

Majority Reaction to Minority Representation in Local Government*

Aly Somani

University of Toronto

March 23, 2021

Abstract

ABSTRACT: This paper studies the racial dynamics of local American politics in a setting with non-partisan elections. I exploit close city council elections in California from 1996 to 2017 to implement a regression discontinuity design, which allows me to study the causal effects of a nonwhite candidate's victory against a white candidate. I find that in cities where the nonwhite candidate won ("treatment"), compared to cities where the nonwhite candidate lost ("control"), more white candidates run in the next election. While this effect is partly driven by the presence of more white runners-up in the treatment group, the electoral success of the nonwhite candidate also leads to the entry of new white candidates in the next election. This effect is driven by cities that have gone through bigger demographic changes over the past few decades, which suggests that changes in the racial composition of the city and the associated perception of threat to the dominant status of whites within the city (e.g. Jardina, 2014) are the likely mechanism behind the observed effect.

*I am grateful to my supervisor, Gustavo Bobonis, for his invaluable guidance, as well as to Arthur Blouin and David Price for their feedback at various stages of the project. I would also like to thank Lauren Brandt, Jordi Mondria, and other participants in ECO4060 and the CEPA brown bag seminar at the University of Toronto.

1 Introduction

The racial composition of the United States is changing, and fast. According to projections from the U.S. Census Bureau, racial minority groups will make up a majority of the U.S. national population by 2044 (Colby and Ortman, 2015). However, the numerical representation of racial minorities in local politics remains low. A 2018 survey conducted by the International City/County Management Association (ICMA) revealed that about 90% of the 21,466 council members across 3,677 cities surveyed were non-Hispanic whites (ICMA, 2019).

This descriptive under-representation of minorities is often blamed in popular media for the substantive under-representation of and, consequently, worse outcomes for minority groups.¹²³ Research in political economy has also shown that increased numerical representation of minorities can improve minority outcomes. For instance, Logan (2018) finds that the election of a black politician during post-U.S. Civil War Reconstruction decreased the black-white literacy gap. Beach et al. (2018) find that a nonwhite candidate's victory in city council elections in California lead to differential gains in housing prices in majority nonwhite neighbourhoods due to increased business activity in minority neighbourhoods and changes in police behaviour. Nye et al. (2014) also find that the election of a black mayor is associated with increased black employment and labour force participation. Outside the U.S., Pande (2003) finds that political reservation for minorities in Indian states increased redistribution of resources towards the minority groups that benefited from the quota. A closely related literature has shown that increasing women's - another traditionally under-represented group - political representation results in policy choices that are more favourable to women (Iyer et al., 2012; Chattopadhyay and Duflo, 2004).

These papers clearly highlight the importance of studying the dynamics of the competition between different identity groups in the political arena. These dynamics could be informative about why minority representation remains low. To that end, this paper studies the the racial dynamics of the electoral competition in a local, non-partisan setting. The paper sets out to answer three questions in particular: Does the electoral success of a nonwhite candidate lead to subsequent increases the participation of nonwhites as political candidates? Does the election of a nonwhite candidate lead to a response by the white citizen-candidates? Does the election of a nonwhite candidate (and the subsequent response by white and nonwhite citizen-candidates) affect electoral outcomes and, consequently, the racial composition of the legislature in the future?

To answer these questions, I use data on city council candidates in California from 1996 to 2017. I combine this with the 2000 and 2010 decennial census surname files to predict the ethnicity of the candidates based on their last names. I then use close elections between a (non-Hispanic)

¹<https://apnews.com/4c6c0cf4d1aa4c8eba374876b8a24533/divided-america-minorities-missing-many-legislatures>

²<https://nextcity.org/daily/entry/anaheim-city-council-vote-latino-district-at-large-california>

³<https://www.demos.org/research/problem-african-american-underrepresentation-city-councils>

white candidate and a nonwhite candidate (non-Hispanic black, non-Hispanic Asian, Hispanic) in a regression discontinuity (RD) design to generate quasi-random assignment to “treatment” and “control” groups, which allows me to isolate the causal effects of a nonwhite candidate’s victory from all other systematic differences between the treated and control groups.

I do not find any evidence of an increase in the number of nonwhites running for city council in the next election. This is in contrast with a closely related set of studies that find that increasing female representation in politics motivates more female candidates to run. For example, Gilardi (2015) in Swiss municipal elections between 1970 and 2010, and Brown et al. (2018) in India between 1977 and 2014.

Next, I also find that in cities where the nonwhite candidate won compared to cities where the nonwhite candidate lost, more white candidates run in the next election. The average number of white candidates per seat in the sample within the chosen RD bandwidth is 1.2. The treatment effect estimated at the cut-off suggests is approximately 0.5 more white candidates per seat. This effect is, in part, due to the presence of more white runners-up in the treatment group (approximately 0.2 out of 0.5). However, I also find that the nonwhite candidate’s election leads to an entry of new white candidates running in the next election (approximately 0.28 out of 0.5).

I explore several potential mechanisms for the observed entry of new white candidates. I do not find evidence that the effect is driven by differences in voter preferences along racial lines or by policy changes occurring after the election of the nonwhite candidate. Neither do I find any evidence of a differential political cycle across the treatment and control groups. Heterogeneity analysis suggests that reaction by white citizen-candidates is primarily driven by cities that have experienced a bigger increase in the nonwhite share of the population. Based on the recent literature on the politicization of white identity, this suggests that the electoral success of the nonwhite candidate could be leading to the entry of new white candidates by making demographic changes more salient to the whites residing in these cities.

Recent political science research employing experimental survey designs have shown that highlighting the salience of changes to the country’s racial composition leads white Americans in the U.S. to express greater political conservatism (Craig and Richeson, 2014a, 2014b), oppose welfare programs (Wetts and Willer, 2018), and express strong exclusionary reactions related to immigration policies (Enos, 2014). In her book, *White Identity Politics*, Jardina (2019) uses survey data to show that as immigration rates in the U.S. grew in the 1990s, presenting a threat to the existing racial hierarchy, white identity became more relevant to opinions of white Americans on immigration policy. Trebbi et al. (2008) also find that anticipation of demographic changes influenced the choice of electoral system within cities. The authors show that after the passage of the Voting Rights Act in 1965, white majorities strategically chose electoral rules to disenfranchise minority groups. Cities with white majorities that expected an increase in black votes adopted at-large elec-

tions when the black minority in the city was relatively small, allowing whites to win more seats on the city councils. On the other hand, when the minority share was large enough so that whites risked losing at-large elections, white majorities chose district-based elections so that minority votes remained confined to predominantly nonwhite districts.

This result also speaks to another stream of the literature that has focused on the reaction of men, traditionally a dominant group in positions of power, to women, traditionally an under-represented group in positions of power. Perhaps the closest study to mine is Gagliarducci and Paserman (2011), who study gender interactions in municipal governments in Italy between 1993 and 2003. The authors find that municipal legislatures headed by women are more likely to be terminated early, and the effect is most notable in councils which are entirely male and in regions with less favourable views towards working women. Bhalotra et al. (2018) study Indian state legislative elections and find that female electoral success *decreases* entry of new female candidates. Their result is most pronounced in states with higher gender inequality, which the authors argue is suggestive of a “backlash” that discourages women from joining politics. Using a public goods behavioural experiment in Indian villages, Gangadharan et al. (2016) find that male participants were significantly less likely to contribute towards the public good when the group leader was revealed to be a woman rather than a man.

This paper contributes to the existing literature by providing new evidence of changes in the actions of white citizen-candidates, measured as an explicit response in terms of the choice to run for election. This goes beyond the current literature that primarily relies on voting patterns or beliefs expressed in surveys.

Lastly, I also find that as a consequence of the increase in white candidates running for election, more whites get elected to the city council. Since I do not find any evidence of a differential political cycle or a difference in voter preferences at the cut-off, or evidence in support of the hypothesis that policy changes enacted after the election of the nonwhite candidate is leading to the election of more whites in the next election, I conclude that the election of more whites to the city council is most likely a result of more white candidates running in the first place.

The rest of the paper is organized as follows. Section 2 provides information on the institutional context in California and the data used in this study. In Section 3, I discuss the empirical strategy, the potential bias introduced in the estimation due to the (imperfect) race assignment using census surname files, and perform several tests to check that the RD design is valid. Section 4 presents my main empirical findings as well as various robustness checks. Section 5 explores potential mechanisms for the findings. Section 6 concludes the paper and outlines several potential steps for future work.

2 Background and Data

In California, most city councils consist of five members who are elected through a municipal election to represent the city. Council members are elected for staggered four-year terms, with an election every two years with multiple seats being contested. Most city council members are elected at-large, i.e. multiple candidates run for multiple seats. For example, five candidates might be running for three open seats on the city council. In such a scenario, the three candidates with the highest vote shares will win the seat. City council elections in California are non-partisan.

Most cities in California employ the council-manager system of government, where the mayor is selected by the council members from amongst themselves. The selected mayor presides over the council meetings in addition to performing certain ceremonial duties, but has powers which are essentially identical to other council members. Some larger cities use a mayor-council system of government, where the mayor is directly elected by the citizens to oversee the city government. City governments have the power to levy taxes and are primarily responsible for determining the city budget, and appointing and removing certain city officials such as city manager or city clerk, and performing other tasks such as overseeing the functioning of police services, maintenance of public utilities, issuance of zoning or building permits, and provision of waste disposal services.

2.1 California Election Data Archive

The primary dataset for this study is the California Election Data Archive (CEDA), which is a joint project of the Center for California Studies, the Institute for Social Research, and the Office of the California Secretary of State. It contains the first and last names of every candidate that ran in any local government election since 1996 and the number of votes won by that candidate.⁴ CEDA also lists the number of open council seats, which makes it possible to identify the candidates that narrowly won and narrowly lost the election. Since city council elections in California are non-partisan, CEDA does not record any party affiliations for the candidates, neither are they observed by the voters.

2.2 Census Surname Files

CEDA does not identify the race/ethnicity of the candidate. Hence, I use the surname files produced by the U.S. Census Bureau using data from the 2000 and 2010 decennial censuses.⁵ The census surname files contain all the last names reported 100 or more times in the decennial cen-

⁴<http://csus-dspace.calstate.edu/handle/10211.3/210187>

⁵<https://www.census.gov/data/developers/data-sets/surnames.html>

sus, as well as the national-level racial demographics associated with that last name. For each last name, the surname file includes the proportion of people with that last name that identify with a particular race group or Hispanic origin. The 2000 surname file lists just over 151,000 last names reported by over 242 million people. The 2010 surname file lists over 162,000 last names reported by almost 295 million people. The 2000 and 2010 files use the same six Hispanic origin and race categories: non-Hispanic white alone, non-Hispanic black or African American alone, non-Hispanic American Indian and Alaska Native alone, non-Hispanic Asian and Native Hawaiian and Other Pacific Islander alone, non-Hispanic Two or More Races, and Hispanic or Latino origin. For the purposes of this paper, going forward I will refer to these groups as white, black, Native American, Asian, multiracial, and Hispanic, respectively. For candidates running in elections between 1996 and 2004, I use the 2000 census surname file, and for candidates running in elections between 2005 and 2017, I use the 2010 census surname file.

2.3 Demographics Data and City Public Finance Data

I also obtain city demographics from the 1990, 2000, and 2010 decennial censuses from the Integrated Public Use Microdata Series (IPUMS) National Historical Geographic Information System (NHGIS) (Manson et al., 2019).⁶ Moreover, I use public finance data based on the cities' Annual Financial Transaction Reports. This data is maintained by the State Controller's Office.⁷

2.4 Sample Restrictions and Summary Statistics

To ensure that I am comparing the consequences of similar elections across all cities, I drop all candidates identified as mayoral candidates in the dataset. I also drop any candidates running for a short term in the council office. Moreover, I drop any cities with multiple elections in the same year. This is most likely due to district-based elections, which are not explicitly coded in the dataset. This leaves me with 4,011 elections between 1996 and 2017.

Table 1a presents the summary statistics of election-related variables for at-large city council elections. The count of the number of candidates running and the number of candidates elected by race is based on a 70% criterion, whereby a candidate is assigned a race only if people in the census with the same last name associate with that race at least 70% of the time.⁸ A candidate remains unassigned if I cannot match her last name to the census surnames file or if the maximum proportion across the race categories is lower than the threshold being used. The first two columns

⁶<https://www.nhgis.org>

⁷<https://bythenumbers.sco.ca.gov/>

⁸I will discuss the implications of this race prediction for the measurement of the RD treatment effect in the next section. For details on the choice of the 70% criterion, refer to Appendix B.

present the means and standard deviations for all elections in the dataset. The next two columns restrict the sample to elections which are classified as white v. nonwhite.⁹ There are 385 such elections in total. The last two columns restrict to the subsample of white v. nonwhite elections with a margin of victory less than or equal to 4.4%, which is the main bandwidth used for the first result of the paper pertaining to the number of white candidates per seat in the next election.

On average, in a city council election around 5.4 candidates run for election to 2.4 open positions on the city council. Out of those 5.4 candidates, I am able to assign an ethnicity to 3.9 candidates, while the remaining candidates remain unassigned. Note that the number of black candidates matched is very low. This is in part due to the inaccuracy of the race assignment using the census surname files, and in part because blacks make up very small proportion of the city populations in California (see Table 1b).¹⁰ The table also shows that white v. nonwhite elections tend to have more nonwhite candidates fewer white candidates on average compared to the full set of elections. Furthermore, elections within the 4.4% bandwidth around the cut-off look very similar to full set of white v. nonwhite elections.

Table 1b presents the summary statistics of the demographics for all cities where I observe city council elections. The first two columns present the summary statistics for all cities in the dataset. The next two focus on cities where I observe white v. nonwhite elections. The last two focus on cities where I observe white v. nonwhite elections with a margin of victory less than or equal to 4.4%. Compared to all cities in the dataset, cities where I observe white v. nonwhite elections tend to be more diverse, with a higher share of nonwhite population. In particular, they have a much higher share of Hispanics. This also explains the higher number of Hispanic candidates observed in the elections in these cities (Table 1a).

3 Methodology

3.1 Empirical Strategy

To assess the impact of increasing nonwhite representation, I employ a sharp RD design using close elections between a white candidate and a nonwhite candidate. As mentioned earlier, my final sample consists of cities with at-large elections. Hence, I observe multiple candidates contesting for multiple seats. I use the vote shares of the marginal candidates, i.e. the last-place winner and the first-place loser, in my RD design if one of the marginal candidates is white and the other is nonwhite. For example, in a given election, five candidates might be competing for three open

⁹Section 3 provides details on this classification.

¹⁰Based on the last names, I also do not identify any candidates as Native American or multiracial. Hence, for the purposes of this paper I will refer to black, Hispanic, and Asian candidates as nonwhite.

seats on the city council. In such a scenario, the three candidates with the highest vote shares will win the seat. Hence, I classify such an election as white v. nonwhite if among the candidates with the third and fourth highest vote shares, one of them is white and the other is nonwhite.

Several authors have recommended using local linear approximations for RD treatment effect estimation since higher-order polynomials tend to produce over-fitting of the data which can lead to unreliable results near boundary points (e.g., Cattaneo et al., [forthcoming](#); Gelman and Imbens, 2019). Hence, I use a linear specification for the running variable, margin of victory, on either side of the cut-off, and estimate the following regression model within a narrow bandwidth:

$$Y_{c,e+1} = \alpha_{e+1} + \beta_1 \mathbb{1}[NW_{ce}] + \beta_2 Margin_{ce} + \beta_3 \mathbb{1}[NW_{ce}] \times Margin_{ce} + \epsilon_{c,e+1} \quad (1)$$

where $Y_{c,e+1}$ is the outcome variable (e.g., number of white candidates per seat running for city council, number of nonwhite candidates per seat, number of new white candidates per seat, number of whites elected per seat, etc.) in city c in election year $e + 1$, $\mathbb{1}[NW_{ce}]$ is an indicator variable equal to one if a nonwhite candidate won in a narrow election against a white candidate in city c in election year e , $Margin_{ce}$ is the margin of victory for the winning candidate, and α_{e+1} denotes year fixed effects. The parameter of interest is β_1 , which identifies the causal effect of electing a nonwhite candidate relative to the election of a white candidate.¹¹

In a working paper, Pei et al. (2018) argue that one could use a data-driven procedure to choose the optimal polynomial order, and that higher-order polynomials are often optimal. While my preferred specification includes a local linear design, I show that my main results are robust to a quadratic or cubic specification.

The choice of bandwidth is fundamental for RD designs as it directly affects the properties of local polynomial estimation and inference procedures. While a narrower bandwidth reduces any misspecification error due to the local polynomial approximation, it also tends to increase the variance of the estimated coefficients because fewer observations are available for estimation. The most popular approach for bandwidth selection seeks to minimize the mean squared error (MSE) of the local polynomial RD estimator for given a choice of the polynomial order, effectively optimizing the bias-variance trade-off.¹² The procedure yields a different bandwidth for each outcome variable.

