

UCLA

UCLA Electronic Theses and Dissertations

Title

Building a Deeper Understanding of Voter Turnout Behavior in Los Angeles County in Preparation for the 2020 Election

Permalink

<https://escholarship.org/uc/item/0nz6n59w>

Author

Beeman, Paul

Publication Date

2019

Peer reviewed|Thesis/dissertation

UNIVERSITY OF CALIFORNIA

Los Angeles

Building a Deeper Understanding
of Voter Turnout Behavior in Los Angeles County
in Preparation for the 2020 Election

A thesis submitted in partial satisfaction
of the requirements for the degree
Master of Applied Statistics

by

Paul Beeman

2019

© Copyright by

Paul Beeman

2019

ABSTRACT OF THE THESIS

Building a Deeper Understanding of Voter Turnout Behavior in Los Angeles County in Preparation for the 2020 Election

by

Paul Beeman

Master of Applied Statistics in

University of California, Los Angeles, 2019

Professor Frederic R. Paik Schoenberg, Chair

In this study, we set out to better understand the voter behavior of different demographic partitions of Los Angeles County (LAC). We developed a methodology to enrich and analyze LAC voter data by joining it with US census data and modelling vote by mail rates and turnout rates using general additive models. In doing so, we were able to forecast where in LAC we expect to see the highest rates of turnout by vote by mail and turnout in person for the 2020 presidential election. These findings are timely: In 2020, LAC will roll out a new and novel voting system called Voting Solutions for All People. The findings of the current study provide guidance on how to best allocate resources to high-turnout populations and target outreach to low-turnout populations in preparation for the first election under the new system. This paper presents a combination of visualizations, predictive models, and summary statistics that can inform LAC of where to focus outreach to expand the electorate and build a more representative electorate while also improving preparation in areas where turnout has been historically high.

The thesis of Paul Beeman is approved.

Hongquan Xu

Yingnian Wu

Frederic R. Paik Schoenberg, Committee Chair

University of California, Los Angeles

2019

TABLE OF CONTENTS

1	Introduction	1
2	Literature Review	5
3	Data Collection: Sources and Methodology	7
4	Summary Statistics	10
4.1	Los Angeles County Voting Statistics	10
4.2	Census Statistics	14
5	General Additive Models (GAMS)	18
6	Vote By Mail Registration Rates	20
6.1	Smoothed Univariate Relationships Between Census Data and Vote By Mail Registration Rates	21
6.2	Forecasting Vote By Mail Registration Percentiles With GAMS	26
6.3	Model Evaluation and Selection	27
7	Voter Turnout Rates	31
7.1	Smoothed Univariate Relationships Between Census Data and Voter Turnout Rates	31
7.2	Forecasting Voter Turnout Percentiles With GAMS	39
8	Projection Maps	41
9	Conclusion	45
A	R Code	47

A.1	Clean and Merge Data	47
A.2	Summary Statistics	60
A.3	VBM GAMs	66
A.4	Voter Turnout GAMs	101

LIST OF FIGURES

4.1	Total votes cast for major elections in LA County dating back to 1990	12
4.2	Total registered voters vs. total eligible voters in Los Angeles County dating back to 1990, filtered for only general elections	13
4.3	Percent turnout of registered voters in Los Angeles County dating back to 1990	14
4.4	Poverty and public transportation percentiles: 2018	16
4.5	Turnout rate percentiles in 2014	16
4.6	Turnout rate percentiles in 2016	17
4.7	Turnout rate percentiles in 2018	17
6.1	Univariate relationship plot of average income and VBM registration percentage	22
6.2	Univariate relationship plot of poverty rates and vote by mail registration percentage	23
6.3	Univariate relationship plot of average age and vote by mail registration percentage	24
6.4	Univariate relationship plot of hispanic population rate and Vote By Mail registration percentage	25
6.5	Univariate relationship plot of bachelors' degree attainment rate and vote by mail registration percentage	26
7.1	Univariate relationship plot of average income and voter turnout rate	32
7.2	Univariate relationship plot of poverty rates and voter turnout rate	33
7.3	Univariate relationship plot of average age and voter turnout rate	34
7.4	Univariate relationship plot of Hispanic population rate and voter turnout rate .	35
7.5	Univariate relationship plot of bachelors' degree attainment rate and voter turnout rate	36
7.6	Univariate relationship plot of non-graduation from high school rate and voter turnout rate	37

7.7	Univariate relationship plot of percentage of population that are children and voter turnout rate	38
8.1	Projected turnout rate percentiles for 2020	41
8.2	Projected vote by mail registration rate percentiles for 2020	42
8.3	Smoothed relationship between vote by mail registration rates and voter turnout rate	43

LIST OF TABLES

3.1	Breakdown of variables pulled from the census	9
4.1	Recent Major Election Turnout in Los Angeles County [4,5,6]	10
4.2	Turnout and Registration for Vote by Mail vs. In-Person [4]	11
4.3	Percent of voters registered to vote by mail	11

ACKNOWLEDGMENTS

I would like to thank the Economic Roundtable and Los Angeles County Registrar-Recorder/County Clerk for giving me the opportunity to expand my skills while working on important projects to the city of Los Angeles. I would especially like to thank Judy Ly who helped me collect my data for this paper and whose input was invaluable to me developing my R programming skills.

CHAPTER 1

Introduction

For the past decade, the Los Angeles County Registrar Recorder/County Clerk (RRCC) office has been working to overhaul the voting system in Los Angeles County. In 2009, when the effort to build a new system called Voting Solutions for all People (VSAP) began, it was impossible to foresee the dramatic shift in the national political landscape and resulting increase in political engagement that would come as a result.

VSAP will shift voting in Los Angeles from a polling place-based system to a voting center-based system. Voters will no longer be bound to their polling place, but can instead vote at any voting center. While polling places are only open for one night, voting centers will be open for multiple days. Beginning eleven days before election day, there will be one open voting center for every 30,000 voters. Beginning three days before election day and on election day, there will be one open voting center for every 7,500 voters. Another important distinction is that voters can register to vote at any voting center and cast provisional ballots through election day.

The convenience of being able to vote and register to vote at any location for over a week leading up to election provides a trade-off with the diminished number of voting locations open on election day. There will be roughly 3,000 less voting locations on election day. This will of course not be an issue if people vote early; however, if a low percentage of voters cast their ballots before election day, there could be lines that stretch for several hours on election day. Los Angeles County is making a significant bet on the likelihood of voters turning out early at voting centers and/or voting by mail.

On top of increasing accessibility by widening the voting window, Los Angeles County has teamed up with the global design firm IDEO to develop a single, user-friendly interface for all voters to use. While the new system is more elegant, more user friendly, and more accessible, voters will be experiencing it for the first time in 2020. The implementation of a new user interface will likely add to the uncertainty in how successful the VSAP system roll-out will be, specifically affecting how long the average voter will remain inside a voting center.

The 2018 midterm election provided a glimpse of what the worst-case scenario could look like. Voters who missed the registration deadline were able to register at the RRCC on election day and cast provisional ballots. The line was hours long and hundreds of people deep; it was an inspiring sight to behold from a civic engagement standpoint, but one couldn't help but wonder how many people saw the line and decided that the wait was not worth it or simply not feasible.

For VSAP to be successful, Los Angeles County voters need to be educated about the new process and made aware of the benefits of turning out early and/or voting by mail. Voters that are not conditioned to the new system will likely default to what they know, which is voting on election day. This will likely result in a similar situation to the 2018 midterm election, which would be an unacceptable outcome for VSAP.

VSAP has dual goals: first, expanding the electorate, and second, diminishing the burden of voting across all socioeconomic groups. These goals can be in conflict, because over-allocating resources to historically low-turnout communities may increase their turnout but simultaneously risks overburdening historically high-turnout communities in the form of long lines.

The worst possible outcome for VSAP would be a dual failure: (1) experiencing low turnout

in historically low-turnout communities and (2) experiencing long lines in historically high-turnout communities; this would be the result of poor allocation of both resources and outreach. This study aims to provide Los Angeles County with guidelines for both allocation of resources and strategic outreach in preparation for the 2020 election. This is achieved building multiple predictive and descriptive models and providing summary statistics using voter and census tract data from midterm and presidential elections going back to 2012.

Two response variables will be explored in informing this study’s suggestions for voter outreach in terms of turnout. The first response variable we will look at is the percentage of registered voters that are registered to vote by mail in each census tract. The second response variable we will look at is the percentage of registered voters that actually vote.

The investigation of the first response variable (i.e., the percentage of registered voters that are registered to vote by mail in each census tract) is an important and novel aspect of the study because, to the best of the author’s knowledge, there has been no prior extensive research done on what demographic predictors are correlated with vote by mail utilization. Many of the arguments relating to social justice and day-of-voting center on a disproportionate opportunity cost faced by lower income individuals in voting. Vote by mail is, theoretically, a “great equalizer” because it can eliminate lines, travel time, child care, and multiple other barriers to voting. Still, as of 2018, less than 50% of LA County registered voters were registered to vote by mail. Therefore, understanding what demographic predictions are correlated to vote by mail utilization is a key aspect of this study.

Not only does vote by mail theoretically achieve what VSAP hopes to achieve (i.e., decreasing the burden of voting across socioeconomic groups), it has been in existence for decades. Understanding which groups of people have been consistently averse to voting by mail may help Los Angeles County narrow down which groups VSAP could benefit the most, enabling Los Angeles County to target outreach to those groups. For example, if elderly individuals have shown a consistent aversion to vote by mail, Los Angeles County should consider fo-

cusing VSAP outreach towards elderly individuals.

The second response variable (i.e., the percentage of registered voters who actually voted) will also be modeled. By understanding the relationship between the registered voter and voting populations, Los Angeles County can develop greater understanding of which census tracts will produce long lines at voting centers. For example, census tracts that have historically high turnout rates and low rates of vote by mail utilization should trigger concerns for potential long wait times at voting centers.

Finally, a predictive model that is informed by the descriptive models will be built to give Los Angeles County a comprehensive understanding of how many non-vote by mail voters they should expect to turn out from each census tract in the 2020 election.

These predictions will inform Los Angeles County as to where they need to prepare for the largest influx of voters, and also where they need to focus their outreach efforts. The results provided in this paper will highlight areas of Los Angeles County that have both high and low expected turnout. Los Angeles County can then utilize this information to encourage higher turnout in census tracts with low expected turnout and to increase allocation of resources to census tracts with high expected turnout.

CHAPTER 2

Literature Review

A key goal of VSAP is to expand the electorate by making the voting experience more convenient through early voting and election day registration (EDR). However, previous studies on the effects of early voting and EDR have provided mixed levels of support for this hypothesis.

In their paper *Election Day Registration's Effect on U.S. Voter Turnout*, after comparing turnout from several states with EDR laws from before and after the implementation of EDR, Brians and Grofman conclude that EDR “exerts a strong and positive influence on turnout.” Furthermore, they conclude that “although turnout inches higher as closing dates shorten, voter turnout still remains higher with the adoption of EDR than with even very short closing dates.” They attribute this to the convenience factor of being able to make a single trip to both register to vote and vote. Their findings supports EDR as a tool to level the playing field for people who experience a higher opportunity cost when waiting in line to vote, which is in line with the goals of VSAP in Los Angeles. Brians and Grofman also find that EDR increased turnout by 4 percentage points on average and the group that saw the greatest increase in turnout (i.e., by 5%) was the middle socioeconomic status (SES) group, which they define as people with a high school education and middle income. Both high and low SES groups had an increase of 3% after EDR implementation. Brians and Grofman attribute this to higher barriers to voting for voters with low SES and previously existing high representation of voters with high SES in the voting population [1].

According to a study done by Gronke and Toffee in 2008, “early voters are older, better educated, are more likely to declare a partisan affiliation on the voter registration form, and

tend to be exposed to party mobilization efforts.” Gronke and Toffee found these differences to be significant when looking at presidential elections, but insignificant when looking at midterm elections. Gronke and Toffee theorized that this discrepancy was due to the self-selecting nature of midterm voters: midterm voters are already more homogenous in their education and age; as a result, differences between early voters and election day voters are more difficult to discern.

A wide range of studies have found that long lines on election night are disproportionately suffered by urban populations with higher rates of minorities [2]. A study performed in 2011 by Bradley and McNulty found that when a single polling place was moved in Los Angeles County, a subsequent transportation and searching effect resulted in a 3.03% drop in turnout. Some of the drop in turnout was offset by an increase in absentee voting; however, overall, there was a diminished turnout of 1.85% [3].

Under VSAP, with the consolidation of polling places into voting centers, there will also likely be a greater transportation and searching strain on voters as they get to their new voting locations. According to the RRCC, in 2016, 79% of people that utilized one of the limited early voting polling stations dropped ballots off at one of the three closest locations to them. This evidence supports the need to consider transportation limitations in the placement of polling locations, as distance traveled has been shown to have a significant impact on turnout and has a disproportionate strain on lower income individuals with limited access to cars.

CHAPTER 3

Data Collection: Sources and Methodology

In the current study, the census tract level was selected to predict voter turnout levels for the 2020 election. The census tract level was chosen because census data provides a rich, granular description of a county. Using the R package `tidycensus` [4] we were able to connect directly to the US Census API¹ and load data from either the decennial US Census or the 5-year American Community Survey.

Anonymous voter data was aggregated at the census tract level and provided to us directly from the Los Angeles County RRCC. After considerable data cleaning² we were able to join the two data sets using the census tract number variable. Thus, the 2012 5-year American Community Survey was joined with the 2012 presidential election voter data and the 2014 midterm election voter data while the 2017 5-year American Community Survey was joined with the 2016 presidential election voter data and 2018 midterm data.

Data collection methods have been consistently improving in recent years at Los Angeles County over the past four major elections. In 2012, 8.2% of the 4,802,287 registered voters in Los Angeles County were unable to be matched to a census tract due to poor data collection practices. By 2018, only 87 out of 5,325,017 total voters were unsuccessfully mapped to a census tract. Due to the higher quality of data in recent years, we focused mainly on the 2016 presidential and the 2018 midterm election; earlier elections were referenced as needed throughout the study. The data that Los Angeles County provided us included

¹Application Programming Interface

²see R code in appendix for details

voter registration counts, total votes for each election broken down by political party, and whether the voters were registered to vote by mail or vote in person. After each major election, Los Angeles County takes a “snapshot” of registered voters, which allowed us to identify which census tract each voter lived in at their time of voting. For the sake of parsimoniousness, we combined several of the less popular political parties into an “other” bin, resulting in categories of Republican, Democrat, No Party Preference (NPP), and other.

The data available to us from the 5-year American Community Survey is extremely expansive. The 2017 survey provides access to over 25,000 data points for each census tract which was quite intimidating, but the `tidycensus` package provides a useful set of tools to search through variables. Nonetheless, the process of sifting through variables to find meaningful data was far from trivial. The majority of the variables were overly specific, limiting their usefulness in building models. From the 5-year American Community Survey, we chose to explore variables pertaining to age, race, gender, earnings, employment, poverty, education, transportation, and childcare.

Categories	Variables
Population	• Total Population • Population of Voting Age
Gender	• Count of females • Count of males
Race/Ethnicity	• Count of African Americans • Count of European Americans • Count of Asian Americans • Count of Native Americans • Count of Pacific Islanders • Count of Hispanic or Latino
Earnings	• Average Yearly Earnings

Employment	• Count of unemployed in the workforce • Count of employed in the workforce
Poverty	• Count of people under the poverty level • Count of people between the poverty level and 1.5x the poverty level • Count of people above 1.5x the poverty level
Education	• Count of people that did not graduate high school • Count of people with high school education • Count of people with bachelor's degree • Count of people with graduate degree
Transportation	• Count of people that drive to work • Count of people that bike to work • Count of people that take the bus to work
Childcare	• Count of children under 6 • Count of children between 6 and 17

Table 3.1: Breakdown of variables pulled from the census

CHAPTER 4

Summary Statistics

4.1 Los Angeles County Voting Statistics

Election	Total Registrations	Voters	Turnout Percentage
2004 Presidential	3,972,738	2,710,280	77%
2006 Midterm	3,914,138	2,033,119	52%
2008 Presidential	4,093,488	3,368,057	82%
2010 Midterm	4,403,046	2,377,105	54%
2012 Presidential	4,802,287	3,137,427	65%
2014 Midterm	4,908,969	1,502,482	31%
2016 Presidential	5,250,700	3,445,829	67%
2018 Midterm	5,325,017	2,918,717	55%

Table 4.1: Recent Major Election Turnout in Los Angeles County [4,5,6]

The 2008 presidential election had the highest turnout percentage of registered voters in Los Angeles County going back to 1990. The 2018 midterm election had the second highest turnout for a midterm election going back to 1990. It is important to note, however, that although the turnout percentage for the 2018 midterm was not unprecedented, the raw number of voters was staggering. Turnout was nearly double the unusually low turnout from the previous midterm election and exceeded the turnout from the 2004 presidential election.

We also see that vote by mail voters are making up an increasingly high percentage of registered voters in Los Angeles County. In 2018, the majority of the votes cast were by vote by

Year	Registration Category	Votes	Registered	Percent Voted
2012	NON-VBM	2072216	3423684	0.60525913
2012	VBM	1065211	1378603	0.77267422
2014	NON-VBM	848533	3392698	0.25010567
2014	VBM	653949	1516271	0.43128768
2016	NON-VBM	1841164	3068727	0.59997647
2016	VBM	1604665	2181973	0.73541927
2018	NON-VBM	1299833	2745468	0.47344679
2018	VBM	1618884	2579549	0.62758412

Table 4.2: Turnout and Registration for Vote by Mail vs. In-Person [4]

Year	VBM Percent
2012	28.7
2014	30.9
2016	41.6
2018	48.4

Table 4.3: Percent of voters registered to vote by mail

mail voters, even though vote by mail registrations still made up a minority of total voting registrations. Furthermore, vote by mail voters consistently voted at significantly higher rates than in-person voters. This could be a result of more engaged people self-selecting as vote by mail voters, but convenience may also be an important factor.

Figure 4.2 shows that since 2006, voter registration has been growing at a faster rate than the population of eligible voters. Also, significant jumps are seen leading up to presidential elections in recent years. If recent presidential elections are an indicator of what to expect in 2020, there will likely be another large influx of newly registered voters in 2020.

Figure 4.3 shows that voter registration in recent years has followed a pattern of large in-

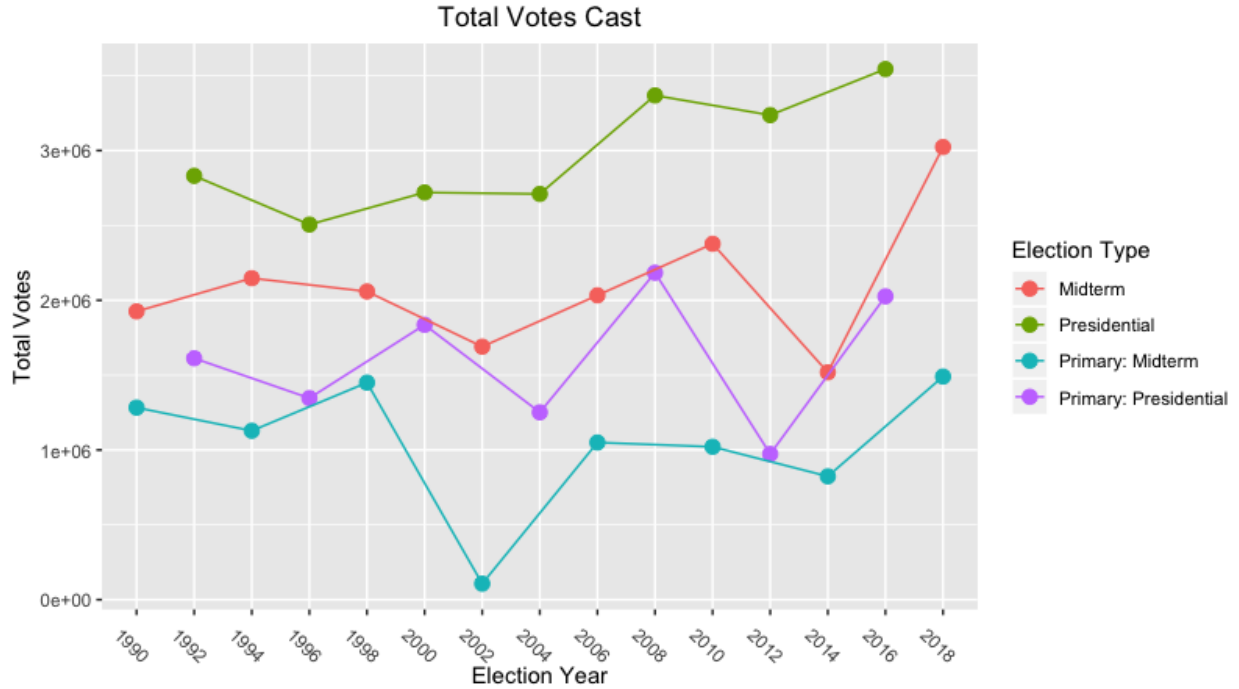


Figure 4.1: Total votes cast for major elections in LA County dating back to 1990

creases leading up to a presidential election followed by small increases in the years leading up to a midterm. The turnout was only 67% in 2016, but thanks to an expanding pool of registered voters, total turnout was the highest on record looking back until 1990.

After President Clinton’s first election, we see a two year stagnation of turnout below 70%. This was followed by a large jump in turnout in the second Bush term. A similar pattern of stagnation can be observed following the first Obama term. If recent history is to repeat itself, we can expect to see a large turnout increase in 2020, similar to the increase in turnout percentage in the second Bush term. In these cases, this may be indicative of the predominantly Democratic-leaning voting population in Los Angeles County being lulled by a elected Democratic president and then energized during the first term of a Republican president. This conclusion is of course based on a small sample size, but in preparation for VSAP implementation, it is important to be prepared for the most extreme possibility.

If we were to see 77% turnout, as we did in the second Bush term, and no further registra-

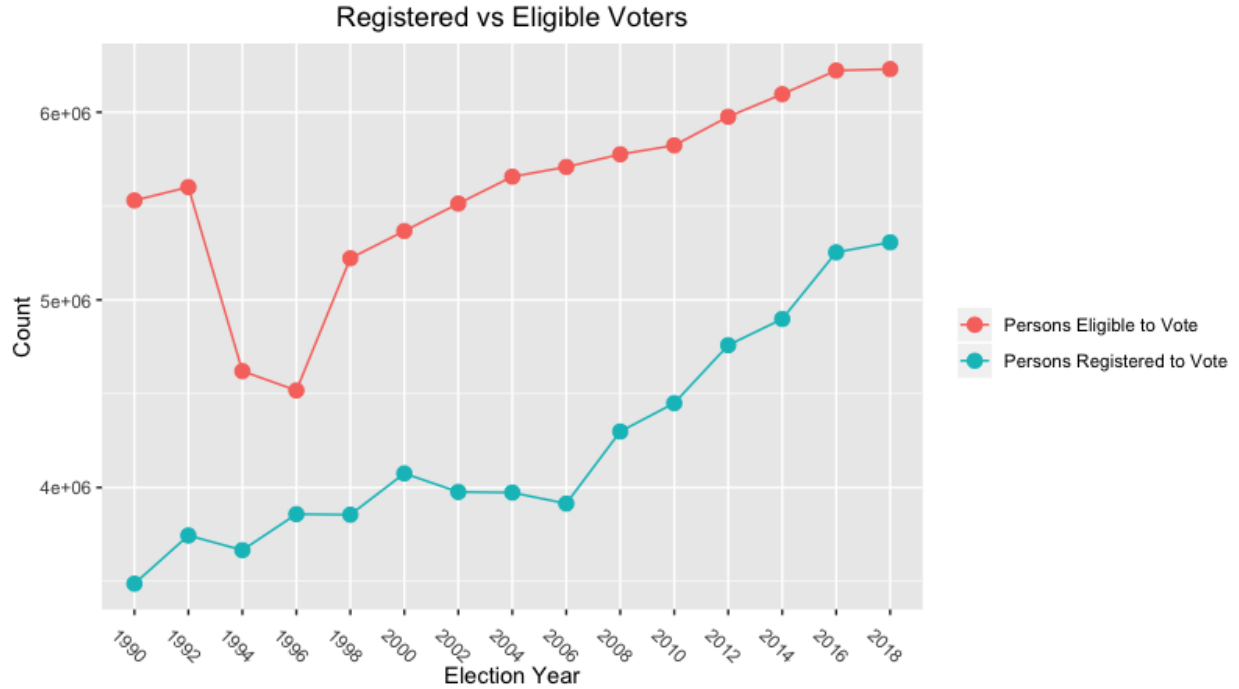


Figure 4.2: Total registered voters vs. total eligible voters in Los Angeles County dating back to 1990, filtered for only general elections

tions¹ from the 2018 registration levels, we would see a turnout of 4,100,263 voters, or over half a million more voters than in the 2016 presidential election. Given the high midterm turnout in 2018 compared to 2014, it seems that political engagement is generally trending up and we should be prepared for a very large turnout in 2020.

¹The assumption of no further registrations is obviously unrealistic but put in place to demonstrate that even for zero growth of Los Angeles County voting population, we should still expect a historically large turnout

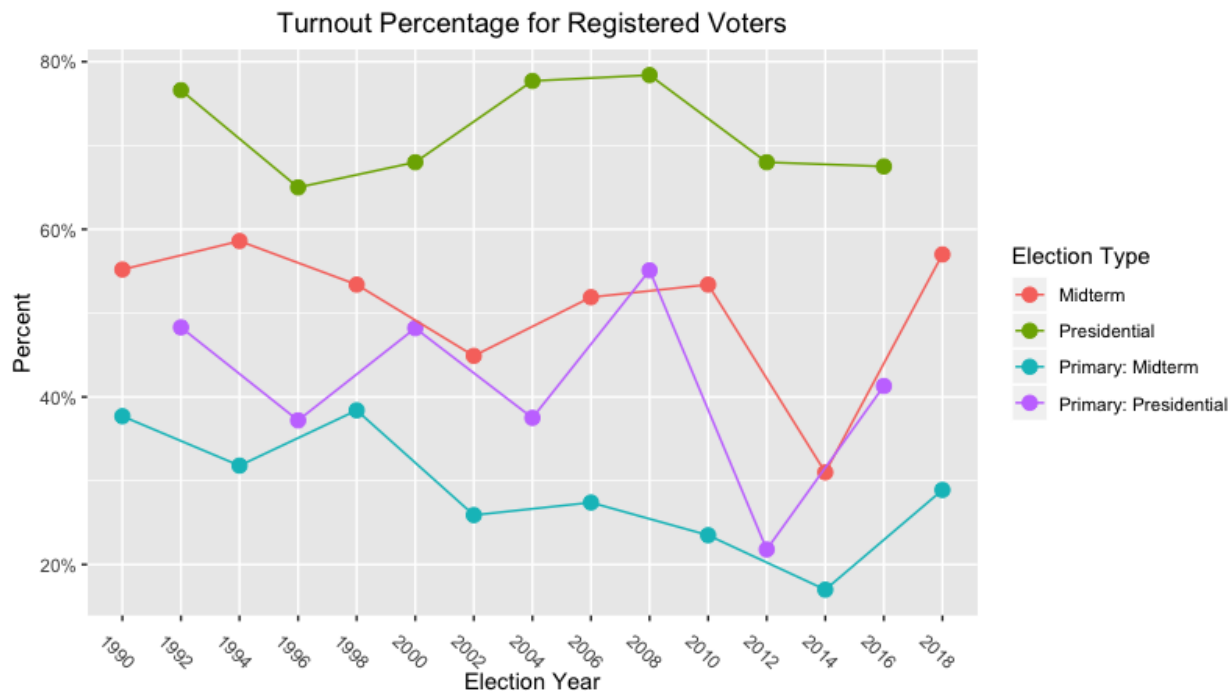


Figure 4.3: Percent turnout of registered voters in Los Angeles County dating back to 1990

4.2 Census Statistics

The R packages `tidycensus` and `ggplot2` were utilized for the purpose of exploring 2017 census tract level data. In performing this preliminary analysis, we begin to see which Angelenos are currently underrepresented in the voting process. Percentiles are mapped for ease of interpretation. In Figures 4.4 through 4.7, we see some visual support for the correlation between poverty, mobility, and turnout.

