing inexpensive goods that appeal to individuals of all income levels.

³⁶We define a household as having dependents if there are children under 18 or if the household has more than two adults. This may induce some measurement error, as we would code three roommates as two adults with a dependent and a single adult taking care of a parent or a sibling as two adults with no dependents.

³⁷An alternative to this approach would be to use an equivalence scale to adjust for the size and the composition of the household. The downside of the alternative approach is that all standard scales (per-capita income, the Oxford scale, the OECD-modified scale, and the square root scale) systematically label households with (more) children as more likely to be poor. Consequently, the ability to predict household income then primarily stems from the ability to predict whether there is a child in the household: tell-tale signs of “being poor” are watching *SpongeBob SquarePants* or buying children’s medications. Under our preferred approach, there is by construction no relationship between poverty and the presence of children, and the relationship between poverty and the number of children is weaker.

Population Survey to identify the distribution of household income for each household type in each year. We focus on the top and the bottom quartile (as opposed to, say, the top and the bottom half or the top and the bottom decile) to balance a desire to make the rich and the poor as different in their income as possible and the pragmatic need to keep our sample sizes sufficiently large.³⁸ Given our definition of rich and poor and our procedure for equalizing sample sizes described in Section 3.4, each prediction of income in a given year is based on 6,394 observations in the MRI, 398 observations in the GSS, and 418 observations in the AHTUS.

We also consider alternative definitions of rich and poor, comparing (i) top half vs. bottom half, (ii) top quartile vs. everyone else, and (iii) bottom quartile vs. everyone else. Under all of these alternative definitions, our qualitative results remain the same (cf: Figure A.15). Throughout the analysis, we use contemporaneous income rather than wealth or lifetime income. While the latter two measures might seem more closely related to what it means to be rich or poor, we do not have data on wealth or lifetime income.

Panel (a) of Figure 1 summarizes our results. There is no evidence of an increasing cultural gap between the rich and the poor based on media consumption, consumer behavior, or time use. The patterns regarding media and consumer behavior, where our sample size is the largest, show that cultural distance is essentially the same in each year. Knowing what TV shows and movies someone watches and what magazines a person reads allows us to correctly predict the person's income group about 80 percent of the time. Knowing what goods and services a person buys, including particular brands, allows us to correctly predict the person's income group between 85 to 89 percent of the time, with no apparent time trend. The gap in how rich and poor spend their time has also been constant; the ability to guess income from time use has been around 60 percent since 1965. We do observe some divergence of attitudes between income groups, mostly between the mid 1970s and the late 1980s; while there have been some year-to-year fluctuations since then, there is no discernible trend over the last quarter-century.

Moreover, if we ignore household types and define rich (poor) as the top (bottom) quartile of household income divided by the square root of the household size, we observe the same temporal trends in cultural distances between the rich and the poor (cf: Figure A.14).

³⁸As income variables available in the GSS, the AHTUS, and the MRI are income brackets, the top and bottom income quartiles obtained from the CPS most often occur within an income bracket rather than at the boundary. Consequently, using income brackets to define top and bottom income quartiles results in some miscategorization. We classify respondents into the top and bottom quartiles to minimize miscategorization (please refer to our Data Appendix for details). The share of miscategorization never exceeds 5 percent, and the extent of miscategorization does not explain almost any of the variance in measured cultural distance. Specifically, if we regress measured cultural distance on a linear time trend and the dummy for the cultural dimension (media consumption, consumer behavior, attitudes, and time use), adding the extent of mismeasurement increases R^2 from 0.389 to 0.390.

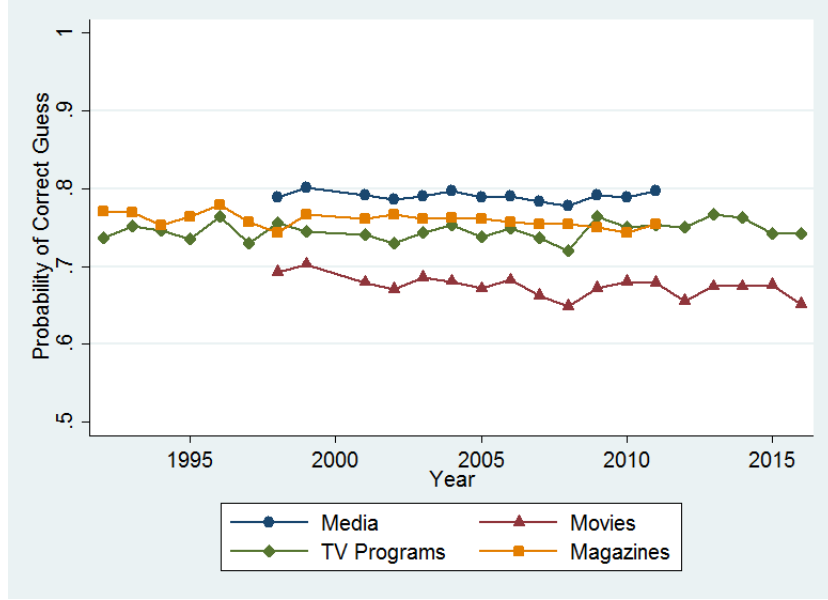


Figure 2: Cultural distance by income over time: media consumption

Note: Data source is the MRI. Sample size each year is 6,394. See text and data appendix for details on sample construction and implementation of machine-learning ensemble method. Presented in the figure is share of correct guesses of respondent's income in the hold-out sample each year. The procedure to guess income in the hold-out sample was repeated 25 times, and the share of guesses reported is the average of these 25 iterations.

4.1.1 Media Consumption

Figure 2 shows how well one can predict income group over time based on the three separate components of media consumption: TV programs, movies, and magazine readership. For reference, we also again report predictability of income based on media consumption overall.

Figure 2 reveals that there has been no divergence between the rich and the poor in any of the sub-components of media consumption. These groups watch different TV shows and different movies and read different magazines, but the extent of that difference has been nearly constant across the last quarter-century. We also see that income differences in consumption of magazines and TV shows is somewhat greater than the difference in consumption of movies.³⁹

It is particularly interesting that predictability of income based on TV shows has been constant over time given that there have been substantial changes both in the number of TV shows available and the number of TV shows watched in each income group.⁴⁰ As we discussed in the introduction,

³⁹This comparison is more meaningful than comparing, say, the cultural distance in TV watching with cultural distance in time use, since those are measured in datasets with very different sample sizes. Of course, even with equal sample sizes, there are important differences in measurement across cultural dimensions. As shown in Figure A.17 in the Online Appendix, the average number of TV shows watched is between 20 and 30 (and similar across income groups), while the average number of movies watched is less than 10 (also similar across income groups).

⁴⁰Figure A.16 plots the number of TV shows in the data over time. Figure A.17 shows the average number of

there is no mechanical reason that necessitates a relationship between the number of options and the size of the cultural gap, but it is nonetheless striking that the cultural gap has been constant even as the number of options has changed substantially.

Of course, the fact that cultural distance in media consumption has been constant over time does not imply that the particular magazines, TV shows, and movies that drive this distance have been the same from year to year. This is most obvious in the case of movies where each year brings a crop of new releases. Panels (a), (b), and (c) of Table 1 report, respectively, the ten movies, TV shows, and magazines that are individually most informative of income. We do this for three separate years spanning the beginning, the middle, and the end of each dataset. Consider movies for example (Panel (b)). If we could ask a person a single question of the form “Did you see movie X” and then guess, based on the answer, whether the person is rich or poor, the best question to ask in 1998 would be “Did you see Jerry Maguire”? Since 35% of rich people and 18% of poor people saw that movie, guessing the person is rich if and only if they say they saw the movie would lead us to guess correctly 57% of the time.⁴¹ All of the ten most informative movies in 1998 are movies that are distinctly rich-people movies. By contrast, the three most informative movies in 2007 are movies whose audiences are distinctly poor. As shown in Panel (a), even though cultural distance in TV consumption has been constant, the specific TV content that is most predictive of income has changed over time. In 1992, watching car racing, bicycle racing, and figure skating – and not watching *Roseanne* – are most indicative of being rich; in 2004, watching football and tennis – and not watching *Cops* – are most indicative of being rich. Panel (c) reveals more stability in the list of magazines whose readership is most indicative of income group. While the rank ordering varies somewhat, the top three magazines are the same: *Newsweek*, *Consumer Reports*, and *Time*.

TV shows watched. The jump in 2009 reflects the addition of cable shows to the data. The secular decrease in the average number of shows watched probably means that consumers increasingly watch a few shows devotedly rather than watching many shows occasionally; the time use data shows an increase rather than a decrease in the total amount of time spent watching television from 1998 to 2012.

⁴¹It does not matter whether we compute the likelihood of a correct guess using Bayes’ rule or estimate the frequency of each binary answer in a training sample and then report the share of correct guesses in a hold-out sample since the sampling variation in the share of people in each group who answer a question a particular way is negligible. For the Jerry Maguire question, for example, the probability that the person is rich conditional on seeing the movie is $\frac{0.35}{0.35+0.18} = 0.66$. The probability that the person is poor conditional on not seeing the movie is $\frac{0.72}{0.72+0.65} = 0.53$. Guessing that the person is rich in the case they have seen the movie and guessing that they are poor otherwise leads us to guess correctly 57% of the time.

Table 1: TV shows, movies, and magazines most indicative of being high-income

Panel (a) TV shows					
1992		2004		2016	
Watched <i>Autoworks 200</i>	57.3%	Watched <i>Super Bowl</i>	58.5%	Watched <i>Super Bowl</i>	57.1%
Watched <i>Busch Clash</i>	57.1%	Watched <i>NFL Monday Night Football</i>	56.1%	Watched <i>Love It Or List It</i>	55.9%
Watched <i>Tour du Pont</i>	56.7%	Watched <i>NFL Regular Season Football</i>	55.9%	Watched <i>Property Brothers</i>	55.7%
Watched <i>US Figure Skating Championship</i>	56.6%	Watched <i>NFL Regular Season Games</i>	55.8%	Watched <i>House Hunters</i>	55.5%
Watched <i>Michigan 500</i>	56.2%	Watched <i>US Open</i>	54.9%	Watched <i>Academy Awards</i>	55.3%
Didn't watch <i>Roseanne</i>	55.8%	Watched <i>College Football Regular Season</i>	54.9%	Watched <i>NCAA Men's Final Four</i>	55.9%
Didn't watch <i>Sunday Night Movie</i>	55.8%	Didn't watch <i>Cops</i>	54.8%	Watched <i>Flip or Flop</i>	54.9%
Watched <i>Miller Genuine Draft 200</i>	55.7%	Watched <i>Academy Awards</i>	54.7%	Watched <i>The Masters</i>	54.8%
Watched <i>Indianapolis 500</i>	55.5%	Watched <i>Wimbledon</i>	54.7%	Watched <i>SNL Specials</i>	54.3%
Watched <i>Fedex St. Jude Classic</i>	55.4%	Watched <i>NCAA Men's Basketball</i>	54.9%	Watched <i>Grammy Awards</i>	53.9%
Panel (b) Movies					
1998		2007		2016	
Watched <i>Jerry Maguire</i>	57.3%	Didn't watch <i>Big Momma's House</i>	54.0%	Watched <i>Gone Girl</i>	54.2%
Watched <i>First Wive's Club</i>	55.1%	Didn't watch <i>Final Destination 3</i>	53.5%	Watched <i>The Hunger Games</i>	52.7%
Watched <i>The English Patient</i>	54.7%	Didn't watch <i>Saw II</i>	53.5%	Didn't watch <i>Teenage Mutant Ninja Turtles</i>	52.7%
Watched <i>Air Force One</i>	53.8%	Watched <i>The Devil Wears Prada</i>	53.1%	Watched <i>Interstellar</i>	52.3%
Watched <i>Michael</i>	53.5%	Watched <i>Walk the Line</i>	53.0%	Didn't watch <i>Annabelle</i>	52.3%
Watched <i>My Best Friend's Wedding</i>	52.6%	Watched <i>Pirates Of The Caribbean 2</i>	52.6%	Didn't watch <i>No Good Deed</i>	52.1%
Watched <i>The Chamber</i>	52.4%	Watched <i>The Da Vinci Code</i>	52.4%	Didn't watch <i>Oujia</i>	52.0%
Watched <i>Evita</i>	52.4%	Watched <i>Syrianna</i>	52.3%	Didn't watch <i>Let's Be Cops</i>	51.9%
Watched <i>Ransom</i>	52.4%	Didn't watch <i>The Exorcism Of Emily Rose</i>	52.2%	Watched <i>The Theory Of Everything</i>	51.6%
Watched <i>One Fine Day</i>	52.2%	Watched <i>Brokeback Mountain</i>	52.2%	Watched <i>Kingsman</i>	51.4%
Panel (c) Magazines					
1992		2002		2011	
Read <i>Newsweek</i>	61.2%	Read <i>Newsweek</i>	60.2%	Read <i>Consumer Reports</i>	57.9%
Read <i>Consumer Reports</i>	60.0%	Read <i>Time</i>	59.6%	Read <i>Newsweek</i>	57.5%
Read <i>Time</i>	59.8%	Read <i>Consumer Reports</i>	58.5%	Read <i>Time</i>	56.7%
Read <i>Business Week</i>	59.1%	Read <i>Business Week</i>	57.4%	Read <i>People</i>	56.7%
Read <i>US News & World Report</i>	58.7%	Read <i>US News & World Report</i>	57.2%	Read <i>Sports Illustrated</i>	55.8%
Read <i>Parade</i>	58.4%	Read <i>Money</i>	57.0%	Read <i>Men's Health</i>	54.8%
Read <i>Money</i>	58.1%	Read <i>Forbes</i>	56.9%	Read <i>Travel & Leisure</i>	54.8%
Read <i>National Geographic</i>	57.4%	Read <i>Fortune</i>	56.8%	Read <i>Forbes</i>	54.7%
Read <i>Forbes</i>	57.1%	Read <i>Architectural Digest</i>	55.7%	Read <i>Economist</i>	54.6%
Read <i>Fortune</i>	57.0%	Read <i>People</i>	55.5%	Read <i>Real Simple</i>	54.6%

Note: Data source is the MRI. Sample size in all panels is 6,394. Reported in each column are the 10 cultural traits most indicative of being rich in that year. The numbers indicate the likelihood of guessing correctly whether an individual is rich or poor based on the answer to the question. For example, in 1992, knowing whether a person watched *Autoworks 200* allows us to guess income correctly 57.3% of the time, whereas knowing whether a person watched *Roseanne* allows us to guess income correctly 55.8% of the time. An affirmative answer to “Did you watch *Autoworks 200*?” and a negative answer to “Did you watch *Roseanne*?” indicate that the person is rich.

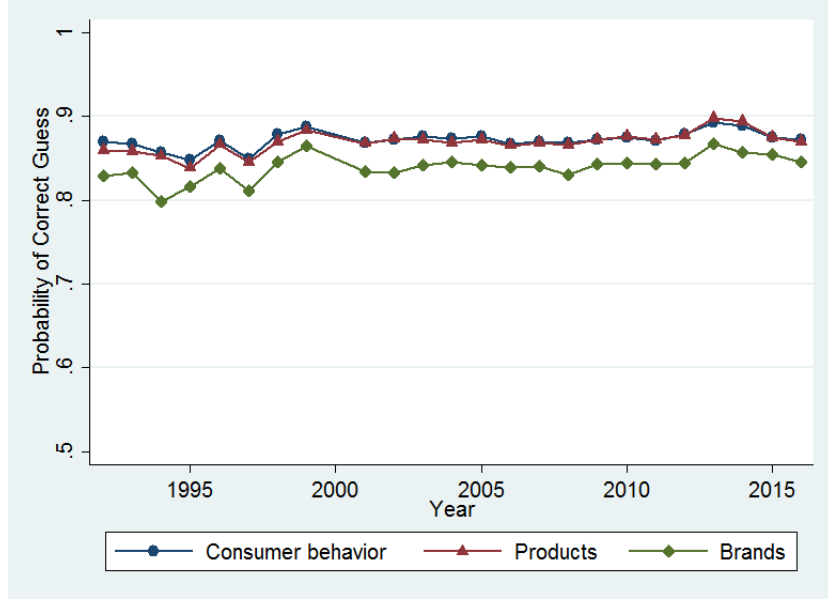


Figure 3: Cultural distance by income over time: consumer behavior

Note: Data source is the MRI. Sample size each year is 6,394. See text and data appendix for details on sample construction and implementation of machine-learning ensemble method. Presented in the figure is share of correct guesses of respondent's income in the hold-out sample each year. The procedure to guess income in the hold-out sample was repeated 5 times, and the share of guesses reported is the average of these 5 iterations.

4.1.2 Consumer Behavior

Figure 3 reports the predictability of income over time separately based on products and brands individuals report buying or owning. We also again report the predictability of income based on the entire consumer behavior data as a benchmark.

We see that the flat trend line previously reported for consumer behavior extends to these two subsets of the consumer data: the probability of correctly guessing someone's income based on the products or brands consumed is essentially the same over the quarter-century of available data. Products and brands consumed have very comparable predictive power. Moreover, the aggregation of the product and brand features does not change much the predictive power of the model compared to using products or brands separately.

Panels (a) and (b) of Table 2 show the variables over time in the product and brand data respectively that are individually most indicative of income group. As in the previous subsection, the fact that cultural distance has not changed over time does not imply that the same features distinguish income groups in each year. Consider panel (a) first. Household goods dominate the 1992 list: owning a dishwasher, a garage door opener, a fireplace, and a telephone answering machine separate rich and poor. Household goods are again present in the 2004 list (dishwasher,

Table 2: Products and brands most indicative of being high-income

Panel (a) Products					
1992		2004		2016	
Own an automatic dishwasher	71.4%	Bought a new vehicle	73.6%	Traveled in the continental US	70.9%
Used dishwasher detergent	70.2%	Used a dishwasher detergent	71.6%	Own a passport	70.3%
Traveled domestically	67.0%	Own a dishwasher	70.8%	Own Bluetooth on vehicle	70.2%
Own a garage door opener	65.8%	Traveled domestically	70.5%	Own heated/cooled seat on vehicle	69.9%
Own a fireplace	65.4%	Own a stereo on vehicle	69.7%	Used dishwasher detergent	69.3%
Own a telephone answering machine	65.3%	Belong to a frequent flier club	69.0%	Own a dishwasher	69.1%
Used dry cleaning services	65.2%	Own a personal computer	68.5%	Belong to a frequent flier club	68.6%
Used overnight delivery services	64.3%	Own an air bag on passenger side	68.5%	Traveled outside of continental US	67.7%
Own a garbage disposer	64.1%	Ordered any item by the Internet	68.4%	Ordered an item by Internet	67.4%
Traveled internationally	64.1%	Own a garage door opener	67.6%	Ordered a plane ticket by Internet	67.3%

Panel (b) Brands					
1992		2004		2016	
Used Grey Poupon Dijon (mustard)	62.2%	Used Land O' Lakes Regular (butter)	59.2%	Own an Iphone	69.1%
Bought Kodak (film)	61.6%	Used Kikkoman (soy sauce)	58.7%	Own an Ipad	66.9%
Used Thomas (English muffin)	61.5%	Did not use BIC (lighter)	58.7%	Used Verizon Wireless	61.0%
Used Cascade - Lemon (dish. detergent)	59.0%	Used Reynold Wrap (aluminum foil)	58.5%	Own an Android phone	59.5%
Used Scotch Magic (transparent tape)	58.7%	Used Bertolli (salad/cooking oil)	58.4%	Used Kikkoman (soy sauce)	59.0%
Used Cut-Rite (waxed paper)	57.7%	Used Scotch Magic (transparent tape)	58.4%	Own HP (printer/fax machine)	58.2%
Used Philadelphia (cream cheese)	57.7%	Own Toshiba (TV set)	58.3%	Used AT&T (cellular network)	58.1%
Used Kikkoman (soy sauce)	57.5%	Used AT&T (long distance call service)	57.5%	Own Samsung (TV set)	58.0%
Used Hellmann's (mayonnaise)	57.4%	Drank Diet Coke (diet cola)	57.5%	Used Cascade Complete	57.6%
Own Sylvania (TV set)	57.4%	Used Kleenex Regular (facial tissue)	57.4%	Used Ziploc (plastic bag)	57.5%

Note: Data source is the MRI. Sample size in all panels is 6,394. Reported in each column are the 10 cultural traits most indicative of being rich in that year. The numbers indicate the likelihood of guessing correctly whether an individual is rich or poor based on the answer to the question. For example, in 1992, knowing whether a person owns an automatic dishwasher allows us to guess income correctly 71.4% of the time, whereas in 2004 knowing whether a person bought a BIC lighter allows us to guess income correctly 58.7% of the time. An affirmative answer to “Do you own an automatic dishwasher?” and a negative answer to “Did you buy BIC lighter?” indicate that the person is rich.

personal computer) as are vehicle-related variables (new car, car stereo, airbags). While travel-related experiences and services are already present in the list of most indicative variables in 1992 and 2004 (domestic and international travel, frequent flyer programs), these travel-related features reach the top of the list in 2016. Half of the most indicative variables of top income in 2016 are related to travel. Beyond this, ownership of car gadgets (bluetooth, heated seats) and dishwashers remain important income class differentiators in 2016.

The brand most predictive of top income in 1992 is Grey Poupon Dijon mustard.⁴² By 2004, the brand most indicative of the rich is Land O’Lakes butter, followed by Kikkoman soy sauce. By the end of the sample, ownership of Apple products (iPhone and iPad) tops the list. Knowing whether someone owns an iPad in 2016 allows us to guess correctly whether the person is in the top or bottom income quartile 69 percent of the time. Across all years in our data, no individual brand is as predictive of being high-income as owning an Apple iPhone in 2016.

We of course acknowledge that some of the differences in consumer behavior between the rich and the poor reflect differences in budgets rather than differences in anything that we should call culture. Presumably, the poor do not own dishwashers because they cannot afford them (or have no space for them) rather than because they have a cultural preference for washing dishes by hand. That said, many of the brands that distinguish the rich and the poor, such as mustard or soy sauce, may reflect cultural influence on food choices (cf. Atkin, 2016). Also, the differences between the rich and the poor in other dimensions of culture, such as media consumption and social attitudes, probably primarily reflect choices rather than opportunities.

4.1.3 Time Use

As panel (a) of Figure 1 shows, the cultural distance in time use between the rich and the poor has been largely constant since 1965. One may argue that there was a slight uptick in cultural distance in 2003, but we urge caution in interpreting this apparent change since 2003 is precisely the year when collection of time use data was taken over by BLS and, while the AHTUS attempts to best harmonize over time the different time use data sources, it is possible that this uptick reflects

⁴²This fact is particularly interesting given the way Grey Poupon has been marketed. In the 1980s and early 1990s, the so-called “Pardon me” TV advertisements aired: a Rolls-Royce pulls up alongside another Rolls-Royce; a passenger in the back seat of one asks a passenger in the back seat of the other: “Pardon me, would you have any Grey Poupon?”; the other passenger responds, “But of course!” Since then, Grey Poupon has often been referenced in hip-hop lyrics as a symbol of status. For example, FM Static has a song with a verse, “And if I had money, then I’d only wear Sean John / Eat my cereal with Grey Poupon.” An analysis by vox.com indicates that almost every year since 1992, at least one hip-hop song has been released referencing Grey Poupon. In 2011, 15 such songs were released.

superior survey quality starting in 2003.

As we already mentioned, we focus our analysis of time use on respondents who are employed full-time, but even if we look at the full sample, we see no trends in cultural distance (cf: Figure A.18). In contrast to media consumption, consumer behavior, and social attitudes, it is less straightforward to identify the individually most informative variables in time use. The problem is that the small sample sizes coupled with a rich potential set of responses (in contrast to the binary responses to questions about media, consumer behavior, and attitudes) drive a wedge between the empirical and the true distribution of responses in a given group.⁴³

4.1.4 Social Attitudes

As discussed above, information on how people answer questions in the GSS was increasingly predictive of income groups until the mid-1980s but with little change for the last 30 years or so. To better understand these patterns, we also consider trends in cultural distance based only on subsets of questions in the GSS. In particular, we separate questions related to: law enforcement; marriage, sex and abortion; life and trust; politics and religion; civil liberties; confidence; and government spending. The Data Appendix reports the complete list of GSS questions included under each theme. Table 3 reports cultural distance based on each of these subsets of GSS questions for the years 1976, 1996, and 2016 (beginning, middle, and end of our sample). We report the full GSS results at the top for reference. The last two columns of Table 3 report the coefficient (and the t-statistic) from an estimate of a linear trend in cultural distance (based on that subset of questions) using all years in the GSS. The topics are listed in order of decreasing estimate of the trend.