Calonico et al. (2018) argue that while the MSE-optimal bandwidth yields an MSE-optimal RD treatment effect estimator, an auxiliary bandwidth should be used to obtain a robust bias-corrected

¹¹The regression model here is written in terms of the margin of victory for the winning candidate. However, this is equivalent to a model in terms of the nonwhite candidate's margin of victory where the treatment effect is estimated at the zero vote margin. In the latter model, the running variable would be negative in cities where the white candidate won and positive in cities where the nonwhite candidate won.

¹²See Cattaneo and Vazquez-Bare (2016) for a recent overview of RD bandwidth selection methods.

(RBC) confidence interval which minimizes the asymptotic coverage error rate. Note that the RBC confidence intervals are not centred at the MSE-optimal point estimate; they are centred at a bias-corrected point estimate. However, Cattaneo et al. (forthcoming) recommend reporting the an MSE-optimal point estimate in conjunction the RBC confidence interval. I will follow this guideline for all my reported results.

The RD estimation procedure developed by Calonico et al. (2017) (CCFT, henceforth) for STATA, called **rdrobust**, implements a data-driven calculation of both the main bandwidth used to obtain an MSE-optimal point estimate and an auxiliary bandwidth to obtain a robust bias-corrected confidence interval for a given choice of the polynomial order. Note that the current implementation of **rdrobust** is not able to automatically drop any redundant dummy variables. Most cities in California conduct city council elections during even-numbered years. This can lead to all elections from some (odd) years getting dropped from the final estimation sample as the algorithm determines the optimal bandwidth, which in turn makes some of the year dummy variables redundant. These redundant dummy variables lead to an error in the bandwidth selection procedure. Hence, for all my outcome variables, I first obtain the bandwidths for a cross-sectional specification without any fixed effects and then include year fixed effects in my final estimation.

While I focus on the CCFT optimal bandwidths for each outcome variable, I also show estimates for the main outcome variables using consistent bandwidths of 4.4% for the main bandwidth and 7.9% for the auxiliary bandwidth. These are the optimal bandwidths derived for the effect of nonwhite candidate's victory on the number of white candidates per seat in the next election.

Lastly, Cattaneo et al. (forthcoming) recommend using the triangular kernel for its desirable asymptotic optimality properties. The triangular kernel function assigns the maximum weight to observations at the cut-off. The weight declines symmetrically and linearly for observations further away from the cut-off. A zero weight is assigned to all observations with the margin of victory outside the chosen bandwidth. Standard errors are clustered at the city level.

3.2 Race Assignment and Measurement Error

Of course, last names are not perfect predictors of race, which could result in a biased estimate of the treatment effect. In particular there are two kinds of measurement errors that need to be considered. The first is related to the selection of the sample. One could use, for instance, a plurality criterion, where a candidate is assigned a race based on which race category do people in the U.S. with the same last name associate most with. However, using such a weak criterion could include elections which are, in fact, not an election between a white and a nonwhite candidate. On the other hand, using a very strict race assignment criterion, say 80%, whereby a candidate is assigned a race only if people with the same last name associate with that race at least 80% of the

time, could lead to a much smaller sample and drop a number of elections which are, in fact, an election between a white and a nonwhite candidate. It is difficult to predict how this should bias my results, since the bias depends on the treatment effect in the (mistakenly) excluded sample, as well as the treatment effect in the (mistakenly) included sample.

Moreover, tightening the race assignment criteria invariably linked with the “whiteness” or “non-whiteness” of the last names of these candidates. Hence, while a higher threshold gives a more accurate sample, it also changes the quality of the sample such that only elections involving candidates with *distinctly* white and nonwhite names will be included in the estimation. Note that this does not negate the finding, but perhaps warrants a narrower interpretation of the result at higher thresholds, that election of candidates with distinctly nonwhite names induces more white candidates to run in the next election.

Second, *for a given sample*, the measurement errors in the outcome variable and the assignment to treatment need to be considered. It is straightforward to work out the bias in the estimate of the treatment effect due to these measurement errors. For simplicity, consider a local randomization RD framework, where the true model is given by:

$$Y^* = \eta + X^* \beta + \epsilon$$

where Y^* is the outcome variable, for instance, the number of white candidates running for election in a city, X^* is a dummy variable for victory of a nonwhite candidate against a white candidate in the previous election in the same city, and ϵ is an error term that is uncorrelated with X^* . I have omitted city and time subscripts.

I do not observe Y^* and X^* . Instead, I observe $Y = Y^* + \eta$, where η represents the measurement error in counting the number of white candidates. This measurement error is due to some nonwhite candidates being counted as white, as well as some white candidates being counted as nonwhite or unassigned. Moreover, I observe $X = X^* + \mu$, where μ represents the measurement error from classifying elections with the nonwhite candidate’s victory as elections with victory for the white candidate, or vice versa.

It can then be shown that:

$$\text{plim } \hat{\beta} = (1 - \gamma^{NW} - \gamma^W)(1 - \alpha - \rho)\beta + \gamma^{NW}\beta_{CX} \quad (2)$$

where

- γ^{NW} is the probability that a nonwhite candidate is misidentified as white (leading to false positives)

- γ^W is the probability that a white candidate is misidentified as either nonwhite or unsigned (leading to false negatives)
- α is the proportion of elections classified as nonwhite victories that are, in fact, nonwhite losses
- ρ is the proportion of elections classified as nonwhite losses that are, in fact, white losses
- C is the total number of candidates running the election

Appendix B1 includes a complete derivation. Note that I cannot sign the bias. The first term shows that any measurement error will bias my estimate downward. However, the second term could lead to an upward bias in my estimate because of counting some nonwhite candidates as white. However, note that β_{CX} in this case refers to the RD estimate with the total number of candidates running in the election as the outcome variable within the sample being used (i.e. X ; not X^*). Since the CEDA dataset includes the total number of candidates running in every election, I can get a consistent estimate of β_{CX} . As part of my main results, I verify that the total number of candidates does not change significantly at the discontinuity.

In Appendix B2, I further assess the quality of race assignment using the census surname files, I use data published along with a recent study by Beach and Jones (2017). The authors also used the CEDA dataset for election years between 2005 and 2011 to study the effects of increasing ethnic diversity in the city council on public spending. The authors collected ethnicity information for city council candidates through various sources. This data allows me to estimate γ^W and γ^{NW} under different assignment criteria (under the assumption that the authors' ethnicity information is accurate). I use the 70% threshold as my preferred race criterion. For details on the choice of the 70% criterion, refer to Appendix B2. Most importantly, under this criterion, γ^{NW} falls to around 10%, severely limiting the upward bias of the estimate. Hence, I am confident that my treatment effect estimates are downward biased.

3.3 A note on NamePrism and Bayesian Updating

Another tool to predict race using names that is gaining popularity in academic research is called NamePrism (Ye et al., 2017).¹³ I did not use NamePrism as my preferred source of race assignment since in the sample of candidates from Beach and Jones (2017) it performed much worse than the census surname files in predicting race. However, as a robustness check, I also show the estimates for my main results using NamePrism's race assignment. For more details on NamePrism and the quality of race assignment with it, see the Appendix C.

¹³See Diamond et al. (2018) for a recent example.

Elliott et al. (2009) show that combining the census surname list with geography-based information using Bayes' rule performs better in predicting ethnicity than methods based on names or geography alone. In my context, using Bayesian updating based on the city's racial make-up to calculate posterior probabilities could lead to endogeneity issues since the assignment to treatment (x-variable) as well as the count of white candidates in the next election (y-variable) will both become functions of the city's racial composition. Hence, I do not combine city demographics with the census surname files. In my robustness checks, I re-estimate the main treatment effects using Bayesian updating either for the treatment assignment (to obtain a better sample) or the counting of candidates running/elected (to reduce measurement error in the outcome variable). The main results hold in both cases.

3.4 Validity of the Regression Discontinuity Design

An important validation test for an RD design involves examining whether the treated and the control groups are similar near the cut-off in terms of observable characteristics determined before the treatment assignment. Thus, except for the treatment status, other observable characteristics should not be changing discontinuously at the cut-off. This requires estimating equation (1) with the available predetermined characteristics as the outcome variables.

In particular, I test for discontinuities in two sets of predetermined characteristics. The first set includes the characteristics of the close elections which are used to determine the treatment status. These include the number of candidates running per seat (by race), number of candidates elected per seat (by race), number of seats being contested, turnout, and the incumbency status of the white and nonwhite candidates used to assign treatment. The second set of characteristics includes the 1990 demographics of the cities where I observe close elections.

The results of these continuity tests are presented in Table 2. The only variables to change at the discontinuity are the number of white candidates getting elected per seat and the number of nonwhite candidates getting elected per seat in the year where I observe the close election. This is the treatment assignment at the cut-off. Appendix Table A1 presents the results of the continuity tests using consistent bandwidths of 4.4% (main) and 7.9% (auxiliary) for all variables.

Next, using the procedures recommended by McCrary (2008) and Cattaneo et al. (2019), I test for manipulation of the running variable - as represented by excess mass on either side of the cut-off - in the close elections involving a white and a nonwhite candidate. Figures 1a and 1b illustrate the continuity of the running variable, nonwhite candidates victory margin. Figure 1a is based on the McCrary test. The dashed lines represent the 95 percent confidence interval around the estimated density. The discontinuity estimate based on McCrary's test (log difference in height) is -0.1355 with a standard error of 0.2394. Figure 1b is based on the density test proposed by Cattaneo et

al. (2019). This figure also plots the estimated density and a 95% confidence interval around the estimated density. The p-value for the test for a discontinuity in the running variable at the cut-off is 0.8901. Hence, I do not find any evidence of a discontinuity in the density at the zero-vote margin.

4 Main Results

4.1 RD Visualization

Figures 2a-2d present the visualizations of my RD design for the four main outcome variables - white candidates per seat, whites elected per seat, white newcomers running per seat, and white newcomers elected per seat. The plot for each dependent variable is restricted to the observations that fall within the MSE-optimal bandwidth (main bandwidth for point estimation). The figures display the linear slopes estimated separately on each side of the cut-off using the election-level data and the procedure developed by Calonico et al. (2017).¹⁴ I also overlay a scatter plot of 0.4pp bin averages of the dependent variable weighted by the number of observations in each bin. While these plots do not account for the year fixed effects used in the final estimation.

4.2 Candidates per Seat

For all my main results in Tables 3-6, I show estimates in Panel A from my preferred specification, which uses optimal data-driven bandwidths to obtain an MSE optimal point estimate and to construct an RBC confidence interval. Panel B presents the results using consistent bandwidths of 4.4% and 7.9% for point estimation and optimal inference, respectively.

Table 3 presents the results for the first set of outcomes pertaining to the number of candidates per seat. Column 1 shows that in cities where the nonwhite candidate won a close election against a white candidate compared to cities where the nonwhite candidate lost, more white candidates run in the next election. The estimated effect is approximately 0.52 candidates per seat. On average, about 2.2 seats are open in an election within the main sample (Table 1a), implying that, on average, about 1.1 more white candidate runs in cities with nonwhite victories compared to cities where the white candidates win. Columns 2 shows a decrease in the number of nonwhite candidates per seat running in the next election after the nonwhite candidate's victory in a close election. However, it is estimated imprecisely. I also find that the number of candidates that were not assigned a race also decreases at the cut-off (column 3). However, the point estimate is much

¹⁴These would be same as slopes estimated in equation (1) without the year fixed effects.

lower than the increase in the number of white candidates per seat. Hence, the increase in the number of white candidates could not simply be driven by a measurement error. Moreover, the point estimate for the discontinuity in the number of unassigned candidates (-0.18) is the same as the point estimate for the balance test conducted for the number of unassigned candidates at $t = 0$ (though, the estimate for the balance test was not statistically significant). Hence, there is practically no change in the number of unassigned candidates due to the victory of the nonwhite candidate.

Column 4 shows that the number of total candidates per seat in the next election does not change significantly at the cut-off. As discussed in Section 3.2, the additive bias in the measurement error is the product of the probability that a nonwhite candidate is classified as white (γ^{NW}) and the change in the total number of candidates at the cut-off (equation (2)). The point estimate of the change in the total number of candidates per seat at the cut-off is only 0.11. Based on the ethnicity data for a sample of candidates from Beach and Jones (2017), I estimated γ^{NW} to be about 0.1 under the 70% criterion (see Appendix B2 for details). This would imply that the upward bias in the treatment effect is only about 0.011, which is negligible compared to the estimated treatment effect on the number of white candidates per seat in the next election.

Table 4 explores the source of the change in the number of candidates by race. Columns 1-4 focus on white candidates per seat, while columns 5-8 focus on nonwhite candidates per seat. Column 1 shows that the increase in white candidates per seat in the next election is partly driven by the number of runners-up from the previous election (the close election used in the RD). This increase is driven by the fact that the cities in the treatment group have an additional white runner-up compared to the cities in the control group. If I remove the first-place runner-up who lost the close election to a nonwhite candidate, I see that the the number of other runners-up does not change at the cut-off. Columns 5 and 6 show analogous results for the nonwhite candidates, though the estimates are not statistically significant from zero. Columns 3 and 7 show that the number of incumbents running per seat also does not change at the cut-off for either whites or nonwhites, respectively.

In column 4, I find that the victory of a nonwhite candidate against a white candidate in a close election leads to the entry of white newcomers - candidates who are neither runners-up, nor incumbents. However, I do not find evidence of a similar response by nonwhite citizen-candidates to the election of a white candidate (column 8). While the point estimate is, in fact, negative, i.e. there are more newcomers in cities where the white candidate won, it is statistically insignificant. Moreover, the point estimate for nonwhite newcomers (≈ -0.12) is much smaller than the point estimate for white newcomers (≈ 0.28).

4.3 Elected per Seat

Next, I study how the composition of those that get elected in the next election differs at the cut-off. Table 5 presents the results for discontinuities in the number of candidates elected per seat. Column 1 shows that in cities where the nonwhite candidate won a close election against a white candidate compared to cities where the nonwhite candidate lost, more white candidates get elected in the next election. Columns 2 and 3 indicate that the higher number of whites getting elected is explained by a lower number of nonwhite and unassigned candidates getting elected, respectively. Though I do not have enough power to reject the null hypothesis of no change in the number of nonwhites elected at the cut-off.

Table 6 explores the source of the change in the elected candidates. Columns 1-4 focus on white candidates per seat, while columns 5-8 focus on nonwhite candidates per seat. Column 1 shows that more white runners-up get elected to the city council in the next election in cities where the nonwhite candidate won the close election. However, when I remove the first-place runner-up, I find no evidence of a change in the number of other runners-up getting elected (column 2). This suggests that the result in column 1 is driven by the “runner-up effect” from Anagol and Fujiwara (2016). I find a symmetric, though smaller, runner-up effect for nonwhite candidates (columns 5 and 6). Columns 3 and 7 show that the number of incumbents (re-)elected per seat also does not change at the cut-off for either whites or nonwhites, respectively.

In column 4, I find that some of the newcomers that entered the race after the victory of a nonwhite candidate are able to get elected to the city council. However there is no change at the cut-off in the number of nonwhite newcomers getting elected (column 8).

4.4 Robustness Checks

I present an extensive set of robustness checks for the main results pertaining to the number of white candidates per seat, number of whites elected per seat, number of white newcomers running per seat, and the number white newcomers elected per seat. Here, I will focus my discussion on the two most important results of the paper: the higher number white newcomers running in the next election after the success of a nonwhite candidate, and the higher number of white newcomers getting elected in the next election.

Appendix Figures A1 (white candidates per seat), A2 (whites elected per seat), A3 (white newcomers running per seat), and A4 (white newcomers elected per seat) show the sensitivity of the main results to changes in bandwidth and the polynomial order. Each figure includes three graphs based on the choice of polynomial order: (a) linear, (b) quadratic, and (c) cubic. In each

graph, I plot the coefficient estimates from RD design for a range of main bandwidths between 2% and 15%. For each chosen main bandwidth (h) used for point estimation, I calculate an auxiliary bandwidth (b) - which is used to construct the RBC confidence interval - using the ratio of the optimal data-driven bandwidths for the given polynomial order. The vertical line in each graph indicates the MSE-optimal data-driven bandwidth for point estimation derived using the procedure developed by Calonico et al. (2017).

Appendix Figure A3a shows how the estimated effect, the 90% standard confidence interval, and the 90% RBC confidence interval vary with changes in the main bandwidth under a linear specification for white newcomers running per seat. While the point estimate decreases with the use of larger bandwidths and the effect is statistically insignificant, note that increasing the bandwidth will increase the bias of the local linear estimate. This is because as we increase the bandwidth, a local linear specification becomes an inappropriate approximation of the underlying data. Appendix Figures A3b and A3c show that the result continues to be significant for larger bandwidths under a quadratic or cubic specification. At smaller bandwidths, the quadratic and cubic specifications give an even higher point estimate, which might be partly due to over-fitting.