The maps of poverty and public transportation rate percentiles (Figure 4.4) are near inverses of that of voter turnout (Figures 4.5 through 4.7). In terms of turnout rate² percentiles, the maps of 2014, 2016, and 2018 look very similar. This suggests that even as absolute turnout goes up or down, the percentile rank does not change by too much.

²number of voters who turned out divided by the number of registered voters

In order to dig deeper, the 2017 census tracts were divided into two groups based on whether they were above or below the 50th percentile for voter turnout rates in the 2018 midterm election. The average yearly earnings of a worker in the higher turnout group were \$24,167 higher than the average yearly earnings of a worker in the lower turnout group. The higher turnout group reported a 10% poverty rate compared to 24% in the lower turnout group. The population of the higher turnout group was 40.1% minority compared to 55.7% in the lower turnout group. The higher turnout group had an average age of 40.6 compared to 33.1 in the lower turnout group. 13.1% of the voting-age population in the lower turnout group report having a college degree, while 40.4% of the voting-age population in the higher turnout group report having at least a bachelor's degree. 56.6% of voters in the higher turnout group utilized vote by mail compared to 52.9% in the lower turnout group. In short, census tracts that turn out in higher numbers tend to be populated by people that make more money, have higher levels of education, are older, and contain a smaller minority population than people that populate census tracts that turn out in lower numbers. The higher turnout group also utilizes vote by mail at a higher rate than the lower turnout group, which is consistent with previous studies that observed that older, more educated people tend to utilize early voting systems like vote by mail at higher rates.

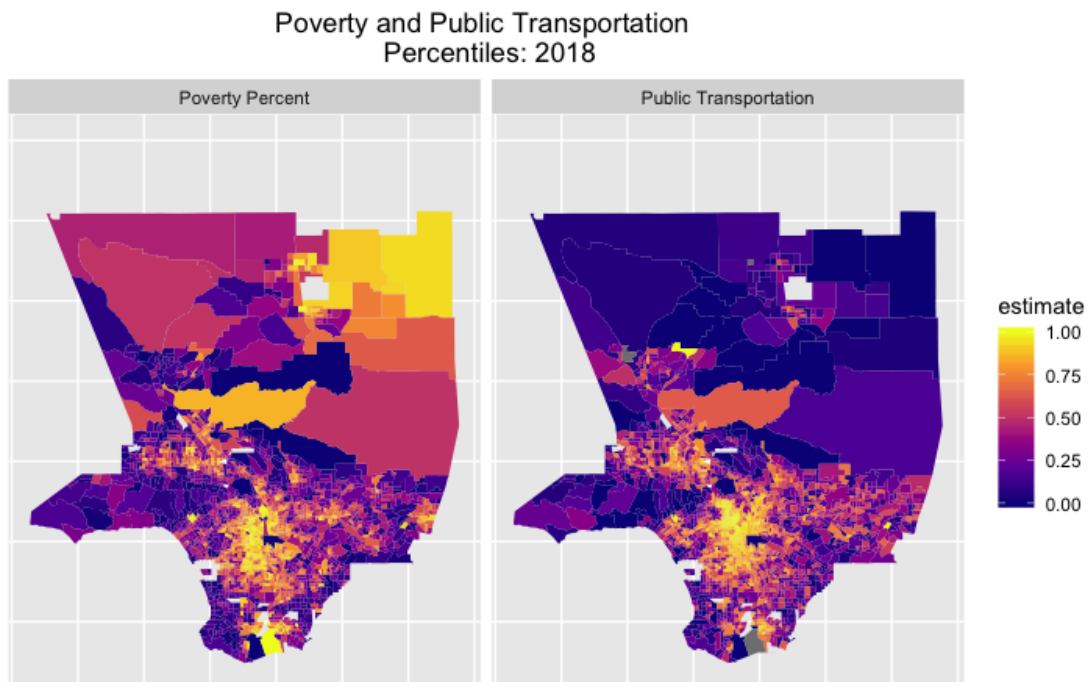


Figure 4.4: Poverty and public transportation percentiles: 2018

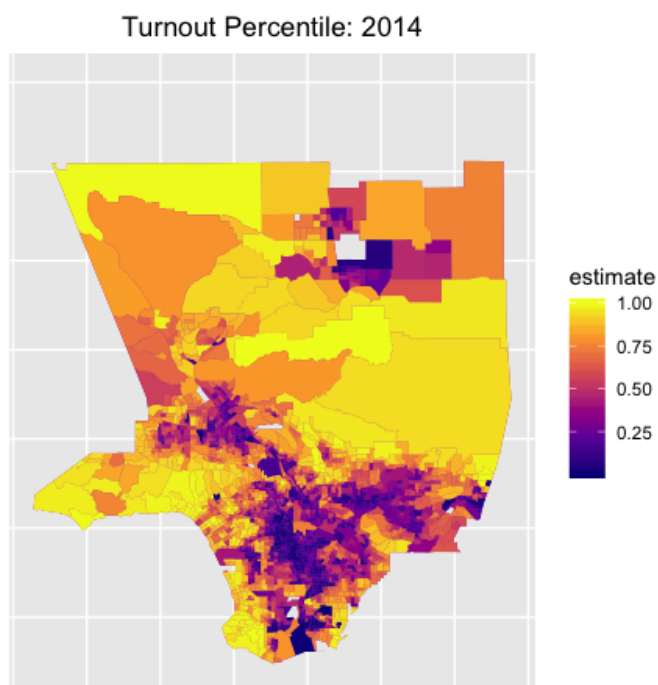


Figure 4.5: Turnout rate percentiles in 2014

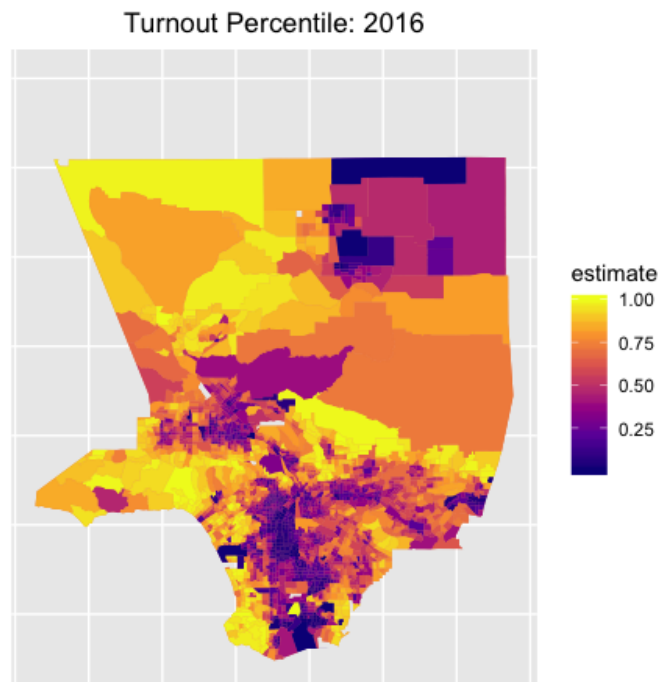


Figure 4.6: Turnout rate percentiles in 2016

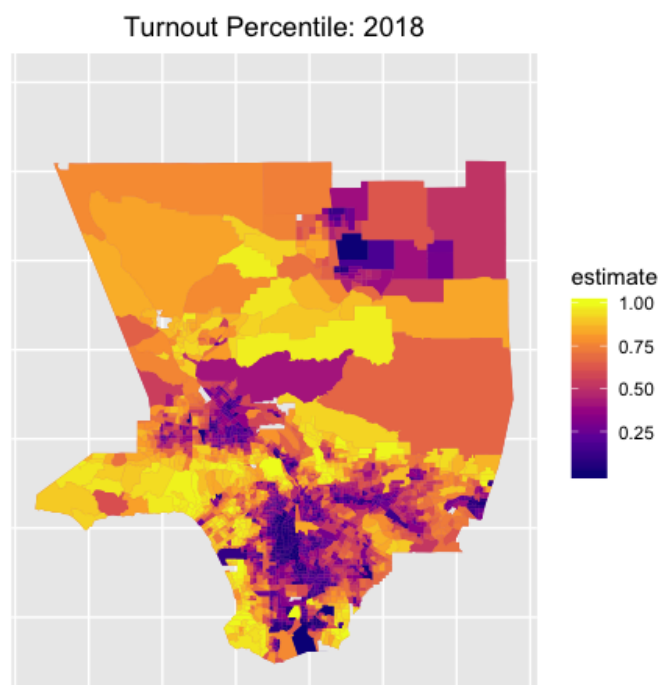


Figure 4.7: Turnout rate percentiles in 2018

CHAPTER 5

General Additive Models (GAMS)

In the current study, we rely on general additive models (GAMs) to model our response variables. GAMs have been chosen in order to take advantage of both their flexible modeling power and simple interpretation. Instead of assuming that the response variables are a linear combination of independent variables, GAMs model the dependent variable as a sum of smoothed functions of the independent variable [7].

$$Y_i = \alpha + \sum_{j=1}^p f_j(X_{ij}) + \epsilon_i$$

The fact that GAMs are additive allows for easy interpretation of the relationships between the predictor variables and the response variable by simply looking at plots of the curved relationships. This is useful for communication to nontechnical audiences. Furthermore, we can adjust the level of smoothness to avoid overly wiggly relationships to increase the interpretability and plausibility of the relationship, assuming that true relationships between variables are relatively smooth [8].

There are a number of methods that can be used to smooth the predictor variables and multiple smoothing methods can be used within the same model. Generally speaking, one can get to very similar results using different methods. In this paper, we use regression splines. Regression splines are built from a linear combination of basis functions that are not reliant on the Y variable. Put simply, the curved approximation of the relationship between our predictor and outcome variable is a weighted sum of k basis functions. The greater k is, the more wiggly the smoother is; therefore, we can use the k parameter to

control for the bias/variance trade-off of the smoother.

$$s(x) = \sum_{l=1}^K B_{l,q}(x)\beta_l = B'\beta$$

For our numeric predictors, we predominantly use thin plate splines. Thin plate splines measure a combination of R-squared and penalize for how much the function must bend through the data. This approach is called penalized regression splines and attempts to strike a compromise between fitting the data well without losing extrapolative powers due to overfitting. Thin plate splines penalize the wiggleness of the regression spline by minimizing the following function¹:

$$E_{\text{tps,smoth}}(f) = \sum_{i=1}^K \|y_i - f(x_i)\|^2 + \lambda \iint \left[\left(\frac{\partial^2 f}{\partial x_1^2} \right)^2 + 2 \left(\frac{\partial^2 f}{\partial x_1 \partial x_2} \right)^2 + \left(\frac{\partial^2 f}{\partial x_2^2} \right)^2 \right] dx_1 dx_2$$

In the case of proportion predictions, such as vote by mail percentages, we use the beta regression, which is useful when the outcome variable is continuous, but bounded between 0 and 1.

¹here, x is two dimensional for the sake of simplicity of explanation

CHAPTER 6

Vote By Mail Registration Rates

Vote by mail is of particular interest because it theoretically already provides an option to nearly completely alleviate the strains caused by day-of voting. It should then follow that as the percentage of registered voters that are registered to vote by mail increases, voter turnout percentage would follow a similar increasing trend. While this does occur, a rather large increase in vote by mail registration percentage corresponded to only a modest increase in voter turnout for the past two presidential elections. For the 2012 and 2016 presidential elections, there was a 12.9% increase in percentage of registered voters who were registered to vote by mail, but turnout increased only 2%, from 65% to 67%.

Voter registration has been steadily increasing along with the eligible population of voters in Los Angeles. At the same time, the total number of people actually voting and the percentage of vote by mail voters have steadily increased. However, turnout percentages have not shown the same upward trend. This is consistent with a study performed by fivethirtyeight[9] that showed that early voting reform has historically not resulted in increased turnout percentages, but instead just shifted when people who already voted went to the polls [REF].

It is possible that Los Angeles County is seeing a similar phenomenon, in that people that already vote are simply shifting from voting at the poll to voting by mail. This would account for the higher vote by mail registration and the higher turnout percentage among vote by mail voters, and it would also explain why we are not seeing large increases in overall turnout percentages as a result of increased vote by mail registration.

In other words, if people who are already voting at the polls migrate to vote by mail, we would expect to see higher turnout among vote by mail voters, but not higher overall turnout because it is not expanding the electorate - just redistributing the electorate to a new method of voting. This is still a positive outcome for Los Angeles County in terms of cost-saving and diminishing lines at the polling place, but there is little evidence to show that it is succeeding in expanding the electorate. Diving into the data to deepen our understanding of the relationships between the demographics of neighborhoods and the vote by mail registration percentage of each neighborhood will allow us to understand who is benefiting from vote by mail and who is averse to vote by mail. We can use this knowledge to craft outreach policies and roll-out strategies in regards to VSAP. It is important to note that the dependent variable we are observing is the percentage of those registered to vote who are registered to vote by mail. This dependent variable was chosen to best answer the question: Who are the registered voters that are hesitant to adopt vote by mail?

6.1 Smoothed Univariate Relationships Between Census Data and Vote By Mail Registration Rates

GAMs are well-suited to looking at social science problems when it is more important to understand the relationships between the predictor variables and the outcome variable. In other words, when the interpretation of the model by a non-statistical audience is just as, if not more, important than a powerful prediction, a GAM is a good choice to model the data. In this case, we can look at the smoothed relationships between different demographic characteristics of a census tract and the percentage of voters in that census tract that are registered to vote by mail to immediately gain visual insights as to what groups of individuals are more likely to vote by mail. From a predictive standpoint, we do not improve substantially beyond simply looking at the lag variable of the previous election, but from a descriptive standpoint, we can learn a sizeable amount of information about what popu-

lations are signing up for vote by mail and which populations could benefit from targeted outreach.

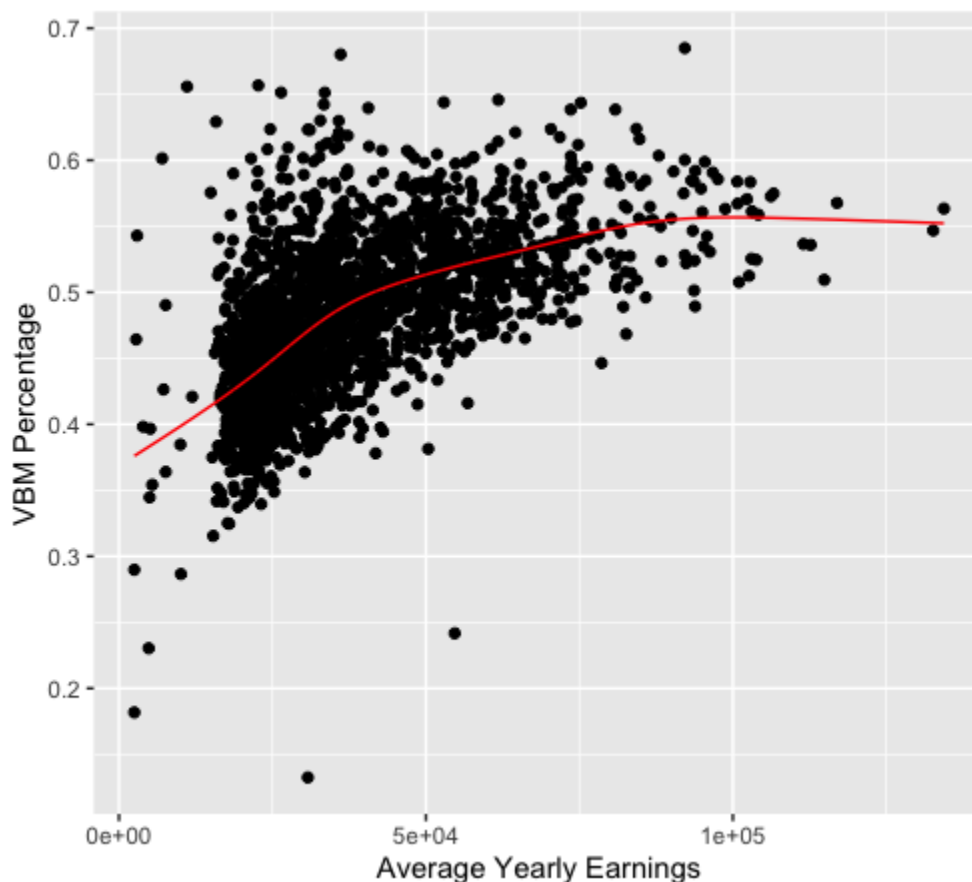


Figure 6.1: Univariate relationship plot of average income and VBM registration percentage

We see in figure 6.1 that as average income increases so does vote by mail registration, up to around \$80,000 at which point vote by mail registration tapers off.

We also see in Figure 6.2 that as poverty rates increase, vote by mail percentage decreases. This is discouraging, because vote by mail has the most potential to ease the burden of voting for impoverished census tracts. However, it appears that adoption of vote by mail has not caught on in these communities.

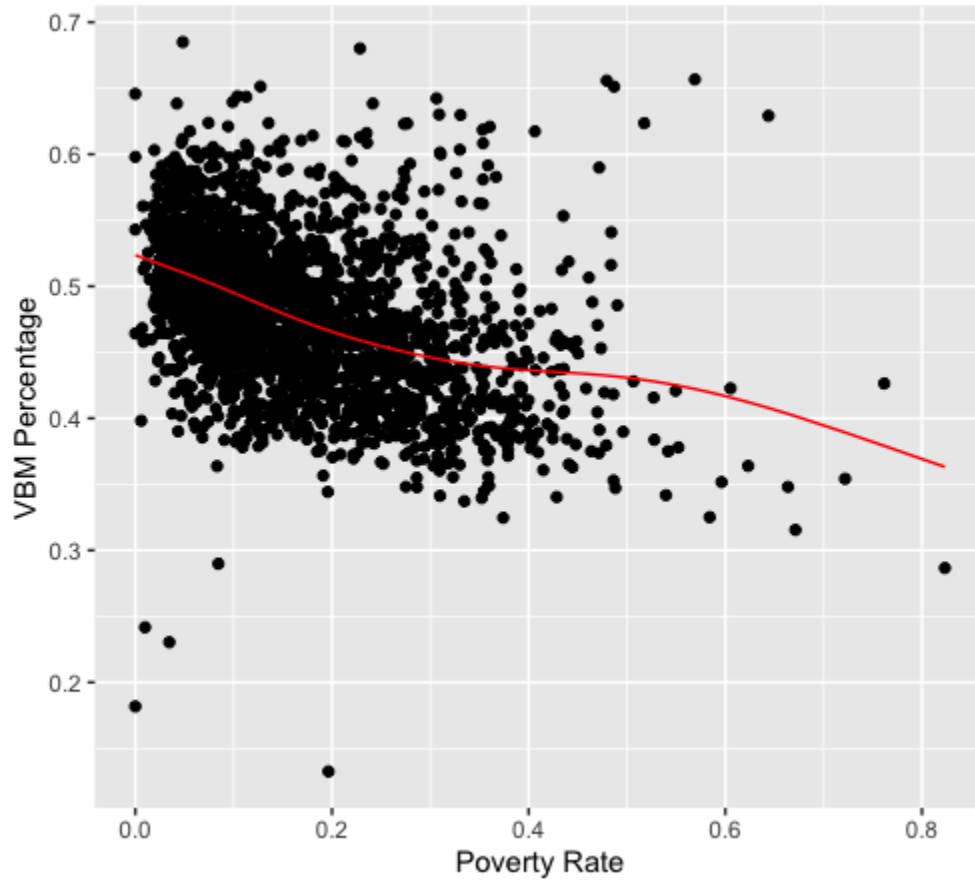


Figure 6.2: Univariate relationship plot of poverty rates and vote by mail registration percentage

Census tracts with an older population also tend to have higher vote by mail registration rates. Figure 6.3 shows a clear linear relationship between older average ages and vote by mail registration percentage.

Figure 6.4 shows that Hispanic communities tend to have lower vote by mail registration rates. The relationship appears to be linear. The relationship between African American population percentage and vote by mail registration does not appear to be strong in any particular direction. This is also the case when looking at Asian population percentages.

Census tracts with a higher rate of people with bachelors' degrees also have higher vote by mail turnout. Generally, census tracts with higher rates of drivers also have higher rates of

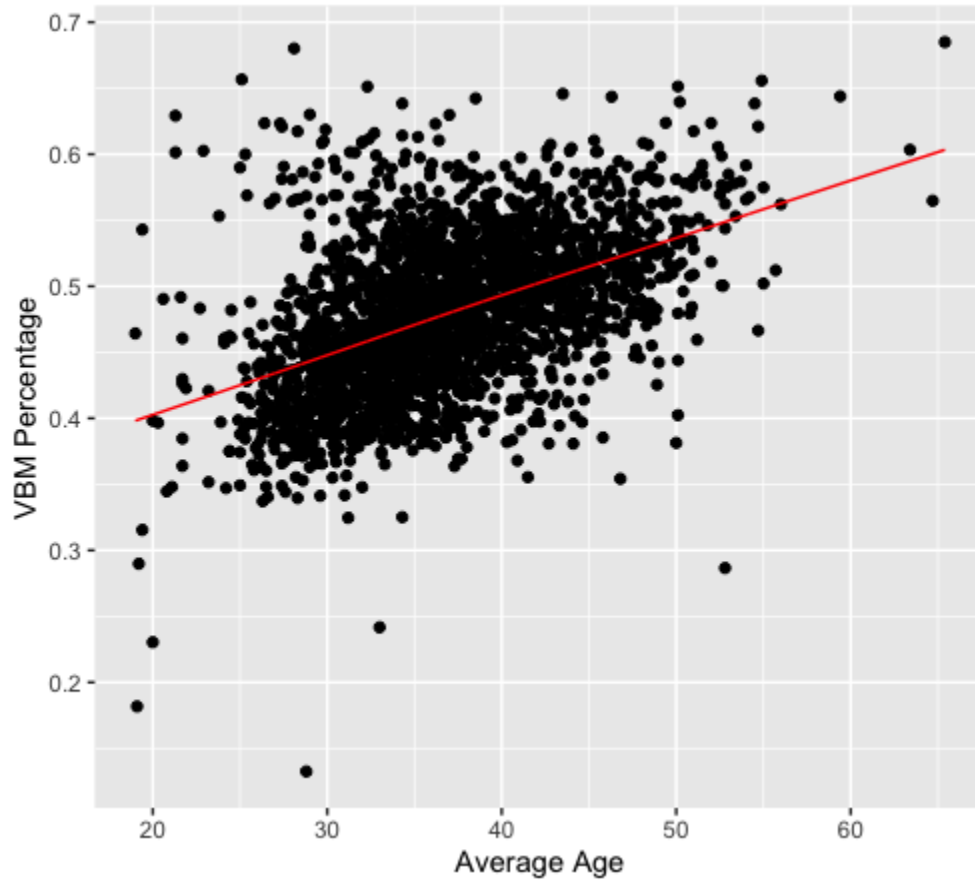


Figure 6.3: Univariate relationship plot of average age and vote by mail registration percentage

vote by mail registration.

This initial exploration supports the theory that vote by mail is predominantly being utilized by more affluent, older populations with higher levels of educational attainment. Census tracts with lower average incomes, lower educational attainment, lower rates of drivers, and younger populations have lower rates of registered voters who are registered to vote by mail.

A possible explanation for this is that people who intend on voting are switching to vote by mail for its convenience, but people who are registered to vote but do not intend to vote do not see any point in taking the extra step to change their registration status. This would

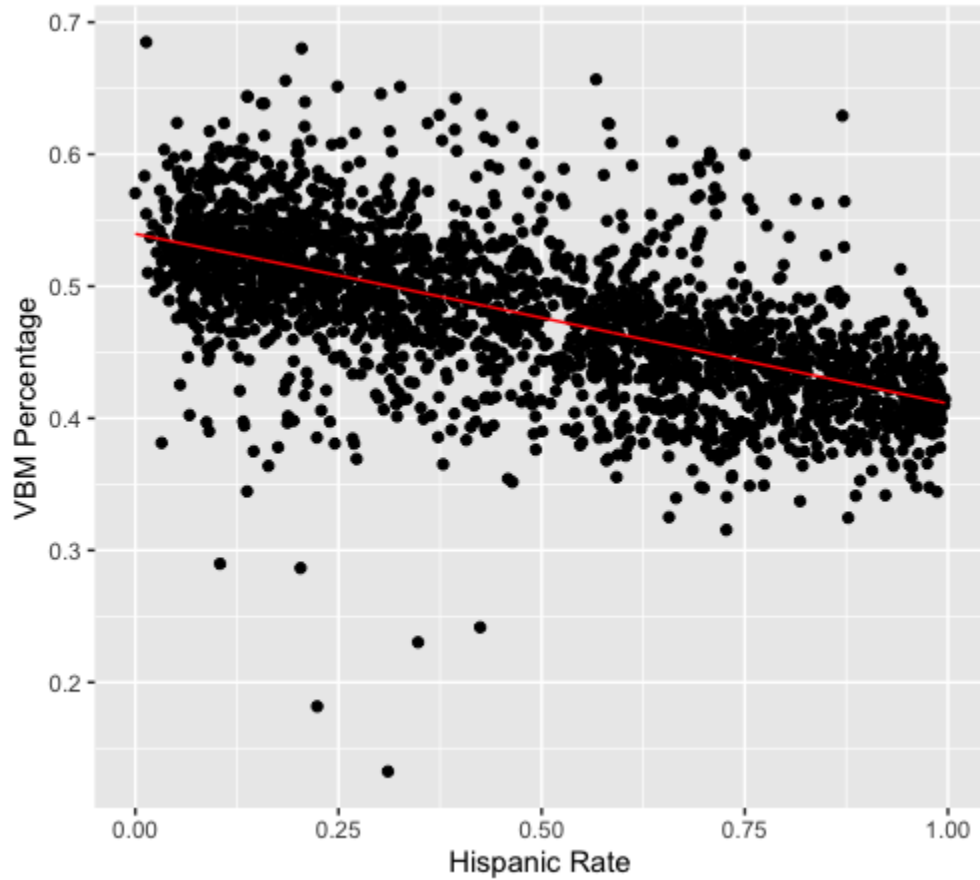


Figure 6.4: Univariate relationship plot of hispanic population rate and Vote By Mail registration percentage

explain why we see similar relationships between our predictor variables and vote by mail rate as we see between our predictor variables and voter turnout rates. Another explanation could be that certain populations are more suspicious of the vote by mail process. Whatever the explanation, if Los Angeles County wants to utilize vote by mail as a tool to expand the voting population, it needs to target outreach to communities that exhibit low turnout rates, which tend to be more impoverished communities.