Table 3 reveals that the rich and the poor have diverged the most in their attitudes toward law enforcement. At the other extreme, there is no indication of any divergence based on confidence in institutions or views about government spending. Table A.1 in the Online Appendix reports the ten social attitudes that are single-handedly most predictive of being rich in 1976, 1996, and 2016. The top of these lists is remarkably stable in each year: voting and trusting people are among the three most individually predictive variables in all three years. Figure 4 presents a more systematic analysis of stability of relative predictability over time. We rank order all of the variables based on

⁴³For example, in 1965, some poor respondents spent 5 hours a week on gardening, none spent 7 hours, but some spent 9; among the rich respondents, some spent 5 hours on gardening, some spent 7, but none spent 9. It would obviously be mistaken to conclude from this pattern of responses that anyone who spent 7 hours on gardening must be rich while anyone who spent 9 hours on gardening must be poor. This issue is the primary problem tackled by Gentzkow et al. (2017) in their analysis of polarization of political speech. We could follow their approach or other methods for dealing with the finite sample bias here, but for ease of exposition we try to avoid customizing our analyses for particular groups or dimensions of cultural distance.

Table 3: Cultural distance by income over time: social attitudes

	1976	1996	2016	Coefficient	T-statistic
All GSS	70.4%	74.0%	76.8%	0.20	3.45
Law Enforcement	54.9%	54.9%	63.9%	0.35	3.79
Politics & Religion	63.9%	64.8%	67.4%	0.12	2.49
Life & Trust	61.2%	61.7%	64.9%	0.11	2.44
Marriage, Sex, Abortion	63.1%	60.2%	64.0%	0.10	1.36
Civil Liberties	62.9%	60.2%	59.4%	0.09	1.06
Confidence	57.7%	61.7%	56.1%	-0.07	-0.83
Government Spending	60.1%	59.5%	57.4%	-0.12	-2.45

Note: Data source is the GSS. Sample size in each year is 398. Rows 2 to 8 present the results of the machine-learning ensemble method when performed only on the subset of GSS variables in that row. See Data Appendix for the list of GSS questions included in each row. See text and Data Appendix for details on sample construction and implementation of the ensemble. Columns 1 to 3 present the share of correct guesses of respondent’s income in the hold-out sample in 1976, 1996, and 2016. The procedure to guess income in the hold-out sample was repeated 100 times. The remaining columns present results from a univariate regression of the share of correct guesses between 1976 and 2016 on a linear year trend, including the estimated coefficient on year (column 4) and the t-statistic associated with that estimated coefficient.

how informative they are about income in 1976. We use this ranking to determine each variable’s vertical position throughout the graph. We then color-code the relative informativeness of each variable in each year, with the most informative variables colored dark red, the least informative ones dark blue, and lighter shades of red and blue in between. If the relative informativeness of variables were perfectly stable, each horizontal line would be uniformly colored. Figure 4 reveals that there are substantial changes in the relative importance of specific questions over time, but a small set of highly predictive variables remain highly predictive each year.

4.2 Education

One concern about the classification of households into rich and poor based on income is that we only observe current income, while permanent income might be far more relevant. In this section, we analyze cultural divergence by education, which can be seen as a proxy for permanent income. Examining cultural divergence across education groups is also informative because of the role that education has played in the rise of income inequality (Katz and Murphy 1992). We classify people as less educated if they have at most completed high school, and more educated otherwise.^{44,45}

Panel (b) of Figure 1 summarizes our results. In short, we find similar patterns as with income

⁴⁴Given this definition and sample size equalization over time, each prediction of education level in a given year is based on 9,674 observations in the MRI, 652 observations in the GSS, and 524 observations in the AHTUS.

⁴⁵The definition of these two groups based on an absolute level of education means that we avoid the issue of potential miscategorization discussed in footnote 38. The fact that we mostly observe similar temporal trends as with income gives us confidence that our results on income were not compromised by miscategorization.

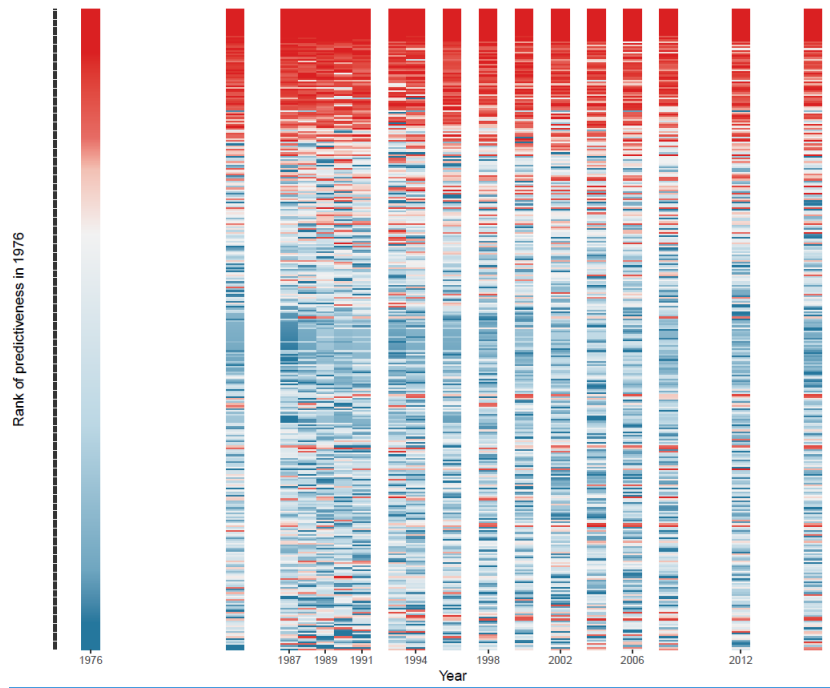


Figure 4: Stability over time of attitudes most indicative income

Note: Data source is the GSS. Sample size is 398. Variables are ranked from bottom to top throughout the graph by increasing order of correctly guessing income in 1976 based on that variable only. Each variable's relative informativeness in subsequent years is color-coded, with the most informative variables in each year color-coded dark red and the least informative color-coded dark blue, and lighter shades of red and blue in between. See Data Appendix for implementation details.

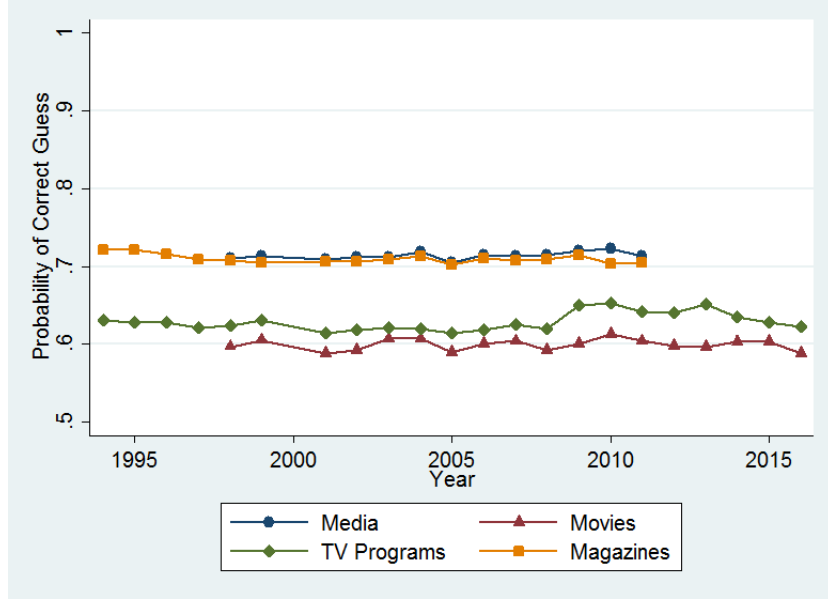


Figure 5: Cultural distance by education over time: media consumption

Note: Data source is the MRI. Sample size each year is 9,674. See text and data appendix for details on sample construction and implementation of machine-learning ensemble method. Presented in the figure is share of correct guesses of respondent's education in the hold-out sample each year. The procedure to guess education in the hold-out sample was repeated 25 times, and the share of guesses reported is the average of these 25 iterations.

with the exception that we see no divergence in social attitudes. The ability to correctly guess one's education group based on social attitudes is essentially the same in 2016 as it was in 1976. Because the results by education mostly match those by income, we do not go through all of them in detail here. More generally, any result reported for some group but not for another in the body of the paper is available in the Online Appendix.

4.2.1 Media Consumption

Figure 5 shows that the cultural distance in media consumption by education has been stable both overall as well as based on any one of the three sub-categories of media (movies, TV shows, and magazines).

The figure also shows that we can predict education based on magazine readership about as well as we can do with the full set of media consumption variables. Table A.2 reports the set of movies, TV shows, and magazines most indicative of higher education. Magazines that are highly indicative of being educated are stable over time, with *Newsweek* and *Time* topping the list throughout the sample. Sports-related TV programs, while also indicative of education, appear less important than they were for income. In 1994, *Rescue 911* and *Unresolved Mysteries* were the two TV shows whose

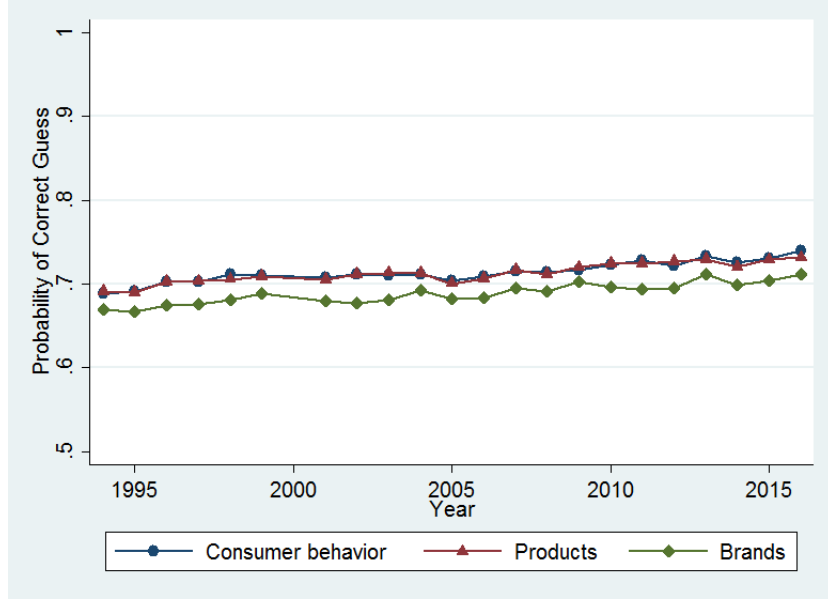


Figure 6: Cultural distance by education over time: consumer behavior

Note: Data source is the MRI. Sample size each year is 9,674. See text and data appendix for details on sample construction and implementation of machine-learning ensemble method. Presented in the figure is share of correct guesses of respondent's education in the hold-out sample each year. The procedure to guess education in the hold-out sample was repeated 5 times, and the share of guesses reported is the average of these 5 iterations.

viewership was most informative of someone education's level, with not watching these shows being indicative of being more educated. In 2005, watching the *Super Bowl* and not watching *Cops* were the most indicative of being educated. By 2016, watching *Love It or List It* and *Property Brothers*, both HGTV shows, were the most indicative of being educated.

4.2.2 Consumer Behavior

As indicated above, our ability to correctly guess one's education based on consumer behavior has remained mostly stable, with maybe a slight increase over time. This is also true when we attempt to predict education using products and brands separately, as shown in Figure 6.

Products and brands most indicative of having attended college are dominated by technological goods (cf: Table A.3). The specific goods reflect waves of technological innovation with the more educated separating themselves from the less educated by being earlier adopters of new technologies. For example, personal computers and PC-related devices are more dominant in the early years, while owning a tablet is most indicative of being educated in 2016. For brands, having bought Kodak film, owning Windows XP, and owning an Apple iPhone are most informative about someone's education in 1994, 2005, and 2016, respectively. Across time, the consumption of travel-related

items also separate the more and the less educated.

4.2.3 Time Use

Aguiar and Hurst (2007) find that less-educated individuals have experienced greater increases in leisure compared to their more-educated counterparts. Given this finding, it might seem surprising that we do not see an increase in cultural differences in time use by education level. The mean number of hours per week spent on leisure⁴⁶ was indeed roughly the same for the two education groups in 1975, but the less-educated were spending relatively more hours per week on leisure by 2003-2012. That said, the overall distribution of time spent on leisure for the two groups is very similar, both in 1975 and 2003-2012 (cf: Figure A.19). There is substantial heterogeneity in the amount of time spent on leisure within each group, and this heterogeneity is much greater than the mean difference across groups. Moreover, as noted by Aguiar and Hurst (2007), the within-group heterogeneity has also increased over time.⁴⁷ Consequently, time use is no more informative about education now than it was in the past.

4.2.4 Social Attitudes

While we saw some evidence of a growing divide between the rich and the less rich based on their social views, we fail to detect any such time trend based on education. Table 4 shows that, just as we observed for income groups, the more- and less-educated have been somewhat diverging over time in their answers to questions related to law enforcement and life and trust. On the other hand, it has become more difficult over time to tell the more- and less-educated apart based on views towards government spending, confidence levels, and views on civil liberties. Aggregating across all topics, the distance in social attitudes has been constant over time.

4.3 Gender

Much has changed for women over the last half century. Educationally, women turned an educational deficit relative to men into an educational surplus (Goldin et al., 2006). Women's labor force participation rate increased, though it seems to have reached a plateau in the mid-1990s; women's

⁴⁶Leisure is defined as time spent watching TV, socializing, playing sports, reading, engaging in hobbies, sleeping, eating, and engaging in personal care.

⁴⁷Aguiar and Hurst (2007) write: "We also document a significant dispersion of leisure within educational categories... while the growing leisure gap between educational groups is substantial, it is more than matched by the growing within-group dispersion."

Table 4: Culture distance by education over time: social attitudes

	1976	1996	2016	Coefficient	T-statistic
All GSS	70.7%	70.2%	71.1%	0.03	0.59
Law Enforcement	57.4%	56.8%	59.9%	0.10	1.98
Life & Trust	60.7%	60.6%	62.2%	0.07	1.83
Politics & Religion	63.7%	61.6%	63.0%	0.03	0.77
Marriage, Sex, Abortion	64.8%	60.7%	61.8%	-0.07	-1.20
Civil Liberties	66.6%	63.5%	62.8%	-0.10	-1.78
Confidence	62.6%	58.8%	57.3%	-0.13	-2.31
Government Spending	62.1%	57.0%	57.1%	-0.17	-4.16

Note: Data source is the GSS. Sample size in each year is 652. Rows 2 to 8 present the results of the ensemble machine-learning method when performed only on the subset of GSS variables in that row. See Data Appendix for the list of GSS questions included in each row. See text and Data Appendix for details on sample construction and implementation of the ensemble. Columns 1 to 3 present the share of correct guesses of respondent’s education in the hold-out sample in 1976, 1996, and 2016. The procedure to guess education in the hold-out sample was repeated 100 times. The remaining columns present results from a univariate regression of the share of correct guesses between 1976 and 2016 on a linear year trend, including the estimated coefficient on year (column 4) and the t-statistic associated with that estimated coefficient.

labor market earnings converged towards those of men, but this convergence also appears to have slowed down in the most recent decade (Bertrand, 2018). While these well-established trends may a priori suggest shrinking cultural divides between the genders, other forces may have pushed in the other direction. First, women’s greater financial independence may have allowed them to also achieve greater cultural independence from their husbands, with the goods, the experiences, and the media they consume becoming more closely aligned with their own personal preferences. Similarly, the decline in marriage and the rise in divorce may have contributed to a cultural divergence between men and women as cultural choices became less likely to take place within the confines of the couple. Moreover, as noted by Edlund and Pande (2002), changes in family structure may have directly affected women’s social attitudes, by strengthening their redistributive preferences, support for greater government spending, and overall support for more democratic political platforms.⁴⁸

Panel (c) of Figure 1 summarizes our results by gender.⁴⁹ There is no evidence of an increasing cultural gap between men and women based on media consumption or social attitudes. Our ability to predict gender based on consumer behavior is nearly perfect in every period, so this is one instance where our approach to measuring cultural distance is ill suited: the presence of a few highly gender-specific goods masks any potential changes in the gender gap in consumption of other goods. Finally, we see that men and women’s time use became much more similar from 1965

⁴⁸See also Montgomery and Stuart (1999) and Box-Steffensmeier, Boef, and Lin (2000).

⁴⁹Given sample size equalization over time, each prediction of gender in a given year is based on 15,036 observations in the MRI, 1,000 observations in the GSS, and 668 observations in the AHTUS.

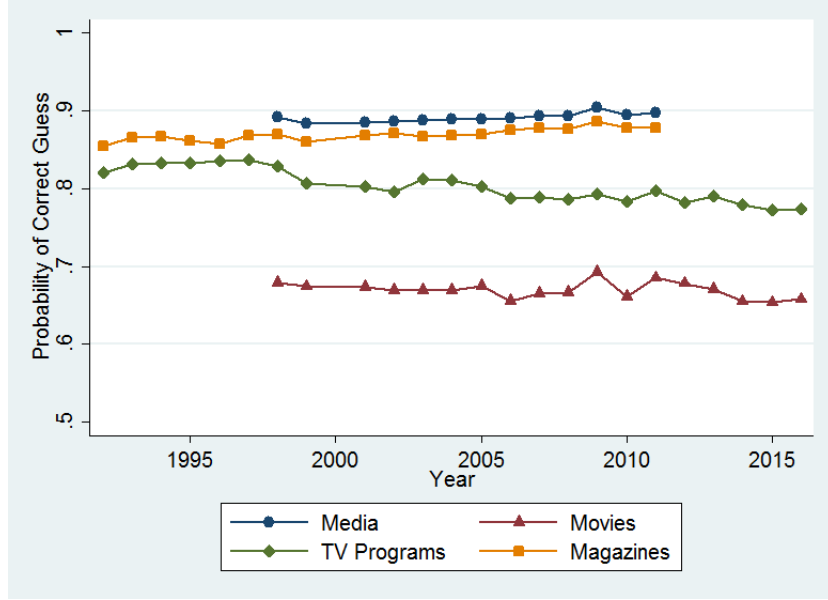


Figure 7: Cultural distance by gender over time: media consumption

Note: Data source is the MRI. Sample size each year is 15,036. See text and data appendix for details on sample construction and implementation of machine-learning ensemble method. Presented in the figure is share of correct guesses of respondent's gender in the hold-out sample each year. The procedure to guess gender in the hold-out sample was repeated 25 times, and the share of guesses reported is the average of these 25 iterations.

to the mid-1990s, but there has been no further change in this dimension of cultural distance since the mid-1990s.

4.3.1 Media Consumption

Figure 7 shows that the cultural distance in media consumption by gender has been broadly stable in all the three sub-categories of media, with perhaps some mild divergence in consumption of magazines and mild convergence in consumption of TV shows. The figure also reveals that magazine readership is most informative about gender, with limited gains in predictive power coming from adding data on TV shows and movies to the ensemble algorithm.

Table A.5 in the Online Appendix shows the movies that men and women sort on. Movies most indicative of males tend to fall into the action, thriller, and sci-fi categories while dramas and romantic comedies are most indicative of females. Table A.5 also reveals that in the early years, the most discriminating movies were those watched primarily by women; in 1998, each of the top ten most predictive movies was more likely to be watched by a woman. In the later years, gender-specific movies are mostly those watched primarily by men. Yet, despite these changes, the overall difference in movie-watching by gender has been constant over time.

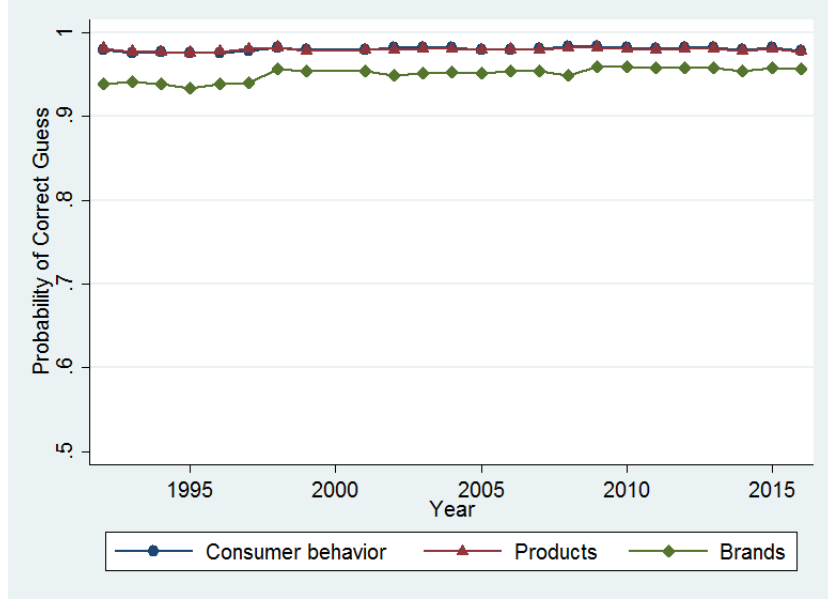


Figure 8: Cultural distance by gender over time: consumer behavior

Note: Data source is the MRI. Sample size each year is 15,036. See text and data appendix for details on sample construction and implementation of machine-learning ensemble method. Presented in the figure is share of correct guesses of respondent's gender in the hold-out sample each year. The procedure to guess gender in the hold-out sample was repeated 5 times, and the share of guesses reported is the average of these 5 iterations.

Table A.5 also shows that, unsurprisingly, fashion and housekeeping magazines are the most distinctive of a female reader, while *Sports Illustrated* is the most distinctive of a male reader. In television, most indicative of a male viewer in all years is watching college and professional American football, with nearly all TV shows predictive of gender in early and middle parts of the sample being football-related. Creeping into the list of the top TV shows most indicative of gender in the final survey year (2016) are 3 HGTV shows, which are disproportionately watched by women.

4.3.2 Consumer Behavior

As indicated above, our approach to measuring cultural distance is not suitable for comparing consumer behavior of men and women since there are some highly gender-specific products that collectively allow us to perfectly predict gender based on consumer behavior in every year.

Table A.6 in the Online Appendix lists the products and brands that are individually most indicative of gender. In particular, the consumption of personal care and makeup is so common and so segmented by gender that knowing whether one bought, used, or own these products provides sufficient information to infer gender nearly perfectly.⁵⁰

⁵⁰We could throw out some "overly gender-specific" products from our data and measure the predictability of the

4.3.3 Time Use

The changes we observe over time in the differences between how men and women spend their time are most striking. The time pattern we document in panel (c) of Figure 1 is reminiscent of the well-known time series on women’s labor force participation in the US, with increases in women’s labor force participation up until the mid-1990s and a subsequent plateau. However, recall that the sample of men and women we study in the time use data is restricted to the full-time employed, so the observed pattern cannot be due to changes in women’s likelihood of being employed. In fact, we observe the same pattern if we predict gender based on shares of non-work time spent on various non-work-related activities (cf: Figure A.20); ways that men and women spend their time outside of work became more similar from 1965 to 1995, but this convergence has stopped since then. This pattern also echoes the fact that attitudes toward gender roles became more egalitarian over time, but this progress stalled in the mid-1990s (Fortin 2015).

4.3.4 Social Attitudes

As shown in panel (c) of Figure 1, we observe no time trend in one’s ability to predict gender based on social views and norms.

Table 5 examines cultural distance by gender based on the seven thematic subsets of GSS questions. We observe only one topic for which the genders appear more divided today than in the past, namely marriage, sex, and abortion. The one module over which men and women have converged over time is views about life and trust. We observe no systematic time trend in cultural distance for the other five themes covered in the GSS survey. The lack of divergence in views about government spending and politics and religion stands in contrast with the work which has argued that women have moved further over time to the political left of men (Edlund and Pande 2002).

4.4 Race

Our motivation for studying racial differences in culture over time is similar to our motivation for studying differences in culture by income groups over time. Just as growing cultural gaps between

other variables but, as we discuss in footnote 43, for ease of exposition we prefer not to customize our approach for particular groups or dimensions of culture. In the Nielsen data, where we cannot perfectly predict gender based on consumer behavior, we observe convergence over time in shopping behavior between men and women (see Figure A.9). The main explanation for why we can perfectly predict gender in MRI but not Nielsen is because Nielsen collects products and brands bought, not used or owned. For example, in 2004, 76% of women in MRI reported using lipstick or lip gloss; in contrast, only 21% of women in Nielsen had bought lipstick.