Appendix Figure A4a-A4c show analogous graphs for the estimated effect on white newcomers elected per seat. Similar to the results for white newcomers running per seat, the estimate from a linear specification becomes insignificant with larger bandwidths, but the result is significant for quadratic and cubic specifications.

Appendix Figures A5 (white candidates per seat), A6 (whites elected per seat), A7 (white newcomers running per seat), and A8 (white newcomers elected per seat) show results from other robustness tests conducted. In each figure, I first show the main point estimate and a 95% RBC confidence interval for comparison. For each robustness test, new optimal bandwidths are chosen. I then plot the coefficient estimated and a 95% RBC confidence interval.

In the first two robustness tests, I use Bayes' rule to either get a more accurate sample of elections or a more accurate count of the candidates. As discussed in section 3.3, in the context of this paper, Bayesian updating can lead to endogeneity concerns as the assignment to treatment and the outcome variables will both become functions of the city's demographics. In my first robustness test, I re-estimate the treatment effect using the sample (x -variable) of elections classified as white v. nonwhite under a 70% criterion with the posterior probabilities derived using the Bayes' rule while keeping the count (y -variable) based on the raw probabilities from the census surname files. In the second robustness test, I re-estimate the effect using the count based on the the posterior probabilities derived using the Bayes' rule while using the sample based on the raw probabilities from the census surname files.

The next two robustness checks related to changing the estimation of the standard errors, first by

clustering at the county level instead of the city level, and then by using the default default procedure for variance-covariance estimation - heteroskedasticity-robust nearest neighbour variance estimator with a minimum of three neighbours used - implemented by Calonico et al. (2017). The next two tests control for white share of the population in 2000 and in 2010, respectively.

The next set of estimates test the sensitivity of the results to observations located very close to the cut-off. The “donut-hole” approach was developed primarily for cases where discontinuity is detected in the density of the running variable, which is not the case here (see Figures 1a and 1b). If any systematic manipulation of running variable occurred, it is natural to assume that the observations closest to the cut-off are those most likely to have engaged in manipulation. Here, I re-estimate the treatment effect after dropping all observations where the margin of victory is less than 0.1, 0.2, 0.3, or 0.4 percentage points. Note that 0.3 and 0.4 percentage points are quite large donut holes compared to other studies that have used this approach (e.g., Barreca et al., 2011). The overall sample of close white v. nonwhite elections used in this paper consists of 385 elections. Approximately 7 percent and 8.8 percent of these observations have a the margin of victory less than 0.3 and 0.4 percentage points, respectively. These would be an even larger share of the observations within any chosen bandwidth around the cut-off. Dropping the observations close to the cut-off inevitably involves (more) extrapolation of the linear functions on both sides of the cut-off. Furthermore, we are dropping observations that are the most important to the estimation of the RD treatment effect. The theoretical properties of a donut RD approach are not well known, and thus the results should be interpreted with caution.

Next, I test the sensitivity of the estimates to different race assignment criteria using the census surname files. The last set of robustness tests involve re-estimating the treatment effect under various race assignment criteria using race predictions from NamePrism.

Appendix Figure A7 shows the robustness tests for white newcomers running per seat. Using a sample based on Bayesian updated probabilities, the estimated effect is insignificant. The point estimate is, in fact, higher than the main estimate. However, the sample size drops significantly and thus the effect is estimated imprecisely. The estimated effect increases when using a Bayesian updated count in the next robustness check. The effect is robust to how the standard errors are estimated or controlling for the white share of the population in 2000 or 2010. The effect does appear to be highly sensitive the observations closest to the cut-off, as shown by the results from the donut RDs. This certainly deserves further investigation in future work. The effect is very robust to different race assignment criteria using both the census surname files or NamePrism. The effect is statistically insignificant under the 80% criterion using NamePrism. However, the point estimate is very similar to the main effect and the insignificance appears to be due to the drop in the number of observations. This could also be, in part, due the an increase in the amount of downward bias due to measurement error at the 80% threshold for white candidates (Appendix Figure C2).

Appendix Figure A8 shows the robustness tests for white newcomers elected per seat. The estimate is robust to using a Bayesian updated sample or count, using different procedures to estimating the standard errors, or controlling for white share of the population in 2000 or 2010. Similar to the result for newcomers running in the next election, the result for newcomers elected also appears to be sensitive to the exclusion of the observations closest to the cut-off. When re-estimating using the different race assignment criteria with the census surname files, the effect is statistically insignificant under the Plurality or 50% criterion. However, this is probably due to a less accurate sample. The estimated effect is much smaller and statically indistinguishable from zero under any race assignment criterion using NamePrism.

5 Interpretation of the Results

In this section, I explore various potential mechanisms behind the entry of new white candidates after the election of a nonwhite candidate to the city council.

5.1 Voter Preferences

One of the reasons behind the entry of new white candidates could be that these candidates are anticipating differences in voter preferences between cities where the nonwhite candidate won the previous close election and cities where the nonwhite candidate lost. These could be preferences along racial lines. For instance, a preference for maintaining a racial balance on the city council. Such a preference would result in voters wanting to elect more white candidates in the following election. This could also be driven by voters' opinions about the particular white or nonwhite candidate elected in the close elections, for instance, because of their policy choices. This could reduce the preference for the incumbent's racial group, potentially benefiting new white candidates.¹⁵

In Figures 3a-3b and Appendix Tables A2a-A2b, I test for differences in voter preferences after the election of the nonwhite candidate. To test for voter differences along racial lines, I restrict myself to elections (in the next election year) where I observe at least one white candidate and at least one nonwhite candidate. Unfortunately, I do not have granular data on voter preferences, actions, or identities. However, I try to explore changes in voter preferences through several measures. I also test for voter preferences for the winner of the RD election by looking at the incumbents probability of running and the probability of winning in the next election when the incumbent is

¹⁵While some of the white newcomers running for elections are successful in winning a seat on the city council, that is not necessarily an indication of a difference in voter preferences. More white newcomers getting elected may simply be the result of more white newcomers running in the first place.

eligible to run (two elections forward).¹⁶ For each variable, I use optimal bandwidths for point estimation and inference. I then plot the estimated coefficient and a 95% RBC confidence interval.

In Figure 3a and Appendix Table 2a, I first test for differences in turnout in the next election. If the preferences of the eligible voters are changing significantly, it might result in an increase or decrease in turnout in the next election. However, I find no significant change in the turnout at the cut-off.¹⁷ Another way I test for a change in voter preferences is using indicator variables for whether the highest-ranked candidate in the election is white or nonwhite. If voters in cities where the nonwhite candidate won the close election have a strong preference for white council members in the next election, it is likely that the most successful (highest-ranked) candidate in the next election will be white. While I do see an increase in the probability that the highest-ranked candidate is white ($\approx 16\%$) and a corresponding decrease in the probability that the highest-ranked candidate is nonwhite ($\approx 20\%$), the estimates are statistically indistinguishable from zero. The next set of results using indicator variables for the lowest-ranked candidate's racial group sheds more light on this. Here, at the cut-off, I observe a large and statistically significant increase in the likelihood that the lowest-ranked is white. This suggests that the previous result where I observe an increase in the likelihood of the highest-ranked candidate being white is likely driven by the compositional changes of the candidates rather than changes in voter preferences; there are more white candidates running, some of whom end up being the highest-ranked candidates, and some end up being the lowest-ranked candidates. Hence, I am reluctant to interpret this as a significant change in voter preferences along racial lines.

In Figure 3b and Appendix Table 2b, I study voter preferences through vote shares in the next election (for highest- and lowest-ranked candidates by race) and the election after that (for winners of the RD elections). First, I consider the vote share of the highest-ranked candidate *within* the racial groups: the vote share of the highest-ranked white candidate and the highest-ranked nonwhite candidate. If voters in cities where the nonwhite candidate won the close election have a strong preference for white council members in the next election, it is likely that the most successful white (nonwhite) candidate in these cities received a higher (lower) vote share than the most successful white (nonwhite) candidate in cities where the nonwhite candidate lost the previous election. I find no evidence of such a difference in the vote shares of the highest-ranked white or nonwhite candidates at the cut-off. Similarly, I test for differences in the vote shares of the lowest-ranked white and nonwhite candidates at the cut-off. Once again the differences are small and statistically insignificant.

The last set of results in Figure 3a relates to the winners of the close election (RD sample). I find

¹⁶I am using the term eligibility loosely here. By eligibility, I simply mean that I am looking at an election two periods after the close election. Candidate's actual eligibility status will depend on the term limit in the city, as well as how many terms the candidate has been in the office.

¹⁷Admittedly, there might still be a change in the composition of the voters. Unfortunately, I do not have data on voter turnout by racial groups.

that, at the cut-off, nonwhite candidates are less likely to run for re-election and also, conditional on running, less likely to be re-elected. However, the results are not statistically significant. Furthermore, I find that, at the cut-off, nonwhite incumbents from the close election, conditional on running, receive higher vote share than the white candidates (Figure 3b and Appendix Table 2b). Again, the estimate is indistinguishable from zero. Overall, evidence of differences in voter preferences for the RD incumbents at the cut-off is rather mixed.

While these are admittedly crude measures, I do not find any evidence to conclude that the entry of new white candidates in the city council elections is driven by a change in the voter preferences along racial lines or due to voters' opinions about the candidates elected in the close election.

5.2 Political Cycles

Another potential reason for observing the entry of new white candidates after the election of a nonwhite candidate is due to a differential political cycle between cities on either side of the cut-off. This is somewhat related to the idea of voter preferences for the racial composition of the city council. It might be the case that the voters elect more nonwhite candidates in one year, then adjust their votes to elect more white candidates in the next election to keep a racially balanced city council. This might also be because seats held by a white incumbent are thought of as "white seats" within the city and thus the voters may have a preference to continue electing a white candidate for open seats previously held by a white council member. Such a preference will also result in a political cycle where the election of more nonwhite candidates is followed by the election of more white candidates.

Figures 4a-4f explore the extended dynamics across the cities on either side of the cut-off. In each graph, I show the results of estimating the effect of the nonwhite candidate's victory in election period e on the number of white candidates per seat in different election years relative to e . For each year (in each graph), I use optimal bandwidths for point estimation and inference. I then plot the estimated coefficient and a 95% RBC confidence interval. The results for election year $e + 1$ represent the main results discussed earlier in Tables 3-6. The figures also serve as an extended robustness check for the continuity of predetermined covariates at the cut-off.

A word of caution is warranted here. The number of observations used to obtain each confidence interval (in each graph) is different. This is primarily because I observe a fix set of elections between 1996 and 2017. Hence, for instance, for elections in year 1996 or 1997 that are classified as close white v. nonwhite elections used in the RD sample, I do not observe the previous elections, which would have been held in 1994 or 1995. Similarly, for elections in 1998 or 1999, I will observe the prior election (years 1996 and 1997), but not the election two cycles prior (years 1994 and 1995). Analogously, for close white v. nonwhite elections in 2014 or 2015, I only observe the election in

the next cycle and not in period beyond that.¹⁸

Figure 4a shows no differential cyclicalities in white candidacy prior to election year $e + 1$. Similarly, Figure 4b shows no differential political cycles in the racial composition of the successful candidates prior to period e . In period e , fewer white candidates get elected, which is precisely the treatment (also shown in Table 2). In the period $e + 1$, I observe that more white candidates get elected (also shown in Table 5). While statistically insignificant, there is a drop in the number of whites that get elected in period $e + 2$. However, this is not an indication of a differential political cycle across the cities in the treatment and control groups. Figures 4c and 4d shed further light on the source of this drop. Figure 4c shows a drop in white incumbents running for election in period $e + 2$. However, in conjunction with Figure 4b, it is clear that this drop is due to fewer whites getting elected in period e , which would naturally result in fewer white incumbents in period $e + 2$. Figure 4d shows that as a result of fewer white incumbents running, fewer white incumbents get elected in period $e + 2$, which drives the estimate of fewer whites elected overall in period $e + 2$.

In Figures 4e and 4f, I study the dynamics of white newcomers running and white newcomers elected over time. I find no evidence of any differential cyclicalities in either case. Moreover, I find that the entry of new white candidates after nonwhite candidate's victory is only a short-term response. After period $e + 1$, the number of new entrants at the cut-off does not change discontinuously. Similarly, I find no discontinuity in the number of white newcomers getting elected after the election year $e + 1$.

Overall, these graphs show no evidence of differential political cycles in the racial composition of the candidates running for or elected to the city council. Hence, it is unlikely that the main results of the paper pertaining to the entry of white candidates and the election of these white newcomers is driven by a differential political cycle.

5.3 Public Policy

In this section, I study whether the entry of the new candidates is driven by policy changes implemented by the new city council after the election of the nonwhite candidate. City councils are most influential in budgetary matters, and thus I use public finance data from the cities' Annual Financial Transaction Reports. Unfortunately, the data is only available starting in the fiscal year 2002-2003. Moreover, beginning in the fiscal year 2016-2017, the State Controller's Office adopted new accounting standards, making the data for that fiscal year onwards incomparable to the previous years.

¹⁸Sample size in each period: $n_{e-3} = 199$, $n_{e-2} = 253$, $n_{e-1} = 321$, $n_e = 385$, $n_{e+1} = 385$, $n_{e+2} = 309$, $n_{e+3} = 236$. Note that these are the full sample sizes. The effective number of observation used in the estimation will vary with the optimal bandwidth chosen for each plotted coefficient.

To study the effect of the nonwhite candidate's victory on city's finances, I take the average of the data available over the two years prior to the elections and compare them with the average of the two years after the election. It is appropriate to issue a caveat here. The fiscal year for cities runs from July to June, while in most cities the elections occur in November. Hence, for a candidate elected in, say, Nov 2006, the period before the elections consists of the fiscal years 2004-05 and 2005-06, and the period after the election consists of the fiscal years 2006-07 and 2007-08. Thus, the elected candidate was only on the city council for part of the first period considered in the post-election period.

Figure 5 and Appendix Table A3 present the analysis of the public finance data. For each variable, I use optimal bandwidths for point estimation and inference. I then plot the estimated coefficient and a 95% RBC confidence interval. Figure 5a focuses on the overall size of the budget, as measured by revenue, expenditure, and the surplus ratio.¹⁹ I find that compared to the cities where the nonwhite candidate lost the close election, cities where the nonwhite candidate was successful increased their revenue by about 12% (significant at 1%) and expenditure by approximately 9% (significant at 10%). This is a significantly large increase in the budget size in a very short period of time.²⁰

Figure 5b shows the changes in budget preferences at the cut-off. I look at the allocation of expenditure across the budget categories defined in the financial reports: general expenditure, public safety, transportation, community development, health, culture and leisure, public utilities, and other expenses. I find a 1.7 percentage points decrease in the share of the expenditure allocated to public utilities (significant at 10%) in cities where the nonwhite candidate won compared to cities where the white candidate won. There appear to be economically significant increases in the share of expenditure allocated to transportation (1.4%) and to community development (2.4%). However, these are statistically indistinguishable from zero.

The significant change in revenue warrants further investigation in future work. There is data available in the Annual Financial Transaction Reports on the various sources of the revenue. For now, to examine whether the changes in the city finances - in revenue to be precise - are driving the entry of white candidates and their electoral success in the next election, I re-estimate my main RD treatment effects by controlling for these changes. The results are presented in Figure 6 and Appendix Table A4. Each estimate is derived using optimal bandwidths for point estimation and a 95% RBC confidence interval.

Once again, I will focus on the results for white newcomers running and white newcomers elected in the next election. I first plot the main estimate from the full sample for comparison. Next,

¹⁹The surplus ratio is computed as the surplus (revenue minus expenditure) divided by the total revenue.

²⁰I have re-estimated the change in revenue across a range of bandwidths and polynomial orders and the result is robust. Results of the robustness test are available upon request.

I restrict myself to the set of elections used in the analysis of public finance data (same as in Figure 5 and Appendix Table A3), and then re-estimate the treatment effect without any controls. Lastly, I use the public finance subsample and control for the change in revenue observed after the election of the nonwhite candidate. The estimates for both white newcomers running and white newcomers elected in the next election are smaller in the public finance subsample compared to the main estimate, and they are statistically insignificant. This is partly due to the loss of power from having fewer observations. However, in both cases, when I control for the change in revenue, my point estimate moves very little, suggesting that the change in revenue has very little influence on the outcome variables.

Thus, I do not find strong evidence in support of the hypothesis that the entry of new white candidates or the success of these entrants in the next election is driven by policy changes implemented after the election of the nonwhite candidate.