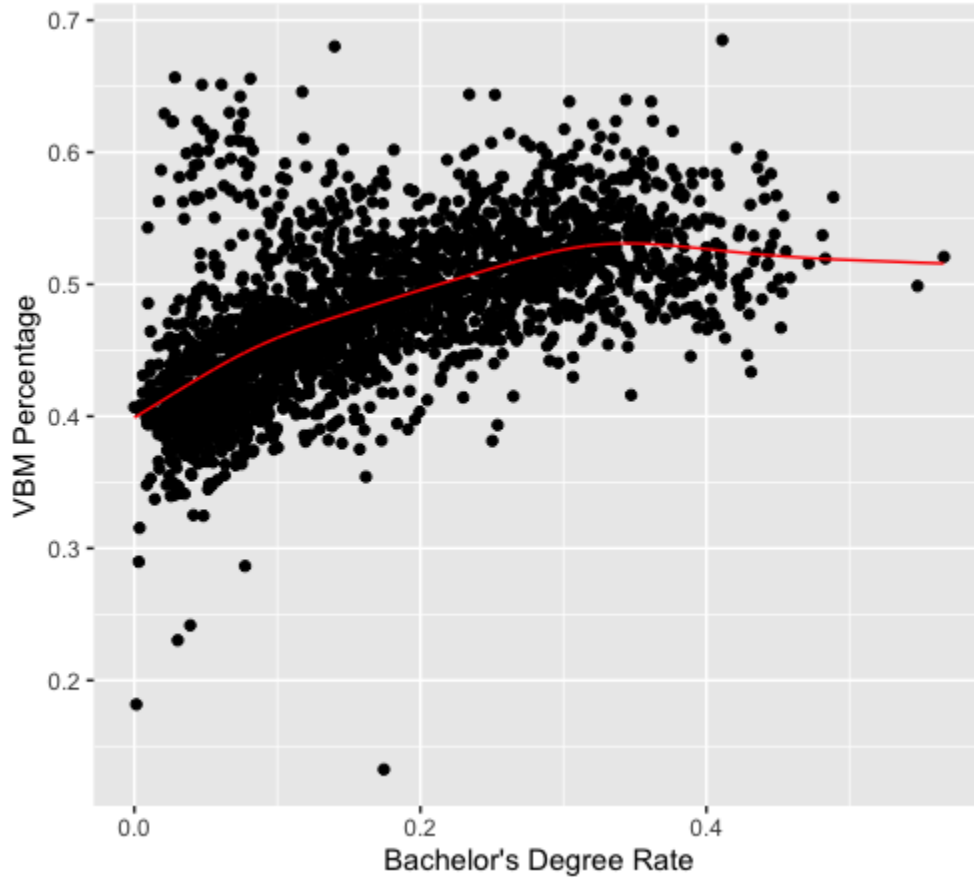


Figure 6.5: Univariate relationship plot of bachelors' degree attainment rate and vote by mail registration percentage

6.2 Forecasting Vote By Mail Registration Percentiles With GAMs

While the original goal of modeling this data was to predict vote by mail registration rates at the census tract level, out of training predictions proved to systematically under or over predict the vote by mail registration rate depending on the year.

Ideally, a multivariate time series approach could have been used to capture the trend of vote by mail registration but since reliable census tract level data only goes back four elections that approach was not a viable option. With this in mind, our first attempt was to build a model to predict vote by mail registration rates. The model was built using 2014 as the training data and a combination of demographic data points and a single lag variable (the

vote by mail registration rate from 2012). Then, the model was used to attempt to predict the 2018 vote by mail registration rates using 2016 as the lag variable. The results were disheartening, but not at all surprising; the model failed to capture the effect of the political shocks of 2016 and the increased voter activity that came after. This approach to modeling the data was taken because 2018 and 2014 were both midterm elections, so it seemed possible that the previous presidential election turnout combined with demographic data points would create a sufficiently good model to predict future voting activity.

In retrospect, the shortcomings of this approach may have been foreseeable; simply looking at the increase in vote by mail registration from 2012 to 2014 compared to the increase in vote by mail registration from 2016 to 2018, while being aware of the evolving political environment over those 6 years, one could assume that a model with only demographic and lag data would not likely capture the political shock of the election of Donald Trump. While the model did not do a good job of predicting absolute vote by mail registration rates, it did perform well when predicting relative vote by mail registration rates, answering the question: Where should we expect the highest and lowest rates of vote by mail registration?

6.3 Model Evaluation and Selection

As previously mentioned, the 2014 data was used with lags from 2012 to model vote by mail registration. The 2014 census tracts were divided into 75% training and 25% test data. Stepwise selection was used, in which a full model was fitted and insignificant variables were removed, comparing root mean squared error (RMSE), Akaike information criterion (AIC), and Bayesian information criterion (BIC) to select the model. Initially, the percentage of voters registered as vote by mail was used as the outcome variable.

The full model, using median age, median earnings, percent Hispanic, percent African American, percent white, percent Asian, percent living in poverty, percent with a bachelors' degree,

percent with a graduate degree, percent with a high school degree, percent of population that are children, and vote by mail percent from the preceding federal election as predictor variables, performed well in modeling the vote by mail percent at the census tract level in 2014. The RMSE for the test data was 0.02302, meaning that on average the predicted vote by mail percent for a census tract was only in error by 2.302%. When insignificant variables (percent white, percent Asian, and percent with a bachelors' degree) were stripped away, the resulting RMSE decreased to 2.3%, indicating that the full model was overfitting noise in the data and doing a worse job of generalizing outside of the training data. After multiple cycles of stripping away variables with the highest p-values, we obtained a model that simply used the single lag variable. This lag-only model achieved a RMSE of 2.4%. The lag only model has an R-squared of 0.816 compared to the full model, which has an R-squared of 0.836. When looking at BIC and AIC, the two criteria disagreed on which model was the best. The lag only model had the lowest BIC while the full model had the lowest AIC.

This leads us to the somewhat unsatisfying conclusion that a model using just the information from the previous election had nearly the same predictive power as a model that included the lag variable along with an array of demographic data points from the US Census. However, the true usefulness of the census data is not in its added predictive power, but instead in its descriptive power. Census data helps us answer the question: what are the demographics of census tracts that are adopting vote by mail and what are the demographics of the census tracts that are not adopting vote by mail? When the lag variable is removed, all of the demographic data points from the full model are significant other than the percent of population with a bachelors' degree, and the R-squared is 0.527. Adding the lag variable makes the demographic data points redundant in the model, but the lag variable alone does not provide us with much insight into what type of person is likely to vote by mail. It is important to look at both the univariate relationship plots and the final model to get a fuller picture of the breakdown of vote by mail voters and how they are distributed throughout LAC.

The major failing of this model is that it did not translate when looking at 2018. The relationship between the lag variable of 2014 and the VBM registration percent of 2014 was not consistent with that of the lag variable of 2018 and the VBM registration percent of 2018. Put simply, this model built from 2014 data systematically underpredicted vote by mail registration in 2018. The RMSE increased to 4.6% and nearly every prediction was an underestimate.

This clearly shows that demographic data alone fails to capture the current political environment, which is obviously a driving force for registration and turnout but is not easily measured by population parameters. The 2018 midterm election saw a 24% increase in turnout rate over the 2014 midterm, which is most likely a reflection of the response to the 2016 presidential election. This begs the question: if so much of registration and turnout hinges on variables that are not currently observable and have no historical reference points for the 2020 election (i.e., who the Democratic candidate will be or what the results of the impeachment investigation will be), what use can these models provide? The answer is that while these models do a poor job of predicting absolute values of vote by mail registration, they do a fairly good job of predicting ordinal values.

In other words, the model does a good job of answering the question: where do we expect to see the highest rates of vote by mail registration? even though it does a bad job of answering the question: what percent of registered voters in each census tract do we expect to be registered to vote by mail? This is a key outcome: understanding where the highest rates of vote by mail registration exist will provide a roadmap of how to best allocate resources for Los Angeles County.

When we changed the response variable to reflect the percentile of vote by mail registration, RMSE was only 0.046, which is to say that the previous vote by mail registration rate is a good indicator of what we will see in the next election. It is worth giving an example at this point to help clarify interpretation of what this model is predicting: If a given census tract

had the highest vote by mail registration percentage, it would be given a value of 1 because it is in the 100th percentile. If the model predicts 0.96 for that census tract the residual is 0.04.

The residuals are also i.i.d. normal with mean close to 0, so there is no systematic under- or over-prediction. The recent past is a good indicator of the future from a relative ranking standpoint, but not from an absolute vote by mail rate standpoint. We should look to the past election to understand where there will be larger vote by mail populations and thus smaller in-person turnout and we can look at census data to understand what these populations look like and what demographic groups and neighborhoods should be the focus of outreach for vote by mail registration.

CHAPTER 7

Voter Turnout Rates

When modelling voter turnout rates, we took the same approach as was used in modelling vote by mail registration percentage. However, for voter turnout rates, adding the census data into our model resulted in a significant improvement over simply using the lag of the previous election.

7.1 Smoothed Univariate Relationships Between Census Data and Voter Turnout Rates

We see in Figure 7.1 that as average yearly earnings increase, voter turnout percent also increases sharply up to an average yearly income of approximately \$75,000; above this income, the increase in voter turnout percent tapers off. This is a similar relationship as was observed in vote by mail rates, but this relationship is more pronounced.

Unsurprisingly, Figure 7.2 shows that voter turnout rates sharply decline as poverty rates of a community increase. This is consistent with previous research that shows a positive correlation between one's financial security and one's probability of voting.

Figure 7.3 shows that older populations have higher turnout rates of registered voters than census tracts with younger average ages.

Similar to vote by mail registration, we see in Figure 7.4 a negative relationship between the

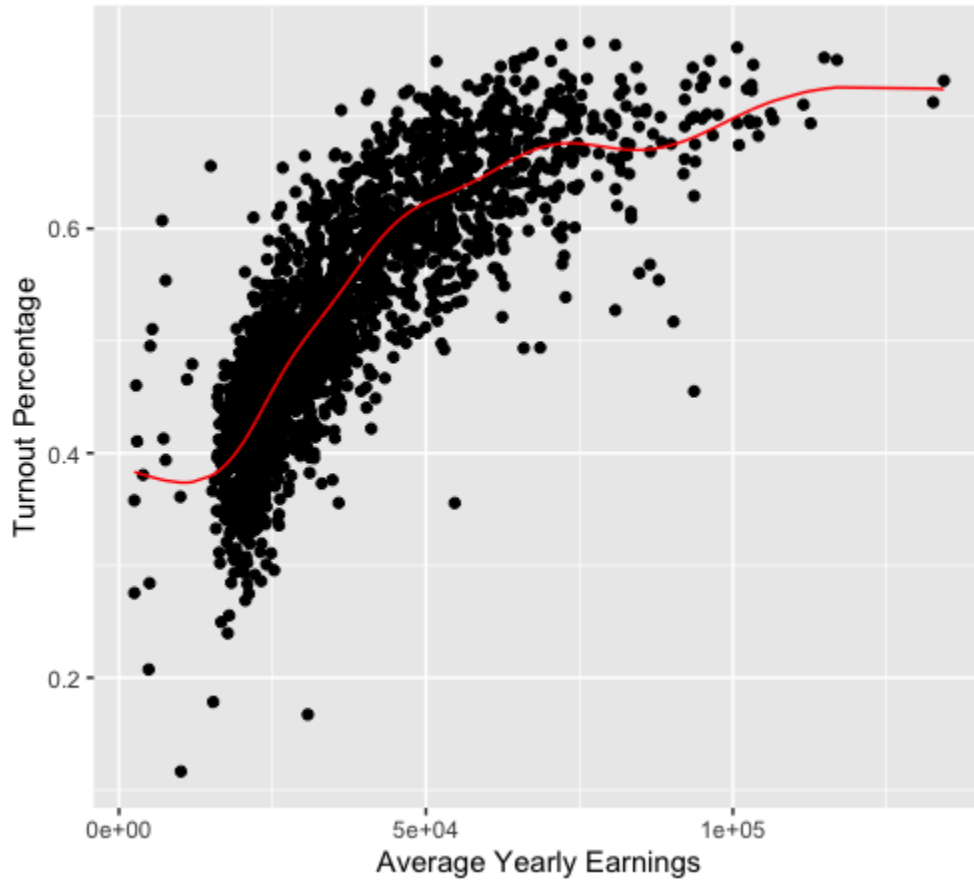


Figure 7.1: Univariate relationship plot of average income and voter turnout rate

density of Hispanic people in a census tract and the voter turnout rate in that census tract.

In Figures 7.5 and 7.6, we see that higher levels of educational attainment correspond to higher voter turnout rates.

Higher rates of children appear to correspond to higher turnout rates up to 10%; above this point, increased rates of children corresponds to lower turnout rates. This is consistent with the theory that finding childcare is a detriment to turning out to vote.

The univariate relationships between census tract data points and voter turnout rates are consistent with previous research in which older, more financially secure people who have attained higher levels of education turn out at higher rates than younger, less financially

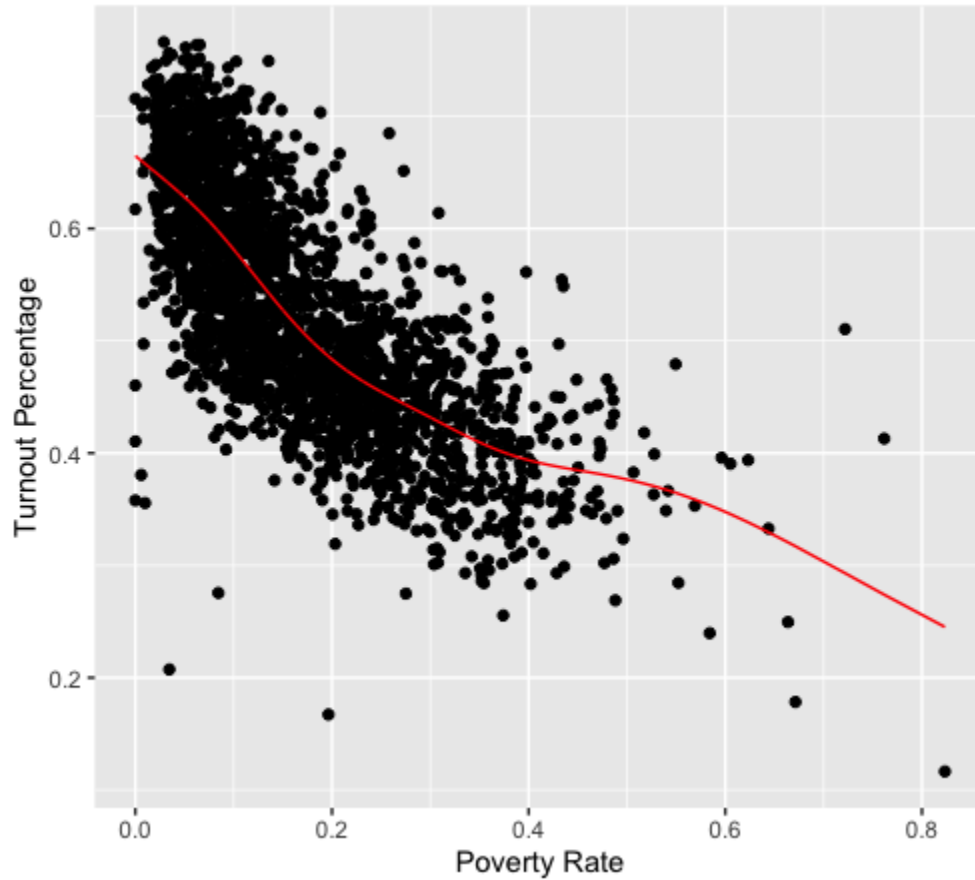


Figure 7.2: Univariate relationship plot of poverty rates and voter turnout rate

secure people who have not attained higher levels of education. It is worth noting that if we look at the correlation matrix between covariates, income is strongly positively correlated to both higher rates of bachelors' degree achievement and higher average ages. At this point a clear picture of the average Los Angeles voter begins to emerge: this voter is older and more financially secure than the average Angeleno, and as a result of their financial security they likely have access to better transportation and are less burdened by the one day voting sytem. These voters are also more likely to have college degrees and more likely to utilize vote by mail.

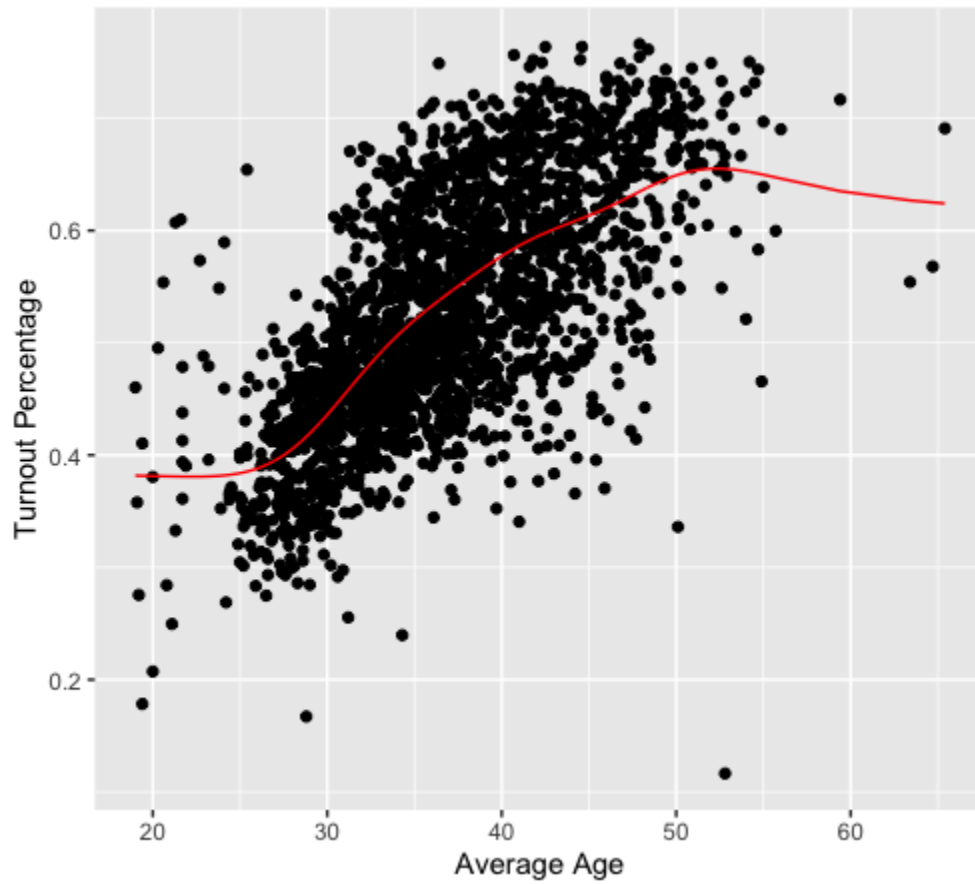


Figure 7.3: Univariate relationship plot of average age and voter turnout rate

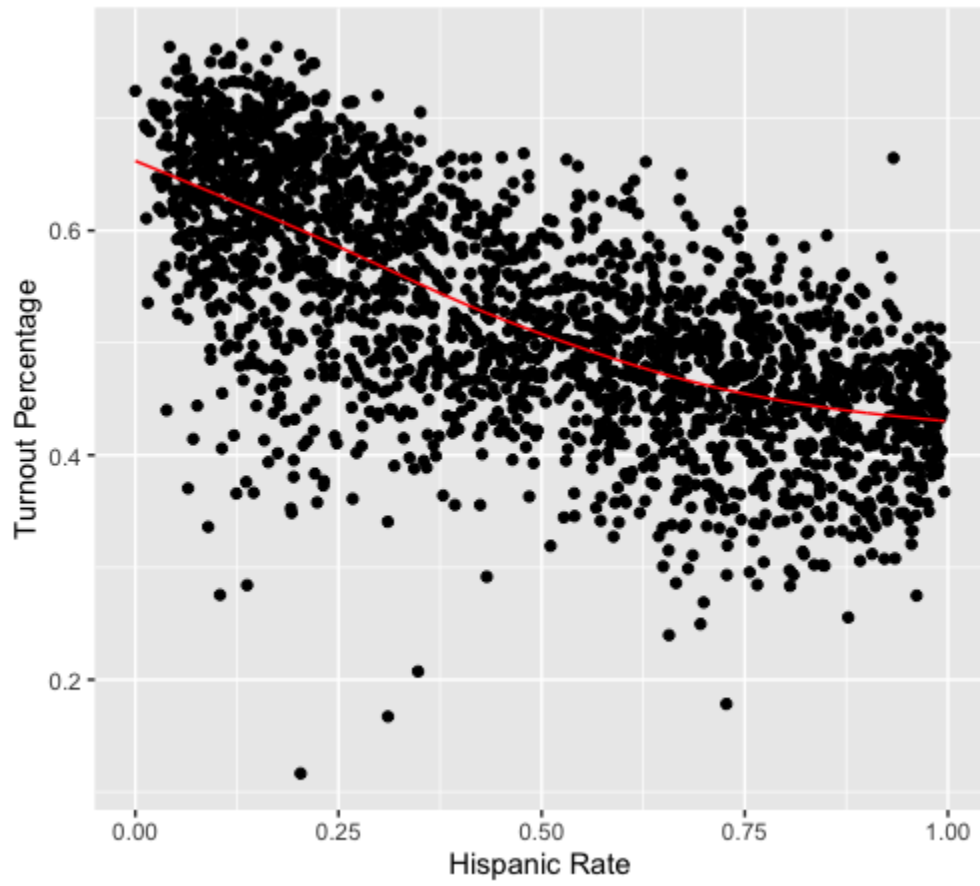


Figure 7.4: Univariate relationship plot of Hispanic population rate and voter turnout rate

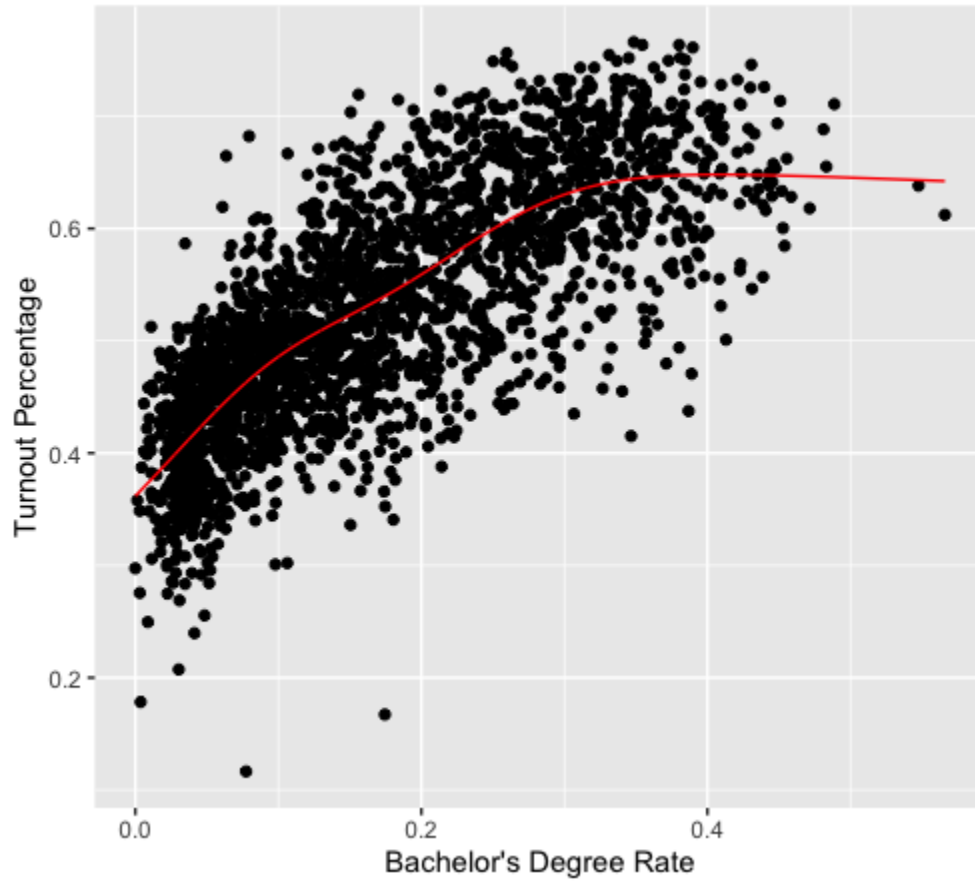


Figure 7.5: Univariate relationship plot of bachelors' degree attainment rate and voter turnout rate