Table 5: Culture distance by gender over time: social attitudes

	1976	1996	2016	Coefficient	T-statistic
All GSS	72.8%	70.3%	69.3%	-0.08	-2.25
Marriage, Sex, Abortion	57.8%	60.4%	61.7%	0.09	2.61
Law Enforcement	57.8%	63.6%	57.8%	0.04	0.73
Confidence	55.3%	55.7%	56.4%	0.03	0.74
Politics & Religion	56.3%	52.0%	53.4%	-0.02	-0.71
Civil Liberties	53.2%	50.9%	51.3%	-0.01	-0.23
Government Spending	61.0%	56.4%	60.0%	-0.09	-1.89
Life & Trust	68.7%	64.6%	59.8%	-0.28	-7.17

Note: Data source is the GSS. Sample size in each year is 1,000. Rows 2 to 8 present the results of the ensemble machine-learning method when performed only on the subset of GSS variables in that row. See Data Appendix for the list of GSS questions included in each row. See text and Data Appendix for details on sample construction and implementation of the ensemble. Columns 1 to 3 present the share of correct guesses of respondent’s gender in the hold-out sample in 1976, 1996, and 2016. The procedure to guess gender in the hold-out sample was repeated 100 times. The remaining columns present results from a univariate regression of the share of correct guesses between 1976 and 2016 on a linear year trend, including the estimated coefficient on year (column 4) and the t-statistic associated with that estimated coefficient.

rich and poor may hinder social mobility, a growing cultural divide between whites and non-whites may cause continued economic struggles for minority groups in the US.

Given our limited sample sizes, we are unable to conduct our analysis across many racial categories. Instead, we focus on comparison of whites and non-whites. Given this grouping and our procedure for equalizing sample sizes described in Section 3.4, each prediction of race in a given year is based on 4,150 observations in the MRI, 234 observations in the GSS, and 298 observations in the AHTUS.

There is a rich literature on cultural differences by race. Fryer and Levitt (2004) discuss some of the prior work on the black-white cultural divide on dimensions such as musical tastes and linguistic patterns. Fryer and Levitt (2004) also highlight anecdotal evidence of racial differences in consumer behavior (e.g., the sharp racial differences in the popularity of different cigarette brands) and media consumption (e.g., *Seinfeld*’s huge following among whites and limited appeal among blacks). However, a systematic documentation of changes over time in cultural differences by race is missing. One important exception is Fryer and Levitt’s (2004) analysis of trends over time in the names whites and blacks give to their children. They show that differences in name choices were relatively small in the 1960s, but a profound shift took place in the early 1970s, especially among blacks living in more racially segregated areas. In this section, we analyze whether such cultural divergence also took place on other cultural dimensions, focusing on differences between whites and all non-whites.

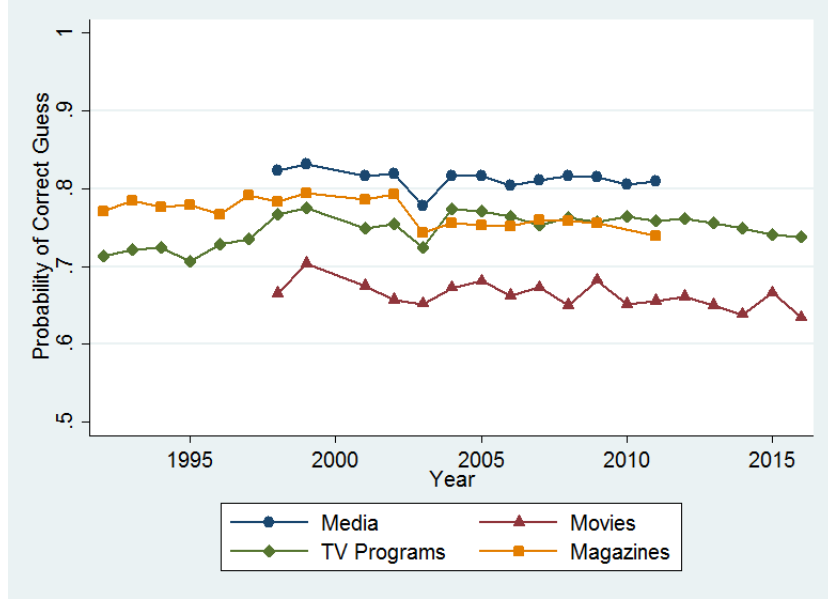


Figure 9: Cultural distance by race over time: media consumption

Note: Data source is the MRI. Sample size each year is 4,150. See text and data appendix for details on sample construction and implementation of machine-learning ensemble method. Presented in the figure is share of correct guesses of respondent's race in the hold-out sample each year. The procedure to guess race in the hold-out sample was repeated 25 times, and the share of guesses reported is the average of these 25 iterations.

Panel (d) of Figure 1 summarizes our results. With the exception of an outlier data point in 1975, there is no apparent trend in our ability to predict one's race based on time use. There is also no evidence of growing racial cultural divides based on patterns of media consumption. Furthermore, there is no evidence of any growing racial divides based on social attitudes; if anything, there has been some weak convergence. The one dimension where we do observe increasing racial differences is with respect to consumption choices, with much of the increase in the racial gap occurring during the 1990s.

4.4.1 Media Consumption

Panel (a) of Figure 9 examines racial differences in consumption of the three media subcategories. While there are no steep trends, the differences in consumption of magazines and movies seem to have gotten somewhat smaller, while the differences in TV programs have increased.

Unlike with the prior results (income, education, gender), in the case of race, the combination of all data on media consumption substantially increases the predictive power of our model compared to focusing on a subcategory (e.g., magazines).

Looking at specific media products that are individually most predictive of race (cf: Table

A.8), a few patterns emerge. Thrillers and horror movies appear distinctively more popular among non-whites than whites; a few movies also emerge that appear a priori targeted at a more black audience (such as *The Preacher's Wife* in 1998 or *Big Momma's House 2* in 2007). In contrast, dramas appear distinctively more popular among whites. The same three magazines (*Ebony*, *Jet*, *Essence*) are most indicative of race throughout our sample period. In contrast, the TV shows most indicative of being white vary quite a lot over time. Some popular sitcoms clearly have differential appeal across racial lines (e.g., *In Living Color* in 1992 was more popular among non-whites while *The Big Bang Theory* in 2016 was more popular among whites). There is also evidence that whites and non-whites sort into different sports on TV, with American football being more popular among whites and basketball more popular among non-whites.

4.4.2 Consumer Behavior

As already indicated, we observe a growing divide between whites and non-whites in terms of consumption choices. The probability of correctly guessing race from consumer behavior grows from roughly 80 percent in 1992 to roughly 90 percent in 2016, with much of the increase taking place during the 1990s and early 2000s. Figure 10 reports on these patterns when we restrict the consumer data to only products or brands. These time series display the same patterns as the full consumer behavior data, with increases in the probability of correctly guessing race concentrated in the first half of the sample period.

The list of products most indicative of being white (cf: Table A.9) is intriguing. A few common items in each year (1992, 2004, and 2016) are pets and flashlights, with whites being distinctive in owning those items.⁵¹ Other items on the list may be a reflection of the systematic income differences between whites and non-whites, such as owning a dishwasher or having a car with cruise control. It also appears that whites are substantially more invested into kitchen appliances than non-whites.

4.4.3 Time Use

We do not see a clear trend in racial differences in time use, though the data show a blip in 1975. We suspect this data point is an accidental outlier, but without a formal take on standard errors (cf: Section 3.5), we cannot take a firm stance. If we restrict our attention to data since 2003,

⁵¹Owning pets might simply reflect living in the suburbs; our data does not allow us to explore this. The MRI provides information on whether the respondent lives in a large core-based statistical area, but does not indicate whether the residence is an urban or a suburban area.

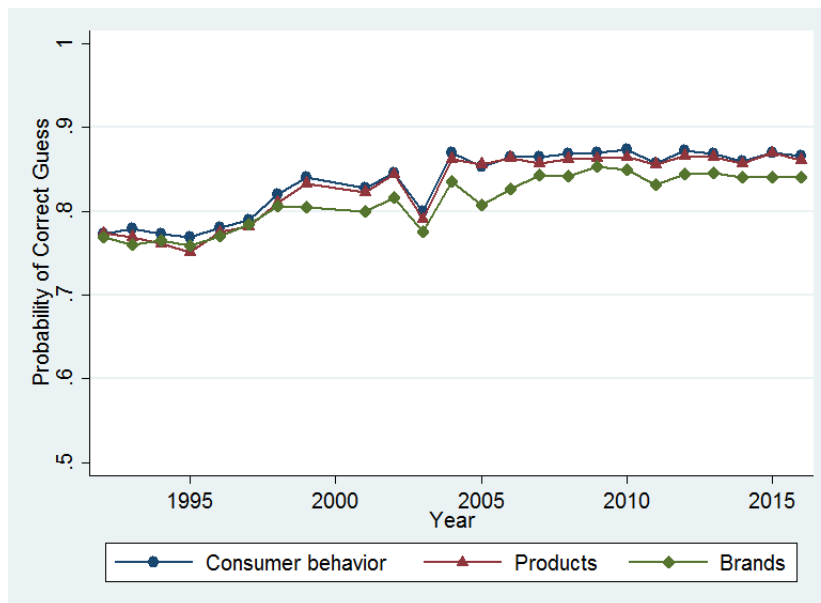


Figure 10: Cultural distance by race over time: consumer behavior

Note: Data source is the MRI. Sample size each year is 4,150. See text and data appendix for details on sample construction and implementation of machine-learning ensemble method. Presented in the figure is share of correct guesses of respondent's race in the hold-out sample each year. The procedure to guess race in the hold-out sample was repeated 5 times, and the share of guesses reported is the average of these 5 iterations.

which gives us much larger sample sizes, we see a very stable difference in time use by race over that decade (cf: Figure A.21).

4.4.4 Social Attitudes

Cultural distance in social attitudes by race shows slight convergence, with the probability of an accurate guess decreasing from 80 percent to 75 percent over the 40 years of the data. This overall pattern, however, masks some sharp differences in trends across sub-categories of the GSS.

As shown in Table 6, whites and non-whites have grown apart on their views on law enforcement. In 1976, one could correctly predict race based on these views 60 percent of the time but by 2016 that number was up to 70 percent. On the other hand, whites and non-whites have sharply converged in their views on life and trust, government spending, and politics and religion. For example, in 1976, one could correctly predict race based on views towards government spending 74 percent of the time but by 2016 this number was down to 56 percent.

Table 6: Culture distance by race over time: GSS

	1976	1996	2016	Coefficient	T-statistic
All GSS	80.5%	79.4%	74.8%	-0.14	-2.60
Law Enforcement	60.9%	65.4%	70.3%	0.25	5.19
Civil Liberties	50.4%	56.6%	55.3%	0.16	2.21
Marriage, Sex, Abortion	60.7%	58.1%	58.9%	0.05	0.76
Confidence	56.2%	58.5%	53.4%	-0.15	-2.42
Politics & Religion	73.0%	69.1%	61.5%	-0.23	-5.41
Government Spending	74.3%	71.8%	56.0%	-0.33	-3.94
Life & Trust	68.8%	61.7%	52.6%	-0.40	-6.53

Note: Data source is the GSS. Sample size in each year is 234. Rows 2 to 8 present the results of the ensemble machine-learning method when performed only on the subset of GSS variables in that row. See Data Appendix for the list of GSS questions included in each row. See text and Data Appendix for details on sample construction and implementation of the ensemble. Columns 1 to 3 present the share of correct guesses of respondent’s race in the hold-out sample in 1976, 1996, and 2016. The procedure to guess race in the hold-out sample was repeated 100 times. The remaining columns present results from a univariate regression of the share of correct guesses between 1976 and 2016 on a linear year trend, including the estimated coefficient on year (column 4) and the t-statistic associated with that estimated coefficient.

4.5 Political Ideology

Of all the cultural divides under study in this paper, the divide that separates Democrats from Republicans, or liberals from conservatives, has received the most prior attention. A large literature in political science documents the rising polarization of Democrats and Republicans in Congress, and the discussions of political polarization is undoubtedly on the rise (Gentzkow 2016), but the literature to date has been far less conclusive on whether Republicans and Democrats in the US population overall have been growing apart. Most of the academic work has focused on differences in social attitudes captured in the GSS or the ANES, with some studies concluding that polarization is on the rise (e.g., Abramowitz and Saunders, 2008; Draca and Schwarz, 2018) and others rejecting this conclusion (e.g., Fiorina and Abrams, 2008; Glaeser and Ward, 2006). Much of the disagreement between these studies is ultimately driven by differences in how polarization is measured. We contribute to this literature by examining what our definition of cultural distance implies for changes in the attitude-differences between liberals and conservatives. Furthermore, we extend the literature by also analyzing the divide between liberals and conservatives in other aspects of culture, namely media consumption and consumer behavior. (Political affiliation is not available in our time use data.)

In the GSS, respondents are categorized as (a) extremely liberal, (b) liberal, (c) slightly liberal, (d) moderate, (e) slightly conservative, (f) conservative, or (g) extremely conservative; we define respondents as liberal if they identify themselves as (a), (b), or (c) and conservative if they identify

themselves as (e), (f), or (g). In the MRI, respondents are categorized as (a) very liberal, (b) somewhat liberal, (c) middle of the road, (d) somewhat conservative, or (e) very conservative; we define respondents as liberal if they identify themselves as (a) or (b) and conservative if they identify themselves as (d) or (e). Given our sample-size-equalization procedure, this yields 4,864 observations per year in the MRI and 566 observations per year in the GSS. We discovered a sharp change in the MRI data in the number of missing observations on ideology after 2009, so based on our principle of keeping the quality of the data constant over time, we analyze media consumption and consumer behavior by ideology only in the 1994-2009 period.

Panel (e) of Figure 1 summarizes our results. Our ability to predict someone’s political ideology based on patterns of media consumption or consumer behavior is essentially constant across years. On the other hand, we document a growing divide between liberals and conservatives based on their stated social attitudes.

4.5.1 Media Consumption

Figure 11 shows that the stability of the cultural distance between liberals and conservatives based on the full media bundle broadly extends to the three separate media sub-components, with one exception. Liberals and conservatives converged in terms of the TV shows they watch during the 1990s; since then, the difference in their TV habits has been mostly stable.

We also observe that TV shows are not only more predictive of ideology than movies or magazines, they are more predictive than all three subcomponents put together. While this might seem puzzling at first glance, it simply reflects the fact that our machine learning algorithm is not fully non-parametric, so inclusion of additional, less predictive variables can decrease predictive power of the estimated model.⁵²

Table A.11 in the Online Appendix lists the TV programs, movies, and magazines that are single-handedly most indicative of political ideology. There is substantial variation over time in the list of most predictive TV shows. The contrast between the list of top shows in 2001 and 2009 is particularly interesting. The three TV programs most indicative of ideology in 2001 are *The Academy Awards*, *Will and Grace*, and *Friends*, with liberals disproportionately watching these shows. In contrast, the three TV programs most indicative of ideology by 2009 are all Fox news programs: *The O’Reilly Factor*, *Fox and Friends*, and *Hannity and Colmes*, with conservatives dis-

⁵²This is especially the case for random forests. If instead of ensemble (which includes random forest), we predict ideology using an elastic net or a regression tree, the predictive power is greater when we use all media variables rather than TV shows alone.

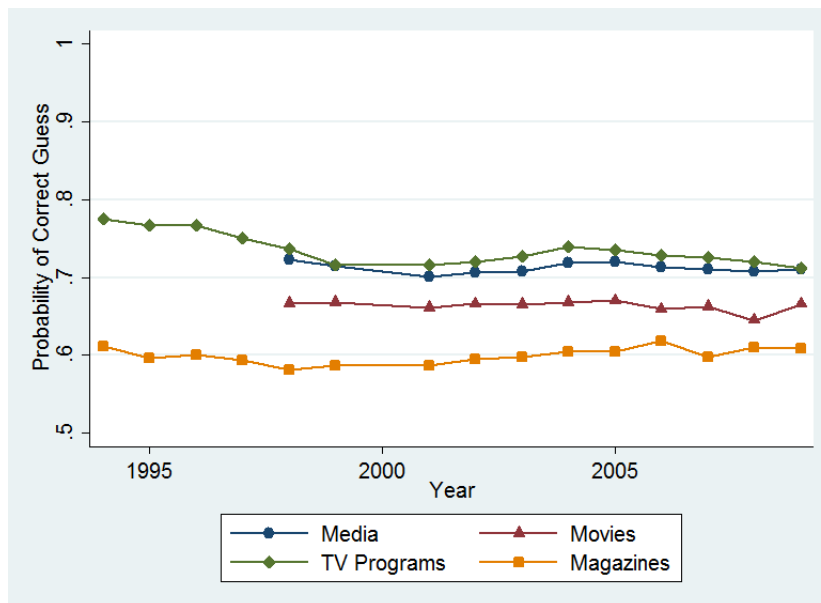


Figure 11: Cultural distance by political ideology over time: media consumption

Note: Data source is the MRI. Sample size each year is 4,864. See text and data appendix for details on sample construction and implementation of machine-learning ensemble method. Presented in the figure is share of correct guesses of respondent's political ideology in the hold-out sample each year. The procedure to guess political ideology in the hold-out sample was repeated 25 times, and the share of guesses reported is the average of these 25 iterations.

proportionately watching all three. This contrast provides a good reminder of one key (potentially undesirable) feature of our measure of cultural distance, which is that the measure does not take any stance on either which elements of cultural behavior are important nor which elements are close to one another. One might argue that the distance between liberals and conservative is greater in 2009 than in 2001 in that *The O'Reilly Factor* or *Fox and Friends* are more different (the vodka in our earlier example) from the rest of what TV has to offer (the water) than *The Academy Awards* or *Friends* are (the Sprite). Moreover, even if one does not take a stance on the “distance” between *Friends* and other sitcoms, we might worry about differences in where people get their news much more than about differences in where people get their non-news entertainment.

4.5.2 Consumer Behavior

As indicated in Figure 1, our ability to correctly predict political ideology based on the basket of goods and brands consumed hovers around the low 70 percent range throughout the sample period. Figure 12 shows this pattern is similar when we restrict the consumer behavior information to either products or brands, with the exception of a noticeable increase in ideological distance based on brands at the end of the 1990s.

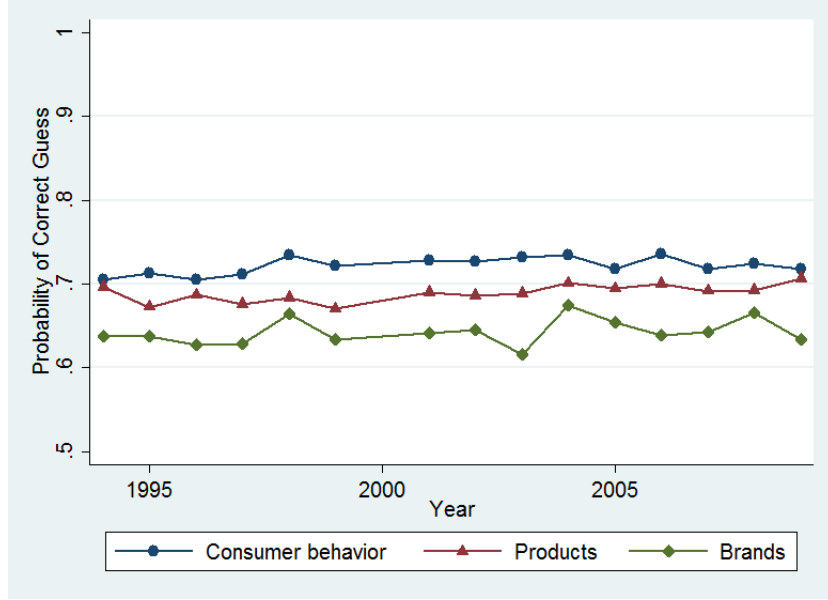


Figure 12: Cultural distance by political ideology over time: consumer behavior

Note: Data source is the MRI. Sample size each year is 4,864. See text and data appendix for details on sample construction and implementation of machine-learning ensemble method. Presented in the figure is share of correct guesses of respondent's political ideology in the hold-out sample each year. The procedure to guess political ideology in the hold-out sample was repeated 5 times, and the share of guesses reported is the average of these 5 iterations.

The list of products individually most distinctive of ideology is interesting (cf: Table A.12). In all years, liberals distinguish themselves from conservatives by drinking alcohol. Conservatives, on the other hand, are much more likely to engage in fishing.

Ideology-specific brands are mostly food, primarily with brands indicative of conservatives who disproportionately buy Jell-O gelatin desserts and eat at *Arby's*.⁵³

4.5.3 Social Attitudes

Differences in social attitudes between liberals and conservatives is the dimension of cultural distance that has received the most prior attention. Based on our measure, we find that liberals and conservatives are more different today in their social attitudes than they have ever been in the last 40 years. Moreover, this divergence is not a recent phenomenon. Furthermore, Table 7 shows that liberals and conservatives have diverged in their views on almost every one of the seven thematic subsets in the GSS.⁵⁴ In particular, while we detect no time trend in our ability to tell liberals

⁵³In 2009, the most predictive TV show primarily watched by liberals is *The Daily Show with Jon Stewart*. In this show, making fun of *Arby's* is perhaps the most commonly repeated gag. At the same time, the brand most predictive of ideology in 2009 was in fact *Arby's*.

⁵⁴Recall, as we mentioned in Footnote 21, that in analyzing ideological differences in social attitudes, we drop questions about the respondent's political affiliation and questions about how the respondent voted in a presidential

Table 7: Culture distance by political ideology over time: GSS

	1976	1996	2016	Coefficient	T-statistic
All GSS	68.4%	74.5%	81.1%	0.41	6.03
Marriage, Sex, Abortion	63.0%	70.5%	77.8%	0.41	5.22
Voting Participation & Religion	50.9%	57.7%	61.8%	0.33	5.21
Confidence	58.6%	62.6%	69.5%	0.31	5.65
Government Spending	64.2%	62.5%	70.7%	0.23	3.69
Law Enforcement	64.1%	62.4%	67.2%	0.19	4.25
Life & Trust	52.9%	50.9%	53.6%	0.13	2.13
Civil Liberties	60.5%	56.7%	57.2%	0.02	0.17

Note: Data source is the GSS. Sample size in each year is 566. Rows 2 to 8 present the results of the ensemble machine-learning method when performed only on the subset of GSS variables in that row. See Data Appendix for the list of GSS questions included in each row. See text and Data Appendix for details on sample construction and implementation of the ensemble. Columns 1 to 3 present the share of correct guesses of respondent’s political ideology in the hold-out sample in 1976, 1996, and 2016. The procedure to guess political ideology in the hold-out sample was repeated 100 times. The remaining columns present results from a univariate regression of the share of correct guesses between 1976 and 2016 on a linear year trend, including the estimated coefficient on year (column 4) and the t-statistic associated with that estimated coefficient.

and conservatives apart based on views towards civil liberties, we see cultural divergence in the remaining six dimensions. Divergence has been greatest in views on marriage, sex and abortion; voting participation and religion; and confidence.

5 Conclusion

We study temporal trends in cultural distances as reflected in media consumption, consumption choices, time use, and social views between groups in the US defined by income, education, gender, race, and political ideology. We use a machine learning approach to measure cultural distance, an approach that is well suited to this particular application given the rich set of features and traits that define someone’s culture. The main take-away of our analysis is that, except for a few noteworthy exceptions, cultural distances have remained broadly constant over time. This take-away runs against the popular narrative of the US becoming an increasingly divided society.

There are, however, a few important caveats to our main finding. First, our approach does not take a stance on what features of culture matter for healthy and productive interactions between groups in society. As we discussed earlier, social frictions may be more affected by whether we get our news from similar sources than by whether we watch different sitcoms. That said, perhaps people primarily connect by talking about sitcoms and sports rather than about the news. Nothing

election; accordingly, the questions in the theme “Politics & Religion” are here replaced by a subset on “Voting participation & Religion”. This theme includes only questions about respondent’s religion, attendance of religious services, and whether he or she voted in the last two presidential elections.

in our data provides guidance on which aspects of culture are most important for our ability to “get along.”

A second limitation of our approach is that it can only analyze cultural distances one dimension at a time, since no single dataset encompasses data on media diet, consumption behavior, time use, and social attitudes. It is possible that there have been changes in the correlation between these components of culture over time and that an analysis that draws on an integrated dataset would come to different conclusions from ours.

Finally, our assessment of the extent of cultural divides within the US has been focused on looking at pre-specified groups (rich vs. poor, more vs. less educated, man vs. woman, white vs. non-white, liberal vs. conservative). In future work, we plan to explore trends in polarization across social groups that are defined endogenously, based on their distinct cultural traits.