5.4 Demographic Change and Perceived Threat

There is a burgeoning literature in political science that studies the politicization of white identity triggered by the change in the racial composition of the U.S. In her 2014 dissertation, *The Demise of Dominance: Group Threat and the New Relevance of White Identity for American Politics*, Ashley Jardina argues that in the past political scientists have concluded that there was no such thing as white identity when applied to politics. Researchers concluded that “the experience of being white in the U.S., and the privileges and advantages white individuals incur as a result of their objective race, make it unlikely that their race comprises a salient identity.” On the other hand, researchers consistently found evidence that racial identity shaped the political attitudes and behaviours of racial minority groups like blacks (e.g Chong and Rogers, 2005), Latinos (e.g Barreto and Pedraza, 2009), and (e.g Junn and Masuoka, 2008).

Jardina (2014) argues that this it is not the case white identity is not salient, but that it is conditional. She argues that in the past white identity has been salient to whites’ political attitudes in times when whites perceive a threat to their dominant status in the American racial hierarchy. Jardina (2014) then uses survey data to show that as immigration rates in the U.S. grew in the 1990s, presenting a threat to the existing racial hierarchy, white identity became more relevant to opinions of white Americans on immigration policy. In summarizing her argument, she writes:

“The key to understanding the formation and import of identity among dominant groups, I argue, is in perceptions of threat; for such groups, identity becomes salient in reaction to beliefs about the relatively threatened or waning status of the group. White Americans, in particular, are responding to the *threat of population changes* and the elec-

toral success of non-white candidates like Barack Obama.” [p. 5, Jardina (2014); emphasis added]

In the quote above, Jardina identifies two particular threats: threat of population changes, and the electoral success of nonwhite candidates. Here, I will test for any possible linkages between these two threats. The idea behind my empirical strategy in this section is to test if the electoral success of the nonwhite candidate is possibly making demographic changes more salient to the whites in these cities. If this is the case, then whites living in cities with bigger demographic changes in recent decades will likely perceive a bigger threat to their dominant status after the election of the nonwhite candidate, compelling more of them to run for city council in an attempt to gain back their dominance in local government.

To test whether the observed entry of new white candidates is driven by the change in the racial composition of the cities and the perceived threat associated with it, I stratify my RD sample of white v. nonwhite elections based on the percentage change in the nonwhite share of the population between 1990 and 2010. I then re-estimate the RD treatment effects within the subsample of cities with below or above median change in the nonwhite share. Note that in these estimates, cities on either side of the cut-off have gone through similar - big or small - changes in the nonwhite share of the population. Hence, the only difference at the cut-off is the treatment, or the electoral success of the nonwhite candidate.

Figure 7 and Appendix Tables 5a-5b present the results of this analysis. In Figure 7, each estimate is derived using optimal bandwidths for point estimation and a 95% RBC confidence interval. On the left, I present the balance tests within each subsample using data from the same year as the close election used in the RD estimation. I do not find any evidence of a discontinuity at the cut-off in the number of white newcomers running in the close election in either the cities with below or above median change in the nonwhite share of the population. Note that the estimate is negative (and statistically significant) for the number of white newcomers elected in the close election. This is simply the consequence of the treatment. Fewer whites get elected and fewer white newcomers get elected.

Now consider the estimates for the number of white newcomers in the next election. I cannot reject the null hypothesis that the coefficient in the subsample with below median change is the equal to the coefficient in the subsample with above median change since their 95% confidence intervals overlap. However, the point estimate suggests that in the subsample with higher than median change, about 0.4 new white candidates ran for every open seat on the city council after the election of the nonwhite candidate. The estimate is significant at the 10% level. On the other hand, the point estimate in the subsample with lower the median change is only 0.07 new white candidates running per seat, and indistinguishable from zero at any conventional significance levels. This suggests that the results in the overall sample are driven by the subsample of cities

that have experienced big changes in their racial composition between 1990 and 2010.

The estimates for the number of white newcomers elected in the next election are insignificant in both subsamples and the difference in coefficients is not as stark as in the case for white newcomers running per seat. The estimates for the total number of white candidates elected in the next election can further shed some light on this. The estimate in the subsample below the median is slightly higher than the estimate in the subsample above the median, though the confidence intervals overlap significantly. Moreover, the estimate in the above median sample is statistically insignificant. Hence, there is likely no significant difference in the voters' preferences at the cut-off for white newcomers within the subsamples below or above median.

The heterogeneity analysis suggests that the entry of new white candidates is, in part, driven by demographic changes and the associated perceived threat by these white citizen-candidates. However, I do not find any evidence that the voters are also responding to such a perception of threat.²¹

6 Conclusion

In this paper, I study the effects of increasing nonwhite representation in local elections. I exploit close elections between the last-place winner and the first-place loser in at-large city council elections across California where one of these candidates is white and the other is nonwhite. In cities where the nonwhite candidate wins the close election compared to cities where the nonwhite candidate loses, I find an increase in the number white candidates per seat running in the next election. I find that this increase is not simply the result of having more runners-up from the previous election who choose to run again. In fact, I find an increase in white newcomers running per seat at the cut-off.

I also study the electoral outcomes in the next election, and some of the white newcomers are successful in their bid for a seat on the city council. However, I do not find that this is a result of a change in voter preferences along racial lines or due to differential political cycles. Thus, it does not appear that the entry of new white candidates is induced by an anticipation of voter preferences for such candidates. Moreover, I find little evidence that the white entrants are motivated by the policy changes occurring after the election of the nonwhite candidate.

In my heterogeneity analysis, I find that the entry of new white candidates is mostly occurring in cities that have gone through bigger demographic changes as measured by the change in the

²¹This also provides further evidence in support of the conclusions in Section 5.1 that the success of the white entrants is not driven by changes in voter preferences.

nonwhite share of the population. This suggests that that the entry of new white candidates is, in part, driven by demographic changes and the associated perceived threat by these white citizen-candidates. However, I do not find any evidence that the voters are also responding to such a perception of threat.

References

- Aigner, Dennis J. 1973. "Regression with a binary independent variable subject to errors of observation." *Journal of Econometrics* 1 (1): 49–59.
- Anagol, Santosh, and Thomas Fujiwara. 2016. "The runner-up effect." *Journal of Political Economy* 124 (4): 927–991.
- Barreca, Alan I, Melanie Guldi, Jason M Lindo, and Glen R Waddell. 2011. "Saving babies? Revisiting the effect of very low birth weight classification." *The Quarterly Journal of Economics* 126 (4): 2117–2123.
- Barreto, Matt A, and Francisco I Pedraza. 2009. "The renewal and persistence of group identification in American politics." *Electoral Studies* 28 (4): 595–605.
- Beach, Brian, and Daniel B Jones. 2017. "Gridlock: Ethnic diversity in government and the provision of public goods." *American Economic Journal: Economic Policy* 9 (1): 112–36.
- Beach, Brian, Daniel B Jones, Tate Twinam, and Randall Walsh. 2018. "Minority Representation in Local Government." Working paper.
- Bhalotra, Sonia, Irma Clots-Figueras, and Lakshmi Iyer. 2018. "Pathbreakers? Women's Electoral Success and Future Political Participation." *The Economic Journal* 128 (613): 1844–1878.
- Brown, Ryan, Hani Mansour, and Stephen D O'Connell. 2018. "Closing the Gender Gap in Leadership Positions: Can Expanding the Pipeline Increase Parity?"
- Calonico, Sebastian, Matias D Cattaneo, and Max H Farrell. 2018. "Optimal Bandwidth Choice for Robust Bias Corrected Inference in Regression Discontinuity Designs." Working paper.
- Calonico, Sebastian, Matias D Cattaneo, Max H Farrell, and Roco Titiunik. 2017. "rdrobust: Software for regression-discontinuity designs." *The Stata Journal* 17 (2): 372–404.
- Cattaneo, Matias D, Nicolás Idrobo, and Roco Titiunik. Forthcoming. "A practical introduction to regression discontinuity designs." *Cambridge Elements: Quantitative and Computational Methods for Social Science-Cambridge University Press I*.
- Cattaneo, Matias D, Michael Jansson, and Xinwei Ma. 2019. "Simple local polynomial density estimators." *Journal of the American Statistical Association*, 1–7.
- Cattaneo, Matias D, and Gonzalo Vazquez-Bare. 2016. "The choice of neighborhood in regression discontinuity designs." *Observational Studies* 2 (134): A146.
- Chattopadhyay, Raghavendra, and Esther Duflo. 2004. "Women as policy makers: Evidence from a randomized policy experiment in India." *Econometrica* 72 (5): 1409–1443.
- Chong, Dennis, and Reuel Rogers. 2005. "Racial solidarity and political participation." *Political Behavior* 27 (4): 347–374.

- Colby, Sandra, and Jennifer M Ortman. 2015. "Projections of the size and composition of the US population: 2014 to 2060." *Current Population Reports*, 25–43.
- Craig, Maureen A, and Jennifer A Richeson. 2014a. "More diverse yet less tolerant? How the increasingly diverse racial landscape affects white Americans racial attitudes." *Personality and Social Psychology Bulletin* 40 (6): 750–761.
- . 2014b. "On the precipice of a majority-minority America: Perceived status threat from the racial demographic shift affects White Americans political ideology." *Psychological Science* 25 (6): 1189–1197.
- Diamond, Rebecca, Timothy McQuade, and Franklin Qian. 2018. *The effects of rent control expansion on tenants, landlords, and inequality: Evidence from San Francisco*. Technical report. National Bureau of Economic Research.
- Elliott, Marc N, Peter A Morrison, Allen Fremont, Daniel F McCaffrey, Philip Pantoja, and Nicole Lurie. 2009. "Using the Census Bureaus surname list to improve estimates of race/ethnicity and associated disparities." *Health Services and Outcomes Research Methodology* 9 (2): 69.
- Enos, Ryan. 2014. "Causal effect of intergroup contact on exclusionary attitudes." *Proceedings of the National Academy of Sciences* 111 (10): 3699–3704.
- Gagliarducci, Stefano, and M Daniele Paserman. 2011. "Gender interactions within hierarchies: Evidence from the political arena." *The Review of Economic Studies* 79 (3): 1021–1052.
- Gangadharan, Lata, Tarun Jain, Pushkar Maitra, and Joseph Vecchi. 2016. "Social identity and governance: The behavioral response to female leaders." *European Economic Review* 90:302–325.
- Gelman, Andrew, and Guido Imbens. 2019. "Why high-order polynomials should not be used in regression discontinuity designs." *Journal of Business & Economic Statistics* 37 (3): 447–456.
- Gilardi, Fabrizio. 2015. "The temporary importance of role models for women's political representation." *American Journal of Political Science* 59 (4): 957–970.
- ICMA. 2019. "2018 Municipal Form of Government Survey Report."
- Iyer, Lakshmi, Anandi Mani, Prachi Mishra, and Petia Topalova. 2012. "The power of political voice: women's political representation and crime in India." *American Economic Journal: Applied Economics* 4 (4): 165–93.
- Jardina, Ashley. 2014. "Demise of Dominance: Group Threat and the New Relevance of White Identity for American Politics."
- . 2019. *White identity politics*. Cambridge University Press.
- Junn, Jane, and Natalie Masuoka. 2008. "Asian American identity: Shared racial status and political context." *Perspectives on Politics* 6 (4): 729–740.

- Logan, Trevon D. 2018. "Do Black Politicians Matter?" Working paper.
- Manson, Steven, Jonathan Schroeder, David Van Riper, and Steven Ruggles. 2019. "IPUMS National Historical Geographic Information System: Version 14.0."
- McCrary, Justin. 2008. "Manipulation of the running variable in the regression discontinuity design: A density test." *Journal of econometrics* 142 (2): 698–714.
- Nye, John VC, Ilia Rainer, and Thomas Stratmann. 2014. "Do black mayors improve black relative to white employment outcomes? Evidence from large US cities." *The Journal of Law, Economics, & Organization* 31 (2): 383–430.
- Pande, Rohini. 2003. "Can mandated political representation increase policy influence for disadvantaged minorities? Theory and evidence from India." *American Economic Review* 93 (4): 1132–1151.
- Pei, Zhuan, David S Lee, David Card, and Andrea Weber. 2018. "Local Polynomial Order in Regression Discontinuity Designs." Working paper.
- Trebbi, Francesco, Philippe Aghion, and Alberto Alesina. 2008. "Electoral rules and minority representation in US cities." *The Quarterly Journal of Economics* 123 (1): 325–357.
- Wetts, Rachel, and Robb Willer. 2018. "Privilege on the precipice: Perceived racial status threats lead white americans to oppose welfare programs." *Social Forces* 97 (2): 793–822.
- Ye, Junting, Shuchu Han, Yifan Hu, Baris Coskun, Meizhu Liu, Hong Qin, and Steven Skiena. 2017. "Nationality classification using name embeddings." In *Proceedings of the 2017 ACM on Conference on Information and Knowledge Management*, 1897–1906. ACM.

Tables and Figures

Table 1a: Summary Statistics of Elections

	All Elections		White v. Nonwhite		Within 4.4% BW	
	Mean	SD	Mean	SD	Mean	SD
Number of seats	2.4	0.6	2.3	0.6	2.2	0.5
Number of candidates	5.4	2.5	5.7	2.4	5.5	2.1
White candidates	3.1	2.0	2.7	1.7	2.7	1.6
Black candidates	0.0	0.1	0.0	0.1	0.0	0.1
Hispanic candidates	0.7	1.3	1.4	1.6	1.4	1.5
Asian candidates	0.1	0.5	0.3	0.7	0.3	0.6
Unassigned candidates	1.4	1.3	1.3	1.2	1.2	1.1
Whites elected	1.4	0.9	1.2	0.8	1.2	0.8
Blacks elected	0.0	0.1	0.0	0.1	0.0	0.1
Hispanics elected	0.3	0.6	0.5	0.7	0.5	0.7
Asians elected	0.1	0.2	0.1	0.3	0.1	0.4
Unassigned elected	0.6	0.7	0.4	0.6	0.4	0.6
Turnout	19,145	32,879	23,038	26,749	22,012	26,515
Observations	4011		385		231	

Notes: Summary statistics for city council elections for a full term in office. The counts of the number of candidates running and the number of candidates elected by race are based on the 70% threshold, i.e. the candidate is assigned a race if at least 70% of the people in the decennial census with the same last name identify themselves as a member of that racial group.

Table 1b: Summary Statistics of City Demographics

	All Cities		Cities with White v. Nonwhite Elections		White v. Nonwhite Elections Within 4.4% BW	
	2000	2010	2000	2010	2000	2010
Total population	52,542 [186,796]	58,204 [191,785]	48,319 [52,620]	54,886 [58,016]	50,415 [55,910]	57,415 [61,607]
White share	0.548 [0.255]	0.478 [0.253]	0.446 [0.232]	0.373 [0.223]	0.422 [0.228]	0.347 [0.216]
Nonwhite share	0.418 [0.261]	0.488 [0.262]	0.522 [0.239]	0.597 [0.232]	0.547 [0.235]	0.623 [0.227]
Black share	0.035 [0.052]	0.035 [0.047]	0.040 [0.053]	0.040 [0.049]	0.042 [0.056]	0.041 [0.051]
Hispanic share	0.299 [0.247]	0.348 [0.253]	0.381 [0.251]	0.432 [0.255]	0.395 [0.254]	0.446 [0.260]
Asian share	0.084 [0.106]	0.105 [0.128]	0.100 [0.126]	0.125 [0.151]	0.110 [0.135]	0.135 [0.160]
Other share	0.033 [0.016]	0.034 [0.017]	0.032 [0.015]	0.030 [0.016]	0.031 [0.014]	0.030 [0.016]
Number of cities	455	461	200	203	152	155

Notes: Standard deviations are reported in square brackets. These summary statistics are based on data from the 2000 and 2010 decennial censuses. Means and standard deviations for the RD sample are based on the sample of observations within the MSE-optimal bandwidth of 4.4% used to estimate the main treatment effect on white candidates per seat.