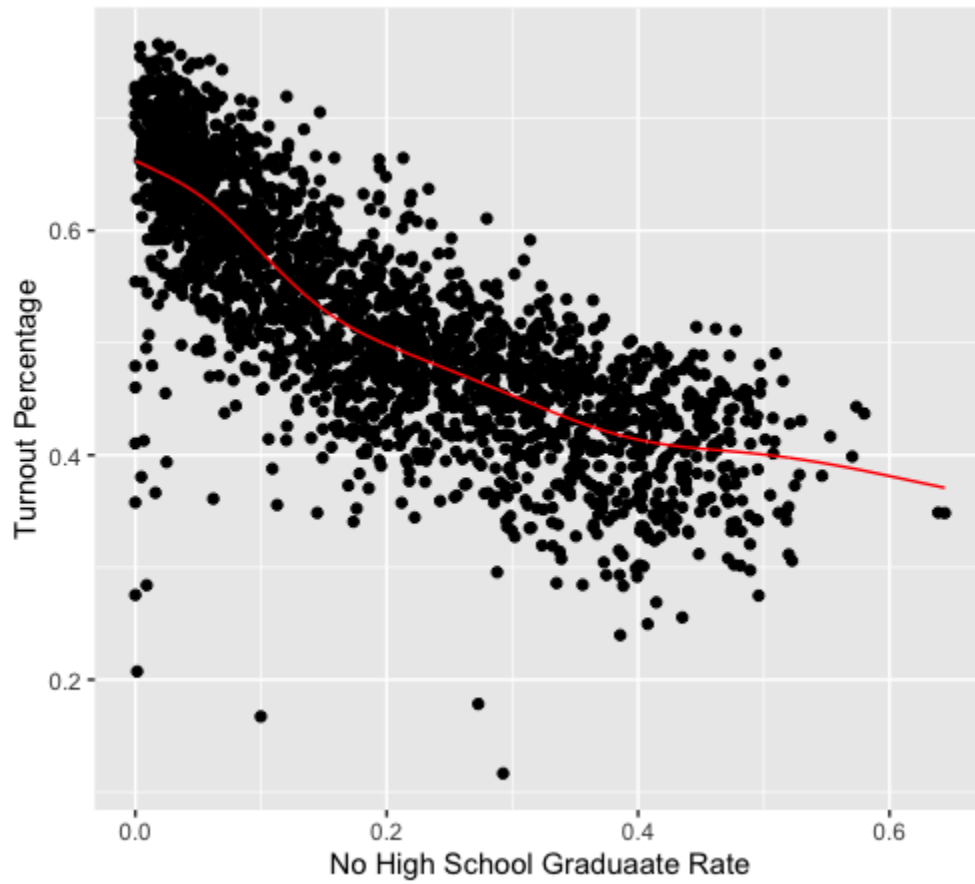


Figure 7.6: Univariate relationship plot of non-graduation from high school rate and voter turnout rate

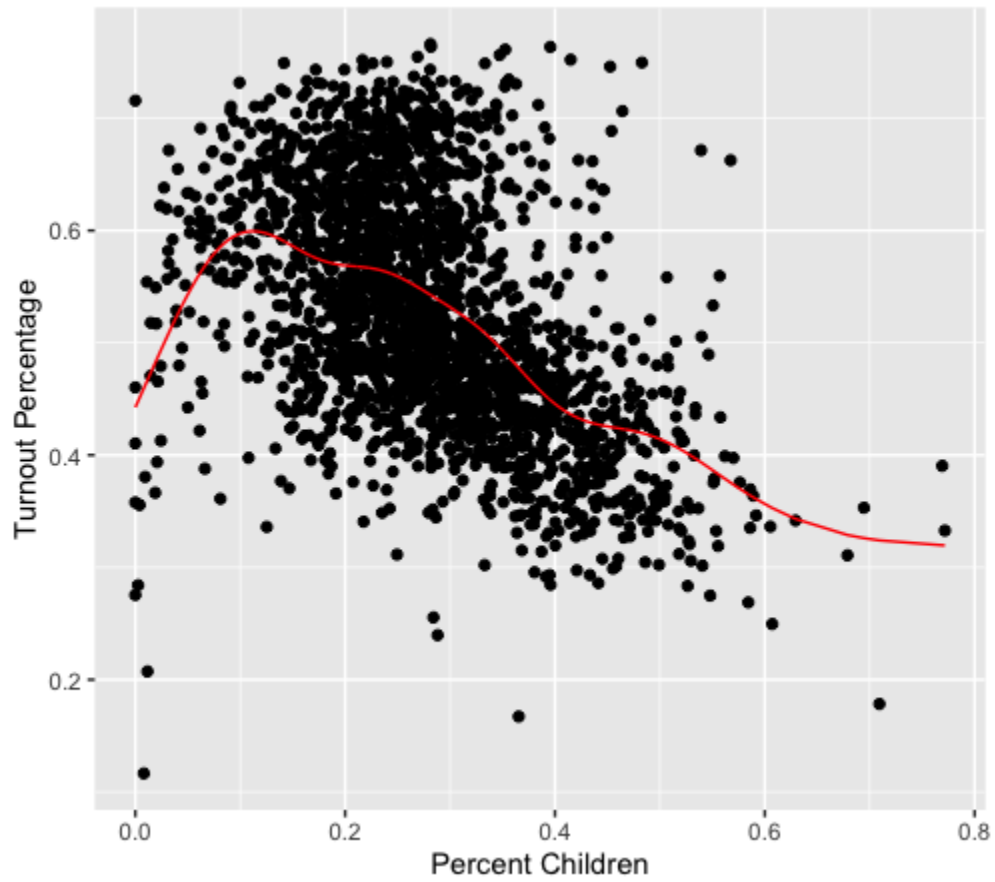


Figure 7.7: Univariate relationship plot of percentage of population that are children and voter turnout rate

7.2 Forecasting Voter Turnout Percentiles With GAMs

When building a model for turnout percent we once again took the approach of modelling 2014 turnout rates using age, earnings, Hispanic percent, African American percent, white percent, poverty percent, Asian percent, bachelors' degree percent, graduate school percent, high school graduation percent, no high school graduation percent, child percent, and turnout percent from the previous federal election. The RMSE of this full model was only 0.02517. We removed insignificant variables in a stepwise fashion and compared AIC, BIC, and RMSE. The result was that many redundant variables were removed (i.e., there were several variables that reflected educational attainment). The final model chosen utilized age, Hispanic percent, African American percent, poverty percent, Asian percent, bachelors' degree percent, graduate degree percent, child percent, and 2012 turnout percent. The final RMSE for the test data was .02519.

However, we once again see that while this approach does a good job of describing the past, it catastrophically fails when faced with shocks that cannot be accounted for in census and voter data. When we use the same model to predict 2018 turnout, the RMSE skyrockets to 0.2353 and systematically underpredicts the turnout rate.

This outcome led us back to taking an ordinal approach to modelling the data. We transformed the response variable from turnout percentage to turnout percentage percentile and we achieved results with normally distributed residuals and respectable RMSE of 0.09. The difference between modelling turnout percentage percentiles and vote by mail percentage percentiles is that in the case of turnout, census data had a meaningful impact on predicting future turnout. When we used the lag only model to predict turnout percentiles the RMSE went up to 0.098.

This result is encouraging because it shows that there are both descriptive and predictive benefits in joining census tract data to voter data. Demographic data, to an extent, not only

captures the relationship between demographics and turnout rates at a static point in time, but also captures information about changes in turnout rates from one election to the next.

CHAPTER 8

Projection Maps

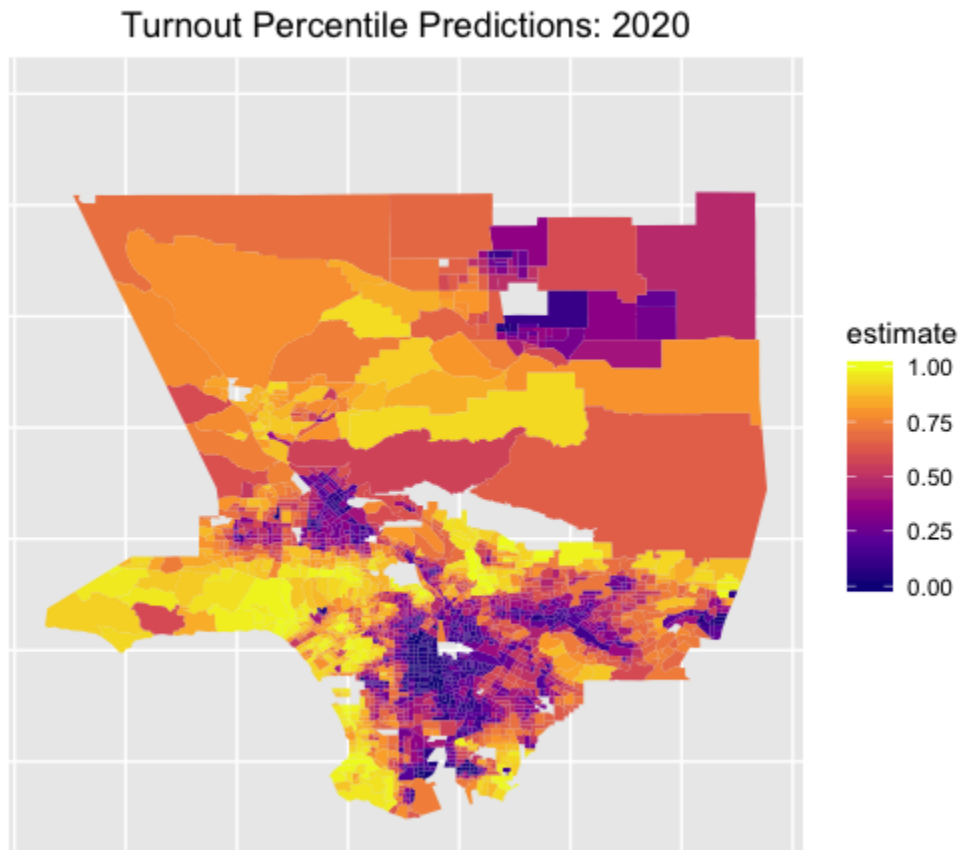


Figure 8.1: Projected turnout rate percentiles for 2020

When looking at our projected percentiles for vote by mail registration rate in Figure 8.2 and voter turnout rate in Figure 8.1 we see two maps that look fairly similar. This indicates that there is a positive correlation between being registered to vote by mail and one's probability of voting.

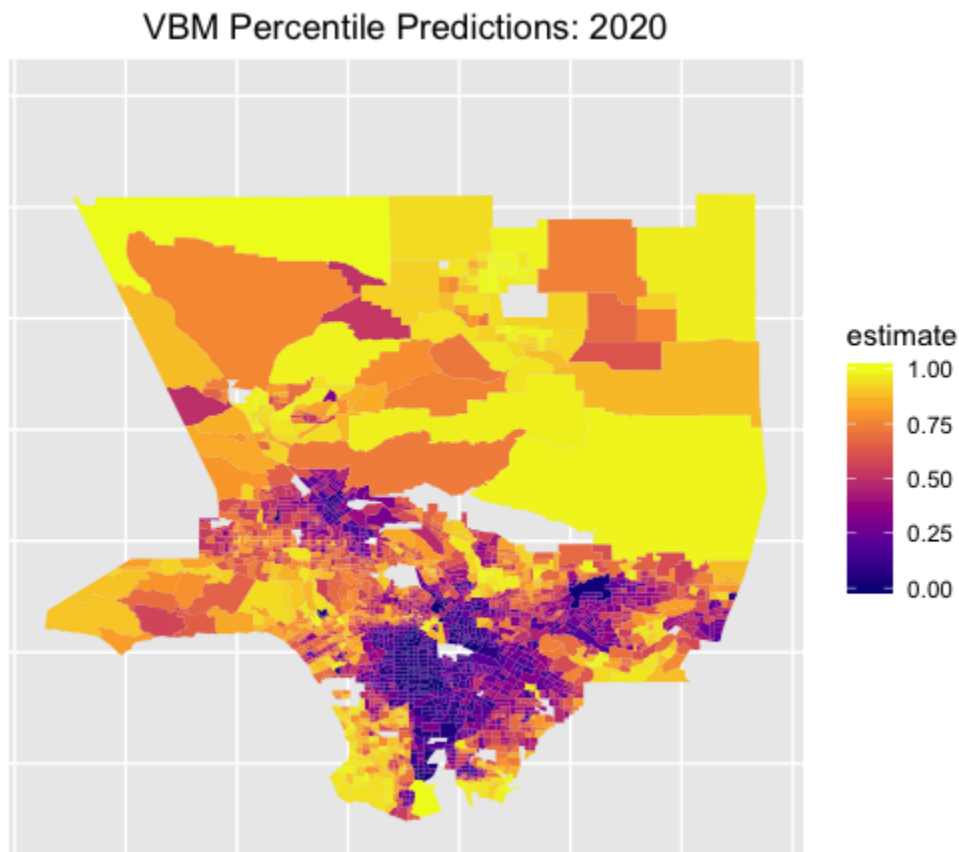


Figure 8.2: Projected vote by mail registration rate percentiles for 2020

When simply plotting the percentage of voters registered to vote by mail in a given census tract against the turnout percent for that census tract in Figure 8.3, we see that as vote by mail percentage increases so does the expected turnout percentage for that census tract. There are two possible interpretations of this: the first being that vote by mail is making it easier for people to vote and thus resulting in higher turnout, or, alternatively, people who are already likely to vote seek out easier ways to vote. Based on our interpretation of the univariate smoothed relationships presented in this paper it seems that the latter explanation is more likely.

Consistent with the earlier studies, we found that census tracts made up of financially secure, older people with higher levels of educational attainment were more likely to have

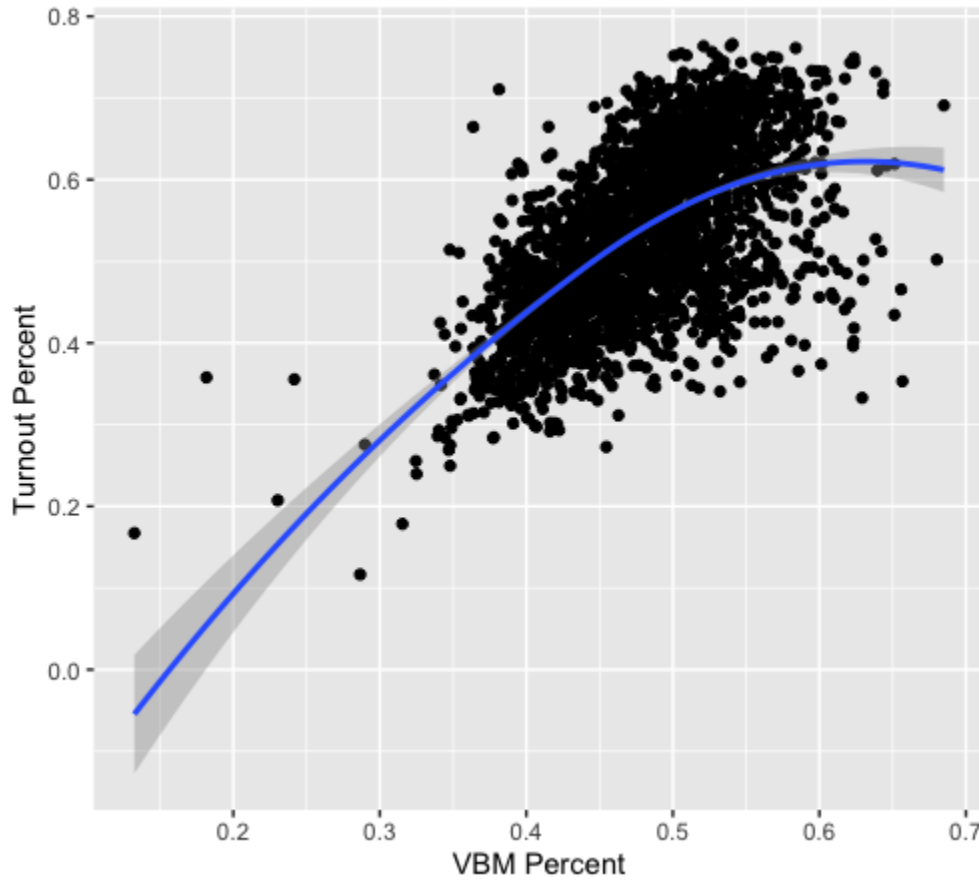


Figure 8.3: Smoothed relationship between vote by mail registration rates and voter turnout rate

higher turnout and vote by mail rates. Looking at our forecast maps for 2020 we see densely populated areas such as south and east Los Angeles that have both low vote by mail rates and low turnout rates. These are the neighborhoods that Los Angeles County should focus outreach on as potential areas to expand voter participation. These neighborhoods tend to be younger, more densely populated, and lower income neighborhoods compared to the rest of Los Angeles County. Also worth noting is that neighborhoods with higher percentages of Hispanic individuals have lower vote by mail registration and voter turnout.

From a preparation and outreach standpoint, Los Angeles County should be concerned about these more densely populated, poorer areas of South LA and East LA. These areas represent

the largest potential in expanding the voting population both in number and in diversity. Reaching out to the young, diverse, less financially secure Angelenos represents an exciting opportunity to have a more representative electorate in Los Angeles County, but it also could spell trouble from a logistical standpoint on election night. This is because studies have shown that people that are already likely to vote are the ones that turn out early. This paper has shown that people that are already likely to vote are more likely to be registered to vote by mail. In other words, getting more people to vote in south and east Los Angeles would also be expanding the population that is likely to show up in person on election night. Los Angeles County needs to make a concerted outreach effort to make these neighborhoods aware of the benefits of turning out to vote early. If this is not done, it is easy to imagine that these voters show up in unprecedented numbers based on the massive turnout rates seen in 2018 across all of Los Angeles County. If this occurs, there will be bottlenecks that disproportionately impact low income neighborhoods, which would be seen as an unequivocal failure in VSAP's rollout.

CHAPTER 9

Conclusion

In this paper, we have shown the power of joining census data to aggregated voter data in helping us understand the driving forces behind Los Angeles County voting behavior. We were able to visualize and understand the relationships between the demographics of a census tract and the voter turnout for that census tract.

We were able to build descriptive models that had low RMSE, but the models failed to predict the actual voter turnout percentage for years that they were not trained on. Put simply, voter data and census tract data did a good job of predicting which census tract would have high or low turnout compared to other census tracts, but failed to predict the magnitude of the turnout due to their inability to capture the effects of political shocks. This paper has successfully laid out a new methodology for building descriptive models of voter behavior using census data and GAMs; still, the failure of current modeling approaches to capture political shocks is significant. In future work, the effects of political shocks must be factored in to truly build accurate models of absolute voter turnout moving forward.

While we were only able to attain aggregate data for use in this study, Los Angeles County has individual level data that can be used to model the likelihood of an individual turning out to vote. As of now, Los Angeles County does not keep demographic data on individual voters, but much can be discerned from just one's name. The R package “gender” uses a large historical dataset of state recorded names and genders to allow the user to infer the gender of an individual based purely from their name. The R package “wru” takes a similar approach to ethnicity inference. Using these methods, along with joining in census data

using the R package “tidycensus”, would give Los Angeles County a rich understanding of an individual’s demographic makeup. Combining these factors with the individual’s voting history would allow Los Angeles County to build models to accurately predict the likelihood of an individual turning out to vote.

The most difficult part of the problem of predicting voter turnout is understanding and quantifying which candidates will generate higher turnout and by how much. In Los Angeles County - and in California at large - a Democrat will, with near certainty, win the state in the 2020 presidential election, but it is easy to imagine a large difference in voter turnout rates depending on the winner of the Democratic primary. It is also easy to imagine large numbers of voters showing up regardless of who the winner of the Democratic primary is simply because of the extremely polar views of the current president. In order to set bounds of expected turnout, polling should be performed to estimate the lower and upper bounds of expected turnout for different Democratic candidates. This information, combined with a rich data set of individual and census level demographic data will give Los Angeles County the best chance to accurately forecast absolute turnout in 2020 and best prepare themselves for the strains that will be placed on VSAP in its inaugural election.

Appendix A

R Code

A.1 Clean and Merge Data

```
1 #####
2 #This script is for collecting, cleaning, and joining data
3 #####
4
5
6
7 #####
8 #Load libraries
9 #####
10
11 library(dplyr)
12 library(tidyr)
13 library(stringr)
14 library(readr)
15 library(readxl)
16 library(tidycensus)
17 library(data.table)
18 library(ggplot2)
19
20 #####
21 #Load Data
22 #####
23 #Voting stats at a census tract level
24 setwd("~/Google Drive/masters thesis copy")
25 CT_votes <- read_csv("Paul_thesis.csv")
```

```

26
27 #Recode election ID
28 CT_votes <- CT_votes %>%
29   mutate(election_type = recode(as.character(election_id), "791" = "
      Presidential",
30                                     "905" = "Midterm",
31                                     "3496" = "Presidential",
32                                     "3861" = "Midterm")) %>%
33   mutate(year = ifelse(election_id == 791, 2012,
34                         ifelse(election_id == 905, 2014,
35                               ifelse(election_id == 3496, 2016,
36                                     ifelse(election_id == 3861, 2018, NA)
37                                     )))) %>%
38   #recode CT_votes party to DEM, REP, NPP, OTHER
39   mutate(party = ifelse(party == "DEM" | party == "REP" | party == "NPP",
40                         party, "Other")) %>%
41   group_by(election_id, ctract, perm_category, party, election_type, year)
42   %>%
43   dplyr::summarise(total.reg = sum(total.reg, na.rm = T), voted = sum(
44     voted, na.rm = T))
45
46 #####
47 #Explore NAs:Some precinct ids could not map to census tracts
48 #####
49
50 CT_NAs <- CT_votes %>%
51   filter(is.na(ctract))
52
53 #we see that data collection has improved over time and most every voter
54   can now be mapped to a
55
56 #ct
57
58 #####
59 #Report summary stats
60 #####

```

```

55 #Voting Percentage by party
56 View(CT_votes %>%
57     mutate(party_recode = ifelse(party == "DEM" | party == "REP" |
58     party == "NPP", party, "other")) %>%
59     group_by(year, party_recode) %>%
60     dplyr::summarise(votes = sum(voted), registered = sum(total.reg),
61     percent_voted = votes/registered))
62
63 #Voting percentage by PERM vs non-PERM
64 View(CT_votes %>%
65     group_by(year, perm_category) %>%
66     dplyr::summarise(votes = sum(voted), registered = sum(total.reg),
67     percent_voted = votes/registered))
68
69 #####
70 #Spread data
71 #####
72 #2018
73 CT_votes_2018 <- CT_votes %>%
74     filter(year == 2018) %>%
75     unite("party_count", c(perm_category, party), sep = "-") %>%
76     select(ctract, party_count, total.reg, voted)
77
78 CT_votes_2018 <- dcast(setDT(CT_votes_2018), ctract ~ paste0(party_count)
79     , value.var = c("total.reg", "voted"), sep = "-") %>%
80     replace(is.na(.), 0)
81 #####
82 #2016
83 CT_votes_2016 <- CT_votes %>%

```

```

84   filter(year == 2016) %>%
85   unite("party_count", c(perm_category, party), sep = "-") %>%
86   select(ctract, party_count, total.reg, voted)
87
88 CT_votes_2016 <- dcast(setDT(CT_votes_2016), ctract ~ paste0(party_count)
    , value.var = c("total.reg", "voted"), sep = "-") %>%
89   replace(is.na(.), 0)
90
91
92 #####
93 #2014
94 CT_votes_2014 <- CT_votes %>%
95   filter(year == 2014) %>%
96   unite("party_count", c(perm_category, party), sep = "-") %>%
97   select(ctract, party_count, total.reg, voted)
98
99 CT_votes_2014 <- dcast(setDT(CT_votes_2014), ctract ~ paste0(party_count)
    , value.var = c("total.reg", "voted"), sep = "-") %>%
100   replace(is.na(.), 0)
101
102
103
104 #####
105 #2012
106 CT_votes_2012 <- CT_votes %>%
107   filter(year == 2012) %>%
108   unite("party_count", c(perm_category, party), sep = "-") %>%
109   select(ctract, party_count, total.reg, voted)
110
111 CT_votes_2012 <- dcast(setDT(CT_votes_2012), ctract ~ paste0(party_count)
    , value.var = c("total.reg", "voted"), sep = "-") %>%
112   replace(is.na(.), 0)
113
114 #do this after joining census data
115 CT_votes_2012 <- CT_votes_2012 %>%

```

```

116 mutate(total_regs = rowSums(CT_votes_2012[,2:9])) %>%
117 mutate(total_votes = rowSums(CT_votes_2012[,10:17]))
118
119
120 #####
121 #Read in census data
122 #####
123
124 #API access
125 #http://api.census.gov/data/key_signup.html
126 census_api_key("b28ca17c1a5ffac1d0c7464d4df0ffa6a912ea7e", install = TRUE,
127               overwrite = T)
128 #Variable lookup 2017
129 v17 <- load_variables(2017, "acs5", cache = TRUE)
130 View(v17)
131 #variable lookup 2012
132 v12 <- load_variables(2012, "acs5", cache = TRUE)
133 View(v17)
134
135 #2017 variables of interest
136 ct_data <- get_acs(geography = "tract", state = "CA", county = "Los
137 Angeles County",
138                   variables = c(population = "B01003_001",
139                                #transportation
140                                #how many people take public
141                                transportation
142                                public_transportation = "B08006_008",
143                                #how many people drive to work
144                                drivers = "B08006_002",
145                                #how many people bike to work
146                                bike = "B08006_014",
147                                #poverty
148                                #below poverty
149                                poverty = "B06012_002",

```

```

148         #1-1.5x poverty
149         poverty_1.5x = "B06012_003",
150         #1.5x or above poverty
151         poverty_2x = "B06012_004",
152         #education
153         #no highschool diploma
154         edu_no_highschool = "B06009_002",
155         #highschool grad
156         edu_highschool = "B06009_003",
157         #bachelors degree
158         edu_bachelors = "B06009_005",
159         #grad school
160         edu_graduate = "B06009_006",
161         #employed 16 and older and in work force
162         employed = "B23025_004",
163         unemployed = "B23025_005",
164         total_labor_force = "B23025_002",
165         #race or ethnicity
166         african_american = "B02009_001",
167         european_american = "C02003_003",
168         asian_american = "C02003_006",
169         native_american = "C02003_005",
170         pacific_islander = "C02003_007",
171         hispanic_latino = "B03002_012",
172         #gender
173         total_male = "B01001_002",
174         total_female = "B01001_026",
175         #age
176         age_median = "B01002_001",
177         #earnings
178         earnings_median = "B24081_001",
179         #childcare
180         under_6 = "B05009_002",
181         over_6_under_18 = "B05009_020",
182         male_under_5 = "B01001_003",

```

```

183         male_5_9 = "B01001_004",
184         male_10_14 = "B01001_005",
185         male_15_17 = "B01001_006",
186         female_under_5 = "B01001_027",
187         female_5_9 = "B01001_028",
188         female_10_14 = "B01001_029",
189         female_15_17 = "B01001_030"),
190     year = 2017,
191     geometry = TRUE,
192     cb = F)
193
194 #2012 variables of interest
195 ct_data_2012 <- get_acs(geography = "tract", state = "CA", county = "Los
    Angeles County",
196     variables = c(population = "B01003_001",
197         #transportation
198         #how many people take public
199         transportation
200         public_transportation = "B08006_008",
201         #how many people drive to work
202         drivers = "B08006_002",
203         #how many people bike to work
204         bike = "B08006_014",
205         #poverty
206         #below poverty
207         poverty = "B06012_002",
208         #1-1.5x poverty
209         poverty_1.5x = "B06012_003",
210         #1.5x or above poverty
211         poverty_2x = "B06012_004",
212         #education
213         #no highschool diploma
214         edu_no_highschool = "B06009_002",
215         #highschool grad
216         edu_highschool = "B06009_003",