References

- Abramowitz, Alan I. and Kyle L. Saunders. 2008. Is polarization a myth? *The Journal of Politics*. 70(2): 542-555.
- Aguiar, Mark and Erik Hurst. 2007. Measuring trends in leisure: The allocation of time over five decades. *Quarterly Journal of Economics*. 122(3): 969-1006.
- Aguiar, Mark and Erik Hurst. 2009. A summary of trends in American time allocation: 1965-2005. *Social Indicators Research*. 93(1): 57-64.
- Alesina, Alberto, Reza Baqir, and William Easterly. 1999. Public goods and ethnic divisions. *Quarterly Journal of Economics*. 114(4): 1234-84.
- Alesina, Alberto and Eliana La Ferrara. 2000. Participation in heterogeneous communities. *Quarterly Journal of Economics*. 115(3): 847-904.
- Alesina, Alberto and Eliana La Ferrara. 2005. Ethnic diversity and economic performance. *Journal of Economic Literature*. 43(3): 762-800.
- Alesina, Alberto, Guido Tabellini, and Francesco Trebbi. 2017. Is Europe an optimal political area? *Brookings Papers on Economic Activity*.
- Autor, David H., Lawrence F. Katz and Melissa S. Kearney. 2008. Trends in U.S. wage inequality: Revising the revisionists. *The Review of Economics and Statistics*. 90(2): 300-323.
- Atkin, David. 2016. The caloric costs of culture: Evidence from Indian migrants. *The American Economic Review*. 106(4): 1144-1181.
- Bertrand, Marianne. 2018. Coase Lecture: the glass ceiling. *Economica*. 85(338): 205-231.
- Bourdieu, Pierre. 1984 [1979]. *Distinction: A Social Critique of the Judgement of Taste*. Translated by Richard Nice. Cambridge, MA: Harvard University Press.
- Box-Steffensmeier, Janet, Suzanna De Boef, and Tse-Min Lin. (2004). The dynamics of the partisan gender gap. *American Political Science Review*. 98(3): 515-528.
- Concordance of Activity Codes. American Heritage Time Use Survey Codebook*. University of Oxford Center for Time Use Research.
- Desmet, Klaus, Ignacio Ortuno-Ortin, and Romain Wacziarg. 2017. Culture, ethnicity, and diversity. *The American Economic Review*. 107(9): 2479-2513.
- Desmet, Klaus and Romain Wacziarg. 2018. The cultural divide. Working paper.
- Draka, Mirko and Carlo Schwarz. 2018. How polarized are citizens? Measuring ideology from the ground-up. Working paper.

- Edlund, Lena and Rohini Pande. 2002. Why have women become left-wing? The political gender gap and the decline in marriage. *The Quarterly Journal of Economics*. 117(3): 917-961.
- Fiorina, Morris P. and Samuel J. Abrams. 2008. Political polarization in the American public. *Annual Review of Political Science*. 11: 563-588.
- Fortin, Nicole. 2015. Gender role attitudes and women's labor market participation: Opting-out, AIDS, and the persistent appeal of housewifery. *Annals of Economics and Statistics*. 117/118: 379-401.
- Fryer, Roland G. Jr. and Steven D. Levitt. 2004. The causes and consequences of distinctively black names. *The Quarterly Journal of Economics*. 119(3): 767-805.
- Gentzkow, Matthew. 2016. Polarization in 2016. Toulouse Network for Information Technology Whitepaper.
- Gentzkow et al. 2017. Measuring polarization in high-dimensional data: Method and application to Congressional speech. Working paper.
- Glaeser, Edward L. and Bryce A. Ward. 2006. Myths and realities of American political geography. *Journal of Economic Perspectives*. 20 (2): 119-144.
- Goldin, Claudia, Lawrence F. Katz and Ilyana Kuziemko. 2006. The homecoming of American college women: The reversal of the college gender gap. *The Journal of Economic Perspectives*. 20(4): 133-156.
- Jaravel, Xavier. 2017. The unequal gains from product innovations: Evidence from the US retail sector. Working paper.
- Katz, Lawrence F. and Kevin M. Murphy. 1992. Changes in relative wages, 1963-1987: Supply and demand factors. *The Quarterly Journal of Economics*. 107(1): 35-78.
- Kaufman, Karen. 2002. Culture wars, secular realignment, and the gender gap in party identification. *Political Behavior*. 24(3): 283-307.
- Meyer, Bruce D. and James X. Sullivan. 2017. Consumption and income inequality in the U.S. since the 1960s. NBER working paper 23655.
- Montalvo, José, G. and Marta Reynal-Querol. 2005. Ethnic polarization, potential conflict, and Civil Wars. *American Economic Review*. 95(3): 796-816.
- Montgomery, Robert and Charles Stuart. 1999. Sex and fiscal desire. Working paper.
- Mossel, Elchanan, Allan Sly, and Omer Tamuz. 2014. Asymptotic learning on Bayesian social networks. *Probability Theory and Related Fields*. 158(1-2): 127-157.
- Mullainathan, Sendhil and Jann Spiess. 2017. Machine learning: An applied econometric approach.

- Journal of Economic Perspectives*. 31(2): 87-106.
- Murray, Charles. 2012. *Coming Apart: The State of White America, 1960-2010*. New York: Crown Forum/Random House.
- Oliver, Eric, Thomas Wood, and Alexandra Bass. 2016. *Liberellas* versus *Konservatives*: Social status, ideology, and birth names in the United States. *Political Behavior*. 38:55–81.
- Politis, Dimitris N., Joseph P. Romano and Michael Wolf. 1999. *Subsampling*. New York: Springer series in statistics.
- Waldfogel, Joel. 2003. Preference externalities: An empirical study of who benefits whom in differentiated product markets. *The RAND Journal of Economics*. 34(3): 557–568.
- Wolfram, Walt, and Erik Thomas. 2002. *The Development of African American English*. Oxford, UK: Blackwell Publishers.
- Zimmerman, Seth. 2017. Making the One Percent: The Role of Elite Universities and Elite Peers. Working paper.

Online Appendix for: Coming apart? Cultural distances in the United States over time

Marianne Bertrand and Emir Kamenica

June 2018

A.1 Data Appendix

A.1.1 Sample Construction

General Social Survey

We use the General Social Survey (GSS) to measure cultural distance for social attitudes. We use 18 interspersed years from 1976 to 2016 (1976, 1984, 1987, 1988, 1989, 1990, 1991, 1993, 1994, 1996, 1998, 2000, 2002, 2004, 2006, 2008, 2012, and 2016). While the GSS is available from 1972 to 2016 (annually from 1972 to 1991, 1993, and bi-annually from 1994 to 2016), we restrict the analysis to the 18 interspersed years above as the preferred trade-off between maximizing the number of years (and time coverage) and maximizing the number of common questions asked in each year.

We use 84 questions from the GSS. We define a variable as a dummy variable for each response to a question. For example, the question “Are you happy?” has five possible responses: 1) very happy, 2) pretty happy, 3) not too happy, 4) don’t know, and 5) no answer. We define a variable “Are you happy - very happy” as a dummy variable that equals 1 for response 1) to the question and 0 otherwise. We do the same for the other responses. We organize the full list of variables in seven themes:

Civil liberties: Allow atheists to teach; allow communists to teach; allow militarists to teach; allow racists to teach; allow homosexuals to teach; allow atheists’ books in library; allow communists’ books in library; allow militarists’ books in library; allow racists’ books in library; allow homosexuals’ books in library; allow atheists to speak; allow communists to speak; allow militarists to speak; allow racists to speak; allow homosexuals to speak.

Confidence: confidence in military; confidence in business; confidence in organized religion; confidence in education; confidence in executive branch; confidence in financial institutions; confidence in US Supreme Court; confidence in organized labor; confidence in Congress; confidence in medicine; confidence in the press; confidence in scientific community; confidence in TV.

Government spending: foreign aid; military & defense; solving problems of large cities; halting crime rate; dealing with drug addiction; education; environment; welfare; health care; affirmative action; space exploration programs; income tax too high/adequate/too low.⁵⁵

⁵⁵For the eleven first questions in the government spending module, the GSS has a “split ballot” design since 1984, where one-third of the respondents were asked the original version of the question and another one-third of the respondents were asked a slightly differently worded version of the question. For these questions, we merge the two questions and treat them as the same despite the slight change in wording. For example, for government spending on education, the original question was worded as: “We are faced with many problems in this country, none of which can be solved easily or inexpensively. I’m going to name some of these problems, and for each one I’d like you to name some of these problems, and for each one I’d like you to tell me whether you think we’re spending too much

Law enforcement and gun control: courts dealing with criminals; should marijuana be legal; approve of police striking citizens if: citizen said vulgar things; citizen attacked policemen with fists; citizen attempted to escape custody; citizen questioned as murder suspect; ever approve of police striking citizen; favor/oppose death penalty for murder; favor/oppose gun permits; have gun/pistol/rifle/shotgun at home.

Life, life outlook, and trust: should aged live with their children; afraid to walk at night in neighborhood; opinion of how people get ahead; general happiness; condition of health; people helpful or looking out for selves; any opposite race in neighborhood; if rich, continue or stop working; can people be trusted.

Marriage, sex, and abortion: approve of legal abortion if: strong change of serious defect; woman's health seriously endangered; married – wants no more children; low income – cannot afford more children; pregnant as result of rape; not married; divorce laws; happiness of marriage; homosexual sex relations; feelings about porn laws; premarital sex; extramarital sex; seen X-rated movie last year.

Politics and religion: political party affiliation; liberal vs. conservative; voted for D, R, I or other presidential candidate; voted in the election; how often attend religious services; religion & denomination; how fundamentalist; belief in life after death.^{56 57 58 59}

We use all respondents of ages 18 to 64.

The income variable available in the GSS is family income, and it is reported in income brackets. The income brackets change across years.⁶⁰

To implement the ensemble algorithm, we equalize sample size across years. For each year,

money on it, too little money, or about the right amount. Are we spending too much, too little, or about the right amount on improving the nation's education system?" The altered version use the word "education" instead of "the nation's education system."

⁵⁶When predicting political ideology, we drop three questions: Liberal vs. conservative; Political party affiliation; Voted for D, R, I or other presidential candidate.

⁵⁷For the question voted for D, R, I or other presidential candidate, we use the following questions in the GSS: PRES72, PRES80, PRES84, PRES88, PRES92, PRES96, PRES00, PRES04, PRES08, PRES12. Each of these questions ask which presidential candidate the respondent voted for in the election in year 19XX or 20XX. These questions were asked only for the four years after the election. For example, VOTE88 exists in the GSS for years 1989-1992 only.

⁵⁸For the question voted in the election, we use the following questions in the GSS: VOTE72, VOTE80, VOTE84, VOTE88, VOTE92, VOTE96, VOTE00, VOTE04, VOTE08, VOTE12. Like the PRESXX questions, each of these variables ask whether they voted in the election in year 19XX or 20XX, and were asked only for the four years after the election.

⁵⁹For the religion and denomination questions, we merged the religion question and the Christian denomination question such that we have a response for each Christian denomination and for each non-Christian religion.

⁶⁰There are 12 brackets for the period 1976 to 1976, 16 brackets for the period 1977 to 1981, 17 brackets for the period 1982 to 1985, 20 brackets for the period 1986 to 1990, 21 brackets for the period 1991 to 1996, 23 brackets for the period 1998 to 2004, and 25 brackets for the period 2006 to 2016.

we choose the sample size for each demographic across years to be equal to the sample size of the smallest demographic-year bucket. The ensemble algorithm first reads in the entire cleaned dataset and then selects a random sample that is equal to the sample size listed below. The ensemble algorithm is iterated over 500 random samples.

Income: 398 (The income-year with the smallest sample size is top quartile in 1990, which has a sample size of 199. Hence, the ensemble sample size is $199 \times 2 = 398$.)

Education: 652 (The education-year with the smallest sample size is some college or more in 1976, which has a sample size of 326. Hence, the ensemble sample size is $326 \times 2 = 652$.)

Gender: 1,000 (The gender-year with the smallest sample size is males in 1990, which has a sample size of 500. Hence, the ensemble sample size is $500 \times 2 = 1,000$.)

Race: 234 (The race-year with the smallest sample size is non-whites in 1976, which has a sample size of 234. Hence, the ensemble sample size is $117 \times 2 = 234$.)⁶¹ ⁶²

Political ideology: 566 (The political ideology-year with the smallest sample size is liberals in 2004, which has a sample size of 283. Hence, the ensemble sample size is $283 \times 2 = 566$.) ⁶³

Urbanicity: 236 (The urbanicity-year with the smallest sample size is rural in 1988, which has a sample size of 118. Hence, the ensemble sample size is $118 \times 2 = 236$.)⁶⁴

Age: 892 (The age-year with the smallest sample size is 40 years or older in 1990, which has a sample size of 446. Hence, the ensemble sample size is $446 \times 2 = 892$.)

In the GSS, for most questions, the data is missing for approximately one-third of the sample. This is because the “sociopolitical attitude and behavior questions are administered using a “split-ballot” design - in which items are assigned to two of three ballots, each of which is answered by a random two-thirds of most GSS samples” (Smith et al. 2014).⁶⁵ We impute the missing data as follows. In each demographic-year, among respondents with non-missing values for each

⁶¹In the GSS, we use the question RACE for our race specification. The responses to this question are “white”, “black”, or “other.” This question is available for all years of the GSS.

⁶²In the GSS, there is a question HISPANIC, which identifies whether or not the respondent is Hispanic and has values for detailed country of origin in the Hispanic world (for example, Mexican, Puerto Rican, Cuban, etc.). This variable is available since year 2000. We do not use this variable for our race specification.

⁶³For political ideology, the GSS question that we use is POLVIEW, which has the following responses: extremely liberal; liberal; slightly liberal; moderate; slightly conservative; conservative; and extremely conservative. We define political ideology as equal to one if the responses are extremely liberal, liberal, or slightly liberal. We define political ideology as equal to zero if the responses are slightly conservative, conservative, and extremely conservative. We drop observations with the response moderate.

⁶⁴For urbanicity, the GSS question that we use is SRCBELT, which has the following responses: 12 largest SMSA’s; 13-100 SMSA’s; suburb of 12 largest SMSA’s; suburb of 13-100 largest SMSA’s; other urban; and other rural. We define urbanicity as equal to one for all responses other than “other rural”, zero otherwise.

⁶⁵Smith, Tom W, Peter Marsden, Michael Hout, and Jibum Kim. 2014. *General Social Surveys: Cumulative Codebook*.

question, we compute the distribution of answers (for example, 40% answer “Republican,” 30% answer “Independent,” and 40% answer “Democrat” to the party affiliation question). Then, for each demographic-year, we use the distribution of answers among respondents with non-missing values to randomly impute the response for respondents with missing responses in the same proportions.⁶⁶ After imputing for missing values, we reshaped the data into dummy variables for each question-response.

American Heritage Time Use Survey (AHTUS)

We use the American Heritage Time Use Survey (AHTUS) to measure cultural distance for time use. We use all available years: 1965, 1975, 1985, 1993, 1995, 1998 and annually from 2003 to 2012. For the analysis using income, however, we dropped 1985, 1993, and 1995 because the available income data were too coarse (only approximate income quartiles are available in those years).

We equalize the set of activities across years using an activities crosswalk that is based on the official documentation published by the University of Oxford Center for Time Use Research. After equalizing the set of activities across years, we use all of the 78 available activities, as well as the 8 aggregates of activities from Aguiar and Hurst (2009).⁶⁷ We define a variable as minutes spent on the activity per day. The full list of variables is: general or other personal care; sleep; naps and rest; wash, dress, personal care; personal medical care; meals at work; other meals and snack; main paid work (not at home); paid work at home; second job, other paid work; work breaks; other time at workplace; time looking for work; regular schooling, education; homework; short course or training; occasional lectures and other education or training; food preparation, cooking; set table, wash/put away dishes; cleaning; laundry, ironing, clothing repair; home repairs, maintain vehicle; other domestic work; purchase routine goods; purchase consumer durables; purchase personal services; purchase medical services; purchase repair, laundry services; financial/government services; purchase other services; general care of older children; medical care of children; play with children; supervise/help with homework; read to/with, talk with children; other child care; adult care; general voluntary acts; political and civic activity; worship and religious acts; general out-of-home leisure; attend sporting event; go to cinema; theater, concert, opera; museums, exhibitions; café, bar, restaurant; parties or receptions; sports and exercise; walking; physical activity/sports with

⁶⁶We note that the above method of imputation uses only the marginal distribution (the distribution of each variable X by demographic group) and not the joint distribution (the joint distribution of variable X, Y, and Z by demographic group).

⁶⁷The 8 aggregates of activities are: market work; home maintenance; obtain goods and services; other home production; non-market work; child care; leisure; and other.

child; hunting, fishing, boating, hiking; gardening; pet care, walk dogs; receive or visit friends; other in-home social, games; artistic activity; crafts; hobbies; relax, think, do nothing; read books, periodicals, newspapers; listen to music; listen to radio; watch television, video; writing by hand; conversation, phone, texting; and use computer.⁶⁸

Travel: travel to or from work; travel related to education; travel related to consumption; travel related to child care; travel related to volunteering and worship; other travel.

We use full-time employed respondents of ages 18 to 64.

The income variable available in AHTUS is family income, and it is available in income brackets. The income brackets change across years.⁶⁹

To implement the ensemble algorithm, we equalize sample size across years. For each year, we choose the sample size for each demographic across years to be equal to the sample size of the smallest demographic-year bucket. The ensemble algorithm first reads in the entire cleaned dataset and then selects a random sample that is equal to the sample size listed below. The ensemble algorithm is iterated over 500 random samples.

Income: 418 (The income-year with the smallest sample size is top quartile in 1998, which has a sample size of 209. Hence, the ensemble sample size is $209 \times 2 = 418$.)

Education: 524 (The education-year with the smallest sample size is some college or more in 1965, which has a sample size of 262. Hence, the ensemble sample size is $262 \times 2 = 524$.)

Gender: 668 (The gender-year with the smallest sample size is females in 1995, which has a sample size of 334. Hence, the ensemble sample size is $334 \times 2 = 668$.)

Race: 298 (The race-year with the smallest sample size is non-whites in 1995, which has a sample size of 149. Hence, the ensemble sample size is $149 \times 2 = 298$.)⁷⁰ ⁷¹

Urbanicity: 756 (The urbanicity-year with the smallest sample size is rural in 1965, which has a sample size of 378. Hence, the ensemble sample size is $378 \times 2 = 756$.)

Age: 1,088 (The age-year with the smallest sample size is 40 years or older in 1985, which has a sample size of 544. Hence, the ensemble sample size is $544 \times 2 = 1,088$.)

⁶⁸The variable “use computer” first appears in the data in 1985. We assign 0 minutes for “use computer” for all observations prior to 1985.

⁶⁹There are 10 brackets for 1965, 18 brackets for 1975, 7 brackets for 1998, and 16 brackets for the period 2003 to 2012.

⁷⁰In AHTUS, we use the variable ETHNIC2 for our race specification. The values of this variable are “white”, “black”, “some other race”, “missing or dirty”, or “not applicable.” We drop observations that have the values “missing or dirty” or “not applicable.” We define the binary race variable as equal to 1 if the value is “white” and 0 if the value is “black” or “some other race.” This variable is available for all years of AHTUS.

⁷¹In AHTUS, there is a variable called HISP which identifies respondent’s Hispanic origin. The variable has values “Yes” or “No” for respondent’s Hispanic origin. This variable is available since year 1995. We do not use this variable for our race specification.

Mediamark Research and Intelligence Survey of the American Consumer (MRI)

We use the Mediamark Research and Intelligence Survey of the American Consumer (MRI) to measure cultural distance for media consumption and consumer behavior. We use all the years that we have access to, which is annually from 1992 to 1999 and annually from 2001 to 2016. The types of variables that we use are:

Movies: “Did you watch movie X in the last 6 months?”

Magazines: “Did you read magazine X in the last 6 months?”⁷²

TV programs: “Did you watch TV program X in the last 7 days / 30 days / 12 months?”

Products: “Do you own product X / Did you use product X / Did you buy product X in the last 30 days / 6 months / 12 months?”⁷³

Brands: “Do you own product from brand X / Did you use product from brand X / Did you buy product from brand X in the last 6 months / 12 months?”

As each question in the MRI has a yes (1) or no (0) answer, we define a variable as a dummy variable equal to 1 for a positive response, 0 otherwise.

MRI includes other variables that we did not use in the analysis. These include: attitudes (political affiliation, health⁷⁴, fashion⁷⁵, general⁷⁶, attitudes towards advertisements⁷⁷, personal attitudes⁷⁸, passionate about topic X⁷⁹), time use (political activity, pets, miles driven on a car, overnight camping trips, visited theme park X in the last year, been to country X in the last 3 years, been to state X in the last year, hours listened to the radio, hours watched TV, interests, hours per week spent on doing X, time spent using the internet, hours spent playing videogame system X/videogame type X, music type X listened to in the last 6 months, hobby X, volunteered for charitable organization, member of an organization or club, leisure activity X), other consumer behavior (shopped in store X in the last 6 months), other media consumption (newspapers⁸⁰, visited social networking site X in the last 30 days, visited website X in the last 30 days).

The number of variables for each module is 83-97 variables each year for movies, 179 to 242

⁷²We did not use magazines which do not require subscription (such as magazines of airlines and retail stores) because exposure to these types of magazines may not capture people’s preferences for reading these magazines.

⁷³We use all products except for financial and insurance products. Same for brands.

⁷⁴An example is “I go to the doctor regularly for check-ups.”

⁷⁵An example is “Comfort is one of the most important factors when selecting fashion products to purchase.”

⁷⁶An example is “Buying American products is important to me.”

⁷⁷An example is “Advertising helps me keep up-to-date about products and services that I need or would like to have.”

⁷⁸An example is “Having material possessions is important.”

⁷⁹Example topics include health care, cooking, and grocery.

⁸⁰Newspapers are not used because of the small number of newspapers included in the dataset; regional newspapers are not included in the US-level data that we have access to.

variables each year for magazines, 507 to 872 variables each year for TV programs, 1,928 to 3,027 variables each year for products, and 5,367 to 6,610 variables each year for brands. We pool movies, magazines, and TV programs together as the media module; there are 871 to 1,186 variables each year for the media module. We also pool products and brands together as the consumer module; there are 7,130 to 9,385 variables each year for the consumer module.

Not all variables are available for all years. While products, brands, and TV programs are available for all years, movies are available for 1998, 1999, and annually from 2001 to 2016. Also, we only use magazines annually from 1992 to 1999 and annually from 2001 to 2011.⁸¹ Hence, for the media module, we only use the overlapping years for movies, magazines, and TV programs, which are 1998, 1999, and annually from 2001 to 2011.

Furthermore, not all demographics are available for all years. While income, gender, and race are available for all years, education and political ideology are available annually from 1994 to 1999 and annually from 2001 to 2016. While we use all available years for education, for political ideology we only use data from 1994 to 1999, and from 2001 to 2009. This is because the share of respondents who do not respond to the political ideology question in the period 2010 to 2013 is substantially higher than in the period 1994 to 2009, while the share in the period 2014 to 2016 is substantially lower than in the period 1994 to 2009. This suggests that the quality of the political ideology question in the period 2010 to 2016 is not the same as in the period 1994 to 2009.

We sample all respondents from ages 20 to 64 instead of all respondents from ages 18 to 64 because age is only available in five-year age groups (20 to 24,..., 60 to 64).

The income variable available in MRI is household income, and it is available in income brackets. The income brackets change across years.⁸²

To implement the ensemble algorithm, we equalize sample size across years. For each year, we choose the sample size for each demographic across years to be equal to the sample size of the smallest demographic-year bucket. The ensemble algorithm first reads in the entire cleaned dataset and then selects a random sample that is equal to the sample size listed below. The ensemble algorithm is iterated over 25 random samples.

Income: 6,394 (The income-year with the smallest sample size is top quartile in 1992, which has a sample size of 3,197. Hence, the ensemble sample size is $3,197 \times 2 = 6,394$.)

⁸¹While magazine data exist in the MRI Media Survey post-2011, the time period was reduced to the last 7 days for the weekly magazines and the last 14 days for the bi-weekly magazines starting in 2012. This makes the “Did you read magazine X” variables in 2012-2016 not comparable to those prior to 2012.

⁸²There are 14 brackets for 1992 and 1993, 15 brackets for the period 1994 to 2001, 16 brackets for the period from 2002 to 2008, and 17 brackets for the period 2009 to 2016.

Education: 9,674 (The education-year with the smallest sample size is high school or less in 2015, which has a sample size of 4,837. Hence, the ensemble sample size is $4,837 \times 2 = 9,674$.)

Gender: 15,036 (The gender-year with the smallest sample size is females in 1996, which has a sample size of 7,518. Hence, the ensemble sample size is $7,518 \times 2 = 15,036$.)

Race: 4,150 (The race-year with the smallest sample size is non-whites in 1992, which has a sample size of 2,075. Hence, the ensemble sample size is $2,075 \times 2 = 4,150$.)⁸³ ⁸⁴

Political ideology: 4,864 (The political ideology-year with the smallest sample size is liberals in 2,432, which has a sample size of 2,432. Hence, the ensemble sample size is $2,432 \times 2 = 4,864$.)

Age: 14,602 (The age-year with the smallest sample size is 40 years or older in 1992, which has a sample size of 7,301. Hence, the ensemble sample size is $7,301 \times 2 = 14,602$.)