Table 2: Balance of Predetermined Covariates

	RD Estimate	Conv. SE	Robust P-value	Main BW	Aux. BW	Obs.	CG Mean
<i>Panel A: Election-related Variables</i>							
<i>Candidates per seat</i>							
Total	-0.092	0.270	0.920	5.3	8.8	253	2.429
White	0.004	0.138	0.920	6.6	11.6	282	1.167
Nonwhite	0.108	0.180	0.492	5.0	8.8	246	0.801
Unassigned	-0.181	0.107	0.148	5.4	9.3	254	0.516
<i>Elected per seat</i>							
White	-0.447***	0.053	< 0.001	7.0	12.9	287	0.504
Nonwhite	0.469***	0.053	< 0.001	7.2	12.1	291	0.310
Unassigned	-0.024	0.046	0.642	6.9	12.4	287	0.188
<i>Other election-related variables</i>							
Num. of seats	-0.112	0.159	0.544	5.4	8.3	254	2.293
Log turnout	0.161	0.290	0.563	5.4	8.4	254	9.236
Winner incumbency status	0.026	0.109	0.676	4.8	8.3	245	0.496
Loser incumbency status	0.003	0.117	0.787	4.1	8.0	218	0.217
<i>Panel B: Demographics in 1990</i>							
Log population	0.076	0.303	0.796	5.4	8.2	247	10.061
White share	0.054	0.048	0.219	7.9	14.4	292	0.520
Nonwhite share	-0.057	0.049	0.210	7.6	13.6	289	0.473
Black share	-0.019	0.012	0.121	7.5	12.5	286	0.042
Hispanic share	-0.021	0.059	0.741	5.9	9.3	260	0.357
Asian share	-0.009	0.027	0.732	5.2	8.0	245	0.079
Other share	0.000	0.001	0.979	4.4	7.6	224	0.008

Notes: Each RD specification uses optimal data-driven bandwidths for a cross-sectional specification and includes year fixed effects. The estimation is based on the default settings in the procedure developed by Calonico et al. (2017) - a local linear fit on either side of the cut-off and a triangular kernel. The main bandwidth is used to construct an MSE-optimal RD point estimator. The auxiliary bandwidth is used to construct the robust bias-corrected (RBC) confidence intervals. Note that the RBC confidence intervals are not centred at the MSE-optimal point estimate; they are centred at the bias-corrected point estimate. Standard errors are clustered at the city level.

***, **, * indicate significance at 1%, 5%, and 10% level, respectively, based on the robust p-value.

Table 3: Effect of a Nonwhite Victory on the Number of Candidates per Seat in the Next Election by Race

	White	Nonwhite	Unassigned	Total
	(1)	(2)	(3)	(4)
<i>Panel A: Optimal Bandwidths</i>				
Nonwhite wins	0.519***	-0.182	-0.182*	0.109
Conventional SE	(0.180)	(0.179)	(0.121)	(0.236)
Robust p-value	0.004	0.238	0.085	0.826
Robust 95% CI	[0.194, 0.996]	[-0.644, 0.160]	[-0.493, 0.032]	[-0.476, 0.597]
Main bandwidth	4.4	5.8	5.0	5.5
Auxiliary bandwidth	7.9	9.4	10.1	9.0
Observations	231	264	246	258
Control Group Mean	1.122	0.771	0.516	2.427
<i>Panel B: Consistent Bandwidths (h = 4.4%, b = 7.9%)</i>				
Nonwhite wins	0.519***	-0.240	-0.214*	0.065
Conventional SE	(0.180)	(0.193)	(0.125)	(0.260)
Robust p-value	0.004	0.189	0.057	0.882
Robust 95% CI	[0.194, 0.996]	[-0.707, 0.140]	[-0.543, 0.008]	[-0.534, 0.621]
Main bandwidth	4.4	4.4	4.4	4.4
Auxiliary bandwidth	7.9	7.9	7.9	7.9
Observations	231	231	231	231
Control Group Mean	1.122	0.763	0.514	2.399

Notes: In Panel A, each RD specification uses optimal data-driven bandwidths for a cross-sectional specification and includes year fixed effects. Panel B presents results using the same bandwidths for all outcome variables. All estimates are based on the default settings in the procedure developed by Calonico et al. (2017) - a local linear fit on either side of the cut-off and a triangular kernel. The main bandwidth is used to construct an MSE-optimal RD point estimator. The auxiliary bandwidth is used to construct the robust bias-corrected (RBC) confidence intervals. Note that the RBC confidence intervals are not centred at the MSE-optimal point estimate; they are centred at the bias-corrected point estimate. Standard errors are clustered at the city level.

***, **, * indicate significance at 1%, 5%, and 10% level, respectively, based on the robust p-value.

Table 4: Who Runs for City Council in the Next Election After a Nonwhite Victory?

	White Candidates per Seat				Nonwhite Candidates per Seat			
	Runners-up	Runners-up Exc. FP Runner-up	Incumbents	Newcomers	Runners-up	Runners-up Exc. FP Runner-up	Incumbents	Newcomers
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
Panel A: Optimal Bandwidths								
Nonwhite wins	0.200**	0.089	0.028	0.283*	-0.081	0.056	-0.069	-0.121
Conventional SE	(0.072)	(0.052)	(0.083)	(0.164)	(0.057)	(0.047)	(0.065)	(0.145)
Robust p-value	0.013	0.115	0.738	0.059	0.252	0.291	0.367	0.253
Robust 95% CI	[0.044, 0.371]	[-0.023, 0.216]	[-0.157, 0.221]	[-0.013, 0.705]	[-0.210, 0.055]	[-0.050, 0.168]	[-0.215, 0.079]	[-0.505, 0.133]
Main bandwidth	5.7	7.5	5.8	4.5	5.9	5.6	7.4	4.4
Auxiliary bandwidth	9.1	12.6	9.6	8.4	9.4	9.3	12.8	7.8
Observations	262	293	264	231	267	260	292	230
Control Group Mean	0.078	0.075	0.349	0.695	0.175	0.039	0.181	0.392
Panel B: Consistent Bandwidths ($h = 4.4\%$, $b = 7.9\%$)								
Nonwhite wins	0.188**	0.063	0.044	0.286**	-0.081	0.067	-0.040	-0.119
Conventional SE	(0.078)	(0.064)	(0.091)	(0.165)	(0.064)	(0.052)	(0.077)	(0.145)
Robust p-value	0.037	0.439	0.683	0.047	0.322	0.203	0.769	0.253
Robust 95% CI	[0.011, 0.359]	[-0.087, 0.201]	[-0.162, 0.248]	[0.005, 0.729]	[-0.217, 0.071]	[-0.041, 0.193]	[-0.196, 0.144]	[-0.503, 0.132]
Main bandwidth	4.4	4.4	4.4	4.4	4.4	4.4	4.4	4.4
Auxiliary bandwidth	7.9	7.9	7.9	7.9	7.9	7.9	7.9	7.9
Observations	231	231	231	231	231	231	231	231
Control Group Mean	0.060	0.060	0.366	0.695	0.179	0.041	0.195	0.389

Notes: In Panel A, each RD specification uses optimal data-driven bandwidths for a cross-sectional specification and includes year fixed effects. Panel B presents results using the same bandwidths for all outcome variables. All estimates are based on the default settings in the procedure developed by Calonico et al. (2017) - a local linear fit on either side of the cut-off and a triangular kernel. The main bandwidth is used to construct an MSE-optimal RD point estimator. The auxiliary bandwidth is used to construct the robust bias-corrected (RBC) confidence intervals. Note that the RBC confidence intervals are not centred at the MSE-optimal point estimate; they are centred at the bias-corrected point estimate. Standard errors are clustered at the city level.

***, **, * indicate significance at 1%, 5%, and 10% level, respectively, based on the robust p-value.

Table 5: Effect of a Nonwhite Victory on the Number of Elected per Seat in the Next Election by Race

	White	Nonwhite	Unassigned
	(1)	(2)	(3)
<i>Panel A: Optimal Bandwidths</i>			
Nonwhite wins	0.224**	-0.116	-0.120**
Conventional SE	(0.089)	(0.082)	(0.053)
Robust p-value	0.011	0.180	0.021
Robust 95% CI	[0.058, 0.456]	[-0.317, 0.059]	[-0.255, -0.021]
Main bandwidth	5.1	6.7	4.4
Auxiliary bandwidth	9.0	11.1	8.5
Observations	251	284	231
Control Group Mean	0.489	0.307	0.187
<i>Panel B: Consistent Bandwidths ($h = 4.4\%$, $b = 7.9\%$)</i>			
Nonwhite wins	0.248***	-0.128	-0.120**
Conventional SE	(0.092)	(0.094)	(0.053)
Robust p-value	0.008	0.196	0.021
Robust 95% CI	[0.072, 0.483]	[-0.346, 0.071]	[-0.259, -0.021]
Main bandwidth	4.4	4.4	4.4
Auxiliary bandwidth	7.9	7.9	7.9
Observations	231	231	231
Control Group Mean	0.503	0.310	0.187

Notes: In Panel A, each RD specification uses optimal data-driven bandwidths for a cross-sectional specification and includes year fixed effects. Panel B presents results using the same bandwidths for all outcome variables. All estimates are based on the default settings in the procedure developed by Calonico et al. (2017) - a local linear fit on either side of the cut-off and a triangular kernel. The main bandwidth is used to construct an MSE-optimal RD point estimator. The auxiliary bandwidth is used to construct the robust bias-corrected (RBC) confidence intervals. Note that the RBC confidence intervals are not centred at the MSE-optimal point estimate; they are centred at the bias-corrected point estimate. Standard errors are clustered at the city level.

***, **, * indicate significance at 1%, 5%, and 10% level, respectively, based on the robust p-value.

Table 6: Who Gets Elected in the Next Election After a Nonwhite Victory?

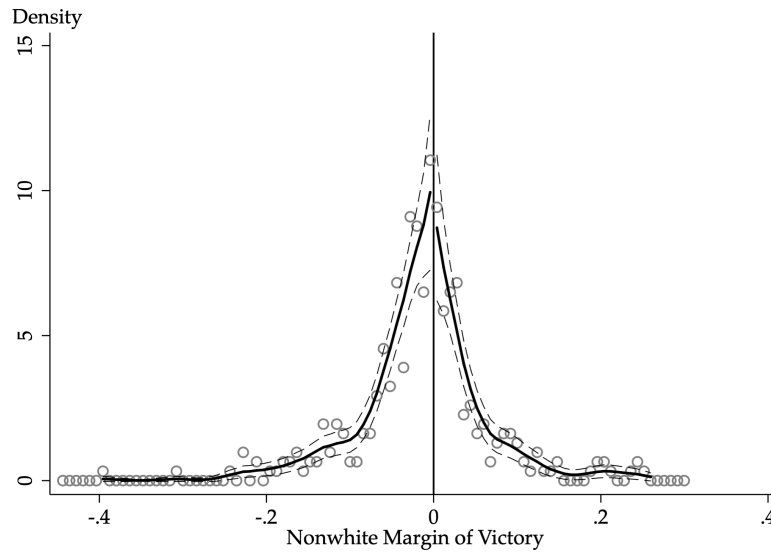
	Whites Elected per Seat				Nonwhites Elected per Seat			
	Runners-up	Runners-up Exc. FP Runner-up	Incumbents	Newcomers	Runners-up	Runners-up Exc. FP Runner-up	Incumbents	Newcomers
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
Panel A: Optimal Bandwidths								
Nonwhite wins	0.107***	0.017	-0.017	0.150**	-0.067*	0.033	-0.039	-0.019
Conventional SE	(0.031)	(0.018)	(0.065)	(0.082)	(0.034)	(0.021)	(0.061)	(0.066)
Robust p-value	0.003	0.430	0.736	0.049	0.093	0.101	0.777	0.506
Robust 95% CI	[0.039, 0.182]	[-0.026, 0.060]	[-0.170, 0.120]	[0.000, 0.363]	[-0.149, 0.012]	[-0.007, 0.083]	[-0.157, 0.118]	[-0.193, 0.095]
Main bandwidth	7.2	7.3	6.7	4.3	6.1	4.9	5.0	4.4
Auxiliary bandwidth	12.5	12.4	12.4	8.2	9.0	8.7	8.8	7.3
Observations	291	291	284	227	274	246	248	231
Control Group Mean	0.012	0.012	0.258	0.219	0.070	0.008	0.152	0.085
Panel B: Consistent Bandwidths ($h = 4.4\%$, $b = 7.9\%$)								
Nonwhite wins	0.111**	0.021	-0.012	0.148**	-0.067	0.034*	-0.042	-0.019
Conventional SE	(0.039)	(0.023)	(0.076)	(0.082)	(0.039)	(0.021)	(0.064)	(0.065)
Robust p-value	0.013	0.400	0.824	0.049	0.124	0.095	0.777	0.499
Robust 95% CI	[0.024, 0.202]	[-0.029, 0.073]	[-0.190, 0.151]	[0.001, 0.367]	[-0.155, 0.019]	[-0.007, 0.085]	[-0.163, 0.122]	[-0.191, 0.093]
Main bandwidth	4.4	4.4	4.4	4.4	4.4	4.4	4.4	4.4
Auxiliary bandwidth	7.9	7.9	7.9	7.9	7.9	7.9	7.9	7.9
Observations	231	231	231	231	231	231	231	231
Control Group Mean	0.010	0.010	0.267	0.226	0.080	0.009	0.145	0.085

Notes: In Panel A, each RD specification uses optimal data-driven bandwidths for a cross-sectional specification and includes year fixed effects. Panel B presents results using the same bandwidths for all outcome variables. All estimates are based on the default settings in the procedure developed by Calonico et al. (2017) - a local linear fit on either side of the cut-off and a triangular kernel. The main bandwidth is used to construct an MSE-optimal RD point estimator. The auxiliary bandwidth is used to construct the robust bias-corrected (RBC) confidence intervals. Note that the RBC confidence intervals are not centred at the MSE-optimal point estimate; they are centred at the bias-corrected point estimate. Standard errors are clustered at the city level.

***, **, * indicate significance at 1%, 5%, and 10% level, respectively, based on the robust p-value.

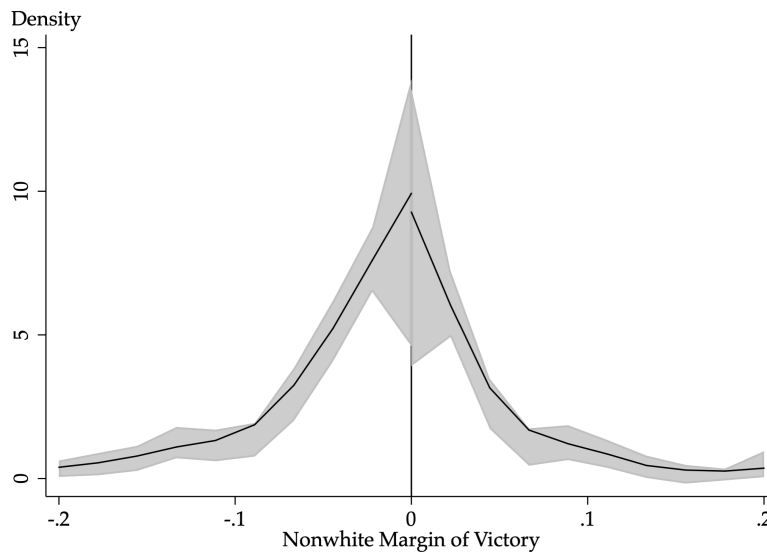
Figure 1: Tests for Discontinuity in Nonwhite Margin of Victory

(a) McCrary Test



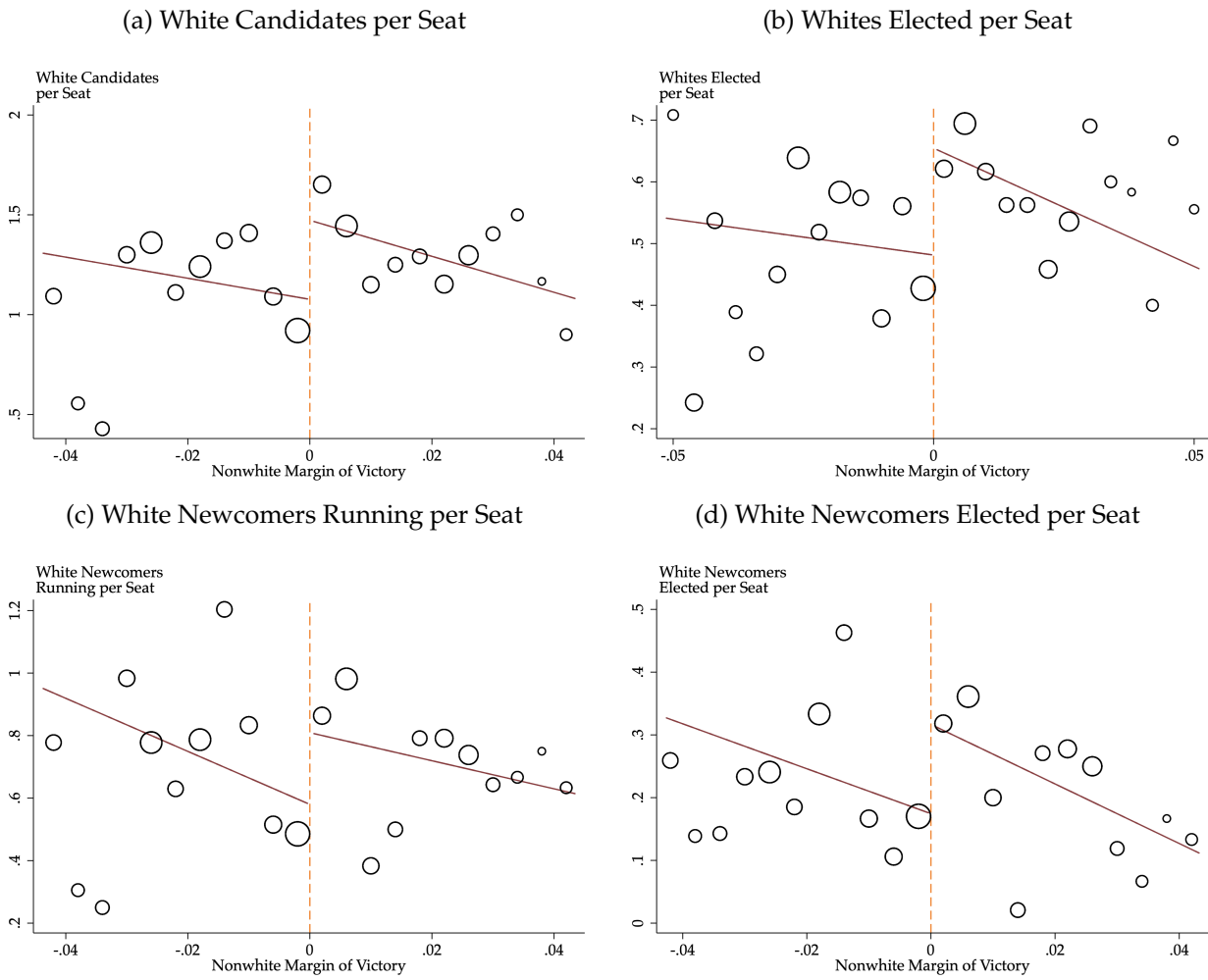
Notes: The figure illustrates the continuity of the running variable, nonwhite candidates victory margin using the methodology developed by McCrary (2008). The dashed lines represent the 95% confidence interval around the estimated density. Discontinuity estimate (log difference in height): -0.1207 (SE: 0.2389).