```



```

216         #bachelors degree
217         edu_bachelors = "B06009_005",
218         #grad school
219         edu_graduate = "B06009_006",
220         #employed 16 and older and in work force
221         employed = "B23025_004",
222         unemployed = "B23025_005",
223         total_labor_force = "B23025_002",
224         #race or ethnicity
225         african_american = "B02009_001",
226         european_american = "C02003_003",
227         asian_american = "C02003_006",
228         native_american = "C02003_005",
229         pacific_islander = "C02003_007",
230         hispanic_latino = "B03002_012",
231         #gender
232         total_male = "B01001_002",
233         total_female = "B01001_026",
234         #age
235         age_median = "B01002_001",
236         #earnings
237         earnings_median = "B24081_001",
238         #childcare
239         under_6 = "B05009_002",
240         over_6_under_18 = "B05009_020",
241         male_under_5 = "B01001_003",
242         male_5_9 = "B01001_004",
243         male_10_14 = "B01001_005",
244         male_15_17 = "B01001_006",
245         female_under_5 = "B01001_027",
246         female_5_9 = "B01001_028",
247         female_10_14 = "B01001_029",
248         female_15_17 = "B01001_030"),
249     year = 2012,
250     geometry = TRUE)

```

```

251
252
253
254
255 #####
256 #clean census data to join to voter data
257 #####
258
259
260 ct_data_clean <- ct_data %>%
261   select(-moe) %>%
262   spread(key = variable, value = estimate) %>%
263   #strip text from census number
264   mutate(ct_number = gsub("[^0-9]", "", NAME)) %>%
265   #add 00 to end of 4 digit numbers to match with county data
266   mutate(ct_number = ifelse(nchar(ct_number) == 4, paste0(ct_number, "00"),
267     ct_number)) %>%
268   #change to numeric
269   mutate(ct_number = as.numeric(ct_number)) %>%
270   #reorder columns
271   #change these numbers when you add variable from census data
272   select(37,1,3:36,38)
273
274 ct_data_2012_clean <- ct_data_2012 %>%
275   select(-moe) %>%
276   spread(key = variable, value = estimate) %>%
277   #strip text from census number
278   mutate(ct_number = gsub("[^0-9]", "", NAME)) %>%
279   #add 00 to end of 4 digit numbers to match with county data
280   mutate(ct_number = ifelse(nchar(ct_number) == 4, paste0(ct_number, "00"),
281     ct_number)) %>%
282   #change to numeric
283   mutate(ct_number = as.numeric(ct_number)) %>%
284   #reorder columns
285   #change these numbers when you add variable from census data

```

```

284   select(37,1,3:36,38)
285 #####
286 #model data
287 #Join census data to 2018 voter data
288 all_data_2018 <- CT_votes_2018 %>%
289   left_join(ct_data_clean, by = c("ctract" = "ct_number")) %>%
290   mutate(ctract = ifelse(ctract == 0, NA, ctract)) %>%
291   mutate(ctract = as.factor(ctract)) %>%
292   mutate(total.reg_PERM = rowSums(CT_votes_2018[,6:9])) %>%
293   mutate(total.reg_NON_PERM = rowSums(CT_votes_2018[,2:6])) %>%
294   mutate(total_regs = rowSums(CT_votes_2018[,2:9])) %>%
295   mutate(total.votes_PERM = rowSums(CT_votes_2018[,14:17])) %>%
296   mutate(total.votes_NON_PERM = rowSums(CT_votes_2018[,10:13])) %>%
297   mutate(total_votes = rowSums(CT_votes_2018[,10:17])) %>%
298   mutate(turnout_percent = total_votes/total_regs) %>%
299   mutate(of_voting_age = population - (female_10_14 + female_15_17 +
300     female_under_5 + female_5_9 +
301     male_10_14 + male_15_17 + male_
302     under_5 + male_5_9)) %>%
303   mutate(reg_percent = total_regs/of_voting_age)
304
305 #Join census data to 2016 voter data
306 all_data_2016 <- CT_votes_2016 %>%
307   left_join(ct_data_clean, by = c("ctract" = "ct_number")) %>%
308   mutate(ctract = ifelse(ctract == 0, NA, ctract)) %>%
309   mutate(ctract = as.factor(ctract)) %>%
310   mutate(total.reg_PERM = rowSums(CT_votes_2016[,6:9])) %>%
311   mutate(total.reg_NON_PERM = rowSums(CT_votes_2016[,2:6])) %>%
312   mutate(total_regs = rowSums(CT_votes_2016[,2:9])) %>%
313   mutate(total.votes_PERM = rowSums(CT_votes_2016[,14:17])) %>%
314   mutate(total.votes_NON_PERM = rowSums(CT_votes_2016[,10:13])) %>%
315   mutate(total_votes = rowSums(CT_votes_2016[,10:17])) %>%
316   mutate(turnout_percent = total_votes/total_regs) %>%
317   mutate(of_voting_age = population - (female_10_14 + female_15_17 +
318     female_under_5 + female_5_9 +

```

```

316                                     male_10_14 + male_15_17 + male_
      under_5 + male_5_9)) %>%
317 mutate(reg_percent = total_regs/of_voting_age)
318
319
320 #Join census data to 2014 voter data
321 all_data_2014 <- CT_votes_2014 %>%
322   left_join(ct_data_2012_clean, by = c("ctract" = "ct_number")) %>%
323   mutate(ctract = ifelse(ctract == 0, NA, ctract)) %>%
324   mutate(ctract = as.factor(ctract)) %>%
325   mutate(total.reg_PERM = rowSums(CT_votes_2014[,6:9])) %>%
326   mutate(total.reg_NON_PERM = rowSums(CT_votes_2014[,2:6])) %>%
327   mutate(total_regs = rowSums(CT_votes_2014[,2:9])) %>%
328   mutate(total.votes_PERM = rowSums(CT_votes_2014[,14:17])) %>%
329   mutate(total.votes_NON_PERM = rowSums(CT_votes_2014[,10:13])) %>%
330   mutate(total_votes = rowSums(CT_votes_2014[,10:17])) %>%
331   mutate(turnout_percent = total_votes/total_regs) %>%
332   mutate(of_voting_age = population - (female_10_14 + female_15_17 +
      female_under_5 + female_5_9 +
333                                     male_10_14 + male_15_17 + male_
      under_5 + male_5_9)) %>%
334   mutate(reg_percent = total_regs/of_voting_age)
335
336 #Join census data to 2014 voter data
337 all_data_2012 <- CT_votes_2012 %>%
338   left_join(ct_data_2012_clean, by = c("ctract" = "ct_number")) %>%
339   mutate(ctract = ifelse(ctract == 0, NA, ctract)) %>%
340   mutate(ctract = as.factor(ctract)) %>%
341   mutate(total.reg_PERM = rowSums(CT_votes_2012[,6:9])) %>%
342   mutate(total.reg_NON_PERM = rowSums(CT_votes_2012[,2:6])) %>%
343   mutate(total_regs = rowSums(CT_votes_2012[,2:9])) %>%
344   mutate(total.votes_PERM = rowSums(CT_votes_2012[,14:17])) %>%
345   mutate(total.votes_NON_PERM = rowSums(CT_votes_2012[,10:13])) %>%
346   mutate(total_votes = rowSums(CT_votes_2012[,10:17])) %>%
347   mutate(turnout_percent = total_votes/total_regs) %>%

```

```

348 mutate(of_voting_age = population - (female_10_14 + female_15_17 +
    female_under_5 + female_5_9 +
349                                     male_10_14 + male_15_17 + male_
    under_5 + male_5_9)) %>%
350 mutate(reg_percent = total_regs/of_voting_age)
351
352
353 #####
354 #LA almanac turnout data
355 #####
356 turnoutdata <- read_excel("turnoutdata.xlsx") %>%
357   #get order right by replacing Primary with "B"
358   mutate(order_reference = gsub("Primary", "B", Election)) %>%
359   mutate(election_year = gsub("[^0-9]+", "", Election)) %>%
360   arrange(order_reference) %>%
361   mutate(election_type = c(rep(c("Primary: Midterm", "Midterm", "Primary:
    Presidential", "Presidential"), 7),
362                           c("Primary: Midterm", "Midterm"))) %>%
363   mutate('Total Votes Cast' = as.numeric(gsub(",", "", 'Total Votes Cast')))
364
365 #####
366 #Data formatted for mapping
367 #####
368 #map friendly data
369 vote_data_2018 <- all_data_2018 %>%
370   select(c(1:17,54:62)) %>%
371   mutate(ctract = as.numeric(as.character(ctract)))
372
373 vote_data_2016 <- all_data_2016 %>%
374   select(c(1:17,54:62)) %>%
375   mutate(ctract = as.numeric(as.character(ctract)))
376
377 vote_data_2014 <- all_data_2014 %>%
378   select(c(1:17,54:62)) %>%
379   mutate(ctract = as.numeric(as.character(ctract)))

```

```

380
381 vote_data_2012 <- all_data_2012 %>%
382   select(c(1:17,54:62)) %>%
383   mutate(ctract = as.numeric(as.character(ctract)))
384
385 map_2018 <- ct_data_clean %>%
386   left_join(vote_data_2018, by = c("ct_number" = "ctract"))
387
388 map_2016 <- ct_data_clean %>%
389   left_join(vote_data_2016, by = c("ct_number" = "ctract"))
390
391 map_2014 <- ct_data_2012_clean %>%
392   left_join(vote_data_2014, by = c("ct_number" = "ctract"))
393
394 map_2012 <- ct_data_2012_clean %>%
395   left_join(vote_data_2012, by = c("ct_number" = "ctract"))
396
397
398 #####
399 #year over year per %
400 #####
401
402 perm_2012 <- all_data_2012 %>%
403   mutate(year = rep("2012"), nrow(all_data_2012)) %>%
404   dplyr::select(year, total_regs, total.reg_PERM)
405
406 perm_2014 <- all_data_2014 %>%
407   mutate(year = rep("2014"), nrow(all_data_2014)) %>%
408   dplyr::select(year, total_regs, total.reg_PERM)
409
410 perm_2016 <- all_data_2016 %>%
411   mutate(year = rep("2016"), nrow(all_data_2016)) %>%
412   dplyr::select(year, total_regs, total.reg_PERM)
413
414 perm_2018 <- all_data_2018 %>%

```

```

415 mutate(year = rep("2018"), nrow(all_data_2018)) %>%
416 dplyr::select(year, total_regs, total.reg_PERM)
417
418 perm_perc <- rbind(perm_2012, perm_2014, perm_2016, perm_2018)
419
420 perm_perc <- perm_perc %>%
421 group_by(year) %>%
422 summarise(percent_perm = sum(total.reg_PERM)/sum(total_regs))

```

A.2 Summary Statistics

```

1 #run thesis.r before running this script
2
3 library(leaflet)
4 library(sf)
5 library(tidyverse)
6
7
8 #####
9 #summary stats
10 #####
11
12 #overall turnout
13
14 #2012
15 sum(all_data_2012$total_votes)/sum(all_data_2012$total_regs)
16
17 #2014
18 sum(all_data_2014$total_votes)/sum(all_data_2014$total_regs)
19
20 #2016
21 sum(all_data_2016$total_votes)/sum(all_data_2016$total_regs)
22
23 #2018
24 sum(all_data_2018$total_votes)/sum(all_data_2018$total_regs)

```

```

25
26 #####
27 #timeseries plots of registration, voters, turnout%
28 #####
29 turnoutdata %>%
30   ggplot(aes(x = election_year, y= 'Turnout of Registered Voters', group =
      election_type)) +
31   geom_point(aes(color = election_type), size = 3) +
32   geom_line(aes(color = election_type)) +
33   labs(x = "Election Year", y= "Percent", title = "Turnout Percentage for
      Registered Voters", color = "Election Type") +
34   theme(plot.title = element_text(hjust = .5),
35         axis.text.x = element_text(angle = -45)) +
36   scale_y_continuous(labels = scales::percent_format(accuracy = 1)) ->
      turnout_percent
37
38 turnoutdata %>%
39   ggplot(aes(x = election_year, y= 'Total Votes Cast', group = election_
      type)) +
40   geom_point(aes(color = election_type), size = 3) +
41   geom_line(aes(color = election_type)) +
42   labs(x = "Election Year", y= "Total Votes", title = "Total Votes Cast",
      color = "Election Type") +
43   theme(plot.title = element_text(hjust = .5),
44         axis.text.x = element_text(angle = -45)) -> total_votes
45
46 turnoutdata %>%
47   filter(!grepl("Primary", election_type)) %>%
48   select(election_year, 'Persons Registered to Vote', 'Persons Eligible to
      Vote') %>%
49   gather(key = "count_type", value = "count", -election_year) %>%
50   arrange(election_year) %>%
51   ggplot() +
52   geom_point(aes(x = election_year, y= count, color = count_type), size =
      3) +

```



```

53 geom_line(aes(x = election_year, y= count, color = count_type, group =
    count_type)) +
54 labs(x = "Election Year", y= "Count", title = "Registered vs Eligible
    Voters", color = "") +
55 theme(plot.title = element_text(hjust = .5),
    axis.text.x = element_text(angle = -45)) -> eligible_registered
57
58 #####
59 #Census Level stats: Use leaflet/ggplot to look at turnout stats and
    demographic stats
60 #we need to look at this as a feature collection instead of a data frame
    for leaflet
61 #####
62
63 #####
64 #ggplot of some basic demographic data
65
66 map_2018 %>%
67   select(geometry,
68     african_american,
69     european_american,
70     hispanic_latino,
71     asian_american) %>%
72   gather(key = "variable", value = "estimate", -geometry) %>%
73   ggplot(aes(fill = estimate)) +
74   facet_wrap(~variable) +
75   geom_sf(color = NA) +
76   coord_sf(crs = 26911) +
77   scale_fill_viridis_c(option = "magma") +
78   theme(axis.text.x = element_blank(),
79     #axis.text.y = element_blank(),
80     axis.ticks = element_blank()) -> ethnicity_plot
81
82 map_2018 %>%
83   mutate('Poverty Percent' = poverty/population,

```

```

84     'Public Transportation' = public_transportation/total_labor_force
      ) %>%
85 filter(population > 0) %>%
86 select(geometry,
87     'Poverty Percent',
88     'Public Transportation') %>%
89 #standardize variables to a percent of the max
90 mutate('Poverty Percent' = percent_rank('Poverty Percent'),
91     'Public Transportation' = percent_rank('Public Transportation'))
      %>%
92 gather(key = "variable", value = "estimate", -geometry) %>%
93 ggplot(aes(fill = estimate)) +
94 facet_wrap(~variable) +
95 geom_sf(color = NA) +
96 coord_sf(ylim = c(33.7,35)) +
97 scale_fill_viridis_c(option = "plasma") +
98 theme(axis.text.x = element_blank(),
99     axis.text.y = element_blank(),
100     axis.ticks = element_blank(),
101     plot.title = element_text(hjust = .5)) +
102 labs(title = "Poverty and Public Transportation \n Percentiles: 2018")
      -> poverty_transportation
103
104
105 map_2018 %>%
106 select(geometry,
107     turnout_percent) %>%
108 mutate(turnout_percentile = percent_rank(turnout_percent)) %>%
109 #filter out wilderness
110 filter(turnout_percent > 0) %>%
111 gather(key = "variable", value = "estimate", -geometry) %>%
112 ggplot(aes(fill = estimate)) +
113 geom_sf(color = NA) +
114 coord_sf(ylim = c(33.7,35)) +
115 scale_fill_viridis_c(option = "plasma") +

```

```

116 theme(axis.text.x = element_blank(),
117         axis.text.y = element_blank(),
118         axis.ticks = element_blank(),
119         plot.title = element_text(hjust = .5)) +
120 labs(title = "Turnout Percentile: 2018") -> turnout_map_2018
121
122
123 map_2016 %>%
124   select(geometry,
125           turnout_percent) %>%
126   mutate(turnout_percentile = percent_rank(turnout_percent)) %>%
127   #filter out wilderness
128   filter(turnout_percent > 0) %>%
129   gather(key = "variable", value = "estimate", -geometry) %>%
130   ggplot(aes(fill = estimate)) +
131   geom_sf(color = NA) +
132   coord_sf(ylim = c(33.7,35)) +
133   scale_fill_viridis_c(option = "plasma") +
134   theme(axis.text.x = element_blank(),
135         axis.text.y = element_blank(),
136         axis.ticks = element_blank(),
137         plot.title = element_text(hjust = .5)) +
138   labs(title = "Turnout Percentile: 2016") -> turnout_map_2016
139
140 map_2014 %>%
141   select(geometry,
142           turnout_percent) %>%
143   mutate(turnout_percentile = percent_rank(turnout_percent)) %>%
144   #filter out wilderness
145   filter(turnout_percent > 0) %>%
146   gather(key = "variable", value = "estimate", -geometry) %>%
147   ggplot(aes(fill = estimate)) +
148   geom_sf(color = NA) +
149   coord_sf(ylim = c(33.7,35)) +
150   scale_fill_viridis_c(option = "plasma") +

```

```

151   theme(axis.text.x = element_blank(),
152         axis.text.y = element_blank(),
153         axis.ticks = element_blank(),
154         plot.title = element_text(hjust = .5)) +
155   labs(title = "Turnout Percentile: 2014") -> turnout_map_2014
156
157
158 #####
159 #summary stats for percentile turnout groups
160 #earnings
161 #age
162 #race
163 #public transport
164 #poverty
165 #PERM %
166 #####
167
168 all_data_2018 %>%
169   filter(population > 100) %>%
170   mutate(turnout_percentile = percent_rank(turnout_percent)) %>%
171   mutate(turnout_categorical = ifelse(turnout_percentile >=.5, "top_50", "
    bottom_50")) %>%
172   group_by(turnout_categorical) %>%
173   dplyr::summarise(avg_earnings = mean(earnings_median, na.rm = T),
174                   poverty_rate = mean(poverty/population),
175                   avg_age = mean(age_median, na.rm = T),
176                   minority_rate = 1 - mean(european_american/population),
177                   college = mean((edu_bachelors + edu_graduate)/ of_
    voting_age),
178                   perm_vote = mean(total_votes_PERM/total_votes))
179
180
181 all_data_2016 %>%
182   filter(population > 100) %>%
183   mutate(turnout_percentile = percent_rank(turnout_percent)) %>%

```

```

184 mutate(turnout_categorical = ifelse(turnout_percentile >=.5, "top_50", "
    bottom_50")) %>%
185 group_by(turnout_categorical) %>%
186 dplyr::summarise(avg_earnings = mean(earnings_median, na.rm = T),
187                 poverty_rate = mean(poverty/population),
188                 avg_age = mean(age_median, na.rm = T),
189                 minority_rate = 1 - mean(european_american/population))
190 all_data_2012 %>%
191 filter(population > 100) %>%
192 mutate(turnout_percentile = percent_rank(turnout_percent)) %>%
193 mutate(turnout_categorical = ifelse(turnout_percentile >=.5, "top_50", "
    bottom_50")) %>%
194 group_by(turnout_categorical) %>%
195 dplyr::summarise(avg_earnings = mean(earnings_median, na.rm = T),
196                 poverty_rate = mean(poverty/population),
197                 avg_age = mean(age_median, na.rm = T),
198                 minority_rate = 1 - mean(european_american/population))
199
200
201 #perm vs non perm
202
203 perm_votes <- CT_votes %>%
204   group_by(year, perm_category) %>%
205   dplyr::summarise(votes = sum(voted),
206                   registered = sum(total.reg),
207                   percent_voted = votes/registered)
208
209 #PERM registration percentage
210 perm_votes %>%
211   group_by(year) %>%
212   summarise(percent_perm = registered[perm_category == "PERM"]/(registered
    [perm_category == "NON-PERM"] +registered[perm_category == "PERM"]))

```

A.3 VBM GAMs

```

1 #####
2 #The goal of this script is to use GAMs to model
3 #the percentage of voter age population that are registered to
4 #vote by mail
5 #We will explore basic assumptions about barriers people face in voting
6 #and see if those are good predictors of if people are using VBM or not.
7 #We will attempt to use both beta regression to model the proportions
      (0,1)
8 #as well as modeling the absolute number.
9 #We will evaluate the strength and usefulness (interpretability) of each
      model.
10 #We will compare mse and aic and gcv
11 #cross validate
12 #####
13 #libraries
14 library(mgcv)
15 library(caret)
16 library(tigris)
17 library(spdep)
18 library(tidyverse)
19 library(PerformanceAnalytics)
20 library(gamclass)
21 library(caTools)
22
23 #load environment
24 load("/Users/pbeeman/Google Drive/masters thesis copy/masterenvironment.
      RData")
25
26 #2014 data manipulation
27 PERM_data <- all_data_2014 %>%
28   mutate(PERM_percent = total.reg_PERM/total_regs,
29           children = (over_6_under_18 + under_6)) %>%
30   filter(total_regs > 100) %>%
31   dplyr::select(african_american,
32                 age_median,

```

```

33         asian_american,
34         earnings_median,
35         edu_bachelors,
36         edu_graduate,
37         edu_highschool,
38         edu_no_highschool,
39         employed,
40         european_american,
41         hispanic_latino,
42         native_american,
43         pacific_islander,
44         population,
45         poverty,
46         total_female,
47         total_male,
48         unemployed,
49         children,
50         PERM_percent,
51         turnout_percent,
52         GEOID,
53         geometry
54     )
55
56 #add it lag variables
57
58 lags_2012 <- all_data_2012 %>%
59   filter(total_regs > 100) %>%
60   mutate(lags = total.reg_PERM/total_regs) %>%
61   dplyr::select(lags, GEOID)
62
63
64 PERM_data <- PERM_data %>%
65   left_join(lags_2012, "GEOID")
66
67

```

```

68 #quick exploration of the linear relationship between PERM registration
    and turnout
69 #we see that perm registration is positively coorelated to voter turnout
70
71 summary(lm(PERM_data$turnout_percent ~ PERM_data$PERM_percent))
72 hist(PERM_data$PERM_percent)
73 hist(PERM_data$turnout_percent)
74 plot(PERM_data$PERM_percent, PERM_data$turnout_percent)
75 abline(.02374, 1.05352, col = "red")
76
77 #latino
78 plot(PERM_data$hispanic_latino/PERM_data$population, PERM_data$PERM_
    percent)
79
80 #start modeling
81 PERM_data <- PERM_data %>%
82   dplyr::select(-turnout_percent) %>%
83   mutate(hispanic_percent = hispanic_latino/population,
84          black_percent = african_american/population,
85          white_percent = european_american/population,
86          pov_percent = poverty/population,
87          asian_percent = asian_american/population,
88          bachelor_percent = edu_bachelors/(population-children),
89          grad_percent = edu_graduate/(population - children),
90          hs_percent = edu_highschool/(population - children),
91          nohs_percent = edu_no_highschool/(population - children),
92          employed_percent = employed/(population - children),
93          unemployed_percent = unemployed/(population - children),
94          child_percent = children/(population - children)
95   ) %>%
96   na.omit()
97
98 #split into training/test data
99 sample_size <- floor(0.75 * nrow(PERM_data))
100 set.seed(123)

```



```

101 train_ind <- sample(seq_len(nrow(PERM_data)), size = sample_size)
102
103 PERM_data <- PERM_data[train_ind, ]
104 PERM_data_test <- PERM_data[-train_ind, ]
105
106
107 #####
108 #Markov Random Fields Geo approach
109 #####
110
111 #read in shap files for CTs
112 LA_shapes <- tracts("CA", "Los Angeles", cb = TRUE)
113
114
115 #check for empty geometries
116 which(is.na(st_dimension(PERM_data$geometry)))
117
118
119 #convert to neighborhood
120 nb <- poly2nb(PERM_data$geometry, row.names = PERM_data$GEOID)
121 names(nb) <- attr(nb, "region.id")
122 str(nb[1:6])
123
124 #fit GAM for PERM using MRF with a beta response
125 #https://stat.ethz.ch/R-manual/R-patched/library/mgcv/html/Beta.html
126
127 ctrl <- gam.control(nthreads = 6) # use 6 parallel threads, reduce if
    fewer physical CPU cores
128 m1 <- gam(PERM_percent ~ s(as.factor(GEOID), bs = 'mrf', k = 20, xt = list
    (nb = nb)), # define MRF smooth
129     data = PERM_data,
130     method = 'REML',
131     # fast version of REML smoothness selection
132     family = betar()) # fit a beta regression
133

```

```

134 viz_data <- PERM_data %>%
135   mutate(new_vals = predict(m1, type = 'response'))
136
137
138 summary(m1)
139 gam.check(m1)
140
141 #visualize
142 #cross validation
143 #https://rdrr.io/cran/gamclass/man/CVgam.html
144
145 #####
146 #Let's start looking at univariate relationships and plotting them out
147 #####
148 #####
149 #avg earnings
150 #####
151 earnings <- gam(PERM_percent ~ s(earnings_median) , data = PERM_data,
152               family = betar())
153 summary(earnings)
154 plot(earnings, residuals = T)
155
156 #partial effects plot
157 p <- predict(earnings, type="lpmatrix")
158 beta <- coef(earnings)[grepl("median", names(coef(earnings)))]
159 s <- p[,grepl("median", colnames(p))] %*% beta
160
161 ggplot(data=cbind.data.frame(s, PERM_data$earnings_median), aes(x=PERM_
162   data$earnings_median, y=s)) + geom_line()
163
164 #fitted vs predictors
165
166 fit_data <- data.frame(cbind(earnings$fitted.values, earnings$y, PERM_data
167   $earnings_median))
168 names(fit_data) <- c("yhat", "y", "earnings")
169

```

```

166 fit_data %>%
167   ggplot() +
168   geom_point(aes(x = earnings, y = y)) +
169   geom_line(aes(x = earnings, y = yhat), color= "red")
170
171 #####
172 #poverty percent
173 #####
174
175
176 pov <- gam(PERM_percent ~ s(pov_percent) , data = PERM_data, family =
      betar())
177 summary(pov)
178 plot(pov)
179
180 #partial effects plot
181 p <- predict(pov, type="lpmatrix")
182 beta <- coef(pov)[grepl("pov", names(coef(pov)))]
183 s <- p[,grepl("pov", colnames(p))] %*% beta
184 ggplot(data=cbind.data.frame(s, PERM_data$pov_percent), aes(x=PERM_data$
      pov_percent, y=s)) + geom_line()
185
186 #fitted vs predictors
187
188 fit_data <- data.frame(cbind(pov$fitted.values, pov$y, PERM_data$pov_
      percent))
189 names(fit_data) <- c("yhat", "y", "pov")
190
191 fit_data %>%
192   ggplot() +
193   geom_point(aes(x = pov, y = y)) +
194   geom_line(aes(x = pov, y = yhat), color= "red")
195
196 #####
197 #unemployed percent