A.1.2 Ensemble Algorithm

We use a machine learning approach to determine how predictable group membership is from a set of variables in a given year. In particular, we use an ensemble method that consists in running multiple separate algorithms and then averaging the prediction of these algorithms with weights chosen by cross-validation (Mullainathan and Spiess, 2017). We use three machine learning algorithms: elastic net regression (tuned by lambda and alpha), regression tree (tuned by the minimal node size of each tree), and random forest (tuned by the minimal node size of each tree and the proportion of variables used in each tree). We “ensemble” across algorithms with weights determined by OLS. The ensemble algorithm yields a prediction (posterior probability) that the respondent is in the given group (top income quartile, some college or more, etc.) for each respondent. We define “guess” as 1 if the prediction is greater than or equal to 0.5, 0 otherwise. We report the share of correct guesses in the hold-out sample (30%). The procedure is as follows.

1. Partition the data into a training sample (70%) and a hold-out sample (30%).
2. Tuning step (general)
 - (a) Divide the training sample randomly into 5 folds.
 - (b) For each fold, fit the algorithm for every tuning parameter value on all 4 other folds and for predictions on the current fold.

⁸³In MRI, the race variable has the following values for the listed years: 1992-1997 - “White,” “African American,” or “Other;” 1998-2002 - “White,” “African American,” “Asian,” or “Other;” 2003-2016 - “White,” “African American,” “American Indian or Alaska Native,” “Asian,” or “Other.”

⁸⁴In MRI, there is a variable that identifies whether the respondent is of Hispanic origin. This variable is available since year 2007. We do not use this variable for our race specification.

- (c) From 2(b), obtain one prediction per tuning parameter for every observation in the full training sample. Now, average the squared-error loss for each tuning parameter value over the full training sample.
 - (d) Based on loss estimates in 2(c), choose the tuning parameters that minimize the squared-error loss.
 - (e) Fit the algorithm with the chosen tuning parameter on the full training sample.
 - (f) Repeat steps 2(b)-2(e) for each algorithm (elastic net regression, regression tree, random forest).
3. Tuning parameters (specific to each algorithm)
- (a) Elastic net regression
 - i. In 2(c), elastic net regression is fit for a grid of values of lambda and alpha for the following objective function: $\min_{\beta_0, \beta} \frac{1}{N} \sum_{i=1}^N w_i l(y_i, \beta_0 + \beta^T x_i) + \lambda[(1 - \alpha)\|\beta\|_2^2 + \alpha\|\beta\|_1]$
 - A. Lambda ranges from e^{-8} to e^{10} , in increments of 0.5 for the exponent (i.e. -8, -7.5, ..., 9.5, 10). Lambda controls the penalty on the coefficients. As lambda grows larger, the penalty grows stronger, and coefficients are forced closer to zero.
 - B. Alpha grid is 0, 0.5, and 1. $\alpha = 1$ case is LASSO, $\alpha = 0$ case is the ridge regression, and $\alpha = 0.5$ is the intermediate case. Alpha specifies the type of penalty applies to the coefficients. When $\alpha = 1$ (LASSO), coefficients are penalized based on the sum of their absolute values (L1 penalty). When $\alpha = 0$ (ridge regression), coefficients are penalized based on the sum of their squared values (L2 penalty). When alpha is between 0 and 1, the coefficients are penalized based on both L1 and L2 penalties, and the weights are determined by alpha.
 - (b) Regression tree
 - i. In 2(c), regression tree is fit for a grid of values of minimum node size (“minbucket”), where node size is the number of observations belonging to a terminal node. The grid for node size is (1, 5, 10, 20, 30, 40, 50, 70, 100, 150, 500). The depth of the tree is determined by the node size: the smaller the node size, the deeper the tree.
 - (c) Random forest

- i. In 2(c), random forest is fit for a grid of values of 1) minimum node size of each tree (“node sizes”) and 2) the proportion of variables used in each tree (“pmtry”). The number of trees is set to 100. The grid for node sizes is (5, 10, 20, 50, 100, 200, 400, 1000) and the grid for pmtrys is (0.1, 0.2, 0.3, 0.4).

4. Ensemble step

- (a) From 2, we have obtained one prediction for each algorithm for every observation in the full training sample.
- (b) Fit weights by running a linear regression (OLS) of the outcome on the predicted values for each algorithm in the full training sample, and store the resulting linear model.
- (c) To predict in the hold-out sample, fit each algorithm on the full training sample, obtain predictions for each algorithm on the hold-out sample, and then ensemble the predictions with the linear model obtained in 4(b).

5. Ensemble algorithm implementation

- (a) For each dataset-year, implement the ensemble algorithm where:
 - i. LHS = Income / Education / Gender / Race / Political Ideology / Urbanicity / Age (dummy variables)
 - ii. RHS = Dataset
- (b) Iterate the ensemble algorithm for X number of random subset of the dataset (X=500 for attitudes and time use, X=25 for media, movies, TV programs, magazines, consumer behavior, products, and brands).
- (c) For each iteration, compute the hold-out sample share of correct guesses.
 - i. The ensemble algorithm outputs the predictability that a respondent is in the income / demographic group for each year.
 - ii. We guess whether the respondent is in that income / demographic group if the predictability is greater than or equal to / less than 0.5.
 - iii. Then, for each respondent, we have the true income / demographic of the respondent and our guess using the RHS variables. We compute the hold-out sample share of correct guesses.

- iv. The ensemble algorithm uses 70% of the data to generate a prediction model (training sample), and designates the remaining 30% as the hold-out sample. We only use the hold-out sample to compute the share of correct guesses.
- (d) For each dataset-year, average the hold-out sample share of correct guesses across the iterations.

A.1.3 Bayesian Algorithm

We use a Bayesian approach to determine how predictable group membership is from a single variable in a given year. We use the results from the Bayesian approach to produce a) the table of top 10 cultural traits that are most indicative of membership in a demographic group and b) the heat map of cultural traits that are indicative of membership in a demographic group (for attitudes only). The procedure is as follows:

1. Partition the data into a training sample (80%) and a hold-out sample (20%).
2. In the training sample, calculate the probability of any response (e.g. watched *Fox and Friends*) conditional on the respondents' membership in a demographic group.
3. Turning to the hold-out sample, guess whether a respondent is in a demographic group conditional on his or her response given the conditional probabilities derived in the training sample.
4. Compute the hold-out sample share of correct guesses.
5. Repeat steps (1) to (4) 100 times for time use and attitudes, and 5 times for media and consumer behavior. Obtain the average hold-out sample share of correct guesses across iterations.
6. Using the full sample, take the average probability of any response conditional on the respondents' membership in a demographic group and store it. (This is needed to record the direction of guess.)

The procedure for producing the tables of the top ten cultural traits that are most indicative of group membership is then as follows. First, we rank each response in decreasing order the average hold-out sample share of correct guesses obtained by the Bayesian procedure. Second, we report the average hold-out sample share of correct guesses for the ten responses with the highest

share of correct guesses. Third, we use the average probability of any response conditional on the respondents' group membership to know the direction of the prediction (e.g., watching *Fox and Friends* is predictive of being conservative).

The procedure for producing the heat map of cultural traits that are indicative of group membership (for attitudes only) is as follows. First, we rank each variable in increasing order of the average hold-out sample share of correct guesses obtained by the Bayesian procedure for the first year (1976 for attitudes). Variables are vertically ranked throughout the heat map figure based on that 1976 order. Second, in each subsequent year, we assign to each variable its rank in increasing order of the average hold-out sample share of correct guesses for that year. We then assign color-code to each variable's relative rank in each year, with the most informative variables being color-coded dark red and the least informative color-coded dark blue, and lighter shades of red and blue in between.

A.1.4 Defining income quartile cutoffs by household groups using the Current Population Survey (CPS)

We use family income for the GSS and AHTUS and household income for the MRI. Note that the income variables in all three of our main datasets are in income brackets, not continuous dollar amounts. As the CPS top / bottom income quartile cutoffs by household groups most often occur within an income bracket, using income brackets does not exactly capture the top / bottom income quartiles in the CPS. We describe below the method we use to minimize this mismeasurement.

First, we define household groups as follows. We define the households with one adult and no children as household group 1, households with two adults and no children as household group 2, households with two adults and children as household group 3, and households with one adult and children. Households with more than two adults were classified into household group 3; adults other than the two primary adults are regarded as dependents.

The procedure for defining the income quartile dummy variable is as follows. For every year-household group, we obtain from the CPS the top and bottom quartile income cutoffs as well the full income distribution. For each of the three datasets (GSS, AHTUS, MRI), we then consider all possible assignments of observations to top and bottom quartiles based on the income brackets available in that dataset-year. For each possible assignment, we count the number of observations that actually are in top / bottom quartile according to the CPS but not assigned as such, as well as the number of observations that actually are not top / bottom quartile according to the CPS but

assigned as such. We call the sum of these two numbers the number of mis-measured observations. For each dataset-year-household group, we then generate the top and bottom quartile variables by choosing the assignment that minimizes the number of mis-measured observations.

The share of mis-measured observations, when averaged across household groups (with weights corresponding to the number of observations in each household group), are summarized below.

1. Top quartile:

- (a) GSS: average - 2.7%, minimum - 1.3%, maximum - 5.1%
- (b) AHTUS: average - 5.0%, minimum - 0.6%, maximum - 8.5%
- (c) MRI: average - 4.0%, minimum - 1.6%, maximum - 6.9%

2. Bottom quartile

- (a) GSS: average - 1.5%, minimum - 0.7%, maximum - 3.6%
- (b) AHTUS: average - 4.0%, minimum - 2.2%, maximum - 7.3%
- (c) MRI: average - 1.2%, minimum - 0.4%, maximum - 2.1%

While the share of mis-measured observations is less than 5% for most dataset-quartiles, the share is larger than 5% (and thus not negligible) for: MRI for years 2007-2013 for the top quartile; AHTUS for years 1965, 1998, and 2006-2012 for the top quartile; and AHTUS for year 1998 for the bottom quartile. To investigate the effect of mismeasurement on our ability to predict, we regress the average hold-out sample share of correct guesses on an intercept, average share of mismeasurement for the top and bottom quartiles, year dummies, and dataset dummies. First, we find that the coefficient on the average share of mismeasurement is not statistically significant (coefficient = -0.21, t-statistic = -0.39). Second, we find that the R-squared increases only minimally when we include the average share of mismeasurement; in fact, the adjusted R-squared decreased. From these two observations, we conclude that while the level of mismeasurement is not negligible, its effect on our ability to predict does not appear to be substantive.

A.2 Main Additional Results

A.2.1 Income

Table A.1: Attitudes and norms most indicative of being high-income

1976		1996		2016	
Trusts people	65.1%	Voted in the election	64.4%	Voted in the election	65.2%
Voted in the election	65.0%	Trusts people	62.0%	Approve of police striking citizens	63.4%
Allow homosexuals' books in library	64.4%	Allow atheists to teach	61.1%	Trusts people	62.2%
Allow communists to speak	63.0%	Allow abortion for single women	60.3%	Allow abortion for single women	61.8%
Spending on space exploration is adequate	62.8%	Allow racists to speak	60.1%	Not afraid to walk at night	61.5%
Homosexual sex is not always wrong	62.4%	Federal income tax is too high	59.7%	Voted for Republican pres. candidate	61.2%
Allow homosexuals to speak	62.3%	Confident in the scientific community	59.7%	Allow abortion for married women	61.2%
Voted for Republican pres. candidate	62.2%	Allow abortion for married women	59.5%	Homosexual sex is not wrong at all	60.6%
Allow communists' books in library	62.0%	Allow racists' books in library	59.1%	Allow abortion for low-income women	60.6%
Allow militarists' books in library	61.7%	Allow atheists' books in library	59.0%	Allow militarists' books in library	59.4%

Note: Data source is the GSS. Sample size is 398. Reported in each column are the 10 cultural traits most indicative of being rich in that year. The numbers indicate the likelihood of guessing correctly whether an individual is rich based on the answer to the question. For example, in 1976, knowing whether a person trusts people allows us to guess income correctly 65.1% of the time, whereas knowing whether a person thinks homosexual sex is not always wrong allows us to guess income correctly 62.4% of the time. An affirmative answer to “Do you trust people?” and a negative answer to “Is homosexual sex always wrong?” indicate that the person is rich.

A.2.2 Education

Table A.2: TV shows, movies, and magazines most indicative of being more educated

Panel (a) TV shows					
1994		2005		2016	
Didn't watch <i>Rescue 911</i>	56.0%	Watched <i>Super Bowl</i>	53.9%	Watched <i>Love It Or List It</i>	53.6%
Didn't watch <i>Unresolved Mysteries</i>	54.6%	Didn't watch <i>Cops</i>	53.8%	Watched <i>Property Brothers</i>	53.2%
Watched <i>Wimbledon</i>	54.2%	Watched <i>Academy Awards</i>	53.1%	Watched <i>House Hunters</i>	53.2%
Didn't watch <i>Oprah Winfrey Show</i>	54.1%	Watched <i>NFL Monday Night Football</i>	52.6%	Watched <i>Academy Awards</i>	53.0%
Watched <i>NCAA Basketball Championship</i>	54.0%	Watched <i>NCAA Men's Basketball</i>	52.6%	Watched <i>Flip or Flop</i>	52.7%
Didn't watch <i>In the Heat of the Night</i>	53.9%	Didn't watch <i>WWE Smackdown!</i>	52.6%	Didn't watch <i>Criminal Minds</i>	52.3%
Didn't watch <i>Country Music Awards</i>	53.9%	Didn't watch <i>Noticiero Univision</i>	52.5%	Watched <i>Grammy Awards</i>	52.1%
Watched <i>Superbowl</i>	53.8%	Didn't watch <i>NASCAR Daytona 500</i>	52.5%	Watched <i>NCAA's Final Four</i>	52.1%
Didn't watch <i>America's Most Wanted</i>	53.6%	Watched <i>The Masters</i>	52.4%	Watched <i>SNL Specials</i>	52.0%
Didn't watch <i>Married with Children</i>	53.3%	Didn't watch <i>Fear Factor</i>	52.4%	Watched <i>Wimbledon</i>	52.0%
Panel (b) Movies					
1998		2007		2016	
Watched <i>Jerry Maguire</i>	54.6%	Didn't watch <i>Big Momma's House 2</i>	52.6%	Watched <i>Gone Girl</i>	53.6%
Watched <i>The English Patient</i>	53.3%	Watched <i>Walk The Line</i>	52.4%	Watched <i>The Hunger Games</i>	53.0%
Watched <i>First Wive's Club</i>	52.7%	Watched <i>The Chronicles Of Narnia</i>	52.4%	Watched <i>Interstellar</i>	52.5%
Watched <i>Star Trek First Contact</i>	52.1%	Watched <i>Pirates of the Caribbean 2</i>	52.0%	Watched <i>Guardians of the Galaxy</i>	51.8%
Watched <i>The Empire Strikes Back</i>	52.0%	Watched <i>The Da Vinci Code</i>	51.9%	Watched <i>Into the Woods</i>	51.7%
Watched <i>Star Wars - Special Edition</i>	51.9%	Watched <i>The Devil Wears Prada</i>	51.9%	Watched <i>Big Hero 6</i>	51.6%
Watched <i>Air Force One</i>	51.8%	Didn't watch <i>Saw II</i>	51.9%	Watched <i>Birdman</i>	51.6%
Watched <i>Michael</i>	51.6%	Didn't watch <i>Scary Movie 4</i>	51.8%	Watched <i>The Theory of Everything</i>	51.5%
Watched <i>Ransom</i>	51.5%	Didn't watch <i>When a Stranger Calls</i>	51.8%	Didn't watch <i>Annabelle</i>	51.4%
Watched <i>One Fine Day</i>	51.5%	Didn't watch <i>Get Rich or Die Tryin'</i>	51.7%	Watched <i>The Hobbit</i>	51.3%
Panel (c) Magazines					
1992		2002		2011	
Read <i>Newsweek</i>	60.3%	Read <i>Time</i>	58.7%	Read <i>Time</i>	57.8%
Read <i>Time</i>	59.1%	Read <i>Newsweek</i>	58.5%	Read <i>Newsweek</i>	57.4%
Read <i>US News & World Report</i>	58.6%	Read <i>People</i>	56.7%	Read <i>Consumer Reports</i>	57.0%
Read <i>Consumer Reports</i>	58.0%	Read <i>US News & World Report</i>	55.6%	Read <i>People</i>	56.2%
Read <i>National Geographic</i>	57.1%	Read <i>Consumer Reports</i>	55.5%	Read <i>National Geographic</i>	55.5%
Read <i>Business Week</i>	56.7%	Read <i>National Geographic</i>	55.0%	Read <i>The New Yorker</i>	55.0%
Read <i>Money</i>	56.4%	Read <i>Business Week</i>	54.1%	Read <i>Forbes</i>	54.8%
Read <i>The New Yorker</i>	55.6%	Read <i>Fortune</i>	53.5%	Read <i>Real Simple</i>	54.8%
Read <i>Forbes</i>	55.6%	Read <i>The New Yorker</i>	53.5%	Read <i>O, The Oprah Magazine</i>	54.8%
Read <i>Smithsonian</i>	55.6%	Didn't read <i>National Enquirer</i>	53.3%	Read <i>Travel & Leisure</i>	54.4%

Note: Data source is the MRI. Sample size in all panels is 9,674. Reported in each column are the 10 cultural traits most indicative of being educated in that year. The numbers indicate the likelihood of guessing correctly whether an individual is educated based on the answer to the question. For example, in 1992, knowing whether a person watched *Wimbledon* allows us to guess education correctly 54.2% of the time, whereas knowing whether a person watched *Rescue 911* allows us to guess education correctly 56.0% of the time. An affirmative answer to “Did you watch *Wimbledon*?” and a negative answer to “Did you watch *Rescue 911*?” indicate that the person is educated.

Table A.3: Products and brands most indicative of being more educated

Panel (a) Products					
1994		2005		2016	
Own an imported car	59.6%	Own a PC software	63.1%	Own a tablet PC	63.8%
Traveled in the continental US	59.3%	Own a personal computer	63.0%	Ordered an item by Internet	63.2%
Own a personal computer	58.9%	Own a PC peripheral device	62.3%	Traveled in the continental US	62.8%
Own a PC peripheral device	58.5%	Ordered an item by Internet	61.9%	Own a PC software	62.3%
Traveled domestically by plane	58.5%	Own a desktop	61.1%	Own a passport	62.1%
Own a PC software	58.2%	Traveled in the continental US	60.9%	Own a laptop	61.8%
Used dishwasher detergent	58.2%	Own a word processing software	60.8%	Own a personal computer	61.4%
Own a passport	58.2%	Own a passport	60.8%	Own a PC peripheral device	61.0%
Own an answering machine	58.2%	Own a CD-ROM drive	60.1%	Own a printer	60.6%
Drank wine	57.8%	Own an Inkjet printer	60.0%	Used dishwasher detergent	60.4%

Panel (b) Brands					
1994		2005		2016	
Bought Kodak (film)	55.7%	Own MS Windows XP (OS)	58.5%	Own an Iphone	62.4%
Bought AT&T (calling card)	55.4%	Own a Dell (personal computer)	56.7%	Own an Ipad	60.6%
Used Grey Poupon Dijon (mustard)	55.1%	Used Kikkoman (soy sauce)	55.4%	Used Verizon Wireless	55.9%
Used Kikkoman (soy sauce)	54.9%	Bought at Starbucks (fast food)	55.1%	Used AT&T	55.4%
Used Sony (CD player)	54.1%	Didn't drink Pepsi (regular cola)	54.4%	Own Amazon Kindle	55.3%
Used Sony (TV set)	53.8%	Used Bertolli (salad/cooking oil)	54.2%	Used Netflix	55.0%
Used Philadelphia (cream cheese)	53.7%	Used Kleenex regular (facial tissue)	54.1%	Own HP (printer/fax machine)	54.9%
Used Fuji (film)	53.6%	Bought at Olive Garden (family rest.)	54.0%	Bought at Starbucks (fast food)	54.8%
Used Cascade - Lemon (dish. detergent)	53.5%	Own Sony (CD player)	53.9%	Bought at Chipotle (fast food)	54.6%
Used Ben & Jerry's (ice cream)	53.4%	Used Ziploc (plastic bag)	53.9%	Didn't use Country Crock (butter)	54.4%

Note: Data source is the MRI. Sample size in all panels is 9,674. Reported in each column are the 10 cultural traits most indicative of being educated in that year. The numbers indicate the likelihood of guessing correctly whether an individual is educated based on the answer to the question. For example, in 1994, knowing whether a person owns an imported car allows us to guess education correctly 59.6% of the time, whereas in 2005, knowing whether a person bought Pepsi regular cola allows us to guess education correctly 54.4% of the time. An affirmative answer to “Do you own an imported car?” and a negative answer to “Did you buy Pepsi regular cola?” indicate that the person is educated.

Table A.4: Attitudes and norms most indicative of being more educated

1976		1996		2016	
Allow communists to speak	65.1%	Voted in the election	64.3%	Voted in the election	62.7%
Allow atheists to teach	64.6%	Allow communists to speak	62.1%	Trusts people	61.9%
Allow militarists to speak	63.8%	Allow militarists to speak	61.8%	Allow communists to speak	60.9%
Allow communists' books in library	63.4%	Allow communists' book in library	61.0%	People are helpful	60.5%
Homosexual sex is not always wrong	63.4%	Confident in the scientific community	60.6%	Allow communists to teach	60.2%
Allow communists to teach	62.6%	Trusts people	60.2%	Approve of police striking citizens	59.9%
Allow atheists to speak	62.5%	Allow atheists to teach	59.6%	Allow communists' book in library	59.7%
Allow homosexuals' book in library	62.5%	Allow communists to teach	59.1%	Allow abortion for single women	59.3%
Allow militarists' book in library	62.4%	Allow abortion for single women	59.1%	Homosexual sex is not wrong at all	59.1%
Allow racists to speak	62.2%	Allow racists to speak	58.9%	Allow militarists to speak	58.7%

Note: Data source is the GSS. Sample size is 652. Reported in each column are the 10 cultural traits most indicative of being educated in that year. The numbers indicate the likelihood of guessing correctly whether an individual is educated based on the answer to the question. For example, in 1976, knowing whether a person thinks communists should be allowed to speak allows us to guess education correctly 65.1% of the time, whereas knowing whether a person thinks homosexual sex is not always wrong allows us to guess education correctly 63.4% of the time. An affirmative answer to “Should communists be allowed to speak?” and a negative answer to “Is homosexual sex always wrong?” indicate that the person is educated.

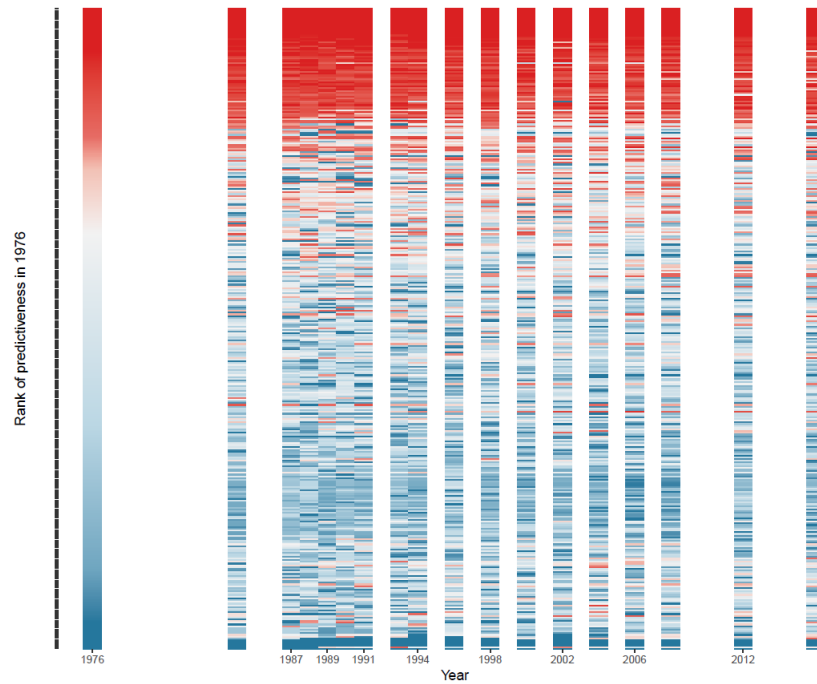


Figure A.1: Stability over time of attitudes most indicative of education

Note: Data source is the GSS. Sample size is 652. Variables are ranked from bottom to top throughout the graph by increasing order of correctly guessing education in 1976 based on that variable only. Each variable's relative informativeness in subsequent years is color-coded, with the most informative variables in each year color-coded dark red and the least informative color-coded dark blue, and lighter shades of red and blue in between. See Data Appendix for implementation details.