(b) RD Density Test by Cattaneo, Jansson and Ma (2018)



Notes: The figure illustrates the continuity of the running variable, nonwhite candidates victory margin using the methodology developed by Cattaneo et al. (2019). The figure plots the estimated density and a 95% confidence interval around the estimated density. P-value for the test for a discontinuity in the running variable at the cut-off: 0.8901 .

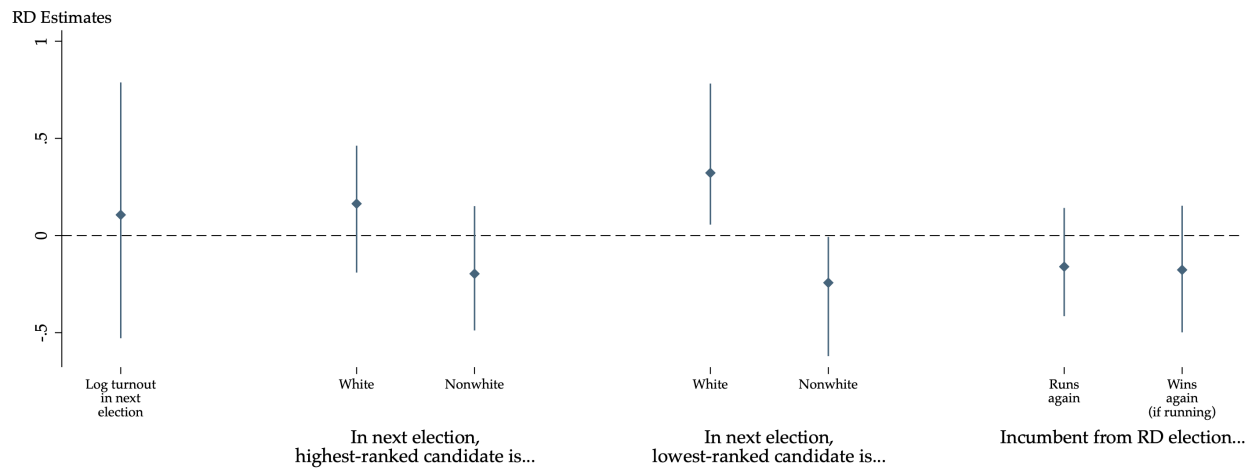
Figure 2: Illustration of the Regression Discontinuity



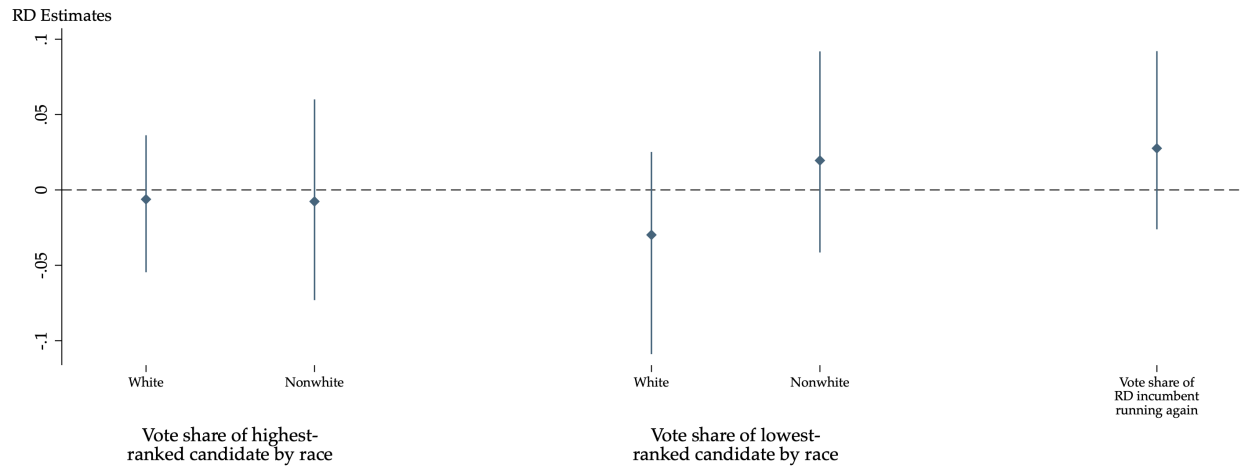
Notes: The figures above display the linear slopes estimated separately on each side of the cut-off using the election-level data and the procedure developed by Calonico et al. (2017). I also overlay a scatter plot of 0.4pp bin averages of the dependent variable weighted by the number of observations in each 0.4pp bin.

Figure 3: Nonwhite Victory and Voter Preferences in the Next Election

(a) Turnout, Candidate Rankings, and Re-election

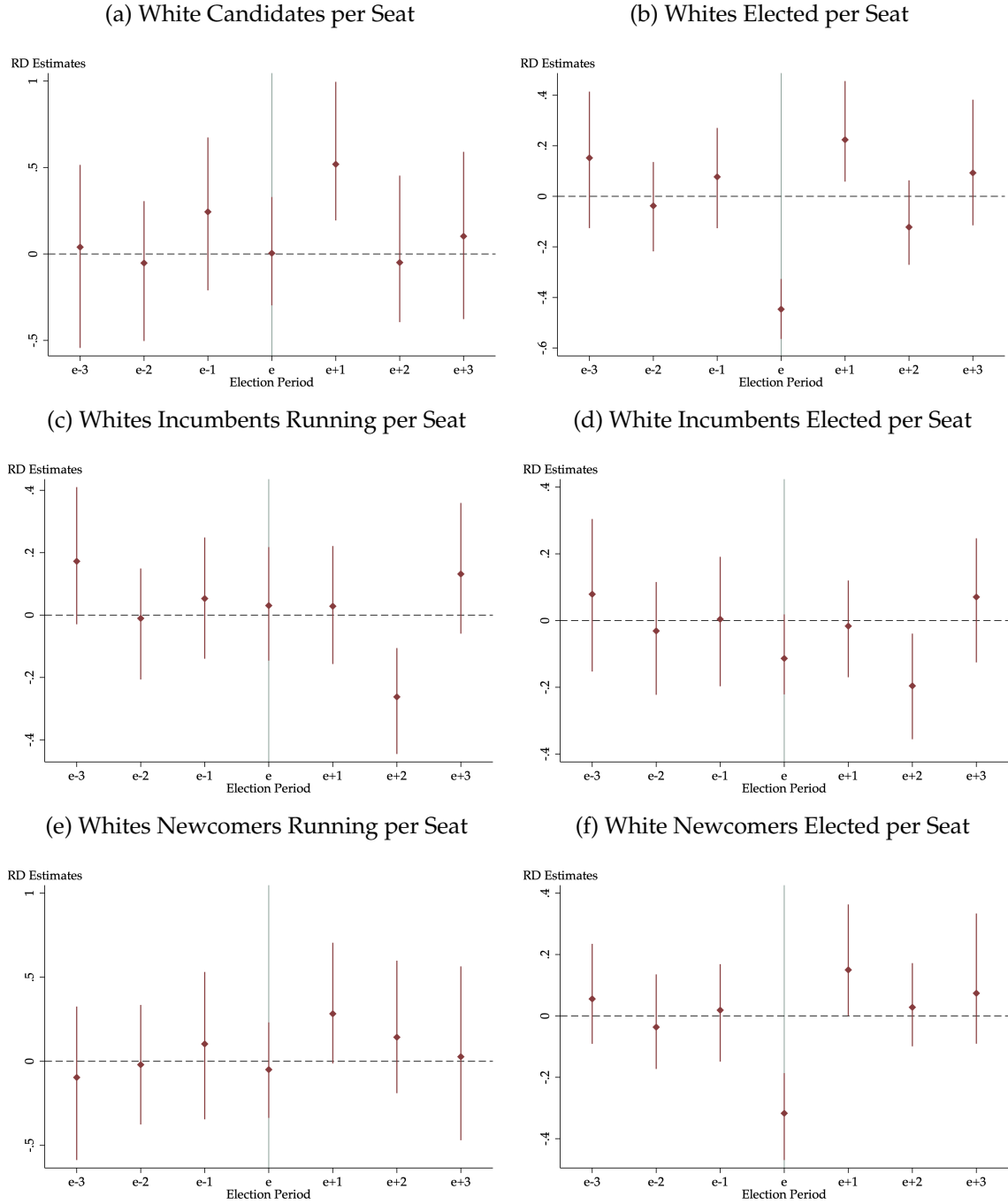


(b) Vote Shares



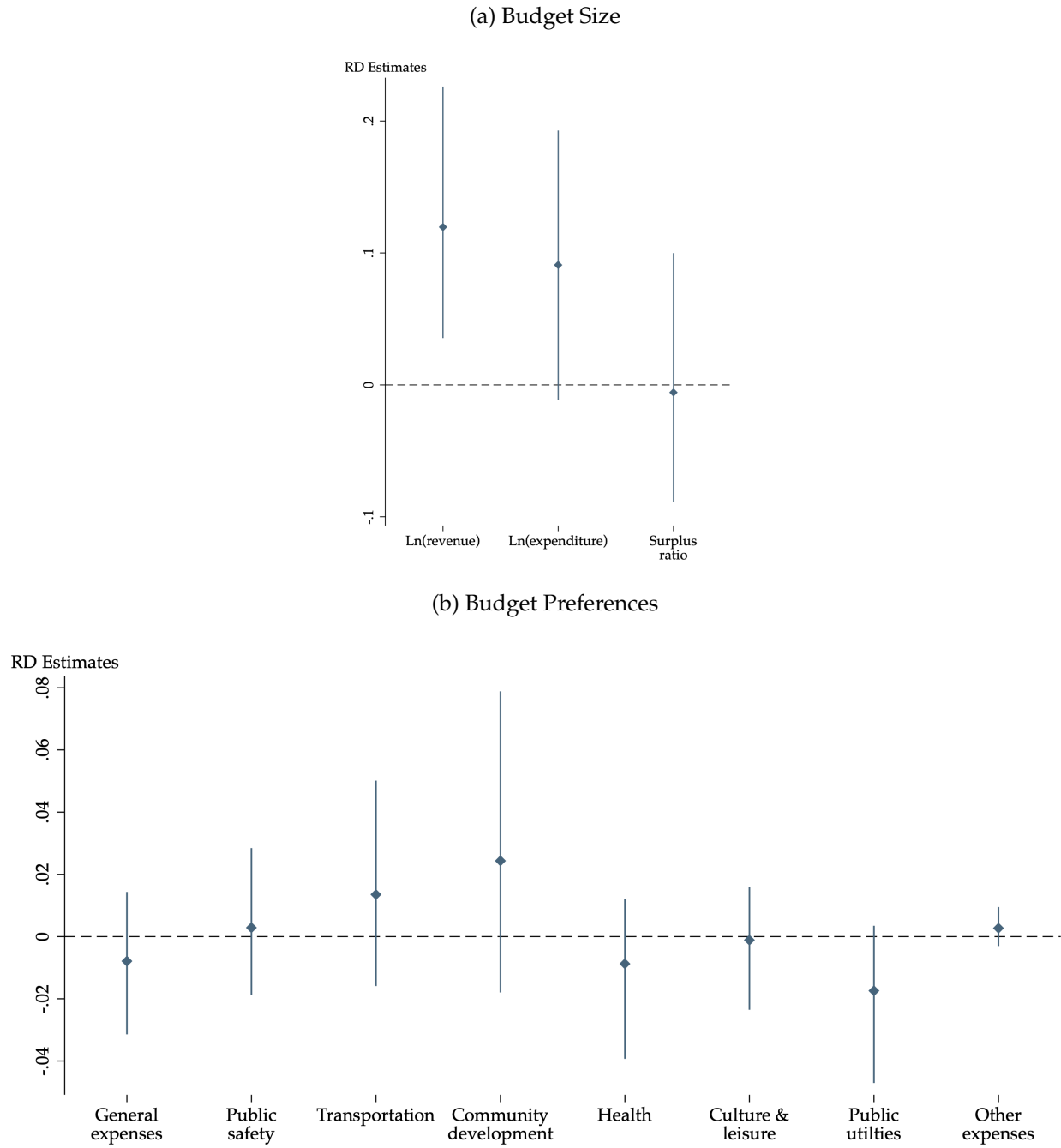
Notes: The figures presents the point estimates and 95% robust bias-corrected confidence intervals. Each RD specification uses optimal data-driven bandwidths for a cross-sectional specification and includes year fixed effects. The estimation is based on the default settings in the procedure developed by Calonico et al. (2017) - a local linear fit on either side of the cut-off and a triangular kernel. The main bandwidth is used to construct an MSE-optimal RD point estimator. The auxiliary bandwidth is used to construct the robust bias-corrected (RBC) confidence intervals. Note that the RBC confidence intervals are not centred at the MSE-optimal point estimate; they are centred at the bias-corrected point estimate. Standard errors are clustered at the city level.

Figure 4: Dynamics Over Time



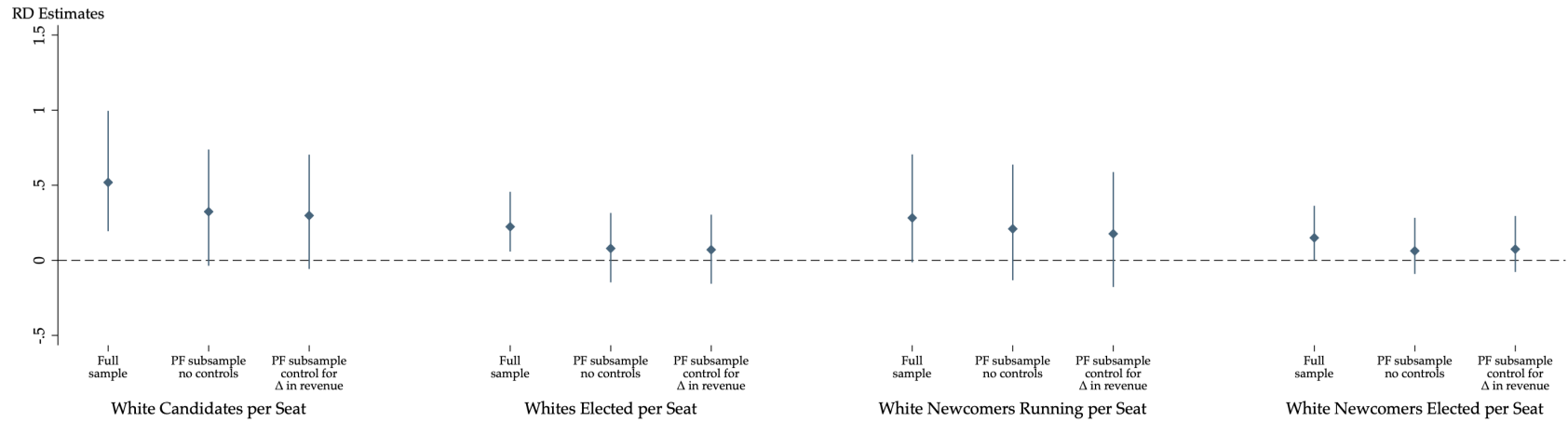
The figure presents the point estimates and 95% robust bias-corrected confidence intervals. In each figure, I present the treatment effect of a nonwhite candidates victory in election year e on the outcome variable in different election years relative to election year e . Each RD specification uses optimal data-driven bandwidths for a cross-sectional specification and includes year fixed effects. The estimation is based on the default settings in the procedure developed by Calonico et al. (2017) - a local linear fit on either side of the cut-off and a triangular kernel. The main bandwidth is used to construct an MSE-optimal RD point estimator. The auxiliary bandwidth is used to construct the robust bias-corrected (RBC) confidence intervals. Note that the RBC confidence intervals are not centred at the MSE-optimal point estimate; they are centred at the bias-corrected point estimate. Standard errors are clustered at the city level.