```

```

198 #####
199
200
201 unemployed <- gam(PERM_percent ~ s(unemployed_percent) , data = PERM_data,
    family = betar())
202 summary(unemployed)
203 plot(unemployed)
204
205 #partial effects plot
206 p <- predict(unemployed, type="lpmatrix")
207 beta <- coef(unemployed)[grepl("unemployed", names(coef(unemployed)))]
208 s <- p[,grepl("unemployed", colnames(p))] %*% beta
209 ggplot(data=cbind.data.frame(s, PERM_data$unemployed_percent), aes(x=PERM_
    data$unemployed_percent, y=s)) + geom_line()
210
211 #fitted vs predictors
212
213 fit_data <- data.frame(cbind(unemployed$fitted.values, unemployed$y, PERM_
    data$unemployed_percent))
214 names(fit_data) <- c("yhat", "y", "unemployed")
215
216 fit_data %>%
217   ggplot() +
218   geom_point(aes(x = unemployed, y = y)) +
219   geom_line(aes(x = unemployed, y = yhat), color= "red")
220
221 #####
222 #age
223 #####
224
225 age <- gam(PERM_percent ~ s(age_median) , data = PERM_data, family = betar
    ())
226 summary(age)
227 plot(earnings)
228

```

```

229 #partial effects plot
230 p <- predict(age, type="lpmatrix")
231 beta <- coef(age)[grepl("median", names(coef(age)))]
232 s <- p[,grepl("median", colnames(p))] %*% beta
233 ggplot(data=cbind.data.frame(s, PERM_data$age_median), aes(x=PERM_data$age_
    _median, y=s)) + geom_line()
234
235 #fitted vs predictors
236
237 fit_data <- data.frame(cbind(age$fitted.values, age$y, PERM_data$age_
    median))
238 names(fit_data) <- c("yhat", "y", "age")
239
240 fit_data %>%
241   ggplot() +
242     geom_point(aes(x = age, y = y)) +
243     geom_line(aes(x = age, y = yhat), color= "red")
244
245 #we see a linear relationship between age and PERM %
246
247 #####
248 #hispanic percent
249 #####
250
251 hispanic <- gam(PERM_percent ~ s(hispanic_percent) , data = PERM_data,
    family = betar())
252 summary(hispanic)
253 plot(hispanic)
254
255 #partial effects plot
256 p <- predict(hispanic, type="lpmatrix")
257 beta <- coef(hispanic)[grepl("hispanic", names(coef(hispanic)))]
258 s <- p[,grepl("hispanic", colnames(p))] %*% beta
259 ggplot(data=cbind.data.frame(s, PERM_data$hispanic_percent), aes(x=PERM_
    data$hispanic_percent, y=s)) + geom_line()

```

```

260
261 #fitted vs predictors
262
263 fit_data <- data.frame(cbind(hispanic$fitted.values, hispanic$y, PERM_data
    $hispanic_percent))
264 names(fit_data) <- c("yhat", "y", "hispanic")
265
266 fit_data %>%
267   ggplot() +
268   geom_point(aes(x = hispanic, y = y)) +
269   geom_line(aes(x = hispanic, y = yhat), color= "red")
270
271 #we see a linear relationship between age and PERM %
272
273 #####
274 #black percent
275 #####
276
277
278 black <- gam(PERM_percent ~ s(black_percent) , data = PERM_data, family =
    betar())
279 summary(black)
280 plot(black)
281
282 #partial effects plot
283 p <- predict(black, type="lpmatrix")
284 beta <- coef(black)[grep1("black", names(coef(black)))]
285 s <- p[,grep1("black", colnames(p))] %*% beta
286 ggplot(data=cbind.data.frame(s, PERM_data$black_percent), aes(x=PERM_data$
    black_percent, y=s)) + geom_line()
287
288 #fitted vs predictors
289
290 fit_data <- data.frame(cbind(black$fitted.values, black$y, PERM_data$black
    _percent))

```

```

291 names(fit_data) <- c("yhat", "y", "black")
292
293 fit_data %>%
294   ggplot() +
295   geom_point(aes(x = black, y = y)) +
296   geom_line(aes(x = black, y = yhat), color= "red")
297
298
299 #####
300 #asian percent
301 #####
302
303
304 asian <- gam(PERM_percent ~ s(asian_percent) , data = PERM_data, family =
      betar())
305 summary(asian)
306 plot(asian)
307
308 #partial effects plot
309 p <- predict(asian, type="lpmatrix")
310 beta <- coef(asian)[grepl("asian", names(coef(asian)))]
311 s <- p[,grepl("asian", colnames(p))] %*% beta
312 ggplot(data=cbind.data.frame(s, PERM_data$asian_percent), aes(x=PERM_data$
      asian_percent, y=s)) + geom_line()
313
314 #fitted vs predictors
315
316 fit_data <- data.frame(cbind(asian$fitted.values, asian$y, PERM_data$asian
      _percent))
317 names(fit_data) <- c("yhat", "y", "asian")
318
319 fit_data %>%
320   ggplot() +
321   geom_point(aes(x = asian, y = y)) +
322   geom_line(aes(x = asian, y = yhat), color= "red")

```

```

323
324
325 #####
326 #hs percent
327 #####
328
329
330 hs <- gam(PERM_percent ~ s(hs_percent) , data = PERM_data, family = betar
      ())
331 summary(hs)
332 plot(hs)
333
334 #partial effects plot
335 p <- predict(hs, type="lpmatrix")
336 beta <- coef(hs)[grepl("hs", names(coef(hs)))]
337 s <- p[,grepl("hs", colnames(p))] %*% beta
338 ggplot(data=cbind.data.frame(s, PERM_data$hs_percent), aes(x=PERM_data$hs_
      percent, y=s)) + geom_line()
339
340 #fitted vs predictors
341
342 fit_data <- data.frame(cbind(hs$fitted.values, hs$y, PERM_data$hs_percent)
      )
343 names(fit_data) <- c("yhat", "y", "hs")
344
345 fit_data %>%
346   ggplot() +
347     geom_point(aes(x = hs, y = y)) +
348     geom_line(aes(x = hs, y = yhat), color= "red")
349
350 #doesn't appear too meaningful
351
352 #####
353 #bachelor percent
354 #####

```



```

355
356
357 bachelor <- gam(PERM_percent ~ s(bachelor_percent) , data = PERM_data,
    family = betar())
358 summary(bachelor)
359 plot(bachelor)
360
361 #partial effects plot
362 p <- predict(bachelor, type="lpmatrix")
363 beta <- coef(bachelor)[grepl("bachelor", names(coef(bachelor)))]
364 s <- p[,grepl("bachelor", colnames(p))] %*% beta
365 ggplot(data=cbind.data.frame(s, PERM_data$bachelor_percent), aes(x=PERM_
    data$bachelor_percent, y=s)) + geom_line()
366
367 #fitted vs predictors
368
369 fit_data <- data.frame(cbind(bachelor$fitted.values, bachelor$y, PERM_data
    $bachelor_percent))
370 names(fit_data) <- c("yhat", "y", "bachelor")
371
372 fit_data %>%
373   ggplot() +
374   geom_point(aes(x = bachelor, y = y)) +
375   geom_line(aes(x = bachelor, y = yhat), color= "red")
376
377 #####
378 #nohs percent
379 #####
380
381
382 nohs <- gam(PERM_percent ~ s(nohs_percent) , data = PERM_data, family =
    betar())
383 summary(nohs)
384 plot(nohs)
385

```

```

386 #partial effects plot
387 p <- predict(nohs, type="lpmatrix")
388 beta <- coef(nohs)[grep1("nohs", names(coef(nohs)))]
389 s <- p[,grep1("nohs", colnames(p))] %*% beta
390 ggplot(data=cbind.data.frame(s, PERM_data$nohs_percent), aes(x=PERM_data$
      nohs_percent, y=s)) + geom_line()
391
392 #fitted vs predictors
393
394 fit_data <- data.frame(cbind(nohs$fitted.values, nohs$y, PERM_data$nohs_
      percent))
395 names(fit_data) <- c("yhat", "y", "nohs")
396
397 fit_data %>%
398   ggplot() +
399     geom_point(aes(x = nohs, y = y)) +
400     geom_line(aes(x = nohs, y = yhat), color= "red")
401
402
403
404 #####
405 #child percent
406 #####
407
408
409 child <- gam(PERM_percent ~ s(child_percent) , data = PERM_data, family =
      betar())
410 summary(child)
411 plot(child)
412
413 #partial effects plot
414 p <- predict(child, type="lpmatrix")
415 beta <- coef(child)[grep1("child", names(coef(child)))]
416 s <- p[,grep1("child", colnames(p))] %*% beta
417 ggplot(data=cbind.data.frame(s, PERM_data$child_percent), aes(x=PERM_data$

```

```

    child_percent, y=s)) + geom_line()
418
419 #fitted vs predictors
420
421 fit_data <- data.frame(cbind(child$fitted.values, child$y, PERM_data$child
    _percent))
422 names(fit_data) <- c("yhat", "y", "child")
423
424 fit_data %>%
425   ggplot() +
426   geom_point(aes(x = child, y = y)) +
427   geom_line(aes(x = child, y = yhat), color= "red")
428
429 #doesn't look like a strong predictor
430
431
432
433 #####
434 #correlation matrix
435 #####
436
437 cor_data <- PERM_data %>%
438   select(20,2,4, 23:31)
439
440 chart.Correlation(cor_data, histogram=TRUE, pch=19)
441
442
443 #####
444 #
445 #
446 #Start modelling
447 #
448 #####
449
450

```

```

451
452 #####
453 #model with just lags
454 #####
455 lag_mod <- gam(PERM_percent ~ s(lags) , data = PERM_data, family = betar
    ())
456 summary(lag_mod)
457 plot(lag_mod)
458
459 #try to better understand relationship with lags by fitting simple linear
    model
460 lag_line <- lm(PERM_percent ~ lags, data = PERM_data)
461 summary(lag_line)
462
463 #####
464 #build a model with lags and predictor variables
465 #####
466 full_gam <- gam(PERM_percent ~ s(age_median) + s(earnings_median) + s(
    hispanic_percent) +
467             s(black_percent) + s(white_percent) + s(pov_percent) + s
    (asian_percent) +
468             s(bachelor_percent) + s(grad_percent) + s(hs_percent) +
    s(nohs_percent),
469             data = PERM_data,
470             family = betar(),
471             select = T)
472
473 summary(full_gam)
474 gam.check(full_gam)
475
476 #lets ad geo data into this
477
478 full_gam_geo <- gam(PERM_percent ~ s(age_median) + s(earnings_median) + s(
    hispanic_percent) +
479             s(black_percent) + s(white_percent) + s(pov_percent) + s

```

```

    (asian_percent) +
480         s(bachelor_percent) + s(grad_percent) + s(hs_percent) +
    s(nohs_percent) +
481         s(child_percent) +
482         s(as.factor(GEOID), bs = 'mrf', k = 20, xt = list(nb =
    nb)),
483         data = PERM_data,
484         family = betar(),
485         select = T)
486
487 summary(full_gam)
488 gam.check(full_gam)
489 plot(full_gam)
490
491 #lets ad lags into this
492 full_gam <- gam(PERM_percent ~ s(age_median) + s(earnings_median) + s(
    hispanic_percent) +
493         s(black_percent) + s(white_percent) + s(pov_percent) + s(
    (asian_percent) +
494         s(bachelor_percent) + s(grad_percent) + s(hs_percent) +
    s(nohs_percent) +
495         s(child_percent) +
496         s(as.factor(GEOID), bs = 'mrf', k = 20, xt = list(nb =
    nb)) +
497         s(lags),
498         data = PERM_data,
499         family = betar(),
500         select = T)
501
502 summary(full_gam)
503 gam.check(full_gam)
504 plot(full_gam)
505
506
507 #remove geo

```

```

508 full_gam <- gam(PERM_percent ~ s(age_median) + s(earnings_median) + s(
    hispanic_percent) +
509         s(black_percent) + s(white_percent) + s(pov_percent) + s
    (asian_percent) +
510         s(bachelor_percent) + s(grad_percent) + s(hs_percent) +
    s(nohs_percent) +
511         s(child_percent) +
512         s(lags),
513         data = PERM_data,
514         family = betar(),
515         select = T)
516
517 summary(full_gam)
518 gam.check(full_gam)
519 plot(full_gam)
520
521 #####
522 #performance
523
524 predict_perm_test <- predict.gam(full_gam, PERM_data_test, se.fit = T,
    type = "response")
525 resid_test <- cbind(predict_perm_test$fit, PERM_data_test$PERM_percent)
526 resid_test <- as.data.frame(resid_test)
527 names(resid_test) <- c("pred", "obs")
528 resid_test <- resid_test %>%
529     mutate(resids = obs - pred)
530
531 rmse_full <- RMSE(obs = resid_test$obs, pred = resid_test$pred)
532 #just for fun I want to double check
533 sqrt(sum(resid_test$resids^2)/nrow(resid_test))
534 #more or less iid normally distributed residual
535 plot(resid_test$obs, resid_test$resids)
536 hist(resid_test$resids)
537 #####
538

```

```

539 #####
540 #strip variables that do not appear to be useful
541 diminished_gam <- gam(PERM_percent ~ s(age_median) + s(earnings_median) +
      s(hispanic_percent) +
542         s(black_percent) + s(pov_percent) +
543         s(bachelor_percent) + s(nohs_percent) +
544         s(child_percent) +
545         s(lags),
546         data = PERM_data,
547         family = betar(),
548         select = T)
549
550 summary(diminished_gam)
551 gam.check(diminished_gam)
552
553 #performance for gam_diminished
554 predict_perm_test2 <- predict.gam(diminished_gam, PERM_data_test, se.fit =
      T, type = "response")
555 resid_test2 <- cbind(predict_perm_test2$fit, PERM_data_test$PERM_percent)
556 resid_test2 <- as.data.frame(resid_test2)
557 names(resid_test2) <- c("pred", "obs")
558 resid_test2 <- resid_test2 %>%
559   mutate(resids = obs - pred)
560
561 rmse_dim <- RMSE(obs = resid_test2$obs, pred = resid_test2$pred)
562 #more or less iid normally distributed residual
563 plot(resid_test2$obs, resid_test2$resids)
564 hist(resid_test2$resids)
565
566 #####
567 #lastly, just check rmse for lag model
568
569 predict_perm_test3 <- predict.gam(lag_mod, PERM_data_test, se.fit = T,
      type = "response")
570 resid_test3 <- cbind(predict_perm_test3$fit, PERM_data_test$PERM_percent)

```

```

571 resid_test3 <- as.data.frame(resid_test3)
572 names(resid_test3) <- c("pred", "obs")
573 resid_test3 <- resid_test3%>%
574   mutate(resid = obs - pred)
575
576 rmse_lag <- RMSE(obs = resid_test3$obs, pred = resid_test3$pred)
577 #more or less iid normally distributed residual
578 plot(resid_test3$obs, resid_test3$resid)
579 hist(resid_test3$resid)
580 #####
581 #
582 #fit on 2018 data and see how it performs out of box
583 #####
584 #2018 data manipulation
585 PERM_data_2018 <- all_data_2018 %>%
586   mutate(PERM_percent = total.reg_PERM/total_regs,
587          children = (over_6_under_18 + under_6)) %>%
588   filter(population > 100) %>%
589   dplyr::select(african_american,
590                 age_median,
591                 asian_american,
592                 earnings_median,
593                 edu_bachelors,
594                 edu_graduate,
595                 edu_highschool,
596                 edu_no_highschool,
597                 employed,
598                 european_american,
599                 hispanic_latino,
600                 native_american,
601                 pacific_islander,
602                 population,
603                 poverty,
604                 total_female,
605                 total_male,

```



```

606         unemployed ,
607         children ,
608         PERM_percent ,
609         turnout_percent ,
610         GEOID ,
611         geometry
612     )
613
614 #add in lag variables
615
616 lags_2016 <- all_data_2016 %>%
617   mutate(lags = total.reg_PERM/total_regs) %>%
618   select(lags, GEOID)
619
620
621 PERM_data_2018 <- PERM_data_2018 %>%
622   left_join(lags_2016, "GEOID")
623
624 #start modeling
625 PERM_data_2018 <- PERM_data_2018 %>%
626   dplyr::select(-turnout_percent) %>%
627   mutate(hispanic_percent = hispanic_latino/population,
628          black_percent = african_american/population,
629          white_percent = european_american/population,
630          pov_percent = poverty/population,
631          asian_percent = asian_american/population,
632          bachelor_percent = edu_bachelors/(population-children),
633          grad_percent = edu_graduate/(population - children),
634          hs_percent = edu_highschool/(population - children),
635          nohs_percent = edu_no_highschool/(population - children),
636          employed_percent = employed/(population - children),
637          unemployed_percent = unemployed/(population - children),
638          child_percent = children/(population - children)
639   ) %>%
640   na.omit()

```

```

641
642 predict_perm_2018 <- predict.gam(diminished_gam, PERM_data_2018, se.fit =
    T, type = "response")
643 resids_2018 <- cbind(predict_perm_2018$fit, PERM_data_2018$PERM_percent,
    as.numeric(PERM_data_2018$GEOID))
644 resids_2018 <- as.data.frame(resids_2018)
645 names(resids_2018) <- c("pred", "obs", "geo")
646 resids_2018 <- resids_2018 %>%
647   mutate(resids = pred - obs)
648
649
650 #systematically underpredicts and RMSE doubles
651 RMSE(obs = resids_2018$obs, pred = resids_2018$pred)
652 #just for fun I want to double check
653 sqrt(sum(resids_2018$resids^2)/nrow(resids_2018))
654 #more or less iid normally distributed residual
655 plot(resids_2018$obs, resids_2018$resids)
656 hist(resids_2018$resids)
657
658 #now let's look at it using the lag model.
659 predict_perm_20182 <- predict.gam(lag_mod, PERM_data_2018, se.fit = T,
    type = "response")
660 resids_20182 <- cbind(predict_perm_20182$fit, PERM_data_2018$PERM_percent,
    as.numeric(PERM_data_2018$GEOID))
661 resids_20182 <- as.data.frame(resids_20182)
662 names(resids_20182) <- c("obs", "pred", "geo")
663 resids_20182 <- resids_20182 %>%
664   mutate(resids = pred - obs)
665
666
667 #systematically underpredicts and RMSE doubles
668 RMSE(obs = resids_20182$obs, pred = resids_20182$pred)
669 #more or less iid normally distributed residual
670 plot(resids_2018$obs, resids_2018$resids)
671 hist(resids_2018$resids)

```

```

672 #here we see that the lag model predicts worse, but not by that much and
      still systematically under predicts
673
674 #####
675 #lets look at 2016 and 2018 as the outcome again, but let's build the
      model with our outcome variable
676 #standardized
677 #
678 #
679 #####
680 #we now have stnadardized variables that we can use are our outcome
      variable
681 #tow different approaches to standardizing
682 #attempt 1
683 PERM_data_standard1 <- PERM_data %>%
684   mutate(PERM_percent2 = PERM_percent/max(PERM_percent)) %>%
685   mutate(lags2 = lags/max(lags, na.rm = T)) %>%
686   dplyr::select(PERM_percent, lags, PERM_percent2, lags2)
687
688 #attempt 2
689 PERM_data_standard2 <- PERM_data %>%
690   mutate(PERM_percent2 = scale(PERM_percent)) %>%
691   mutate(lags2 = scale(lags)) %>%
692   dplyr::select(PERM_percent, lags, PERM_percent2, lags2)
693
694 #this approach ended up not being terribly convincing.
695
696 #####
697 #Let's look at how the model performed if we interpret the out put as an
      ordinal output
698 #we don't need to rerun the model , we just need to order the predictions
      and the actual numbers and see how we did
699 #####
700
701

```

```

702 #let's look at resids_2018 but change obs, pred to ordinal and see what
      the resids are
703
704 ordinal_resids_2018 <- resids_2018 %>%
705   mutate(obs2 = percent_rank(obs)) %>%
706   mutate(pred2 = percent_rank(pred)) %>%
707   mutate(resids2 = obs2-pred2)
708
709 #just lag
710 ordinal_resids_20182 <- resids_20182 %>%
711   mutate(obs2 = percent_rank(obs)) %>%
712   mutate(pred2 = percent_rank(pred)) %>%
713   mutate(resids2 = obs2-pred2)
714
715 #just lag does a decent job of predicting the percentile of the next
      election.
716 RMSE(pred = ordinal_resids_20182$pred, obs = ordinal_resids_20182$obs)
717 #when estimating the percentile rank its RMSE is .045
718 #residuals are roughly normally distributed
719 plot(ordinal_resids_20182$resids2)
720 hist(ordinal_resids_20182$resids2)
721
722
723 #####
724 #winning model diagnostics
725 #####
726 #original
727 #partial effects interpretation: https://www.researchgate.net/figure/Generalized-additive-model-GAM-plots-showing-the-partial-effects-of-selected\_fig4\_324422761
728 gam.check(lag_mod)
729 plot(lag_mod)
730 #ordinal
731 plot(ordinal_resids_20182$resids2)
732 hist(ordinal_resids_20182$resids2)

```

```

733
734 #####
735 #univariate relationships for 2018
736 #####
737 #####
738 #avg earnings
739 #####
740 earnings <- gam(PERM_percent ~ s(earnings_median) , data = PERM_data_2018,
    family = betar())
741 summary(earnings)
742 plot(earnings, residuals = T)
743
744 #partial effects plot
745 p <- predict(earnings, type="lpmatrix")
746 beta <- coef(earnings)[grepl("median", names(coef(earnings)))]
747 s <- p[,grepl("median", colnames(p))] %*% beta
748 ggplot(data=cbind.data.frame(s, PERM_data_2018$earnings_median), aes(x=
    PERM_data_2018$earnings_median, y=s)) + geom_line()
749
750 #fitted vs predictors
751
752 fit_data <- data.frame(cbind(earnings$fitted.values, earnings$y, PERM_data
    _2018$earnings_median))
753 names(fit_data) <- c("yhat", "y", "earnings")
754
755 fit_data %>%
756   ggplot() +
757   geom_point(aes(x = earnings, y = y)) +
758   geom_line(aes(x = earnings, y = yhat), color= "red") +
759   ylab("VBM Percentage") +
760   xlab("Average Yearly Earnings")
761
762 #####
763 #poverty percent
764 #####

```

```

765
766
767 pov <- gam(PERM_percent ~ s(pov_percent) , data = PERM_data_2018, family =
      betar())
768 summary(pov)
769 plot(pov)
770
771 #partial effects plot
772 p <- predict(pov, type="lpmatrix")
773 beta <- coef(pov)[grep1("pov", names(coef(pov)))]
774 s <- p[,grep1("pov", colnames(p))] %*% beta
775 ggplot(data=cbind.data.frame(s, PERM_data_2018$pov_percent), aes(x=PERM_
      data_2018$pov_percent, y=s)) + geom_line()
776
777 #fitted vs predictors
778
779 fit_data <- data.frame(cbind(pov$fitted.values, pov$y, PERM_data_2018$pov_
      percent))
780 names(fit_data) <- c("yhat", "y", "pov")
781
782 fit_data %>%
783   ggplot() +
784   geom_point(aes(x = pov, y = y)) +
785   geom_line(aes(x = pov, y = yhat), color= "red") +
786   ylab("VBM Percentage") +
787   xlab("Poverty Rate")
788
789
790 #####
791 #unemployed percent
792 #####
793
794
795 unemployed <- gam(PERM_percent ~ s(unemployed_percent) , data = PERM_data_
      2018, family = betar())

```

```

796 summary(unemployed)
797 plot(unemployed)
798
799 #partial effects plot
800 p <- predict(unemployed, type="lpmatrix")
801 beta <- coef(unemployed)[grepl("unemployed", names(coef(unemployed)))]
802 s <- p[,grepl("unemployed", colnames(p))] %*% beta
803 ggplot(data=cbind.data.frame(s, PERM_data_2018$unemployed_percent), aes(x=
      PERM_data_2018$unemployed_percent, y=s)) + geom_line()
804
805 #fitted vs predictors
806
807 fit_data <- data.frame(cbind(unemployed$fitted.values, unemployed$y, PERM_
      data_2018$unemployed_percent))
808 names(fit_data) <- c("yhat", "y", "unemployed")
809
810 fit_data %>%
811   ggplot() +
812   geom_point(aes(x = unemployed, y = y)) +
813   geom_line(aes(x = unemployed, y = yhat), color= "red") +
814   ylab("VBM Percentage") +
815   xlab("Unemployment Rate")
816
817
818 #####
819 #age
820 #####
821
822 age <- gam(PERM_percent ~ s(age_median) , data = PERM_data_2018, family =
      betar())
823 summary(age)
824 plot(earnings)
825
826 #partial effects plot
827 p <- predict(age, type="lpmatrix")

```

```

828 beta <- coef(age)[grepl("median", names(coef(age)))]
829 s <- p[,grepl("median", colnames(p))] %*% beta
830 ggplot(data=cbind.data.frame(s, PERM_data_2018$age_median), aes(x=PERM_
      data_2018$age_median, y=s)) + geom_line()
831
832 #fitted vs predictors
833
834 fit_data <- data.frame(cbind(age$fitted.values, age$y, PERM_data_2018$age_
      median))
835 names(fit_data) <- c("yhat", "y", "age")
836
837 fit_data %>%
838   ggplot() +
839     geom_point(aes(x = age, y = y)) +
840     geom_line(aes(x = age, y = yhat), color= "red") +
841     ylab("VBM Percentage") +
842     xlab("Average Age")
843
844
845 #we see a linear relationship between age and PERM %
846
847 #####
848 #hispanic percent
849 #####
850
851 hispanic <- gam(PERM_percent ~ s(hispanic_percent) , data = PERM_data_
      2018, family = betar())
852 summary(hispanic)
853 plot(hispanic)
854
855 #partial effects plot
856 p <- predict(hispanic, type="lpmatrix")
857 beta <- coef(hispanic)[grepl("hispanic", names(coef(hispanic)))]
858 s <- p[,grepl("hispanic", colnames(p))] %*% beta
859 ggplot(data=cbind.data.frame(s, PERM_data_2018$hispanic_percent), aes(x=