A.2.3 Gender

Table A.5: TV shows, movies, and magazines most indicative of being male

Panel (a) TV shows					
1992		2004		2016	
Watched <i>CBS NFL Football Playoffs</i>	62.5%	Watched <i>NFL Regular Season</i>	61.3%	Watched <i>Super Bowl</i>	58.2%
Watched <i>NBC NFL Football Playoffs</i>	62.5%	Watched <i>Super Bowl</i>	61.3%	Watched <i>NFL Live</i>	55.7%
Watched <i>Superbowl</i>	58.3%	Watched <i>NFL Monday Night Football</i>	60.8%	Watched <i>NCAA Men's Final Four</i>	55.1%
Watched <i>Rose Bowl</i>	57.0%	Watched <i>NFL Regular Season Football</i>	60.3%	Watched <i>Sports Center</i>	54.8%
Watched <i>All-Star Basketball Game</i>	55.8%	Watched <i>College Football Regular Season</i>	58.4%	Didn't watch <i>Love It or List It</i>	54.6%
Watched <i>NCAA Men's Championship</i>	55.7%	Watched <i>Super Bowl Pre-game Show</i>	58.1%	Watched <i>College Football Playoffs</i>	54.6%
Watched <i>Pro Football Playoffs</i>	55.6%	Watched <i>College Football</i>	57.5%	Watched <i>Super Bowl Pre-game Show</i>	54.6%
Didn't watch <i>Barbara Walters</i>	55.1%	Watched <i>Rose Bowl</i>	57.4%	Didn't watch <i>House Hunters</i>	54.5%
Watched <i>World League of American Football</i>	55.0%	Watched <i>Fiesta Bowl</i>	57.3%	Watched <i>NCAA Men's Championship</i>	54.2%
Watched <i>Pro Bowl</i>	54.7%	Watched <i>Super Bowl Post-game Show</i>	57.3%	Didn't watch <i>Property Brothers</i>	54.2%
Panel (b) Movies					
1998		2007		2016	
Didn't watch <i>First Wives Club</i>	56.0%	Watched <i>King Kong</i>	52.4%	Watched <i>John Wick</i>	52.7%
Didn't watch <i>The Mirror Has Two Faces</i>	54.0%	Watched <i>Transporter 2</i>	52.3%	Watched <i>Interstellar</i>	52.6%
Didn't watch <i>The Preacher's Wife</i>	53.7%	Didn't watch <i>In Her Shoes</i>	52.3%	Watched <i>Fury</i>	52.2%
Didn't watch <i>Dalmatians</i>	53.6%	Watched <i>Underworld: Evolution</i>	52.1%	Didn't watch <i>Gone Girl</i>	51.9%
Didn't watch <i>One Fine Day</i>	53.4%	Watched <i>X-Men: The Last Stand</i>	51.9%	Didn't watch <i>Annie</i>	51.9%
Didn't watch <i>My Best Friend's Wedding</i>	53.0%	Watched <i>The Legend of Zero</i>	51.8%	Watched <i>The Hobbit</i>	51.8%
Didn't watch <i>Jerry Maguire</i>	53.0%	Watched <i>A History of Violence</i>	51.7%	Watched <i>Guardians of the Galaxy</i>	51.8%
Didn't watch <i>Fly Away Home</i>	53.0%	Watched <i>Firewall</i>	51.7%	Watched <i>The Equalizer</i>	51.7%
Didn't watch <i>The English Patient</i>	52.0%	Watched <i>Mission Impossible 3</i>	51.6%	Didn't watch <i>Into the Woods</i>	51.7%
Didn't watch <i>Michael</i>	52.0%	Didn't watch <i>The Family Stone</i>	51.6%	Watched <i>Mad Max</i>	51.5%
Panel (c) Magazines					
1992		2002		2011	
Didn't read <i>Family Circle</i>	67.6%	Didn't read <i>Woman's Day</i>	65.4%	Didn't read <i>Better Homes & Gardens</i>	65.5%
Didn't read <i>Woman's Day</i>	67.6%	Didn't read <i>Better Homes & Gardens</i>	64.6%	Didn't read <i>Woman's Day</i>	64.8%
Didn't read <i>Good Housekeeping</i>	66.6%	Didn't read <i>Good Housekeeping</i>	64.6%	Didn't read <i>Good Housekeeping</i>	63.8%
Didn't read <i>Ladies' Home Journal</i>	64.6%	Didn't read <i>Family Circle</i>	64.2%	Didn't read <i>Family Circle</i>	62.6%
Didn't read <i>Better Homes & Gardens</i>	64.1%	Read <i>Sports Illustrated</i>	62.3%	Read <i>Sports Illustrated</i>	62.4%
Didn't read <i>McCall's</i>	63.1%	Didn't read <i>Ladies' Home Journal</i>	61.7%	Didn't read <i>People</i>	62.4%
Read <i>Sports Illustrated</i>	62.2%	Didn't read <i>Glamour</i>	60.7%	Didn't read <i>O, The Oprah Magazine</i>	62.3%
Didn't read <i>Redbook</i>	61.6%	Didn't read <i>Martha Stewart Living</i>	60.6%	Didn't read <i>Glamour</i>	60.6%
Didn't read <i>Glamour</i>	59.1%	Didn't read <i>Cosmopolitan</i>	60.3%	Didn't read <i>Martha Stewart Living</i>	60.1%
Didn't read <i>Cosmopolitan</i>	58.7%	Didn't read <i>People</i>	59.3%	Didn't read <i>Ladies' Home Journal</i>	60.0%

Note: Data source is the MRI. Sample size in all panels is 15,036. Reported in each column are the 10 cultural traits most indicative of being male in that year. The numbers indicate the likelihood of guessing correctly whether an individual is male based on the answer to the question. For example, in 1992, knowing whether a person watched *CBS NFL Football Playoffs* allows us to guess gender correctly 62.5% of the time, whereas knowing whether a person watched *Barbara Walters* allows us to guess gender correctly 55.1% of the time. An affirmative answer to “Did you watch *CBS NFL Football Playoffs*?” and a negative answer to “Did you watch *Barbara Walters*?” indicate that the person is male.

Table A.6: Products and brands most indicative of being male

Panel (a) Products					
1992		2004		2016	
Didn't use perfume/cologne for women	90.5%	Didn't use lipstick/lip gloss	88.5%	Didn't use hair care products for women	88.4%
Didn't use lipstick/lip gloss	89.6%	Didn't use perfume/cologne for women	87.4%	Didn't use perfume/cologne for women	85.0%
Didn't use hair care products for women	87.5%	Didn't use hair care products for women	87.0%	Didn't use mascara	83.4%
Didn't use blusher	86.7%	Didn't use facial moisturizers	84.6%	Didn't use lipstick/lip gloss	83.4%
Used aftershave lotion/cologne for men	84.3%	Didn't buy women's clothing	82.6%	Didn't buy women's clothing	83.1%
Didn't use mascara	83.9%	Didn't use mascara	82.0%	Didn't buy women's lingerie/undergarments	81.8%
Didn't buy stockings/pantyhose	82.6%	Didn't use foundation/make-up	80.3%	Didn't use foundation/make-up	80.9%
Didn't use foundation/make-up	82.4%	Did not use blusher	78.1%	Didn't buy eye liner	79.6%
Didn't use face cream/lotion	82.2%	Did not use eye shadow	77.9%	Didn't use eye shadow	77.2%
Didn't use eye shadow	81.5%	Used aftershave lotion/cologne for men	77.2%	Didn't use nail care products/polish	74.7%

Panel (b) Brands					
1992		2004		2016	
Didn't buy L'eggs (stockings)	63.3%	Didn't use Cutex (nail polish remover)	62.6%	Didn't buy Victoria Secret (lingerie)	60.8%
Didn't buy No Nonsense (stocking)	57.3%	Didn't use Lady BIC (disposable razor)	59.0%	Didn't use Bath & Body Works (perfume)	59.0%
Own a Subaru (truck/van/SUV)	56.4%	Didn't use Bath & Body Works (h/b cream)	58.0%	Didn't use Cutex (nail polish remover)	58.2%
Didn't buy Hanes Silk (stockings)	53.8%	Didn't use Bath & Body Works (bath add.)	57.1%	Didn't use Bath & Body Works (h/b cream)	57.4%
Didn't use Philadelphia (cr. cheese)	53.0%	Used Philips Norelco (electric shaver)	56.7%	Didn't use Opi (nail care products)	57.3%
Didn't buy Kodak (film)	53.0%	Didn't use Tampax (tampon)	56.5%	Didn't use Secret invisible (deodorant)	56.9%
Didn't use Murphy's oil soap	52.7%	Didn't use Skintimate (shave gel)	56.1%	Didn't buy Hanes (lingerie)	56.7%
Didn't use Playtex (rubber gloves)	52.7%	Didn't use Maybelline (mascara)	55.9%	Didn't use Dove solid (deodorant)	56.5%
Didn't use Gold Medal (flour)	52.7%	Used Gillette (razor blades)	55.9%	Didn't use Victoria Secret (perfume)	56.4%
Didn't use Heinz (vinegar)	52.6%	Didn't use Bath & Body Works (body wash)	55.9%	Used Degree Men solid (deodorant)	56.2%

Note: Data source is the MRI. Sample size in all panels is 15,036. Reported in each column are the 10 cultural traits most indicative of being male in that year. The numbers indicate the likelihood of guessing correctly whether an individual is male based on the answer to the question. For example, in 1992, knowing whether a person bought aftershave lotion/cologne for men allows us to guess gender correctly 84.3% of the time, whereas knowing whether a person bought perfume/cologne for women allows us to guess gender correctly 90.5% of the time. An affirmative answer to “Did you buy aftershave lotion/cologne for men?” and a negative answer to “Did you buy perfume/cologne for women?” indicate that the person is male.

Table A.7: Attitudes and norms most indicative of being male

1976		1996		2016	
Not afraid to walk at night in neigh.	69.3%	Not afraid to walk at night in neigh.	64.9%	Seen x-rated movie in last year	61.8%
Spending on space expl. is adequate	59.9%	Own gun in home	60.4%	Not afraid to walk at night in neigh.	60.2%
Seen x-rated movie in last year	58.3%	Porn should not be illegal to all	59.1%	Porn should not be illegal to all	57.9%
Favor gun permits	57.8%	Own a pistol/revolver in home	58.0%	Spending on space expl. is too little	57.8%
Porn should not be illegal to all	57.7%	Approve of police striking citizen	57.6%	Not confident in banks/fin. institutions	57.0%
Favor death penalty for murder	57.4%	Own shotgun in home	57.5%	Trusts people	57.5%
Spending on defense is too little	57.4%	Own rifle in home	57.1%	Approve of police striking citizens	55.8%
Not moderate (political view)	55.8%	Favor gun permits	57.1%	Spending on health care is adequate	55.6%
Not a Democrat	55.4%	Spending on space expl. is adequate	56.8%	Own pistol/revolver in home	55.4%
Not allow atheists to teach	55.4%	Seen x-rated movie in last year	56.0%	Favor gun permits	55.4%

Note: Data source is the GSS. Sample size is 1,000. Reported in each column are the 10 cultural traits most indicative of being male in that year. The numbers indicate the likelihood of guessing correctly whether an individual is male based on the answer to the question. For example, in 1976, knowing whether a person is afraid to walk at night in the neighborhood allows us to guess gender correctly 69.3% of the time, whereas knowing whether a person thinks porn should be illegal to all allows us to guess gender correctly 57.7% of the time. An affirmative answer to “Are you afraid to walk at night in the neighborhood?” and a negative answer to “Should porn be illegal to all?” indicate that the person is male.

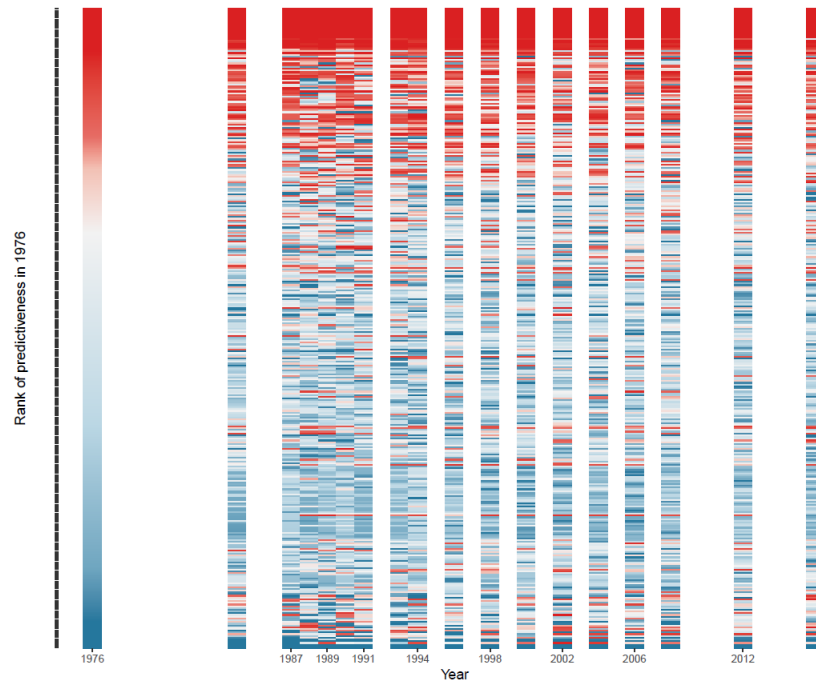


Figure A.2: Stability over time of attitudes most indicative of gender

Note: Data source is the GSS. Sample size is 1,000. Variables are ranked from bottom to top throughout the graph by increasing order of correctly guessing gender in 1976 based on that variable only. Each variable’s relative informativeness in subsequent years is color-coded, with the most informative variables in each year color-coded dark red and the least informative color-coded dark blue, and lighter shades of red and blue in between. See Data Appendix for implementation details.

A.2.4 Race

Table A.8: TV shows, movies, and magazines most indicative of being white

Panel (a) TV shows					
1992		2004		2016	
Didn't watch <i>Arsenio Hall</i>	58.5%	Watched <i>Super Bowl</i>	56.7%	Watched <i>Rudolph the Red-Nosed Reindeer</i>	55.8%
Didn't watch <i>In Living Color</i>	57.2%	Didn't watch <i>The Parkers</i>	55.7%	Watched <i>Macy's Thanksgiving Day Parade</i>	55.6%
Didn't watch <i>Cosby Show</i>	57.1%	Didn't watch <i>NBA Regular Season Games</i>	55.0%	Watched <i>American Pickers</i>	55.5%
Didn't watch <i>A Different World</i>	56.7%	Didn't watch <i>Soul Train Music Awards</i>	54.7%	Watched <i>The Big Bang Theory</i>	54.8%
Watched <i>National Geographic Specials</i>	55.5%	Watched <i>Dick Clark's New Years Rockin'</i>	54.3%	Watched <i>How the Grinch Stole Christmas</i>	54.6%
Didn't watch <i>Motown 30th Anniversary</i>	55.3%	Watched <i>NASCAR Daytona 500</i>	54.3%	Watched <i>SNL Specials</i>	54.6%
Watched <i>Disney Specials</i>	54.6%	Watched <i>Macy's Thanksgiving Day Parade</i>	54.2%	Watched <i>Dick Clark's New Year Rockin'</i>	54.3%
Watched <i>Tournament of Roses Parade</i>	54.3%	Didn't watch <i>Essence Awards</i>	54.2%	Watched <i>Charlie Brown Specials</i>	54.2%
Didn't watch <i>NAACP Image Awards</i>	54.2%	Didn't watch <i>Girlfriends</i>	54.1%	Didn't watch <i>NBA All Star Game</i>	54.2%
Watched <i>Barbara Walters</i>	54.0%	Watched <i>NFL Monday Night Football</i>	54.1%	Watched <i>Kentucky Derby</i>	54.1%
Panel (b) Movies					
1998		2007		2016	
Didn't watch <i>The Preacher's Wife</i>	55.1%	Watched <i>Walk The Line</i>	55.6%	Watched <i>Gone Girl</i>	53.3%
Watched <i>Jerry Maguire</i>	54.8%	Didn't watch <i>Big Momma's House 2</i>	55.5%	Watched <i>The Hunger Games</i>	53.1%
Watched <i>Michael</i>	54.7%	Didn't watch <i>Tyler Perry's Madea's Reunion</i>	53.8%	Watched <i>No Good Deed</i>	52.7%
Watched <i>The English Patient</i>	53.8%	Didn't watch <i>Saw II</i>	53.4%	Watched <i>Inside Out</i>	52.2%
Watched <i>First Wive's Club</i>	53.6%	Didn't watch <i>Final Destination 3</i>	53.2%	Watched <i>St. Vincent</i>	52.0%
Didn't watch <i>Space Jam</i>	53.2%	Didn't watch <i>Transporter 2</i>	53.1%	Watched <i>Birdman</i>	52.0%
Watched <i>That Thing You Do!</i>	52.6%	Didn't watch <i>Get Rich or Die Tryin'</i>	53.0%	Watched <i>Interstellar</i>	51.9%
Didn't watch <i>How to Be a Player</i>	52.5%	Watched <i>Chronicles of Narnia</i>	52.8%	Didn't watch <i>The Equalizer</i>	51.9%
Watched <i>One Fine Day</i>	52.5%	Didn't watch <i>Hostel</i>	52.7%	Watched <i>The Judge</i>	51.7%
Watched <i>My Fellow Americans</i>	52.5%	Watched <i>Brokeback Mountain</i>	52.7%	Watched <i>Wild</i>	51.7%
Panel (c) Magazines					
1992		2002		2011	
Didn't read <i>Ebony</i>	69.5%	Didn't read <i>Ebony</i>	70.7%	Didn't read <i>Ebony</i>	64.1%
Didn't read <i>Jet</i>	68.4%	Didn't read <i>Jet</i>	70.5%	Didn't read <i>Essence</i>	61.9%
Didn't read <i>Essence</i>	62.8%	Didn't read <i>Essence</i>	66.9%	Didn't read <i>Jet</i>	61.8%
Didn't read <i>Black Enterprise</i>	57.3%	Didn't read <i>Black Enterprise</i>	61.3%	Didn't read <i>Black Enterprise</i>	57.8%
Read <i>Modern Maturity</i>	55.7%	Didn't read <i>Vibe</i>	60.3%	Didn't read <i>TV Guide</i>	55.1%
Read <i>National Geographic</i>	54.8%	Didn't read <i>The Source</i>	56.6%	Didn't read <i>ESPN The Magazine</i>	54.9%
Read <i>Consumer Reports</i>	54.3%	Didn't read <i>TV Guide</i>	54.2%	Didn't read <i>National Enquirer</i>	54.8%
Read <i>Reader's Digest</i>	54.0%	Didn't read <i>National Enquirer</i>	54.1%	Didn't read <i>Life & Style Weekly</i>	54.7%
Didn't read <i>Star</i>	53.7%	Didn't read <i>Gentlemen's Quarterly</i>	53.7%	Didn't read <i>Seventeen</i>	54.4%
Read <i>Parade</i>	53.7%	Read <i>People</i>	53.7%	Didn't read <i>Gentlemen's Quarterly</i>	54.3%

Note: Data source is the MRI. Sample size in all panels is 4,150. Reported in each column are the 10 cultural traits most indicative of being white in that year. The numbers indicate the likelihood of guessing correctly whether an individual is white based on the answer to the question. For example, in 1992, knowing whether a person watched *National Geographic Specials* allows us to guess race correctly 55.5% of the time, whereas knowing whether a person watched *Arsenio Hall* allows us to guess race correctly 58.5% of the time. An affirmative answer to “Did you watch *National Geographic Specials*?” and a negative answer to “Did you watch *Arsenio Hall*?” indicate that the person is white.

Table A.9: Products and brands most indicative of being white

Panel (a) Products					
1992		2004		2016	
Own a dishwasher	62.5%	Own cruise control (automobile)	64.5%	Own a pet	63.4%
Own a shovel	62.4%	Own a pet	64.4%	Own a flashlight	63.3%
Own a smoke/fire detector	62.1%	Own a dishwasher	63.8%	Own a dishwasher	62.5%
Own a pet	62.0%	Own a coffee maker	63.7%	Own a sport/recreation equipment	62.4%
Own a microwave	61.7%	Own a smoke/fire detector	63.6%	Own glass ovenware/bakeware	61.9%
Own a flashlight	61.5%	Own a flashlight	63.4%	Own a gas grill	61.6%
Used suntan/sunscreen products	61.4%	Own power locks (automobile)	63.3%	Own a smoke/fire detector	61.5%
Own a hand-held electric mixer	61.3%	Own a hot water heater	63.2%	Own a hot water heater	61.4%
Own a coffee maker	61.2%	Own a hand-held electric mixer	63.0%	Own an air conditioner	61.4%
Own a hose	61.2%	Own air bags on driver side (automobile)	62.9%	Own a built-in dishwasher	60.8%
Panel (b) Brands					
1992		2004		2016	
Used Scotch (transparent tape)	59.8%	Used Scotch (transparent tape)	61.0%	Used Verizon Wireless	60.2%
Bought Kodak (film)	58.7%	Used Cut-rite (waxed paper)	57.7%	Used Thomas' (English muffin)	58.6%
Used Arm & Hammer (baking soda)	56.8%	Own Ford (automobile)	57.4%	Used Shout (laundry pre-treatment)	56.8%
Used Philadelphia (cream cheese)	56.5%	Used Shout (laundry pre-treatment)	57.2%	Used Sweet Baby Ray's (barbecue sauce)	56.5%
Used Cut-rite (waxed paper)	56.4%	Used Arm & Hammer (baking soda)	56.9%	Used Vlasic (pickles)	56.4%
Used Nestle (baking chips)	56.3%	Used Thomas' (English muffin)	56.8%	Used Arm & Hammer (baking soda)	56.0%
Used Pam Regular (cooking product)	56.1%	Used Pam Regular (cooking product)	56.8%	Used Scotch (transparent tape)	55.9%
Used Murphy's oil soap (hh. cleaner)	55.9%	Used Bertolli (salad/cooking oil)	56.7%	Used French's Classic Yellow (mustard)	55.8%
Used Elmer's glue	55.7%	Used Bush's Best (canned beans)	56.5%	Used Windex (glass/surface cleaner)	55.8%
Bought Duracell (batteries)	55.6%	Used JIF (peanut butter)	56.4%	Used Stove Top (stuffing mix/product)	55.8%

Note: Data source is the MRI. Sample size in all panels is 4,150. Reported in each column are the 10 cultural traits most indicative of being white in that year. The numbers indicate the likelihood of guessing correctly whether an individual is white based on the answer to the question. For example, in 1992, knowing whether a person owns a dishwasher allows us to guess race correctly 62.5% of the time. An affirmative answer to “Do you own a dishwasher?” indicates that the person is white.

Table A.10: Attitudes and norms most indicative of being white

1976		1996		2016	
Spending on blacks isn't too little	75.3%	Spending on blacks isn't too little	73.0%	Approve of police striking citizens	65.6%
Not a Baptist	71.5%	Spending on space expl. is adequate	64.4%	Own gun in home	62.3%
Not a fundamentalist	70.3%	Spending on welfare is too much	64.4%	Favor death penalty for murder	61.0%
Trusts people	67.9%	Approve of police striking citizen	64.0%	Own rifle in home	60.5%
Voted for Republican pres. candidate	63.9%	Favor death penalty for murder	61.7%	Spending on blacks isn't too little	60.3%
People are helpful	63.9%	Voted for Republican pres. candidate	61.4%	Voted for Republican pres. candidate	60.3%
Approve of police striking citizens	63.3%	Own gun in home	60.8%	Homosexual sex isn't wrong at all	60.1%
Favor death penalty for murder	61.5%	Divorce laws should not be made easier	60.5%	Own shotgun in home	60.0%
Spending on welfare is too much	60.9%	Trusts people	60.5%	Premarital sex isn't wrong at all	59.3%
Spending on space expl. is adequate	60.9%	Not a Democrat	60.1%	Not confident in the executive branch	58.9%

Note: Data source is the GSS. Sample size is 234. Reported in each column are the 10 cultural traits most indicative of being white in that year. The numbers indicate the likelihood of guessing correctly whether an individual is white based on the answer to the question. For example, in 1976, knowing whether a person trusts people allows us to guess race correctly 67.9% of the time, whereas knowing whether a person thinks spending on blacks is too little allows us to guess race correctly 75.3% of the time. An affirmative answer to “Do you trust people?” and a negative answer to “Is spending on blacks too little?” indicate that the person is white.