Figure 5: Nonwhite Victory and Budget Changes Over the Next Two Years



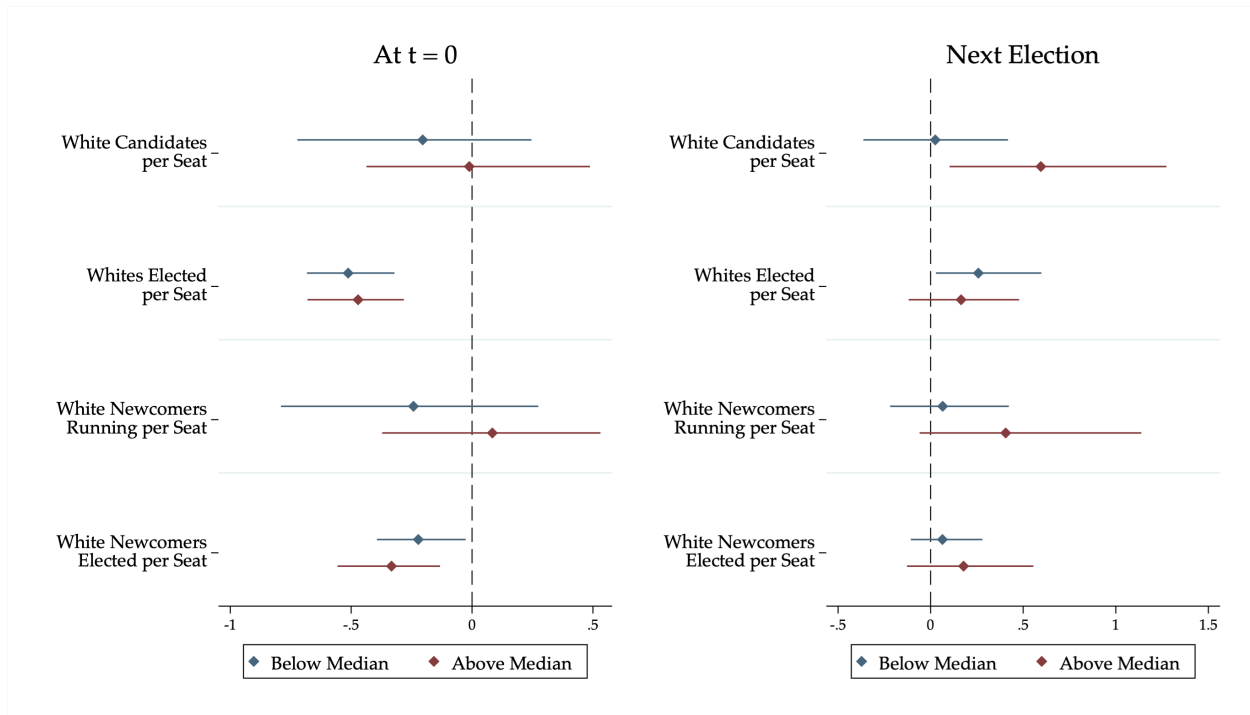
Notes: The figures presents the point estimates and 95% robust bias-corrected confidence intervals. Each RD specification uses optimal data-driven bandwidths for a cross-sectional specification and includes year fixed effects. The estimation is based on the default settings in the procedure developed by Calonico et al. (2017) - a local linear fit on either side of the cut-off and a triangular kernel. The main bandwidth is used to construct an MSE-optimal RD point estimator. The auxiliary bandwidth is used to construct the robust bias-corrected (RBC) confidence intervals. Note that the RBC confidence intervals are not centred at the MSE-optimal point estimate; they are centred at the bias-corrected point estimate. Standard errors are clustered at the city level.

Figure 6: Controlling for Changes in Revenue After a Nonwhite Victory



Notes: The figure presents the point estimates and 95% robust bias-corrected confidence intervals. Each RD specification uses optimal data-driven bandwidths for a cross-sectional specification and includes year fixed effects. The estimation is based on the default settings in the procedure developed by Calonico et al. (2017) - a local linear fit on either side of the cut-off and a triangular kernel. The main bandwidth is used to construct an MSE-optimal RD point estimator. The auxiliary bandwidth is used to construct the robust bias-corrected (RBC) confidence intervals. Note that the RBC confidence intervals are not centred at the MSE-optimal point estimate; they are centred at the bias-corrected point estimate. Standard errors are clustered at the city level.

Figure 7: Heterogeneity by 1990-2010 Change in Nonwhite Share of Population



Notes: The figure presents the point estimates and 95% robust bias-corrected confidence intervals. Each RD specification uses optimal data-driven bandwidths for a cross-sectional specification and includes year fixed effects. The estimation is based on the default settings in the procedure developed by Calonico et al. (2017) - a local linear fit on either side of the cut-off and a triangular kernel. The main bandwidth is used to construct an MSE-optimal RD point estimator. The auxiliary bandwidth is used to construct the robust bias-corrected (RBC) confidence intervals. Note that the RBC confidence intervals are not centred at the MSE-optimal point estimate; they are centred at the bias-corrected point estimate. Standard errors are clustered at the city level.

Appendix A: Additional Analyses

Table A1: Balance of Pre-determined Covariates Using Consistent Bandwidths

	RD Estimate	Conv. SE	Robust P-value	Main BW	Aux. BW	Obs.	CG Mean
<i>Panel A: Election-Related Variables</i>							
<i>Candidates per seat</i>							
Total	-0.098	0.288	0.851	4.4	7.9	231	2.399
White	-0.043	0.156	0.788	4.4	7.9	231	1.122
Nonwhite	0.107	0.189	0.531	4.4	7.9	231	0.763
Unassigned	-0.163	0.112	0.242	4.4	7.9	231	0.514
<i>Elected per seat</i>							
White	-0.435***	0.060	< 0.001	4.4	7.9	231	0.503
Nonwhite	0.462***	0.059	< 0.001	4.4	7.9	231	0.310
Unassigned	-0.029	0.054	0.614	4.4	7.9	231	0.187
<i>Other election-related variables</i>							
Num. of seats	-0.103	0.169	0.593	4.4	7.9	231	2.275
Log turnout	0.177	0.306	0.585	4.4	7.9	231	9.227
Winner incumbency status	0.044	0.113	0.611	4.4	7.9	231	0.516
Loser incumbency status	0.018	0.114	0.840	4.4	7.9	231	0.214
<i>Panel B: Demographics in 1990</i>							
Log population	0.095	0.319	0.780	4.4	7.9	224	10.052
White share	0.063	0.054	0.249	4.4	7.9	224	0.508
Nonwhite share	-0.063	0.055	0.251	4.4	7.9	224	0.484
Black share	-0.031**	0.013	0.013	4.4	7.9	224	0.042
Hispanic share	-0.025	0.063	0.704	4.4	7.9	224	0.365
Asian share	-0.007	0.028	0.820	4.4	7.9	224	0.077
Other share	0.000	0.001	0.942	4.4	7.9	224	0.008

Notes: Each RD specification uses optimal data-driven bandwidths for a cross-sectional specification and includes year fixed effects. The estimation is based on the default settings in the procedure developed by Calonico et al. (2017) - a local linear fit on either side of the cut-off and a triangular kernel. The main bandwidth is used to construct an MSE-optimal RD point estimator. The auxiliary bandwidth is used to construct the robust bias-corrected (RBC) confidence intervals. Note that the RBC confidence intervals are not centred at the MSE-optimal point estimate; they are centred at the bias-corrected point estimate. Standard errors are clustered at the city level.

***, **, * indicate significance at 1%, 5%, and 10% level, respectively, based on the robust p-value.

Table A2a: Nonwhite Victory and Voter Preferences - Turnout, Candidate Rankings, and Re-election

	Next election (e+1)					Re-election (e+2)	
	Log turnout	Top candidate is...		Bottom candidate is...		Incumbent runs again	(If running)
		White	Nonwhite	White	Nonwhite		Incumbent wins again
	(1)	(2)	(3)	(4)	(5)	(6)	(7)
Panel A: Optimal Bandwidths							
Nonwhite wins	0.107	0.164	-0.197	0.322**	-0.243**	-0.160	-0.177
Conventional SE							
Robust p-value	0.699	0.415	0.302	0.024	0.045	0.337	0.300
Robust 95% CI	[-0.528, 0.788]	[-0.191, 0.462]	[-0.488, 0.151]	[0.056, 0.782]	[-0.621, -0.007]	[-0.415, 0.142]	[-0.498, 0.153]
Main bandwidth	5.7	5.1	5.7	3.8	4.0	4.9	5.3
Auxiliary bandwidth	8.8	9.9	10.9	7.6	8.0	8.2	8.1
Observations	263	169	176	148	149	205	118
Control Group Mean	9.261	0.434	0.423	0.411	0.322	0.612	0.718
Panel B: Consistent Bandwidths ($h = 4.4\%$, $b = 7.9\%$)							
Nonwhite wins	0.124	0.169	-0.177	0.248*	-0.215*	-0.153	-0.141
Conventional SE							
Robust p-value	0.729	0.383	0.429	0.055	0.058	0.369	0.437
Robust 95% CI	[-0.568, 0.812]	[-0.197, 0.512]	[-0.496, 0.211]	[-0.008, 0.707]	[-0.600, 0.010]	[-0.411, 0.152]	[-0.464, 0.201]
Main bandwidth	4.4	4.4	4.4	4.4	4.4	4.4	4.4
Auxiliary bandwidth	7.9	7.9	7.9	7.9	7.9	7.9	7.9
Observations	231	156	156	156	156	191	108
Control Group Mean	9.227	0.432	0.453	0.421	0.316	0.604	0.730

Notes: In Panel A, each RD specification uses optimal data-driven bandwidths for a cross-sectional specification and includes year fixed effects. Panel B presents results using the same bandwidths for all outcome variables. All estimates are based on the default settings in the procedure developed by Calonico et al. (2017) - a local linear fit on either side of the cut-off and a triangular kernel. The main bandwidth is used to construct an MSE-optimal RD point estimator. The auxiliary bandwidth is used to construct the robust bias-corrected (RBC) confidence intervals. Note that the RBC confidence intervals are not centred at the MSE-optimal point estimate; they are centred at the bias-corrected point estimate. Standard errors are clustered at the city level.

***, **, * indicate significance at 1%, 5%, and 10% level, respectively, based on the robust p-value.

Table A2b: Nonwhite Victory and Voter Preferences - Vote Shares

	Next election (e+1)				(e+2)
	Vote share of top candidate by race		Vote share of bottom candidate by race		(If running)
	White	Nonwhite	White	Nonwhite	Incumbent vote share
	(1)	(2)	(3)	(4)	(5)
Panel A: Optimal Bandwidths					
Nonwhite wins	-0.006	-0.008	-0.030	0.020	0.028
Conventional SE					
Robust p-value	0.692	0.848	0.221	0.459	0.274
Robust 95% CI	[-0.055, 0.036]	[-0.073, 0.060]	[-0.109, 0.025]	[-0.041, 0.092]	[-0.026, 0.092]
Main bandwidth	4.7	5.3	4.4	6.0	4.8
Auxiliary bandwidth	8.5	8.8	8.2	10.6	7.0
Observations	161	169	156	179	117
Control Group Mean	0.247	0.229	0.156	0.169	0.230
Panel B: Consistent Bandwidths ($h = 4.4\%$, $b = 7.9\%$)					
Nonwhite wins	-0.005	-0.007	-0.029	0.021	0.030
Conventional SE					
Robust p-value	0.739	0.807	0.204	0.502	0.277
Robust 95% CI	[-0.054, 0.038]	[-0.076, 0.059]	[-0.111, 0.024]	[-0.047, 0.096]	[-0.026, 0.091]
Main bandwidth	4.4	4.4	4.4	4.4	4.4
Auxiliary bandwidth	7.9	7.9	7.9	7.9	7.9
Observations	156	156	156	156	108
Control Group Mean	0.248	0.228	0.156	0.168	0.230

Notes: In Panel A, each RD specification uses optimal data-driven bandwidths for a cross-sectional specification and includes year fixed effects. Panel B presents results using the same bandwidths for all outcome variables. All estimates are based on the default settings in the procedure developed by Calonico et al. (2017) - a local linear fit on either side of the cut-off and a triangular kernel. The main bandwidth is used to construct an MSE-optimal RD point estimator. The auxiliary bandwidth is used to construct the robust bias-corrected (RBC) confidence intervals. Note that the RBC confidence intervals are not centred at the MSE-optimal point estimate; they are centred at the bias-corrected point estimate. Standard errors are clustered at the city level.

***, **, * indicate significance at 1%, 5%, and 10% level, respectively, based on the robust p-value.

Table A3: Nonwhite Victory and Budget Changes Over the Next Two Years

	RD Estimate	Conv. SE	Robust P-value	Main BW	Aux. BW	Obs.	CG Mean
Panel A: Optimal Bandwidths							
<i>Budget Summary</i>							
Ln(Revenue)	0.120***	0.043	0.007	4.9	9.2	157	0.020
Ln(Expenditure)	0.091*	0.047	0.082	5.6	9.6	168	0.029
Surplus ratio	-0.006	0.043	0.910	6.7	12.4	182	0.006
<i>Budget Preferences</i>							
General exp. share	-0.008	0.010	0.465	5.2	8.1	161	0.002
Public safety share	0.003	0.011	0.691	4.8	8.9	156	0.011
Transportation share	0.014	0.015	0.310	4.5	7.6	147	-0.011
Community dev. share	0.024	0.021	0.218	4.0	7.4	142	-0.002
Health exp. share	-0.009	0.012	0.301	3.8	8.1	141	-0.004
Culture & leisure share	-0.001	0.009	0.704	4.4	8.0	146	-0.002
Public utilities share	-0.017*	0.012	0.091	6.6	12.0	181	0.002
Other exp. share	0.003	0.003	0.312	4.9	9.6	157	0.000
Panel B: Consistent Bandwidths ($h = 4.4\%$, $b = 7.9\%$)							
<i>Budget Summary</i>							
Ln(Revenue)	0.123***	0.044	0.008	4.4	7.9	147	0.027
Ln(Expenditure)	0.082	0.047	0.128	4.4	7.9	147	0.027
Surplus ratio	0.013	0.049	0.671	4.4	7.9	147	-0.003
<i>Budget Preferences</i>							
General exp. share	-0.011	0.010	0.338	4.4	7.9	147	0.003
Public safety share	0.003	0.011	0.650	4.4	7.9	147	0.010
Transportation share	0.014	0.015	0.292	4.4	7.9	147	-0.011
Community dev. share	0.022	0.020	0.232	4.4	7.9	147	0.000
Health exp. share	-0.008	0.012	0.289	4.4	7.9	147	-0.005
Culture & leisure share	-0.001	0.009	0.707	4.4	7.9	147	-0.002
Public utilities share	-0.023*	0.012	0.055	4.4	7.9	147	0.004
Other exp. share	0.003	0.003	0.250	4.4	7.9	147	0.000

Notes: In Panel A, each RD specification uses optimal data-driven bandwidths for a cross-sectional specification and includes year fixed effects. Panel B presents results using the same bandwidths for all outcome variables. All estimates are based on the default settings in the procedure developed by Calonico et al. (2017) - a local linear fit on either side of the cut-off and a triangular kernel. The main bandwidth is used to construct an MSE-optimal RD point estimator. The auxiliary bandwidth is used to construct the robust bias-corrected (RBC) confidence intervals. Note that the RBC confidence intervals are not centred at the MSE-optimal point estimate; they are centred at the bias-corrected point estimate. Standard errors are clustered at the city level.