```



```

    PERM_data_2018$hispanic_percent, y=s)) + geom_line()
860
861 #fitted vs predictors
862
863 fit_data <- data.frame(cbind(hispanic$fitted.values, hispanic$y, PERM_data
    _2018$hispanic_percent))
864 names(fit_data) <- c("yhat", "y", "hispanic")
865
866 fit_data %>%
867   ggplot() +
868   geom_point(aes(x = hispanic, y = y)) +
869   geom_line(aes(x = hispanic, y = yhat), color= "red") +
870   ylab("VBM Percentage") +
871   xlab("Hispanic Rate")
872
873
874 #we see a linear relationship between age and PERM %
875
876 #####
877 #black percent
878 #####
879
880
881 black <- gam(PERM_percent ~ s(black_percent) , data = PERM_data_2018,
    family = betar())
882 summary(black)
883 plot(black)
884
885 #partial effects plot
886 p <- predict(black, type="lpmatrix")
887 beta <- coef(black)[grep1("black", names(coef(black)))]
888 s <- p[,grep1("black", colnames(p))] %*% beta
889 ggplot(data=cbind.data.frame(s, PERM_data_2018$black_percent), aes(x=PERM_
    data_2018$black_percent, y=s)) + geom_line()
890

```

```

891 #fitted vs predictors
892
893 fit_data <- data.frame(cbind(black$fitted.values, black$y, PERM_data_2018$
      black_percent))
894 names(fit_data) <- c("yhat", "y", "black")
895
896 fit_data %>%
897   ggplot() +
898   geom_point(aes(x = black, y = y)) +
899   geom_line(aes(x = black, y = yhat), color= "red") +
900   ylab("VBM Percentage") +
901   xlab("African American Rate")
902
903
904
905 #####
906 #asian percent
907 #####
908
909
910 asian <- gam(PERM_percent ~ s(asian_percent) , data = PERM_data_2018,
      family = betar())
911 summary(asian)
912 plot(asian)
913
914 #partial effects plot
915 p <- predict(asian, type="lpmatrix")
916 beta <- coef(asian)[grepl("asian", names(coef(asian)))]
917 s <- p[,grepl("asian", colnames(p))] %*% beta
918 ggplot(data=cbind.data.frame(s, PERM_data_2018$asian_percent), aes(x=PERM_
      data_2018$asian_percent, y=s)) + geom_line()
919
920 #fitted vs predictors
921
922 fit_data <- data.frame(cbind(asian$fitted.values, asian$y, PERM_data_2018$

```

```

    asian_percent))
923 names(fit_data) <- c("yhat", "y", "asian")
924
925 fit_data %>%
926   ggplot() +
927   geom_point(aes(x = asian, y = y)) +
928   geom_line(aes(x = asian, y = yhat), color= "red") +
929   ylab("VBM Percentage") +
930   xlab("Asian Rate")
931
932
933
934 #####
935 #hs percent
936 #####
937
938
939 hs <- gam(PERM_percent ~ s(hs_percent) , data = PERM_data_2018, family =
    betar())
940 summary(hs)
941 plot(hs)
942
943 #partial effects plot
944 p <- predict(hs, type="lpmatrix")
945 beta <- coef(hs)[grepl("hs", names(coef(hs)))]
946 s <- p[,grepl("hs", colnames(p))] %*% beta
947 ggplot(data=cbind.data.frame(s, PERM_data_2018$hs_percent), aes(x=PERM_
    data_2018$hs_percent, y=s)) + geom_line()
948
949 #fitted vs predictors
950
951 fit_data <- data.frame(cbind(hs$fitted.values, hs$y, PERM_data_2018$hs_
    percent))
952 names(fit_data) <- c("yhat", "y", "hs")
953

```

```

954 fit_data %>%
955   ggplot() +
956   geom_point(aes(x = hs, y = y)) +
957   geom_line(aes(x = hs, y = yhat), color= "red") +
958   ylab("VBM Percentage") +
959   xlab("High School Graduate Rate")
960
961
962 #doesn't appear too meaningful
963
964 #####
965 #bachelor percent
966 #####
967
968
969 bachelor <- gam(PERM_percent ~ s(bachelor_percent) , data = PERM_data_
    2018, family = betar())
970 summary(bachelor)
971 plot(bachelor)
972
973 #partial effects plot
974 p <- predict(bachelor, type="lpmatrix")
975 beta <- coef(bachelor)[grepl("bachelor", names(coef(bachelor)))]
976 s <- p[,grepl("bachelor", colnames(p))] %*% beta
977 ggplot(data=cbind.data.frame(s, PERM_data_2018$bachelor_percent), aes(x=
    PERM_data_2018$bachelor_percent, y=s)) + geom_line()
978
979 #fitted vs predictors
980
981 fit_data <- data.frame(cbind(bachelor$fitted.values, bachelor$y, PERM_data
    _2018$bachelor_percent))
982 names(fit_data) <- c("yhat", "y", "bachelor")
983
984 fit_data %>%
985   ggplot() +

```

```

986   geom_point(aes(x = bachelor, y = y)) +
987   geom_line(aes(x = bachelor, y = yhat), color= "red") +
988   ylab("VBM Percentage") +
989   xlab("Bachelor's Degree Rate")
990
991
992 #####
993 #nohs percent
994 #####
995
996
997 nohs <- gam(PERM_percent ~ s(nohs_percent) , data = PERM_data_2018, family
          = betar())
998 summary(nohs)
999 plot(nohs)
1000
1001 #partial effects plot
1002 p <- predict(nohs, type="lpmatrix")
1003 beta <- coef(nohs)[grepl("nohs", names(coef(nohs)))]
1004 s <- p[,grepl("nohs", colnames(p))] %*% beta
1005 ggplot(data=cbind.data.frame(s, PERM_data_2018$nohs_percent), aes(x=PERM_
          data_2018$nohs_percent, y=s)) + geom_line()
1006
1007 #fitted vs predictors
1008
1009 fit_data <- data.frame(cbind(nohs$fitted.values, nohs$y, PERM_data_2018$
          nohs_percent))
1010 names(fit_data) <- c("yhat", "y", "nohs")
1011
1012 fit_data %>%
1013   ggplot() +
1014   geom_point(aes(x = nohs, y = y)) +
1015   geom_line(aes(x = nohs, y = yhat), color= "red")
1016
1017

```

```

1018
1019 #####
1020 #child percent
1021 #####
1022
1023
1024 child <- gam(PERM_percent ~ s(child_percent) , data = PERM_data_2018,
              family = betar())
1025 summary(child)
1026 plot(child)
1027
1028 #partial effects plot
1029 p <- predict(child, type="lpmatrix")
1030 beta <- coef(child)[grep1("child", names(coef(child)))]
1031 s <- p[,grep1("child", colnames(p))] %*% beta
1032 ggplot(data=cbind.data.frame(s, PERM_data_2018$child_percent), aes(x=PERM_
              data_2018$child_percent, y=s)) + geom_line()
1033
1034 #fitted vs predictors
1035
1036 fit_data <- data.frame(cbind(child$fitted.values, child$y, PERM_data_2018$
              child_percent))
1037 names(fit_data) <- c("yhat", "y", "child")
1038
1039 fit_data %>%
1040   ggplot() +
1041     geom_point(aes(x = child, y = y)) +
1042     geom_line(aes(x = child, y = yhat), color= "red")
1043
1044
1045
1046
1047 #####
1048 #prediction map for 2020
1049 #####

```

```

1050
1051 PERM_data_2020 <- PERM_data_2018 %>%
1052   mutate(lags = PERM_percent)
1053
1054 #it seems reasonable to rerun the model built on 2016 and 2018 to try and
1055   capture the most recent political
1056
1057 #atmosphere
1058
1059 lag_gam_2018 <- gam(PERM_percent ~ s(lags),
1060                    data = PERM_data_2018,
1061                    family = betar(),
1062                    select = T)
1063
1064 summary(diminished_gam)
1065 gam.check(diminished_gam)
1066
1067 #predict
1068 predict_perm_2020 <- predict.gam(lag_gam_2018, PERM_data_2020, se.fit = T,
1069                                  type = "response")
1070
1071 PERM_data_2020_full <- PERM_data_2020 %>%
1072   mutate(predicts_percentiles = percent_rank(predict_perm_2020$fit))
1073
1074 #map
1075 PERM_data_2020_full %>%
1076   select(geometry,
1077           predicts_percentiles) %>%
1078   gather(key = "variable", value = "estimate", -geometry) %>%
1079   ggplot(aes(fill = estimate)) +
1080   geom_sf(aes(geometry = geometry), color = NA) +
1081   coord_sf(ylim = c(33.7, 35)) +
1082   scale_fill_viridis_c(option = "plasma") +
1083   theme(axis.text.x = element_blank(),
1084         axis.text.y = element_blank(),
1085         axis.ticks = element_blank(),

```

```

1083     plot.title = element_text(hjust = .5)) +
1084     labs(title = "VBM Percentile Predictions: 2020") -> perm_map_2020
1085
1086 #END SCRIPT
1087
1088 plot_data <- all_data_2018 %>%
1089   select(total.reg_PERM, total_regs, turnout_percent, population) %>%
1090   mutate(perm_percent = total.reg_PERM/total_regs) %>%
1091   filter(population >100)
1092
1093 plot_data %>%
1094   ggplot(aes(x= perm_percent, y = turnout_percent)) +
1095   geom_point() +
1096   geom_smooth(method = "loess", span = 2) +
1097   xlab("VBM Percent") +
1098   ylab("Turnout Percent")

```

A.4 Voter Turnout GAMs

```

1 #####
2 #The goal of this script is to use GAMs to model
3 #the percentage of registered voters that turnout to vote
4 #We will explore basic assumptions about barriers people face in voting
5 #and see if those are good predictors of if people are voting.
6 #We will attempt to use both beta regression to model the proportions
   (0,1)
7 #as well as modeling the absolute number.
8 #We will evaluate the strength and usefulness (interpretability) of each
   model.
9 #We will compare mse and aic and gcv
10 #cross validate
11 #####
12 #libraries
13 library(mgcv)
14 library(caret)

```



```

15 library(tigris)
16 library(spdep)
17 library(tidyverse)
18 library(PerformanceAnalytics)
19 library(gamclass)
20 library(caTools)
21
22 #load environment
23 load("/Users/pbeeman/Google Drive/masters thesis copy/masterenvironment.
    RData")
24
25 #2014 data manipulation
26 Votes_data <- all_data_2014 %>%
27   mutate(turnout_percent = total_votes/total_regs,
28           children = (over_6_under_18 + under_6)) %>%
29   filter(total_regs > 100) %>%
30   dplyr::select(african_american,
31                 age_median,
32                 asian_american,
33                 earnings_median,
34                 edu_bachelors,
35                 edu_graduate,
36                 edu_highschool,
37                 edu_no_highschool,
38                 employed,
39                 european_american,
40                 hispanic_latino,
41                 native_american,
42                 pacific_islander,
43                 population,
44                 poverty,
45                 total_female,
46                 total_male,
47                 unemployed,
48                 children,

```

```

49         turnout_percent,
50         GEOID,
51         geometry
52     )
53
54 #add it lag variables
55
56 lags_2012 <- all_data_2012 %>%
57   filter(total_regs > 100) %>%
58   mutate(lags = total_votes/total_regs) %>%
59   dplyr::select(lags, GEOID)
60
61
62 Votes_data <- Votes_data %>%
63   left_join(lags_2012, "GEOID")
64
65
66 #latino
67 plot(Votes_data$hispanic_latino/Votes_data$population, Votes_data$turnout_
      percent)
68
69 #start modeling
70 Votes_data <- Votes_data %>%
71   mutate(hispanic_percent = hispanic_latino/population,
72          black_percent = african_american/population,
73          white_percent = european_american/population,
74          pov_percent = poverty/population,
75          asian_percent = asian_american/population,
76          bachelor_percent = edu_bachelors/(population-children),
77          grad_percent = edu_graduate/(population - children),
78          hs_percent = edu_highschool/(population - children),
79          nohs_percent = edu_no_highschool/(population - children),
80          employed_percent = employed/(population - children),
81          unemployed_percent = unemployed/(population - children),
82          child_percent = children/(population - children)

```

```

83   ) %>%
84   na.omit()
85
86 #split into training/test data
87 sample_size <- floor(0.75 * nrow(Votes_data))
88 set.seed(123)
89 train_ind <- sample(seq_len(nrow(Votes_data)), size = sample_size)
90
91 Votes_data <- Votes_data[train_ind, ]
92 Votes_data_test <- Votes_data[-train_ind, ]
93
94
95 #####
96 #Markov Random Fields Geo approach
97 #####
98
99 #read in shap files for CTs
100 LA_shapes <- tracts("CA", "Los Angeles", cb = TRUE)
101
102
103 #check for empty geometries
104 which(is.na(st_dimension(Votes_data$geometry)))
105
106
107 #convert to neighborhood
108 nb <- poly2nb(Votes_data$geometry, row.names = Votes_data$GEOID)
109 names(nb) <- attr(nb, "region.id")
110 str(nb[1:6])
111
112 #fit GAM for PERM using MRF with a beta response
113 #https://stat.ethz.ch/R-manual/R-patched/library/mgcv/html/Beta.html
114
115 ctrl <- gam.control(nthreads = 6) # use 6 parallel threads, reduce if
    fewer physical CPU cores
116 m1 <- gam(turnout_percent ~ s(as.factor(GEOID), bs = 'mrf', k = 20, xt =

```

```

    list(nb = nb)), # define MRF smooth
117     data = Votes_data,
118     method = 'REML',
119     # fast version of REML smoothness selection
120     family = betar()) # fit a beta regression
121
122 viz_data <- Votes_data %>%
123   mutate(new_vals = predict(m1, type = 'response'))
124
125
126 summary(m1)
127 gam.check(m1)
128
129 #visualize
130 #cross validation
131 #https://rdr.io/cran/gamclass/man/CVgam.html
132
133 #####
134 #Let's start looking at univariate relationships and plotting them out
135 #####
136 #####
137 #avg earnings
138 #####
139 earnings <- gam(turnout_percent ~ s(earnings_median) , data = Votes_data,
    family = betar())
140 summary(earnings)
141 plot(earnings, residuals = T)
142
143 #partial effects plot
144 p <- predict(earnings, type="lpmatrix")
145 beta <- coef(earnings)[grep("median", names(coef(earnings)))]
146 s <- p[,grep("median", colnames(p))] %*% beta
147 ggplot(data=cbind.data.frame(s, Votes_data$earnings_median), aes(x=Votes_
    data$earnings_median, y=s)) + geom_line()
148

```

```

149 #fitted vs predictors
150
151 fit_data <- data.frame(cbind(earnings$fitted.values, earnings$y, Votes_
    data$earnings_median))
152 names(fit_data) <- c("yhat", "y", "earnings")
153
154 fit_data %>%
155   ggplot() +
156   geom_point(aes(x = earnings, y = y)) +
157   geom_line(aes(x = earnings, y = yhat), color= "red")
158
159 #####
160 #poverty percent
161 #####
162
163
164 pov <- gam(turnout_percent ~ s(pov_percent) , data = Votes_data, family =
    betar())
165 summary(pov)
166 plot(pov)
167
168 #partial effects plot
169 p <- predict(pov, type="lpmatrix")
170 beta <- coef(pov)[grepl("pov", names(coef(pov)))]
171 s <- p[,grepl("pov", colnames(p))] %*% beta
172 ggplot(data=cbind.data.frame(s, Votes_data$pov_percent), aes(x=Votes_data$
    pov_percent, y=s)) + geom_line()
173
174 #fitted vs predictors
175
176 fit_data <- data.frame(cbind(pov$fitted.values, pov$y, Votes_data$pov_
    percent))
177 names(fit_data) <- c("yhat", "y", "pov")
178
179 fit_data %>%

```

```

180   ggplot() +
181   geom_point(aes(x = pov, y = y)) +
182   geom_line(aes(x = pov, y = yhat), color= "red")
183
184 #####
185 #unemployed percent
186 #####
187
188
189 unemployed <- gam(turnout_percent ~ s(unemployed_percent) , data = Votes_
    data, family = betar())
190 summary(unemployed)
191 plot(unemployed)
192
193 #partial effects plot
194 p <- predict(unemployed, type="lpmatrix")
195 beta <- coef(unemployed)[grepl("unemployed", names(coef(unemployed)))]
196 s <- p[,grepl("unemployed", colnames(p))] %% beta
197 ggplot(data=cbind.data.frame(s, Votes_data$unemployed_percent), aes(x=
    Votes_data$unemployed_percent, y=s)) + geom_line()
198
199 #fitted vs predictors
200
201 fit_data <- data.frame(cbind(unemployed$fitted.values, unemployed$y, Votes
    _data$unemployed_percent))
202 names(fit_data) <- c("yhat", "y", "unemployed")
203
204 fit_data %>%
205   ggplot() +
206   geom_point(aes(x = unemployed, y = y)) +
207   geom_line(aes(x = unemployed, y = yhat), color= "red")
208
209 #####
210 #age
211 #####

```

```

212
213 age <- gam(turnout_percent ~ s(age_median) , data = Votes_data, family =
      betar())
214 summary(age)
215 plot(earnings)
216
217 #partial effects plot
218 p <- predict(age, type="lpmatrix")
219 beta <- coef(age)[grepl("median", names(coef(age)))]
220 s <- p[,grepl("median", colnames(p))] %*% beta
221 ggplot(data=cbind.data.frame(s, Votes_data$age_median), aes(x=Votes_data$
      age_median, y=s)) + geom_line()
222
223 #fitted vs predictors
224
225 fit_data <- data.frame(cbind(age$fitted.values, age$y, Votes_data$age_
      median))
226 names(fit_data) <- c("yhat", "y", "age")
227
228 fit_data %>%
229   ggplot() +
230   geom_point(aes(x = age, y = y)) +
231   geom_line(aes(x = age, y = yhat), color= "red")
232
233 #we see a linear relationship between age and PERM %
234
235 #####
236 #hispanic percent
237 #####
238
239 hispanic <- gam(turnout_percent ~ s(hispanic_percent) , data = Votes_data,
      family = betar())
240 summary(hispanic)
241 plot(hispanic)
242

```

```

243 #partial effects plot
244 p <- predict(hispanic, type="lpmatrix")
245 beta <- coef(hispanic)[grep1("hispanic", names(coef(hispanic)))]
246 s <- p[,grep1("hispanic", colnames(p))] %*% beta
247 ggplot(data=cbind.data.frame(s, Votes_data$hispanic_percent), aes(x=Votes_
    data$hispanic_percent, y=s)) + geom_line()
248
249 #fitted vs predictors
250
251 fit_data <- data.frame(cbind(hispanic$fitted.values, hispanic$y, Votes_
    data$hispanic_percent))
252 names(fit_data) <- c("yhat", "y", "hispanic")
253
254 fit_data %>%
255   ggplot() +
256   geom_point(aes(x = hispanic, y = y)) +
257   geom_line(aes(x = hispanic, y = yhat), color= "red")
258
259 #we see a linear relationship between age and PERM %
260
261 #####
262 #black percent
263 #####
264
265
266 black <- gam(turnout_percent ~ s(black_percent) , data = Votes_data,
    family = betar())
267 summary(black)
268 plot(black)
269
270 #partial effects plot
271 p <- predict(black, type="lpmatrix")
272 beta <- coef(black)[grep1("black", names(coef(black)))]
273 s <- p[,grep1("black", colnames(p))] %*% beta
274 ggplot(data=cbind.data.frame(s, Votes_data$black_percent), aes(x=Votes_

```



```

    data$black_percent, y=s)) + geom_line()
275
276 #fitted vs predictors
277
278 fit_data <- data.frame(cbind(black$fitted.values, black$y, Votes_data$
    black_percent))
279 names(fit_data) <- c("yhat", "y", "black")
280
281 fit_data %>%
282   ggplot() +
283   geom_point(aes(x = black, y = y)) +
284   geom_line(aes(x = black, y = yhat), color= "red")
285
286
287 #####
288 #asian percent
289 #####
290
291
292 asian <- gam(turnout_percent ~ s(asian_percent) , data = Votes_data,
    family = betar())
293 summary(asian)
294 plot(asian)
295
296 #partial effects plot
297 p <- predict(asian, type="lpmatrix")
298 beta <- coef(asian)[grepl("asian", names(coef(asian)))]
299 s <- p[,grepl("asian", colnames(p))] %*% beta
300 ggplot(data=cbind.data.frame(s, Votes_data$asian_percent), aes(x=Votes_
    data$asian_percent, y=s)) + geom_line()
301
302 #fitted vs predictors
303
304 fit_data <- data.frame(cbind(asian$fitted.values, asian$y, Votes_data$
    asian_percent))

```

```

305 names(fit_data) <- c("yhat", "y", "asian")
306
307 fit_data %>%
308   ggplot() +
309   geom_point(aes(x = asian, y = y)) +
310   geom_line(aes(x = asian, y = yhat), color= "red")
311
312
313 #####
314 #hs percent
315 #####
316
317
318 hs <- gam(turnout_percent ~ s(hs_percent) , data = Votes_data, family =
      betar())
319 summary(hs)
320 plot(hs)
321
322 #partial effects plot
323 p <- predict(hs, type="lpmatrix")
324 beta <- coef(hs)[grepl("hs", names(coef(hs)))]
325 s <- p[,grepl("hs", colnames(p))] %*% beta
326 ggplot(data=cbind.data.frame(s, Votes_data$hs_percent), aes(x=Votes_data$
      hs_percent, y=s)) + geom_line()
327
328 #fitted vs predictors
329
330 fit_data <- data.frame(cbind(hs$fitted.values, hs$y, Votes_data$hs_percent
      ))
331 names(fit_data) <- c("yhat", "y", "hs")
332
333 fit_data %>%
334   ggplot() +
335   geom_point(aes(x = hs, y = y)) +
336   geom_line(aes(x = hs, y = yhat), color= "red")

```

```

337
338 #doesn't appear too meaningful
339
340 #####
341 #bachelor percent
342 #####
343
344
345 bachelor <- gam(turnout_percent ~ s(bachelor_percent) , data = Votes_data,
    family = betar())
346 summary(bachelor)
347 plot(bachelor)
348
349 #partial effects plot
350 p <- predict(bachelor, type="lpmatrix")
351 beta <- coef(bachelor)[grep("bachelor", names(coef(bachelor)))]
352 s <- p[,grep("bachelor", colnames(p))] %*% beta
353 ggplot(data=cbind.data.frame(s, Votes_data$bachelor_percent), aes(x=Votes_
    data$bachelor_percent, y=s)) + geom_line()
354
355 #fitted vs predictors
356
357 fit_data <- data.frame(cbind(bachelor$fitted.values, bachelor$y, Votes_
    data$bachelor_percent))
358 names(fit_data) <- c("yhat", "y", "bachelor")
359
360 fit_data %>%
361   ggplot() +
362   geom_point(aes(x = bachelor, y = y)) +
363   geom_line(aes(x = bachelor, y = yhat), color= "red")
364
365 #####
366 #nohs percent
367 #####
368

```

```

369
370 nohs <- gam(turnout_percent ~ s(nohs_percent) , data = Votes_data, family
      = betar())
371 summary(nohs)
372 plot(nohs)
373
374 #partial effects plot
375 p <- predict(nohs, type="lpmatrix")
376 beta <- coef(nohs)[grep1("nohs", names(coef(nohs)))]
377 s <- p[,grep1("nohs", colnames(p))] %*% beta
378 ggplot(data=cbind.data.frame(s, Votes_data$nohs_percent), aes(x=Votes_data
      $nohs_percent, y=s)) + geom_line()
379
380 #fitted vs predictors
381
382 fit_data <- data.frame(cbind(nohs$fitted.values, nohs$y, Votes_data$nohs_
      percent))
383 names(fit_data) <- c("yhat", "y", "nohs")
384
385 fit_data %>%
386   ggplot() +
387   geom_point(aes(x = nohs, y = y)) +
388   geom_line(aes(x = nohs, y = yhat), color= "red")
389
390
391
392 #####
393 #child percent
394 #####
395
396
397 child <- gam(turnout_percent ~ s(child_percent) , data = Votes_data,
      family = betar())
398 summary(child)
399 plot(child)

```

```

400
401 #partial effects plot
402 p <- predict(child, type="lpmatrix")
403 beta <- coef(child)[grepl("child", names(coef(child)))]
404 s <- p[,grepl("child", colnames(p))] %*% beta
405 ggplot(data=cbind.data.frame(s, Votes_data$child_percent), aes(x=Votes_
    data$child_percent, y=s)) + geom_line()
406
407 #fitted vs predictors
408
409 fit_data <- data.frame(cbind(child$fitted.values, child$y, Votes_data$
    child_percent))
410 names(fit_data) <- c("yhat", "y", "child")
411
412 fit_data %>%
413   ggplot() +
414   geom_point(aes(x = child, y = y)) +
415   geom_line(aes(x = child, y = yhat), color= "red")
416
417 #doesn't look like a strong predictor
418
419
420
421 #####
422 #correlation matrix
423 #####
424
425 cor_data <- Votes_data %>%
426   select(20,2,4, 23:31)
427
428 chart.Correlation(cor_data, histogram=TRUE, pch=19)
429
430
431 #####
432 #

```

```

433 #
434 #Start modelling
435 #
436 #####
437
438 #####
439 #model with just lags
440 #####
441 lag_mod <- gam(turnout_percent ~ s(lags) , data = Votes_data, family =
      betar())
442 summary(lag_mod)
443 plot(lag_mod)
444
445 #try to better understand relationship with lags by fitting simple linear
      model
446 lag_line <- lm(turnout_percent ~ lags, data = Votes_data)
447 summary(lag_line)
448
449 #####
450 #build a model with lags and predictor variables
451 #####
452 full_gam <- gam(turnout_percent ~ s(age_median) + s(earnings_median) + s(
      hispanic_percent) +
453           s(black_percent) + s(white_percent) + s(pov_percent) + s
      (asian_percent) +
454           s(bachelor_percent) + s(grad_percent) + s(hs_percent) +
      s(nohs_percent),
455           data = Votes_data,
456           family = betar(),
457           select = T)
458
459 summary(full_gam)
460 gam.check(full_gam)
461
462 #lets ad geo data into this

```

```

463
464 full_gam_geo <- gam(turnout_percent ~ s(age_median) + s(earnings_median) +
    s(hispanic_percent) +
465         s(black_percent) + s(white_percent) + s(pov_percent)
    + s(asian_percent) +
466         s(bachelor_percent) + s(grad_percent) + s(hs_percent
    ) + s(nohs_percent) +
467         s(child_percent) +
468         s(as.factor(GEOID), bs = 'mrf', k = 20, xt = list(nb
    = nb))),
469         data = Votes_data,
470         family = betar(),
471         select = T)
472
473 summary(full_gam)
474 gam.check(full_gam)
475 plot(full_gam)
476
477 #lets ad lags into this
478 full_gam <- gam(turnout_percent ~ s(age_median) + s(earnings_median) + s(
    hispanic_percent) +
479         s(black_percent) + s(white_percent) + s(pov_percent) + s
    (asian_percent) +
480         s(bachelor_percent) + s(grad_percent) + s(hs_percent) +
    s(nohs_percent) +
481         s(child_percent) +
482         s(as.factor(GEOID), bs = 'mrf', k = 20, xt = list(nb =
    nb)) +
483         s(lags),
484         data = Votes_data,
485         family = betar(),
486         select = T)
487
488 summary(full_gam)
489 gam.check(full_gam)