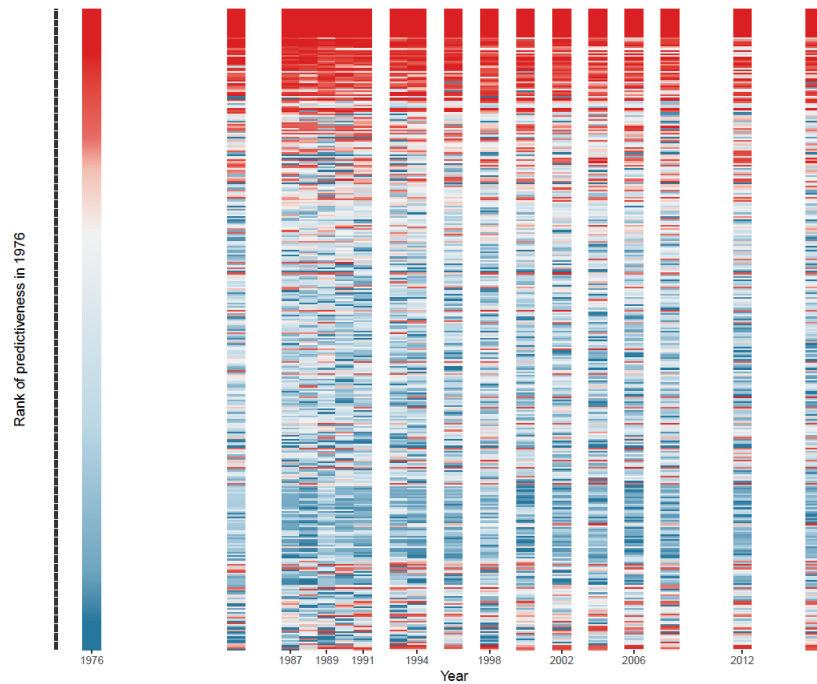


Figure A.3: Stability over time of attitudes most indicative of race

Note: Data source is the GSS. Sample size is 234. Variables are ranked from bottom to top throughout the graph by increasing order of correctly guessing race in 1976 based on that variable only. Each variable's relative informativeness in subsequent years is color-coded, with the most informative variables in each year color-coded dark red and the least informative color-coded dark blue, and lighter shades of red and blue in between. See Data Appendix for implementation details.

A.2.5 Political Ideology

Table A.11: TV shows, movies, and magazines most indicative of being liberal

Panel (a) TV shows					
1994		2001		2009	
Watched <i>Academy Awards</i>	55.4%	Watched <i>Academy Awards</i>	53.0%	Didn't watch <i>The O'Reilly Factor</i>	56.0%
Didn't watch <i>Bob Hope Specials</i>	55.4%	Watched <i>Will & Grace</i>	52.9%	Didn't watch <i>Fox and Friends</i>	55.6%
Didn't watch <i>Rush Limbaugh</i>	55.2%	Watched <i>Friends</i>	52.9%	Didn't watch <i>Hannity & Colmes</i>	55.2%
Didn't watch <i>Country Music Awards</i>	54.5%	Didn't watch <i>Country Music Awards</i>	52.8%	Didn't watch <i>NASCAR Nextel Cup Series</i>	54.4%
Didn't watch <i>Bob Hope Chrysler Classic</i>	53.9%	Watched <i>Ally McBeal</i>	52.8%	Watched <i>The Daily Show</i>	54.2%
Watched <i>SNL Anniv. Specials</i>	53.8%	Didn't watch <i>Wheel of Fortune</i>	52.8%	Watched <i>SNL Specials</i>	54.2%
Didn't watch <i>The Second Half</i>	53.8%	Didn't watch <i>Touched by an Angel</i>	52.7%	Didn't watch <i>Fox Report</i>	53.9%
Watched <i>Phenom</i>	53.6%	Didn't watch <i>Tournament of Roses Parade</i>	52.6%	Didn't watch <i>NASCAR Daytona 500</i>	53.7%
Didn't watch <i>Academy of Country Music</i>	53.4%	Didn't watch <i>US Open</i>	52.6%	Watched <i>Academy Awards</i>	53.7%
Didn't watch <i>Thanksgiving Day Parade</i>	53.4%	Watched <i>The Simpsons</i>	52.6%	Didn't watch <i>Super Bowl</i>	53.7%
Panel (b) Movies					
1998		2004		2009	
Watched <i>Jerry Maguire</i>	55.3%	Watched <i>Chicago</i>	54.8%	Watched <i>Juno</i>	55.4%
Watched <i>The People vs. Larry Flynt</i>	53.4%	Watched <i>The Hours</i>	53.9%	Watched <i>No Country for Old Men</i>	54.0%
Watched <i>Ransom</i>	52.8%	Watched <i>About Schmidt</i>	53.8%	Watched <i>Michael Clayton</i>	53.2%
Watched <i>The English Patient</i>	52.4%	Watched <i>8 Mile</i>	53.6%	Didn't watch <i>Nat'l Treasure: Book of Secrets</i>	52.6%
Watched <i>Michael</i>	52.3%	Watched <i>Adaptation</i>	53.6%	Watched <i>Sweeney Todd</i>	52.4%
Watched <i>Mars Attacks</i>	52.2%	Watched <i>Lord of the Rings 2</i>	53.3%	Watched <i>Sex and the City</i>	52.3%
Watched <i>Donnie Brasco</i>	52.0%	Watched <i>Catch Me If You Can</i>	53.1%	Watched <i>American Vincent</i>	52.1%
Watched <i>First Wives' Club</i>	51.9%	Watched <i>The Pianist</i>	52.9%	Watched <i>The Other Boleyn Girl</i>	52.1%
Watched <i>Sleepers</i>	51.9%	Watched <i>Harry Potter 2</i>	52.8%	Watched <i>Atonement</i>	52.0%
Watched <i>The Long Kiss Goodnight</i>	51.7%	Watched <i>Red Dragon</i>	52.6%	Didn't watch <i>The Chronicles of Narnia</i>	51.9%
Panel (c) Magazines					
1994		2001		2009	
Read <i>Cosmopolitan</i>	54.7%	Read <i>Cosmopolitan</i>	54.1%	Read <i>Vanity Fair</i>	54.9%
Read <i>Rolling Stone</i>	53.8%	Read <i>People</i>	54.0%	Read <i>Rolling Stone</i>	54.6%
Read <i>Vogue</i>	53.2%	Read <i>Rolling Stone</i>	53.8%	Read <i>Vogue</i>	54.0%
Didn't read <i>Reader's Digest</i>	53.2%	Read <i>Entertainment Weekly</i>	53.6%	Read <i>The New Yorker</i>	53.8%
Read <i>TV Guide</i>	53.1%	Didn't read <i>Reader's Digest</i>	53.5%	Didn't read <i>Reader's Digest</i>	53.0%
Read <i>Newsweek</i>	53.1%	Read <i>Vogue</i>	53.2%	Didn't read <i>Field & Stream</i>	52.9%
Read <i>The New Yorker</i>	52.8%	Read <i>The New Yorker</i>	53.0%	Read <i>Time</i>	52.9%
Read <i>Vanity Fair</i>	52.8%	Didn't read <i>Southern Living</i>	52.9%	Read <i>People</i>	52.8%
Read <i>New York Times Magazine</i>	52.8%	Read <i>Newsweek</i>	52.5%	Read <i>Glamour</i>	52.7%
Read <i>Entertainment Weekly</i>	52.8%	Read <i>Vanity Fair</i>	52.3%	Read <i>O, The Oprah Magazine</i>	52.7%

Note: Data source is the MRI. Sample size in all panels is 4,864. Reported in each column are the 10 cultural traits most indicative of being liberal in that year. The numbers indicate the likelihood of guessing correctly whether an individual is liberal based on the answer to the question. For example, in 1994, knowing whether a person watched *Academy Awards* allows us to guess political ideology correctly 55.4% of the time, whereas knowing whether a person watched *Bob Hope Specials* allows us to guess political ideology correctly 55.4% of the time. An affirmative answer to “Did you watch *Academy Awards*?” and a negative answer to “Did you watch *Bob Hope Specials*?” indicate that the person is liberal.

Table A.12: Products and brands most indicative of being liberal

Panel (a) Products					
1994		2001		2009	
Drank any alcoholic beverage	56.0%	Drank any alcoholic beverage	57.9%	Not own a fishing rod	56.9%
Drank bottled water/seltzer	55.5%	Drank imported beer	57.7%	Not own fishing lures/hooks	56.8%
Drank beer	55.5%	Drank any distilled liquor	57.1%	Not own a fishing reel	56.7%
Not own a lawn mower	55.5%	Drank any beer	57.0%	Own any vehicle	56.5%
Not own a portable circular saw	55.3%	Drank any mixed drinks	55.5%	Didn't use frozen bread/dough	56.3%
Bought pre-recorded audio records/tapes/discs	55.2%	Didn't buy religious books (ex-Bible)	55.4%	Drank any alcoholic beverage	56.2%
Drank wine	55.2%	Bought toiletries at a drug store	55.4%	Bought a novel	56.2%
Used tampons	55.0%	Not own a fishing rod	55.4%	Didn't use ranch salad dressing	56.2%
Didn't use gelatin/gelatin desserts	54.9%	Not own a fishing reel	55.4%	Didn't use disposable plates	56.0%
Not own a separate freezer	54.8%	Drank wine	55.2%	Not own other fishing equipments	55.8%

Panel (b) Brands					
1994		2001		2009	
Didn't use Jell-o regular	54.7%	Drank Poland Springs (bottled water)	54.2%	Didn't buy at Arby's (fast food)	55.6%
Didn't use Minute original (rice)	54.1%	Used Celestial Seasonings (tea)	53.7%	Didn't use JIF (peanut butter)	54.4%
Didn't use Kellogg's rice krispies	54.0%	Drank Corona Extra (beer)	53.5%	Didn't buy at Applebee's (family rest.)	54.4%
Didn't use Crisco regular (shortening)	53.7%	Did not buy at Arby's (fast food)	53.5%	Not own a Chevrolet (automobile)	54.2%
Didn't use Heinz (vinegar)	53.5%	Not own a Chevrolet (automobile)	53.4%	Didn't use Tyson (chicken/turkey)	54.2%
Didn't buy Kodak (film)	53.5%	Did not buy at Burger King (fast food)	53.3%	Didn't buy at Sonic (fast food)	54.1%
Didn't use Cut-rite (waxed paper)	53.4%	Used Ben & Jerry's (ice cream)	53.3%	Didn't buy Wrangler (men's clothing)	54.0%
Didn't use Raid (outdoor insecticide)	53.4%	Drank Sam Adams (beer)	53.2%	Didn't use Little Debbie (snack cake)	54.0%
Didn't use Calumet (baking soda)	53.3%	Didn't use Betty Crocker (dry cake mix)	53.2%	Didn't buy Dockers (men's clothing)	54.0%
Didn't use Morton (salt)	53.3%	Didn't use Cracker Barrel (family rest.)	53.2%	Didn't use Cool Whip (whip. topping)	54.0%

Note: Data source is the MRI. Sample size in all panels is 4,864. Reported in each column are the 10 cultural traits most indicative of being liberal in that year. The numbers indicate the likelihood of guessing correctly whether an individual is liberal based on the answer to the question. For example, in 1994, knowing whether a person bought any alcoholic beverage allows us to guess political ideology correctly 56.0% of the time, whereas knowing whether a person owns a lawn mower allows us to guess political ideology correctly 55.5% of the time. An affirmative answer to “Did you buy any alcoholic beverage?” and a negative answer to “Do you own a lawn mower?” indicate that the person is liberal.

Table A.13: Attitudes and norms most indicative of being liberal

1976		1996		2016	
Marijuana should be made legal	66.6%	Homosexual sex is not always wrong	66.1%	Allow abortion for single women	71.8%
Extramarital sex isn't always wrong	63.3%	Premarital sex isn't wrong at all	64.7%	Allow abortion for low income women	71.2%
Porn shouldn't be illegal to all	62.2%	Allow abortion for low income women	63.5%	Homosexual sex isn't wrong at all	69.7%
Homosexual sex isn't always wrong	62.1%	Allow abortion for single women	62.7%	Allow abortion for married women	68.8%
Allow atheists to teach	62.1%	Allow abortion for married women	62.0%	Spending on defense is too much	66.1%
Allow communists to teach	62.0%	Spending on welfare is adequate	61.5%	Allow abortion for rape victims	65.0%
Oppose death penalty for murder	62.0%	Spending on defense is too much	61.3%	Spending on the environment is too little	64.7%
Spending on blacks is too little	61.7%	Spending on the environment is too little	61.1%	Premarital sex isn't wrong at all	64.2%
Divorce laws shouldn't be more difficult	61.7%	Spending on health care is too little	60.2%	Spending on blacks is too little	64.0%
Allow militarists to speak	61.7%	Spending on blacks is too little	60.2%	Confident in the executive branch	63.4%

Note: Data source is the GSS. Sample size is 566. Reported in each column are the 10 cultural traits most indicative of being liberal in that year. The numbers indicate the likelihood of guessing correctly whether an individual is liberal based on the answer to the question. For example, in 1976, knowing whether a person thinks marijuana should be made legal allows us to guess political ideology correctly 66.6% of the time, whereas knowing whether a person thinks porn should be illegal to all allows us to guess political ideology correctly 62.2% of the time. An affirmative answer to “Should marijuana be made legal?” and a negative answer to “Should porn be illegal to all?” indicate that the person is liberal.

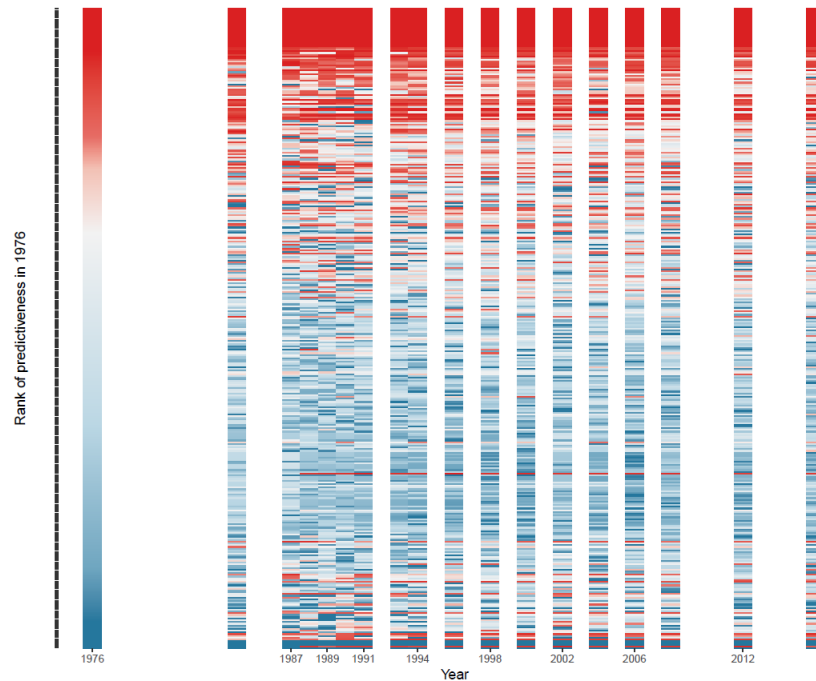


Figure A.4: Stability over time of attitudes most indicative of political ideology

Note: Data source is the GSS. Sample size is 566. Variables are ranked from bottom to top throughout the graph by increasing order of correctly guessing political ideology in 1976 based on that variable only. Each variable's relative informativeness in subsequent years is color-coded, with the most informative variables in each year color-coded dark red and the least informative color-coded dark blue, and lighter shades of red and blue in between. See Data Appendix for implementation details.

A.3 Robustness

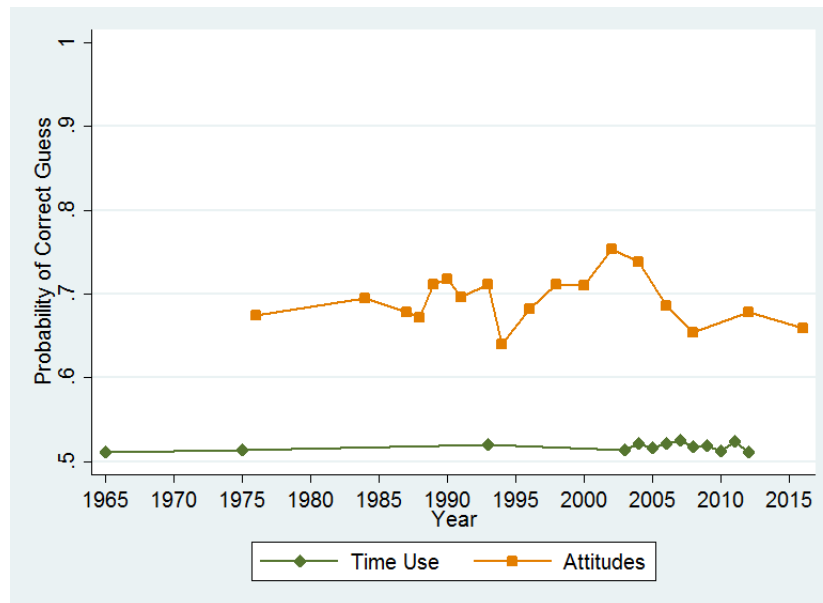


Figure A.5: Cultural distance by urbanicity

Note: Data sources are the GSS and the AHTUS. Sample sizes each year are 756 for time use and 236 for attitudes. See text and data appendix for details on sample construction and implementation of machine-learning ensemble method. Presented in the figure is share of correct guesses of respondent's urbanicity in the hold-out sample each year. The procedure to guess urbanicity in the hold-out sample was repeated 500 times, and the share of guesses reported is the average of these 500 iterations.

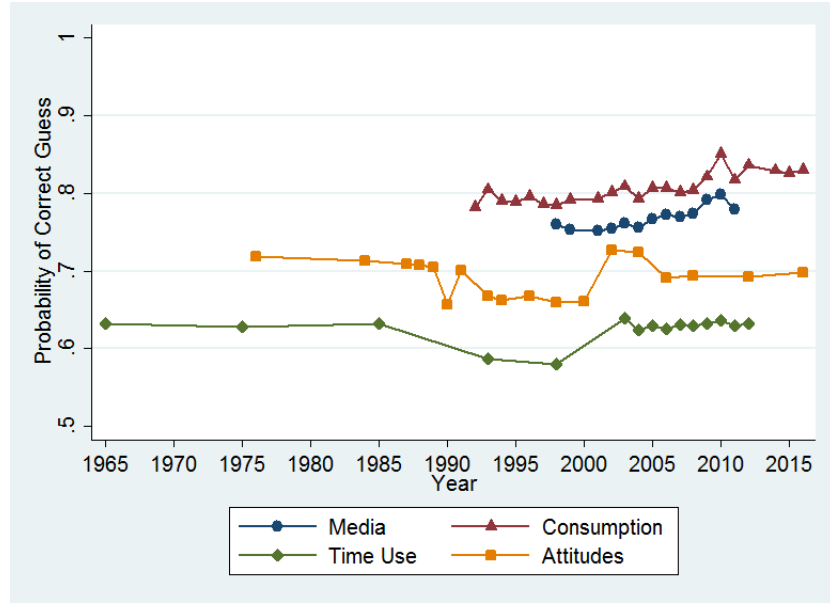


Figure A.6: Cultural distance by age

Note: Data sources are the GSS, the AHTUS, and the MRI. Sample sizes each year are 14,602 for media and consumption, 1,088 for time use, and 892 for attitudes. See text and data appendix for details on sample construction and implementation of machine-learning ensemble method. Presented in the figure is share of correct guesses of respondent's age in the hold-out sample each year. The procedure to guess age in the hold-out sample was repeated 5 times for consumption, 25 times for media, and 500 times for time use and attitudes, and the share of guesses reported is the average of these iterations.

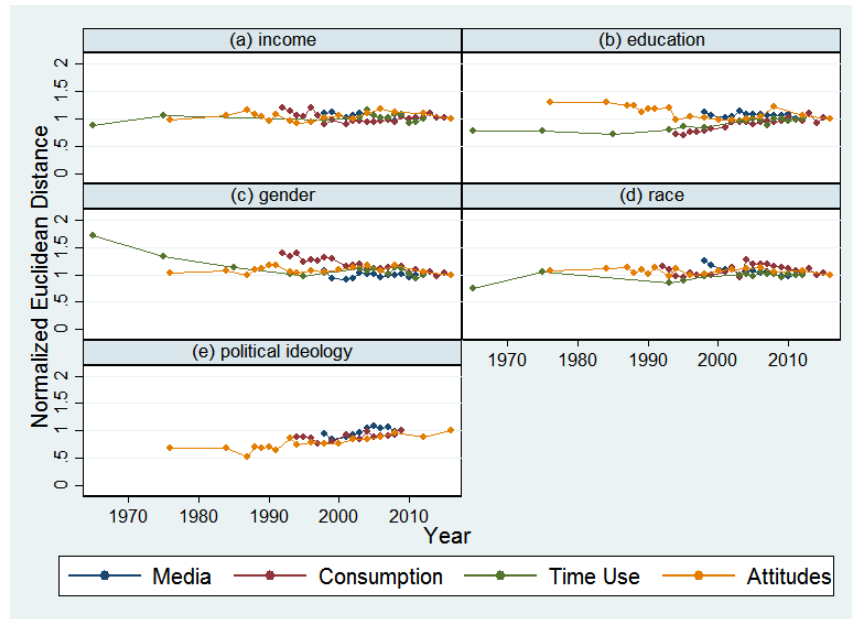


Figure A.7: Cultural distance over time: Euclidean distance

Note: Figure reports normalized Euclidean distances between groups in each year based on media diet, consumer behavior, time use, or social attitudes.

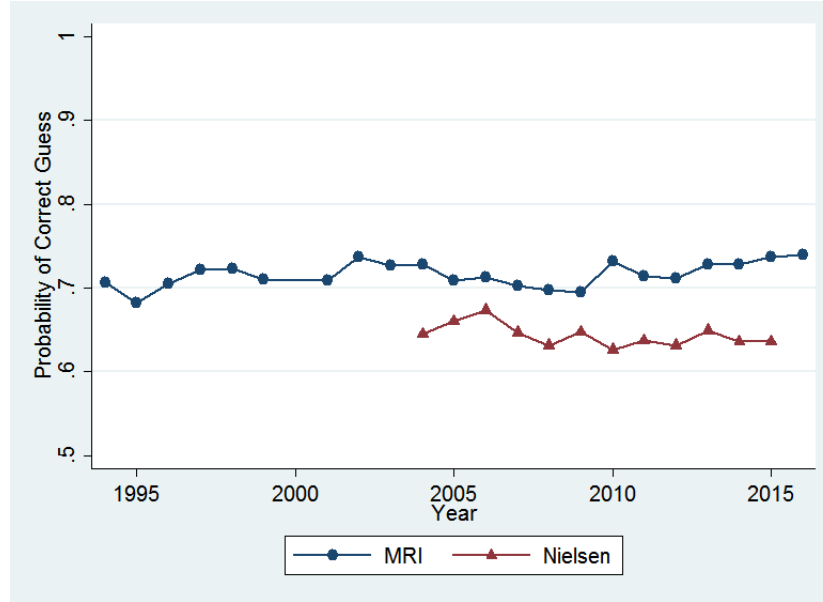


Figure A.8: Cultural distance by education over time: consumer behavior

Note: Data sources are the MRI and Nielsen. Sample sizes each year are 1,628 for MRI and 2,164 for Nielsen. See text and data appendix for details on sample construction and implementation of machine-learning ensemble method. Presented in the figure is share of correct guesses of respondent's education in the hold-out sample each year. The procedure to guess education in the hold-out sample was repeated 5 times, and the share of guesses reported is the average of these iterations.

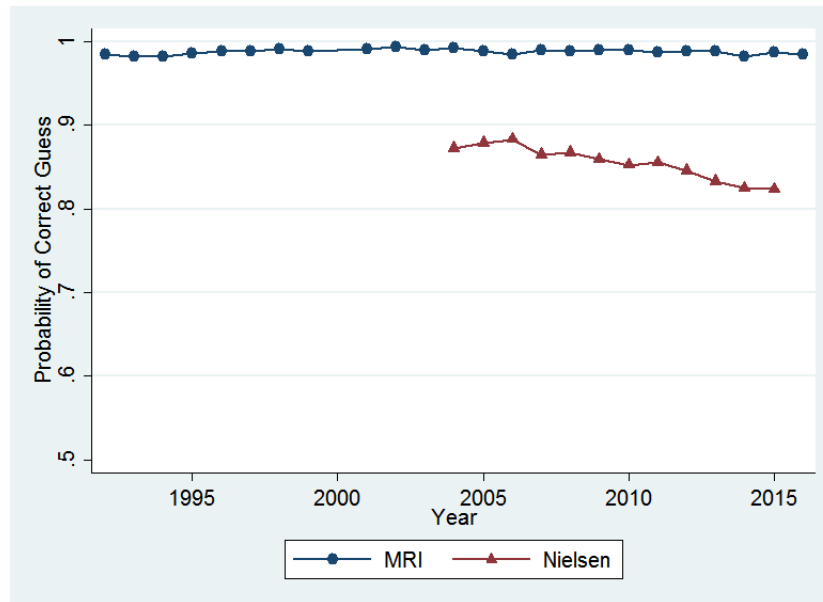


Figure A.9: Cultural distance by gender over time: consumer behavior

Note: Data sources are the MRI and Nielsen. Sample sizes each year are 2,242 for MRI and 4,566 for Nielsen. See text and data appendix for details on sample construction and implementation of machine-learning ensemble method. Presented in the figure is share of correct guesses of respondent's gender in the hold-out sample each year. The procedure to guess gender in the hold-out sample was repeated 5 times, and the share of guesses reported is the average of these iterations.