```

```

490 plot(full_gam)
491
492
493 #remove geo
494 full_gam <- gam(turnout_percent ~ s(age_median) + s(earnings_median) + s(
    hispanic_percent) +
495             s(black_percent) + s(white_percent) + s(pov_percent) + s
    (asian_percent) +
496             s(bachelor_percent) + s(grad_percent) + s(hs_percent) +
    s(nohs_percent) +
497             s(child_percent) +
498             s(lags),
499             data = Votes_data,
500             family = betar(),
501             select = T)
502
503 summary(full_gam)
504 gam.check(full_gam)
505 plot(full_gam)
506
507 #####
508 #performance
509
510 predict_perm_test <- predict.gam(full_gam, Votes_data_test, se.fit = T,
    type = "response")
511 resid_test <- cbind(predict_perm_test$fit, Votes_data_test$turnout_
    percent)
512 resid_test <- as.data.frame(resid_test)
513 names(resid_test) <- c("pred", "obs")
514 resid_test <- resid_test %>%
515     mutate(resids = obs - pred)
516
517 rmse_full <- RMSE(obs = resid_test$obs, pred = resid_test$pred)
518 #just for fun I want to double check
519 sqrt(sum(resid_test$resids^2)/nrow(resid_test))

```



```

520 #more or less iid normally distributed residual
521 plot(resids_test$obs, resids_test$resids)
522 hist(resids_test$resids)
523 #####
524
525 #####
526 #strip variables that do not appear to be useful
527 diminished_gam <- gam(turnout_percent ~ s(age_median) + s(hispanic_
    percent) +
528
    s(black_percent) + s(pov_percent) + s(asian_
    percent) +
529
    s(bachelor_percent) + s(grad_percent) +
530
    s(child_percent) +
531
    s(lags),
532
    data = Votes_data,
533
    family = betar(),
534
    select = T)
535
536 summary(diminished_gam)
537 gam.check(diminished_gam)
538
539 #performance for gam_diminished
540 predict_perm_test2 <- predict.gam(diminished_gam, Votes_data_test, se.fit
    = T, type = "response")
541 resids_test2 <- cbind(predict_perm_test2$fit, Votes_data_test$turnout_
    percent)
542 resids_test2 <- as.data.frame(resids_test2)
543 names(resids_test2) <- c("pred", "obs")
544 resids_test2 <- resids_test2 %>%
545   mutate(resids = obs - pred)
546
547 rmse_dim <- RMSE(obs = resids_test2$obs, pred = resids_test2$pred)
548 #more or less iid normally distributed residual
549 plot(resids_test2$obs, resids_test2$resids)
550 hist(resids_test2$resids)

```

```

551
552 #####
553 #lastly, just check rmse for lag model
554
555 predict_perm_test3 <- predict.gam(lag_mod, Votes_data_test, se.fit = T,
    type = "response")
556 resid_test3 <- cbind(predict_perm_test3$fit, Votes_data_test$turnout_
    percent)
557 resid_test3 <- as.data.frame(resid_test3)
558 names(resid_test3) <- c("pred", "obs")
559 resid_test3 <- resid_test3%>%
560   mutate(resid = obs - pred)
561
562 rmse_lag <- RMSE(obs = resid_test3$obs, pred = resid_test3$pred)
563 #more or less iid normally distributed residual
564 plot(resid_test3$obs, resid_test3$resid)
565 hist(resid_test3$resid)
566 #####
567 #
568 #fit on 2018 data and see how it performs out of box
569 #####
570 #2018 data manipulation
571 Votes_data_2018 <- all_data_2018 %>%
572   mutate(turnout_percent = total_votes/total_regs,
573     children = (over_6_under_18 + under_6)) %>%
574   filter(population > 100) %>%
575   dplyr::select(african_american,
576     age_median,
577     asian_american,
578     earnings_median,
579     edu_bachelors,
580     edu_graduate,
581     edu_highschool,
582     edu_no_highschool,
583     employed,

```

```

584         european_american,
585         hispanic_latino,
586         native_american,
587         pacific_islander,
588         population,
589         poverty,
590         total_female,
591         total_male,
592         unemployed,
593         children,
594         turnout_percent,
595         GEOID,
596         geometry
597     )
598
599
600 #add in lag variables
601
602 lags_2016 <- all_data_2016 %>%
603   mutate(lags = total_votes/total_regs) %>%
604   select(lags, GEOID)
605
606
607 Votes_data_2018 <- Votes_data_2018 %>%
608   left_join(lags_2016, "GEOID")
609
610 #start modeling
611 Votes_data_2018 <- Votes_data_2018 %>%
612   mutate(hispanic_percent = hispanic_latino/population,
613          black_percent = african_american/population,
614          white_percent = european_american/population,
615          pov_percent = poverty/population,
616          asian_percent = asian_american/population,
617          bachelor_percent = edu_bachelors/(population-children),
618          grad_percent = edu_graduate/(population - children),

```

```

619     hs_percent = edu_highschool/(population - children),
620     nohs_percent = edu_no_highschool/(population - children),
621     employed_percent = employed/(population - children),
622     unemployed_percent = unemployed/(population - children),
623     child_percent = children/(population - children)
624   ) %>%
625   na.omit()
626
627 predict_perm_2018 <- predict.gam(diminished_gam, Votes_data_2018, se.fit =
    T, type = "response")
628 resids_2018 <- cbind(predict_perm_2018$fit, Votes_data_2018$turnout_
    percent, as.numeric(Votes_data_2018$GE0ID))
629 resids_2018 <- as.data.frame(resids_2018)
630 names(resids_2018) <- c("pred", "obs", "geo")
631 resids_2018 <- resids_2018 %>%
632   mutate(resids = pred - obs)
633
634
635 #systematically underpredicts and RMSE doubles
636 RMSE(obs = resids_2018$obs, pred = resids_2018$pred)
637 #just for fun I want to double check
638 sqrt(sum(resids_2018$resids^2)/nrow(resids_2018))
639 #more or less iid normally distributed residual
640 plot(resids_2018$obs, resids_2018$resids)
641 hist(resids_2018$resids)
642
643 #now let's look at it using the lag model.
644 predict_perm_20182 <- predict.gam(lag_mod, Votes_data_2018, se.fit = T,
    type = "response")
645 resids_20182 <- cbind(predict_perm_20182$fit, Votes_data_2018$turnout_
    percent, as.numeric(Votes_data_2018$GE0ID))
646 resids_20182 <- as.data.frame(resids_20182)
647 names(resids_20182) <- c("obs", "pred", "geo")
648 resids_20182 <- resids_20182 %>%
649   mutate(resids = pred - obs)

```

```

650
651
652 #systematically underpredicts and RMSE doubles
653 RMSE(obs = resid_2018$obs, pred = resid_2018$pred)
654 #more or less iid normally distributed residual
655 plot(resid_2018$obs, resid_2018$resid)
656 hist(resid_2018$resid)
657 #here we see that the lag model predicts worse, but not by that much and
    still systematically under predicts
658
659 #####
660 #lets look at 2016 and 2018 as the outcome again, but let's build the
    model with our outcome variable
661 #standardized
662 #####
663 #Let's look at how the model performed if we interpret the out put as an
    ordinal output
664 #we don't need to rerun the model , we just need to order the predictions
    and the actual numbers and see how we did
665 #####
666
667
668 #let's look at resid_2018 but change obs, pred to ordinal and see what
    the resid are
669
670 ordinal_resid_2018 <- resid_2018 %>%
671   mutate(obs2 = percent_rank(obs)) %>%
672   mutate(pred2 = percent_rank(pred)) %>%
673   mutate(resid2 = obs2-pred2)
674
675
676 RMSE(pred = ordinal_resid_2018$pred2, obs = ordinal_resid_2018$obs2)
677 #when estimating the percentile rank its RMSE is .09
678 #residuals are roughly normally distributed
679 plot(ordinal_resid_2018$resid2)

```

```

680 hist(ordinal_resids_2018$resids2)
681
682 #the diminished model does a decent job of predicting the percentiles for
    voter turnout for the next election.
683 #let's see how the lag model did
684 ordinal_resids_20182 <- resid_20182 %>%
685   mutate(obs2 = percent_rank(obs)) %>%
686   mutate(pred2 = percent_rank(pred)) %>%
687   mutate(resids2 = obs2-pred2)
688
689
690 RMSE(pred = ordinal_resids_20182$pred2, obs = ordinal_resids_20182$obs2)
691 #when estimating the percentile rank its RMSE is .09
692 #residuals are roughly normally distributed
693 plot(ordinal_resids_2018$resids2)
694 hist(ordinal_resids_2018$resids2)
695
696
697 #the full model does a better job of predicting the percentiles for voter
    turnout for the next election
698 #than just using the lag.
699 #We should be able to use this to get a reasonable estimate for 2020 voter
    turnout percentage percentiles
700
701 #####
702 #winning model diagnostics
703 #####
704 #original
705 #partial effects interpretation: https://www.researchgate.net/figure/Generalized-additive-model-GAM-plots-showing-the-partial-effects-of-selected\_fig4\_324422761
706 gam.check(full_gam)
707 plot(full_gam)
708 #ordinal
709 plot(ordinal_resids_2018$resids2)

```

```

710 hist(ordinal_resids_2018$resids2)
711
712 #What happens if we build a model where percentiles are the outcome
    variable?
713 Votes_data_percentile <- Votes_data %>%
714   mutate(turnout_percent = percent_rank(turnout_percent))
715
716 Votes_data_2018_percentile <- Votes_data_2018 %>%
717   mutate(turnout_percent = percent_rank(turnout_percent))
718
719 full_gam_percentils <- gam(turnout_percent ~ s(age_median) + s(earnings_
    median) + s(hispanic_percent) +
720     s(black_percent) + s(white_percent) + s(pov_percent) + s
    (asian_percent) +
721     s(bachelor_percent) + s(grad_percent) + s(hs_percent) +
    s(nohs_percent) +
722     s(child_percent) +
723     s(lags),
724     data = Votes_data_percentile,
725     family = betar(),
726     select = T)
727
728 summary(full_gam)
729 gam.check(full_gam)
730 plot(full_gam)
731
732 #predict 2018
733 predict_perm_2018_percentile <- predict.gam(full_gam_percentils, Votes_
    data_2018_percentile, se.fit = T, type = "response")
734 resids_2018 <- cbind(predict_perm_2018$fit, Votes_data_2018$turnout_
    percent, as.numeric(Votes_data_2018$GEOID))
735 resids_2018 <- as.data.frame(resids_2018)
736 names(resids_2018) <- c("pred", "obs", "geo")
737 resids_2018 <- resids_2018 %>%
738   mutate(resids = pred - obs)

```

```

739 RMSE(obs = resid_2018$obs, pred = resid_2018$pred)
740
741 #####
742 #univariate relationships for 2018
743 #####
744 #####
745 #avg earnings
746 #####
747 earnings <- gam(turnout_percent ~ s(earnings_median) , data = Votes_data_
      2018, family = betar())
748 summary(earnings)
749 plot(earnings, residuals = T)
750
751 #partial effects plot
752 p <- predict(earnings, type="lpmatrix")
753 beta <- coef(earnings)[grep("median", names(coef(earnings)))]
754 s <- p[,grep("median", colnames(p))] %*% beta
755 ggplot(data=cbind.data.frame(s, Votes_data_2018$earnings_median), aes(x=
      Votes_data_2018$earnings_median, y=s)) + geom_line()
756
757 #fitted vs predictors
758
759 fit_data <- data.frame(cbind(earnings$fitted.values, earnings$y, Votes_
      data_2018$earnings_median))
760 names(fit_data) <- c("yhat", "y", "earnings")
761
762 fit_data %>%
763   ggplot() +
764     geom_point(aes(x = earnings, y = y)) +
765     geom_line(aes(x = earnings, y = yhat), color= "red") +
766     ylab("Turnout Percentage") +
767     xlab("Average Yearly Earnings")
768
769 #####
770 #poverty percent

```



```

771 #####
772
773
774 pov <- gam(turnout_percent ~ s(pov_percent) , data = Votes_data_2018,
      family = betar())
775 summary(pov)
776 plot(pov)
777
778 #partial effects plot
779 p <- predict(pov, type="lpmatrix")
780 beta <- coef(pov)[grepl("pov", names(coef(pov)))]
781 s <- p[,grepl("pov", colnames(p))] %*% beta
782 ggplot(data=cbind.data.frame(s, Votes_data_2018$pov_percent), aes(x=Votes_
      data_2018$pov_percent, y=s)) + geom_line()
783
784 #fitted vs predictors
785
786 fit_data <- data.frame(cbind(pov$fitted.values, pov$y, Votes_data_2018$pov
      _percent))
787 names(fit_data) <- c("yhat", "y", "pov")
788
789 fit_data %>%
790   ggplot() +
791   geom_point(aes(x = pov, y = y)) +
792   geom_line(aes(x = pov, y = yhat), color= "red") +
793   ylab("Turnout Percentage") +
794   xlab("Poverty Rate")
795
796
797 #####
798 #unemployed percent
799 #####
800
801
802 unemployed <- gam(turnout_percent ~ s(unemployed_percent) , data = Votes_

```

```

    data_2018, family = betar())
803 summary(unemployed)
804 plot(unemployed)
805
806 #partial effects plot
807 p <- predict(unemployed, type="lpmatrix")
808 beta <- coef(unemployed)[grep1("unemployed", names(coef(unemployed)))]
809 s <- p[,grep1("unemployed", colnames(p))] %*% beta
810 ggplot(data=cbind.data.frame(s, Votes_data_2018$unemployed_percent), aes(x
    =Votes_data_2018$unemployed_percent, y=s)) + geom_line()
811
812 #fitted vs predictors
813
814 fit_data <- data.frame(cbind(unemployed$fitted.values, unemployed$y, Votes
    _data_2018$unemployed_percent))
815 names(fit_data) <- c("yhat", "y", "unemployed")
816
817 fit_data %>%
818   ggplot() +
819   geom_point(aes(x = unemployed, y = y)) +
820   geom_line(aes(x = unemployed, y = yhat), color= "red")
821
822 #####
823 #age
824 #####
825
826 age <- gam(turnout_percent ~ s(age_median) , data = Votes_data_2018,
    family = betar())
827 summary(age)
828 plot(earnings)
829
830 #partial effects plot
831 p <- predict(age, type="lpmatrix")
832 beta <- coef(age)[grep1("median", names(coef(age)))]
833 s <- p[,grep1("median", colnames(p))] %*% beta

```

```

834 ggplot(data=cbind.data.frame(s, Votes_data_2018$age_median), aes(x=Votes_
      data_2018$age_median, y=s)) + geom_line()
835
836 #fitted vs predictors
837
838 fit_data <- data.frame(cbind(age$fitted.values, age$y, Votes_data_2018$age
      _median))
839 names(fit_data) <- c("yhat", "y", "age")
840
841 fit_data %>%
842   ggplot() +
843     geom_point(aes(x = age, y = y)) +
844     geom_line(aes(x = age, y = yhat), color= "red") +
845     ylab("Turnout Percentage") +
846     xlab("Average Age")
847
848 #we see a linear relationship between age and PERM %
849
850 #####
851 #hispanic percent
852 #####
853
854 hispanic <- gam(turnout_percent ~ s(hispanic_percent) , data = Votes_data_
      2018, family = betar())
855 summary(hispanic)
856 plot(hispanic)
857
858 #partial effects plot
859 p <- predict(hispanic, type="lpmatrix")
860 beta <- coef(hispanic)[grepl("hispanic", names(coef(hispanic)))]
861 s <- p[,grepl("hispanic", colnames(p))] %*% beta
862 ggplot(data=cbind.data.frame(s, Votes_data_2018$hispanic_percent), aes(x=
      Votes_data_2018$hispanic_percent, y=s)) + geom_line()
863
864 #fitted vs predictors

```

```

865
866 fit_data <- data.frame(cbind(hispanic$fitted.values, hispanic$y, Votes_
    data_2018$hispanic_percent))
867 names(fit_data) <- c("yhat", "y", "hispanic")
868
869 fit_data %>%
870   ggplot() +
871   geom_point(aes(x = hispanic, y = y)) +
872   geom_line(aes(x = hispanic, y = yhat), color= "red") +
873   ylab("Turnout Percentage") +
874   xlab("Hispanic Rate")
875
876 #we see a linear relationship between age and PERM %
877
878 #####
879 #black percent
880 #####
881
882
883 black <- gam(turnout_percent ~ s(black_percent) , data = Votes_data_2018,
    family = betar())
884 summary(black)
885 plot(black)
886
887 #partial effects plot
888 p <- predict(black, type="lpmatrix")
889 beta <- coef(black)[grep1("black", names(coef(black)))]
890 s <- p[,grep1("black", colnames(p))] %*% beta
891 ggplot(data=cbind.data.frame(s, Votes_data_2018$black_percent), aes(x=
    Votes_data_2018$black_percent, y=s)) + geom_line()
892
893 #fitted vs predictors
894
895 fit_data <- data.frame(cbind(black$fitted.values, black$y, Votes_data_2018
    $black_percent))

```

```

896 names(fit_data) <- c("yhat", "y", "black")
897
898 fit_data %>%
899   ggplot() +
900   geom_point(aes(x = black, y = y)) +
901   geom_line(aes(x = black, y = yhat), color= "red")
902
903
904 #####
905 #asian percent
906 #####
907
908
909 asian <- gam(turnout_percent ~ s(asian_percent) , data = Votes_data_2018,
1000   family = betar())
910 summary(asian)
911 plot(asian)
912
913 #partial effects plot
914 p <- predict(asian, type="lpmatrix")
915 beta <- coef(asian)[grepl("asian", names(coef(asian)))]
916 s <- p[,grepl("asian", colnames(p))] %*% beta
917 ggplot(data=cbind.data.frame(s, Votes_data_2018$asian_percent), aes(x=
1000   Votes_data_2018$asian_percent, y=s)) + geom_line()
918
919 #fitted vs predictors
920
921 fit_data <- data.frame(cbind(asian$fitted.values, asian$y, Votes_data_2018
1000   $asian_percent))
922 names(fit_data) <- c("yhat", "y", "asian")
923
924 fit_data %>%
925   ggplot() +
926   geom_point(aes(x = asian, y = y)) +
927   geom_line(aes(x = asian, y = yhat), color= "red")

```

```

928
929
930 #####
931 #hs percent
932 #####
933
934
935 hs <- gam(turnout_percent ~ s(hs_percent) , data = Votes_data_2018, family
      = betar())
936 summary(hs)
937 plot(hs)
938
939 #partial effects plot
940 p <- predict(hs, type="lpmatrix")
941 beta <- coef(hs)[grepl("hs", names(coef(hs)))]
942 s <- p[,grepl("hs", colnames(p))] %*% beta
943 ggplot(data=cbind.data.frame(s, Votes_data_2018$hs_percent), aes(x=Votes_
      data_2018$hs_percent, y=s)) + geom_line()
944
945 #fitted vs predictors
946
947 fit_data <- data.frame(cbind(hs$fitted.values, hs$y, Votes_data_2018$hs_
      percent))
948 names(fit_data) <- c("yhat", "y", "hs")
949
950 fit_data %>%
951   ggplot() +
952   geom_point(aes(x = hs, y = y)) +
953   geom_line(aes(x = hs, y = yhat), color= "red") +
954   ylab("Turnout Percentage") +
955   xlab("High School Graduate Rate")
956
957 #doesn't appear too meaningful
958
959 #####

```

```

960 #bachelor percent
961 #####
962
963
964 bachelor <- gam(turnout_percent ~ s(bachelor_percent) , data = Votes_data_
      2018, family = betar())
965 summary(bachelor)
966 plot(bachelor)
967
968 #partial effects plot
969 p <- predict(bachelor, type="lpmatrix")
970 beta <- coef(bachelor)[grep1("bachelor", names(coef(bachelor)))]
971 s <- p[,grep1("bachelor", colnames(p))] %*% beta
972 ggplot(data=cbind.data.frame(s, Votes_data_2018$bachelor_percent), aes(x=
      Votes_data_2018$bachelor_percent, y=s)) + geom_line()
973
974 #fitted vs predictors
975
976 fit_data <- data.frame(cbind(bachelor$fitted.values, bachelor$y, Votes_
      data_2018$bachelor_percent))
977 names(fit_data) <- c("yhat", "y", "bachelor")
978
979 fit_data %>%
980   ggplot() +
981     geom_point(aes(x = bachelor, y = y)) +
982     geom_line(aes(x = bachelor, y = yhat), color= "red") +
983     ylab("Turnout Percentage") +
984     xlab("Bachelor's Degree Rate")
985
986 #####
987 #nohs percent
988 #####
989
990
991 nohs <- gam(turnout_percent ~ s(nohs_percent) , data = Votes_data_2018,

```

```

    family = betar())
992 summary(nohs)
993 plot(nohs)
994
995 #partial effects plot
996 p <- predict(nohs, type="lpmatrix")
997 beta <- coef(nohs)[grep1("nohs", names(coef(nohs)))]
998 s <- p[,grep1("nohs", colnames(p))] %*% beta
999 ggplot(data=cbind.data.frame(s, Votes_data_2018$nohs_percent), aes(x=Votes
    _data_2018$nohs_percent, y=s)) + geom_line()
1000
1001 #fitted vs predictors
1002
1003 fit_data <- data.frame(cbind(nohs$fitted.values, nohs$y, Votes_data_2018$
    nohs_percent))
1004 names(fit_data) <- c("yhat", "y", "nohs")
1005
1006 fit_data %>%
1007   ggplot() +
1008   geom_point(aes(x = nohs, y = y)) +
1009   geom_line(aes(x = nohs, y = yhat), color= "red") +
1010   ylab("Turnout Percentage") +
1011   xlab("No High School Graduaate Rate")
1012
1013 #####
1014 #child percent
1015 #####
1016
1017
1018 child <- gam(turnout_percent ~ s(child_percent) , data = Votes_data_2018,
    family = betar())
1019 summary(child)
1020 plot(child)
1021
1022 #partial effects plot

```



```

1023 p <- predict(child, type="lpmatrix")
1024 beta <- coef(child)[grepl("child", names(coef(child)))]
1025 s <- p[,grepl("child", colnames(p))] %*% beta
1026 ggplot(data=cbind.data.frame(s, Votes_data_2018$child_percent), aes(x=
      Votes_data_2018$child_percent, y=s)) + geom_line()
1027
1028 #fitted vs predictors
1029
1030 fit_data <- data.frame(cbind(child$fitted.values, child$y, Votes_data_2018
      $child_percent))
1031 names(fit_data) <- c("yhat", "y", "child")
1032
1033 fit_data %>%
1034   ggplot() +
1035   geom_point(aes(x = child, y = y)) +
1036   geom_line(aes(x = child, y = yhat), color= "red") +
1037   ylab("Turnout Percentage") +
1038   xlab("Percent Children")
1039
1040 #####
1041 #prediction map for 2020
1042 #####
1043
1044 Votes_data_2020 <- Votes_data_2018 %>%
1045   mutate(lags = turnout_percent)
1046
1047 #it seems reasonable to rerun the model built on 2016 and 2018 to try and
      capture the most recent political
1048 #atmosphere
1049
1050 diminished_gam_2018 <- gam(turnout_percent ~ s(age_median) + s(hispanic_
      percent) +
1051
      s(black_percent) + s(pov_percent) + s(asian_
      percent) +
1052
      s(bachelor_percent) + s(grad_percent) +

```

```

1053         s(child_percent) +
1054         s(lags),
1055         data = Votes_data_2018,
1056         family = betar(),
1057         select = T)
1058
1059 summary(diminished_gam)
1060 gam.check(diminished_gam)
1061
1062 #predict
1063 predict_perm_2020 <- predict.gam(diminished_gam_2018, Votes_data_2020, se.
      fit = T, type = "response")
1064
1065 Votes_data_2020_full <- Votes_data_2020 %>%
1066   mutate(predicts_percentiles = percent_rank(predict_perm_2020$fit))
1067
1068 #map
1069 Votes_data_2020_full %>%
1070   select(geometry,
1071           predicts_percentiles) %>%
1072   gather(key = "variable", value = "estimate", -geometry) %>%
1073   ggplot(aes(fill = estimate)) +
1074   geom_sf(aes(geometry = geometry), color = NA) +
1075   coord_sf(ylim = c(33.7, 35)) +
1076   scale_fill_viridis_c(option = "plasma") +
1077   theme(axis.text.x = element_blank(),
1078         axis.text.y = element_blank(),
1079         axis.ticks = element_blank(),
1080         plot.title = element_text(hjust = .5)) +
1081   labs(title = "Turnout Percentile Predictions: 2020") -> turnout_map_2020
1082
1083 #END SCRIPT

```

References

- [1] Brians, Craig Leonard, and Bernard Grofman. “Election Day Registration’s Effect on U.S. Voter Turnout.” *Social Science Quarterly*, vol. 82, no. 1, 2001, pp. 170–183. JSTOR, www.jstor.org/stable/42955710.
- [2] Famighetti, Christopher. ”Long Voting Lines: Explained.” *Brenon Center for Justice*. 4 November, 2016. <https://www.brennancenter.org/our-work/research-reports/long-voting-lines-explained>
- [3] BRADY, HENRY E., and JOHN E. MCNULTY. “Turning Out to Vote: The Costs of Finding and Getting to the Polling Place.” *The American Political Science Review*, vol. 105, no. 1, 2011, pp. 115–134. JSTOR, www.jstor.org/stable/41480830.
- [4] Judy Ly, 2019. ”Anonymous Voter Data: Los Angeles County 2012-2018.” Los Angeles County Registrar Recorder/County Clerk.
- [5] Los Angles Almanac 2019. ”Eligible & Registered Voters, 1990-2018 Including Turnout & Votes by Mail.” <http://www.laalmanac.com/election/el02.php>
- [6] ”Vote By Mail.” Alex Padilla California Secretary of State. <https://www.sos.ca.gov/elections/voter-registration/vote-mail/#hist>
- [7] Marx, Brian D., and Paul HC Eilers. ”Direct generalized additive modeling with penalized likelihood.” *Computational Statistics & Data Analysis* 28.2 (1998): 193-209.
- [8] Larsen, Kim. “GAM : The Predictive Modeling Silver Bullet.” (2015). <https://pdfs.semanticscholar.org/dea4/adaaf06e6fc99179a2620b7a031188c6e532.pdf>
- [9] Rakich, Nathaniel. ”Early-Voting Laws Probably Don’t Boost Turnout.” *Fivethirtyeight*. 30 January, 2019. <https://fivethirtyeight.com/features/early-voting-laws-probably-dont-boost-turnout/>