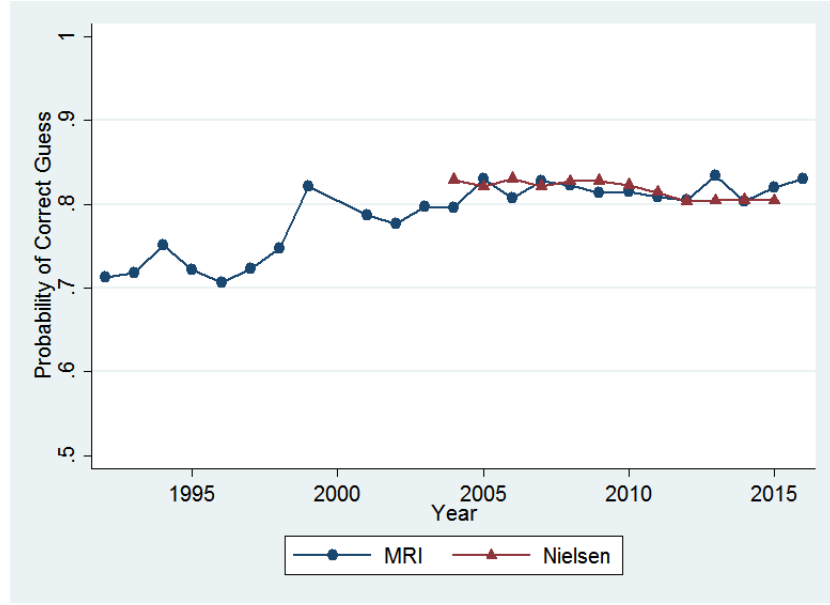


Figure A.10: Cultural distance by race over time: consumer behavior

Note: Data sources are the MRI and Nielsen. Sample sizes each year are 594 for MRI and 2,450 for Nielsen. See text and data appendix for details on sample construction and implementation of machine-learning ensemble method. Presented in the figure is share of correct guesses of respondent's race in the hold-out sample each year. The procedure to guess race in the hold-out sample was repeated 5 times, and the share of guesses reported is the average of these iterations.

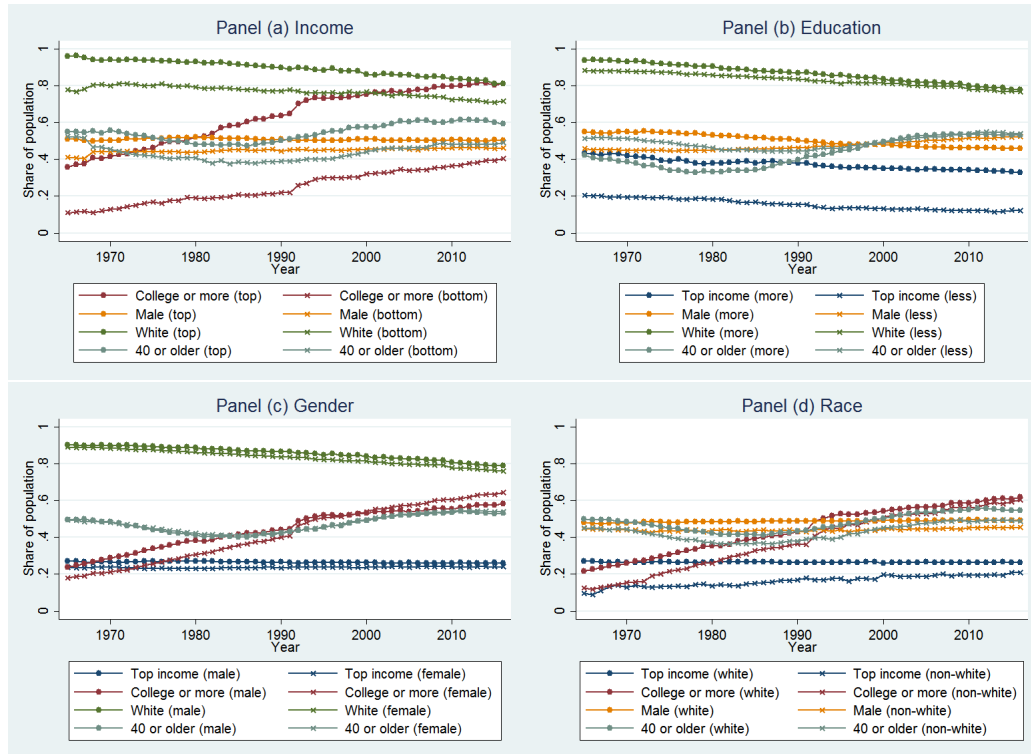


Figure A.11: Compositional changes in income, education, gender, and race

Note: Income defined by top vs. bottom quartile of household income by type.

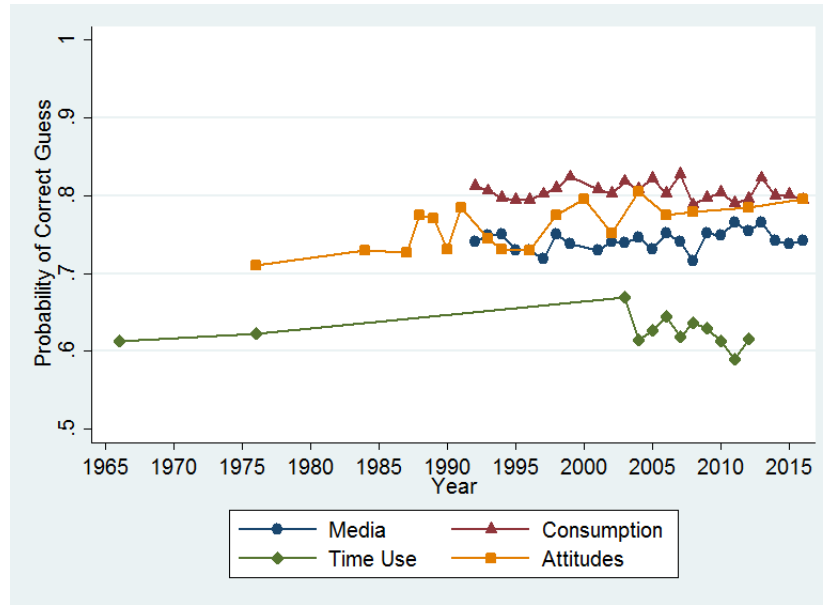


Figure A.12: Cultural distance by income controlling for age

Note: Data sources are the GSS, the AHTUS, and the MRI. Sample sizes each year are 6,026 for media and consumption, 590 for time use, and 372 for attitudes. See text and data appendix for details on sample construction and implementation of machine-learning ensemble method. Presented in the figure is share of correct guesses of respondent's income in the hold-out sample each year. The procedure to guess income in the hold-out sample was repeated 5 times for consumption, 25 times for media, and 500 times for time use and attitudes, and the share of guesses reported is the average of these iterations.

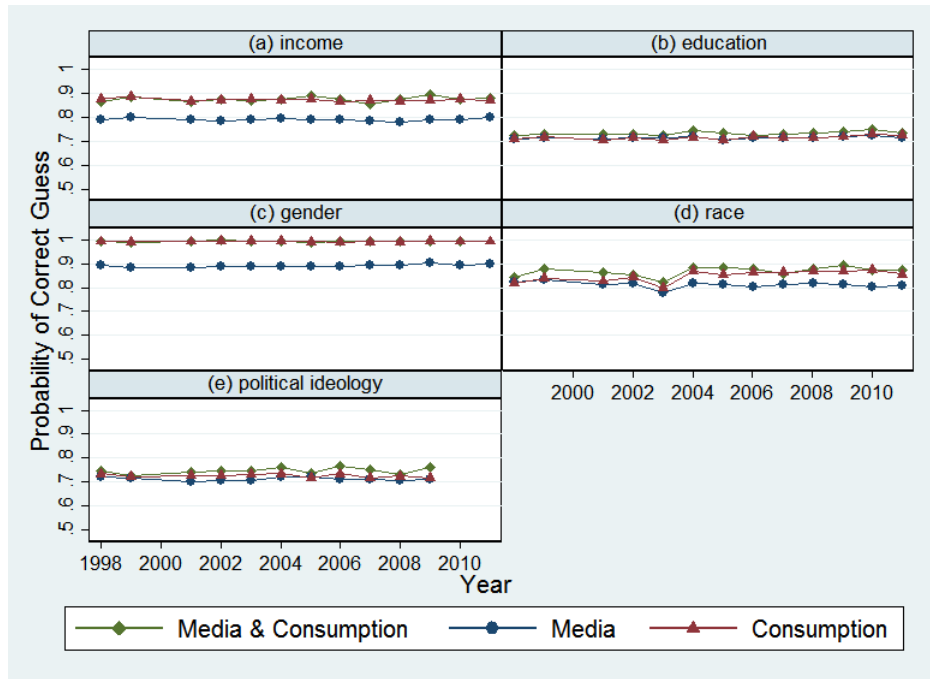


Figure A.13: Cultural distance in both media consumption and consumer behavior

Note: Data source is the MRI. Sample size each year is 6,394. See text and data appendix for details on sample construction and implementation of machine-learning ensemble method. Presented in the figure is share of correct guesses of respondent's income in the hold-out sample each year. The procedure to guess income in the hold-out sample was repeated 25 times, and the share of guesses reported is the average of these 25 iterations.

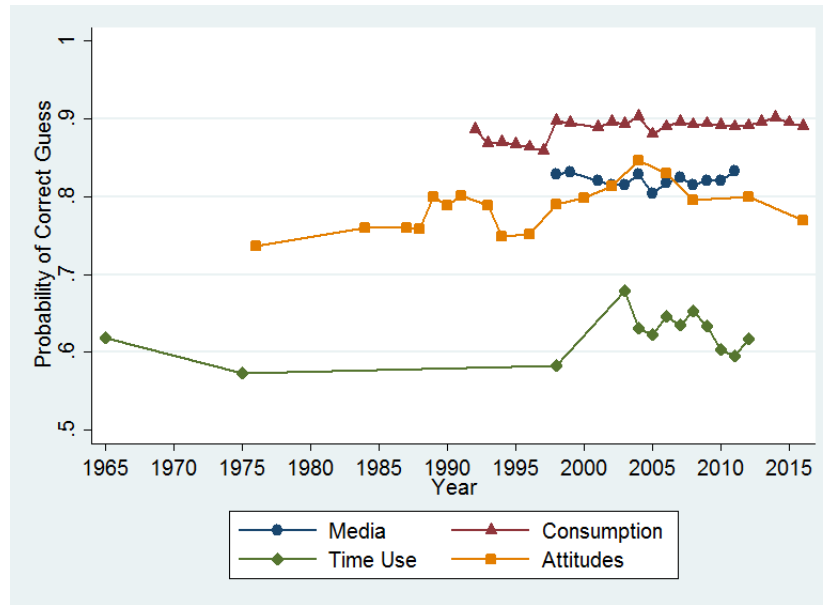


Figure A.14: Cultural distance by income, controlling for household size

Note: Data sources are the GSS, the AHTUS, and the MRI. Sample sizes each year are 4,952 for media and consumption, 312 for time use, and 274 for attitudes. See text and data appendix for details on sample construction and implementation of machine-learning ensemble method. Presented in the figure is share of correct guesses of respondent's income in the hold-out sample each year. The procedure to guess income in the hold-out sample was repeated 5 times for consumption, 25 times for media, and 500 times for time use and attitudes, and the share of guesses reported is the average of these iterations.

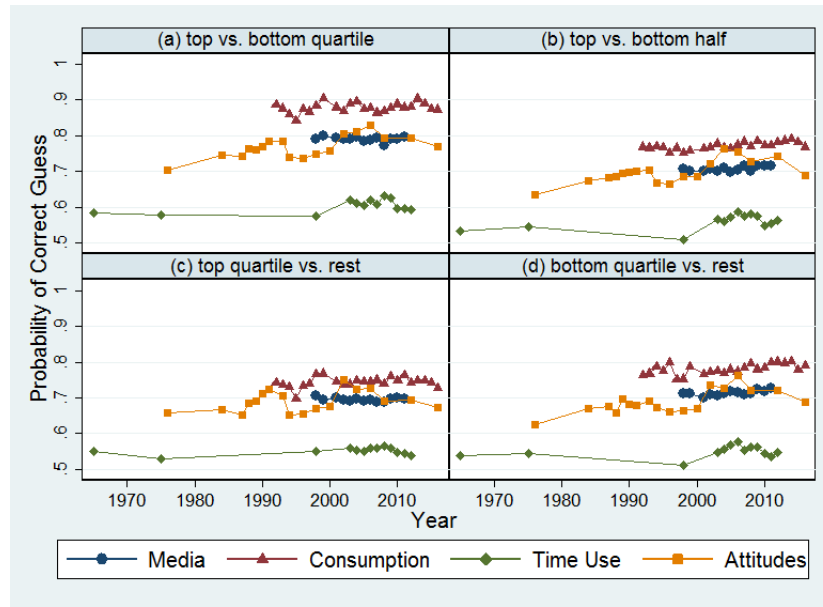


Figure A.15: Alternative income groups

Note: Figure shows the likelihood, in each year, of correctly guessing an individual's group membership based on his/her media diet, consumer behavior, time use, or social attitudes. Panel (a) is equivalent to panel (a) in 1. Panel (b) measures the cultural distance between the top half and the bottom half of the income distribution. Panel (c) measures the distance between top quartile and the rest (second, third, and fourth quartiles), and panel (d) measures the distance between the bottom quartile and the rest (first, second, and third quartiles). See text and data appendix for details on sample construction and implementation of machine-learning ensemble method.

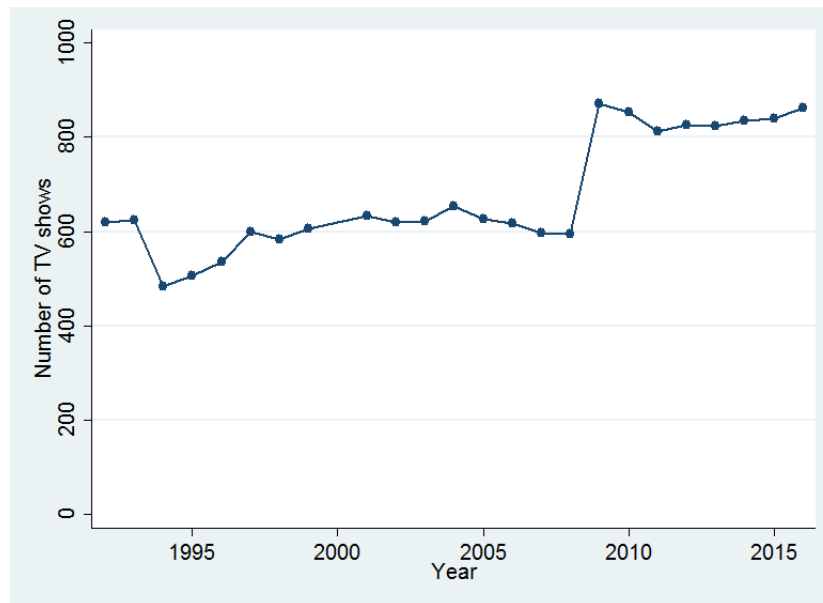


Figure A.16: Number of TV shows in the MRI data

Note: Data source is MRI. The increase in 2009 reflects addition of cable shows.

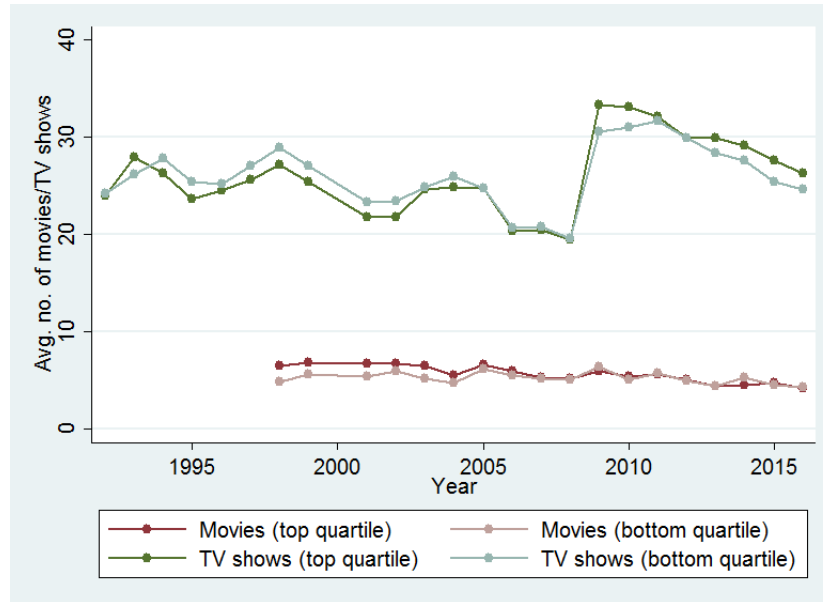


Figure A.17: Average no. of movies and TV shows watched by income in the MRI data
 Note: Data source is MRI. The increase in 2009 reflects addition of cable shows.

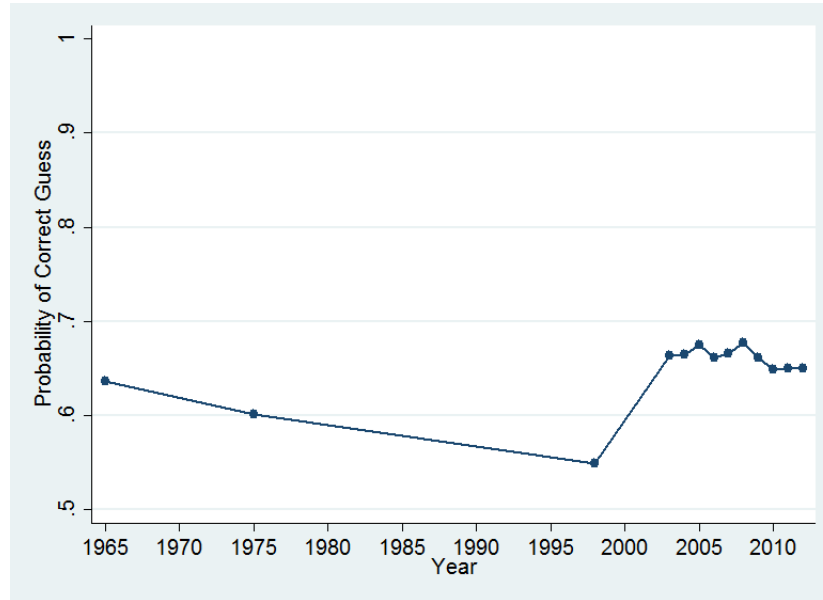


Figure A.18: Cultural distance by income in time use for the full sample
 Note: Data source is the AHTUS. Sample size each year is 706. See text and data appendix for details on sample construction and implementation of machine-learning ensemble method. Presented in the figure is share of correct guesses of respondent's income in the hold-out sample each year. The procedure to guess income in the hold-out sample was repeated 500 times, and the share of guesses reported is the average of these 500 iterations.

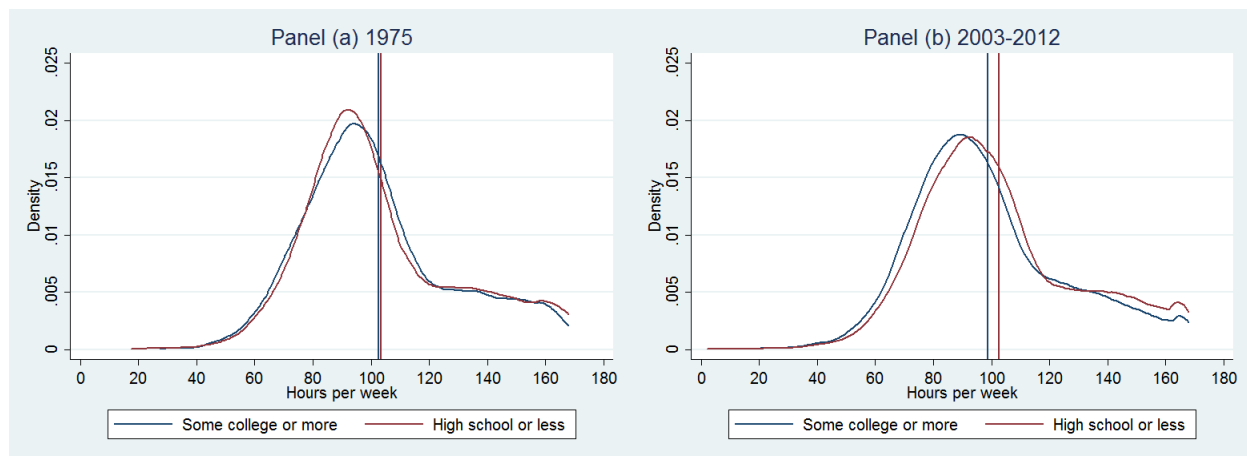


Figure A.19: Distribution of time spent on leisure by education level, 1975 vs. 2003-2012
 Note: Data source is the AHTUS.

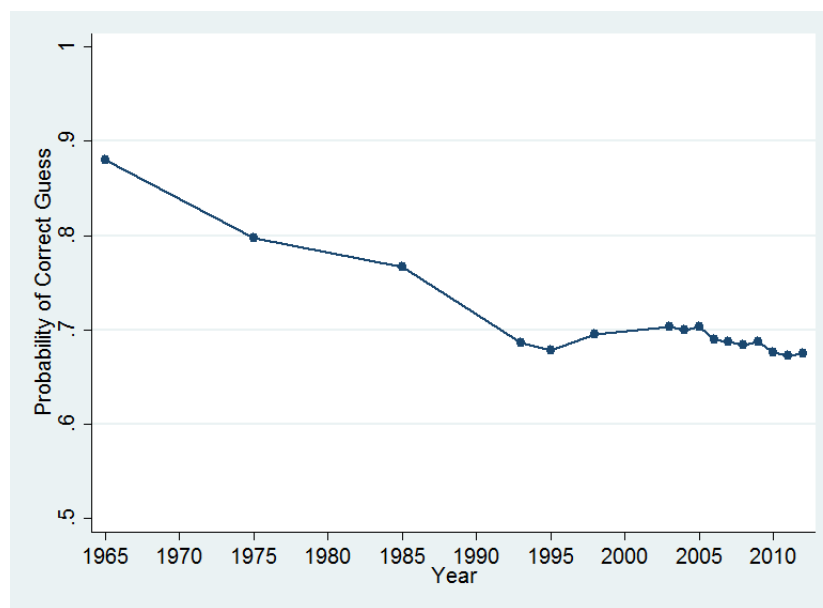


Figure A.20: Gender differences over time in allocation of non-work time
 Note: Data source is the AHTUS. Sample size each year is 812. See text and data appendix for details on sample construction and implementation of machine-learning ensemble method. Presented in the figure is share of correct guesses of respondent's gender in the hold-out sample each year. The procedure to guess gender in the hold-out sample was repeated 500 times, and the share of guesses reported is the average of these 500 iterations.

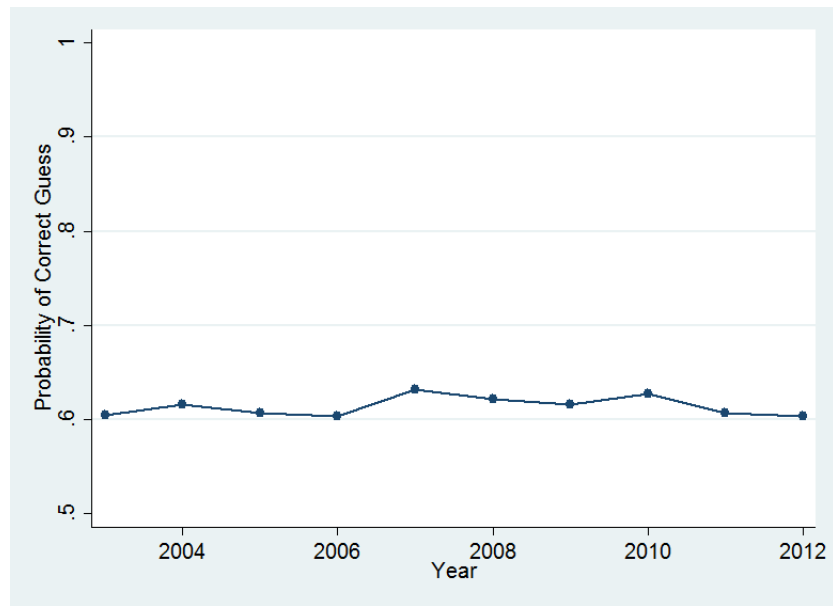


Figure A.21: Cultural distance by race in time use for the 2003-2012 sample

Note: Data source is the AHTUS. Sample size each year is 2,052. See text and data appendix for details on sample construction and implementation of machine-learning ensemble method. Presented in the figure is share of correct guesses of respondent's race in the hold-out sample each year. The procedure to guess race in the hold-out sample was repeated 500 times, and the share of guesses reported is the average of these 500 iterations.