***, **, * indicate significance at 1%, 5%, and 10% level, respectively, based on the robust p-value.

Table A4: Controlling for Changes in Revenue After a Nonwhite Victory

	White Candidates per Seat			Whites Elected per Seat			White Newcomers Running per Seat			White Newcomers Elected per Seat		
	Full Sample	PF Subsample No Controls	PF Subsample Control for Δ in Revenue	Full Sample	PF Subsample No Controls	PF Subsample Control for Δ in Revenue	Full Sample	PF Subsample No Controls	PF Subsample Control for Δ in Revenue	Full Sample	PF Subsample No Controls	PF Subsample Control for Δ in Revenue
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)	(11)	(12)
<i>Panel A: Optimal Bandwidth</i>												
Nonwhite wins	0.519***	0.324*	0.298*	0.224**	0.079	0.071	0.283*	0.210	0.177	0.150**	0.063	0.075
Conventional SE	(0.180)	(0.178)	(0.175)	(0.089)	(0.103)	(0.102)	(0.164)	(0.178)	(0.176)	(0.082)	(0.085)	(0.085)
Robust p-value	0.004	0.076	0.096	0.011	0.475	0.528	0.059	0.199	0.294	0.049	0.313	0.252
Robust 95% CI	[0.194, 0.996]	[-0.037, 0.738]	[-0.058, 0.704]	[0.058, 0.456]	[-0.147, 0.316]	[-0.156, 0.304]	[-0.013, 0.705]	[-0.133, 0.637]	[-0.178, 0.588]	[0.000, 0.363]	[-0.091, 0.283]	[-0.078, 0.295]
Main bandwidth												
Auxiliary bandwidth												
Observations	231	185	192	251	170	170	231	161	174	227	145	145
Control Group Mean	1.122	1.151	1.155	0.489	0.499	0.499	0.695	0.709	0.712	0.219	0.213	0.213
<i>Panel B: Consistent Bandwidths (h = 4.4%, b = 7.9%)</i>												
Nonwhite wins	0.519***	0.389**	0.354*	0.248***	0.097	0.086	0.286**	0.249	0.236	0.148**	0.058	0.070
Conventional SE	(0.180)	(0.191)	(0.192)	(0.092)	(0.108)	(0.107)	(0.165)	(0.180)	(0.180)	(0.082)	(0.085)	(0.085)
Robust p-value	0.004	0.040	0.063	0.008	0.402	0.458	0.047	0.111	0.130	0.049	0.326	0.263
Robust 95% CI	[0.194, 0.996]	[0.021, 0.846]	[-0.021, 0.810]	[0.072, 0.483]	[-0.136, 0.338]	[-0.146, 0.325]	[0.005, 0.729]	[-0.072, 0.699]	[-0.088, 0.686]	[0.001, 0.367]	[-0.093, 0.280]	[-0.080, 0.293]
Main bandwidth												
Auxiliary bandwidth												
Observations	231	147	147	231	147	147	231	147	147	231	147	147
Control Group Mean	1.122	1.122	1.122	0.503	0.503	0.503	0.695	0.695	0.695	0.226	0.226	0.226

Notes: In Panel A, each RD specification uses optimal data-driven bandwidths for a cross-sectional specification and includes year fixed effects. Panel B presents results using the same bandwidths for all outcome variables. All estimates are based on the default settings in the procedure developed by Calonico et al. (2017) - a local linear fit on either side of the cut-off and a triangular kernel. The main bandwidth is used to construct an MSE-optimal RD point estimator. The auxiliary bandwidth is used to construct the robust bias-corrected (RBC) confidence intervals. Note that the RBC confidence intervals are not centred at the MSE-optimal point estimate; they are centred at the bias-corrected point estimate. Standard errors are clustered at the city level.

***, **, * indicate significance at 1%, 5%, and 10% level, respectively, based on the robust p-value.

**Table A5a: Heterogeneity by 1990-2010 Change in Nonwhite Share of Population
White Candidates per Seat & Whites Elected per Seat**

	White Candidates per Seat				Whites Elected per Seat			
	At t = 0		Next Election		At t = 0		Next Election	
	Below Median	Above Median	Below Median	Above Median	Below Median	Above Median	Below Median	Above Median
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
Nonwhite wins	-0.204	-0.011	0.025	0.596**	-0.512***	-0.471***	0.258**	0.165
Conventional SE	(0.207)	(0.202)	(0.171)	(0.252)	(0.077)	(0.086)	(0.129)	(0.126)
Robust p-value	0.334	0.913	0.891	0.021	< 0.001	< 0.001	0.031	0.237
Robust 95% CI	[-0.723, 0.246]	[-0.437, 0.488]	[-0.363, 0.418]	[0.102, 1.274]	[-0.683, -0.321]	[-0.681, -0.282]	[0.029, 0.598]	[-0.118, 0.478]
Main bandwidth	4.7	5.6	5.1	4.5	5.1	5.1	4.3	7.2
Auxiliary bandwidth	8.6	10.0	8.8	7.2	8.2	8.1	7.5	12.7
Observations	120	123	124	108	113	119	104	142
Control Group Mean	1.048	1.451	0.883	1.313	0.436	0.698	0.403	0.595

Notes: Each RD specification uses optimal data-driven bandwidths for a cross-sectional specification and includes year fixed effects. All estimates are based on the default settings in the procedure developed by Calonico et al. (2017) - a local linear fit on either side of the cut-off and a triangular kernel. The main bandwidth is used to construct an MSE-optimal RD point estimator. The auxiliary bandwidth is used to construct the robust bias-corrected (RBC) confidence intervals. Note that the RBC confidence intervals are not centred at the MSE-optimal point estimate; they are centred at the bias-corrected point estimate. Standard errors are clustered at the city level.

***, **, * indicate significance at 1%, 5%, and 10% level, respectively, based on the robust p-value.

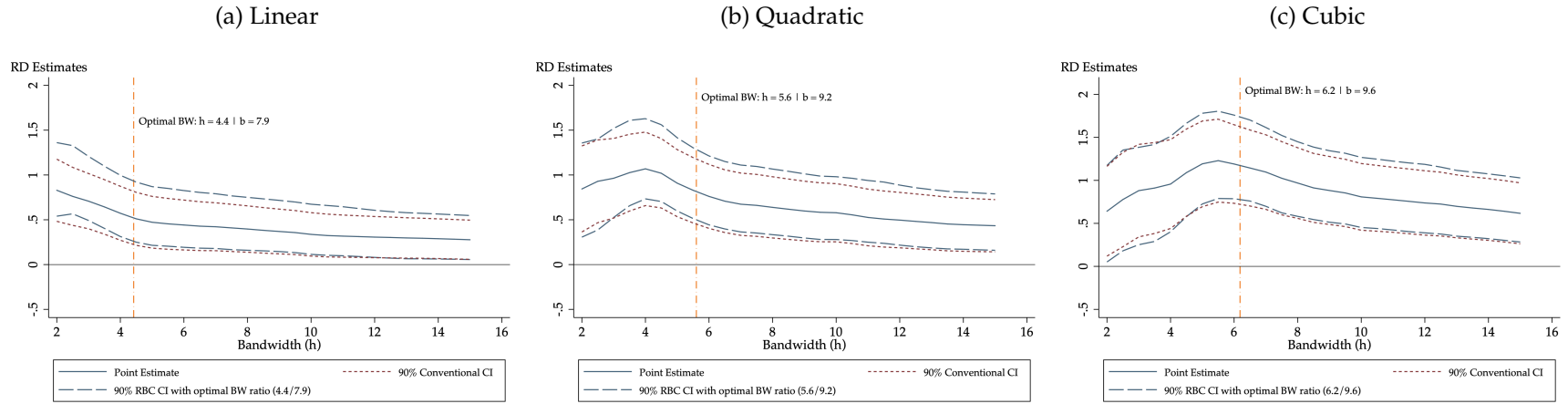
**Table A5b: Heterogeneity by 1990-2010 Change in Nonwhite Share of Population
White Newcomers Running per Seat & White Newcomers Elected per Seat**

	White Newcomers Running per Seat				White Newcomers Elected per Seat			
	At t = 0		Next Election		At t = 0		Next Election	
	Below Median	Above Median	Below Median	Above Median	Below Median	Above Median	Below Median	Above Median
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
Nonwhite wins	-0.242	0.084	0.065	0.406*	-0.222**	-0.333***	0.064	0.178
Conventional SE	(0.223)	(0.194)	(0.136)	(0.263)	(0.078)	(0.090)	(0.087)	(0.145)
Robust p-value	0.341	0.729	0.533	0.078	0.025	0.001	0.382	0.219
Robust 95% CI	[-0.790, 0.274]	[-0.372, 0.531]	[-0.219, 0.422]	[-0.060, 1.138]	[-0.393, -0.026]	[-0.557, -0.132]	[-0.107, 0.280]	[-0.128, 0.555]
Main bandwidth	5.8	7.4	4.3	4.2	6.4	6.2	5.6	4.5
Auxiliary bandwidth	9.5	12.8	7.3	6.9	9.6	10.2	9.1	7.1
Observations	120	144	104	104	126	131	118	107
Control Group Mean	0.513	0.958	0.533	0.867	0.077	0.353	0.169	0.285

Notes: Each RD specification uses optimal data-driven bandwidths for a cross-sectional specification and includes year fixed effects. All estimates are based on the default settings in the procedure developed by Calonico et al. (2017) - a local linear fit on either side of the cut-off and a triangular kernel. The main bandwidth is used to construct an MSE-optimal RD point estimator. The auxiliary bandwidth is used to construct the robust bias-corrected (RBC) confidence intervals. Note that the RBC confidence intervals are not centred at the MSE-optimal point estimate; they are centred at the bias-corrected point estimate. Standard errors are clustered at the city level.

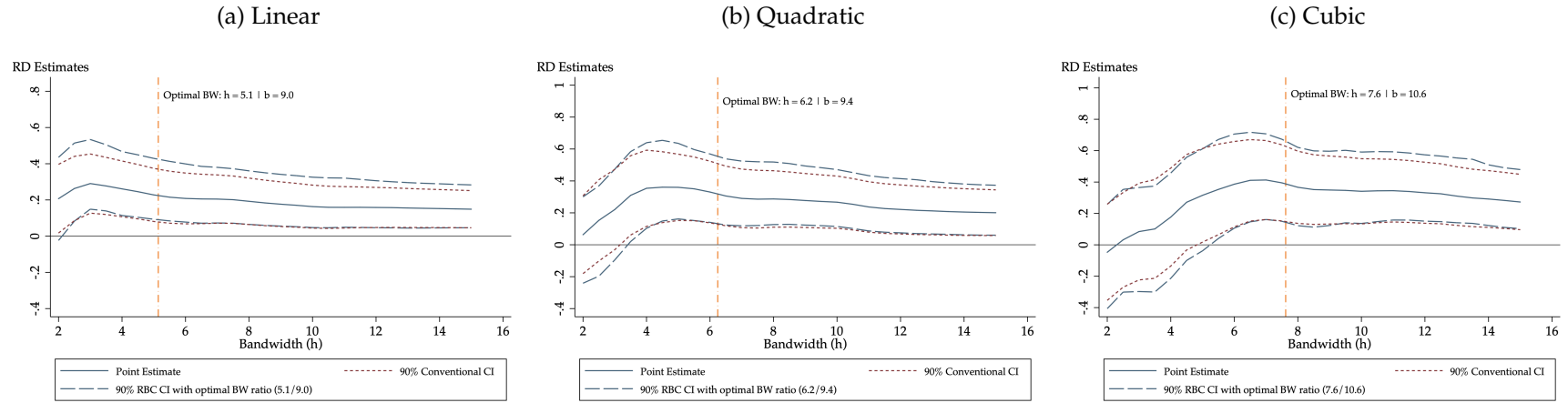
***, **, * indicate significance at 1%, 5%, and 10% level, respectively, based on the robust p-value.

Figure A1: Changing Bandwidths and Polynomial Order - White Candidates per Seat



Notes: The figures above display the coefficient estimates from regressions with bandwidths ranging from 2% to 15%. For each chosen bandwidth (h) used for point estimation, I calculate an auxiliary bandwidth (b) using the the ratio of the optimal data-driven bandwidths for the given polynomial order. The auxiliary bandwidth is then used to create the robust bias-corrected (RBC) confidence interval. Note that the RBC confidence intervals are not centred at the MSE-optimal point estimate; they are centred at the bias-corrected point estimate. Each figure includes a vertical line indicating the MSE-optimal data-driven bandwidth for point estimation derived using the procedure developed by Calonico et al. (2017). Figures (a), (b), and (c) use a linear, quadratic, and cubic specification on either side of the cut-off, respectively. The estimation includes year fixed effects and uses a triangular kernel.

Figure A2: Changing Bandwidths and Polynomial Order - Whites Elected per Seat



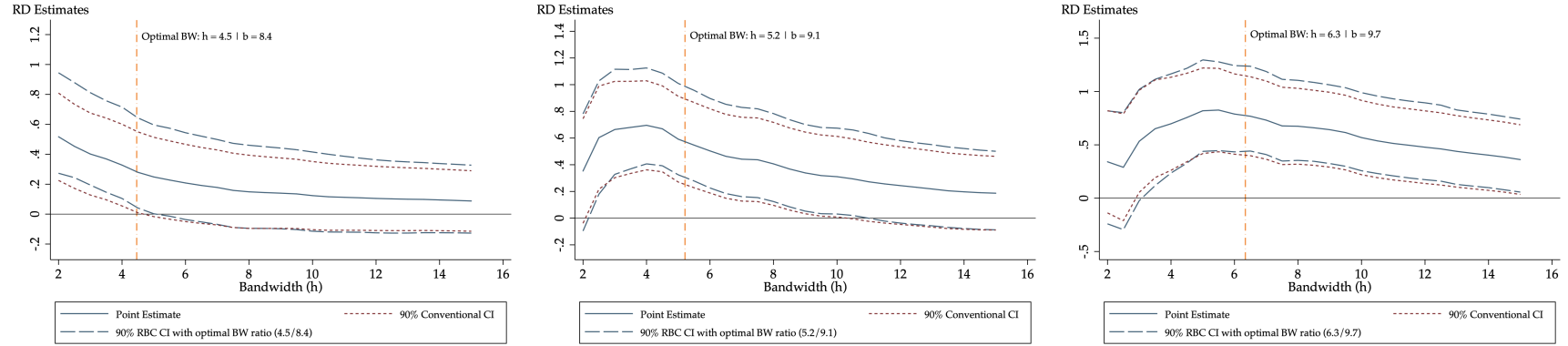
Notes: The figures above display the coefficient estimates from regressions with bandwidths ranging from 2% to 15%. For each chosen bandwidth (h) used for point estimation, I calculate an auxiliary bandwidth (b) using the the ratio of the optimal data-driven bandwidths for the given polynomial order. The auxiliary bandwidth is then used to create the robust bias-corrected (RBC) confidence interval. Note that the RBC confidence intervals are not centred at the MSE-optimal point estimate; they are centred at the bias-corrected point estimate. Each figure includes a vertical line indicating the MSE-optimal data-driven bandwidth for point estimation derived using the procedure developed by Calonico et al. (2017). Figures (a), (b), and (c) use a linear, quadratic, and cubic specification on either side of the cut-off, respectively. The estimation includes year fixed effects and uses a triangular kernel.

Figure A3: Changing Bandwidths and Polynomial Order - White Newcomers Running per Seat

(a) Linear

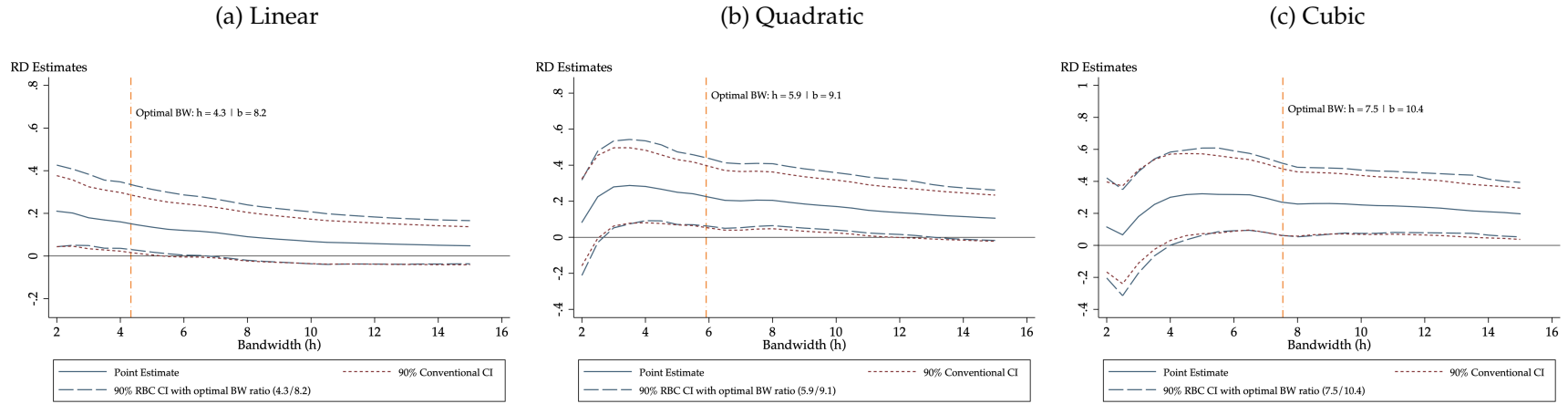
(b) Quadratic

(c) Cubic



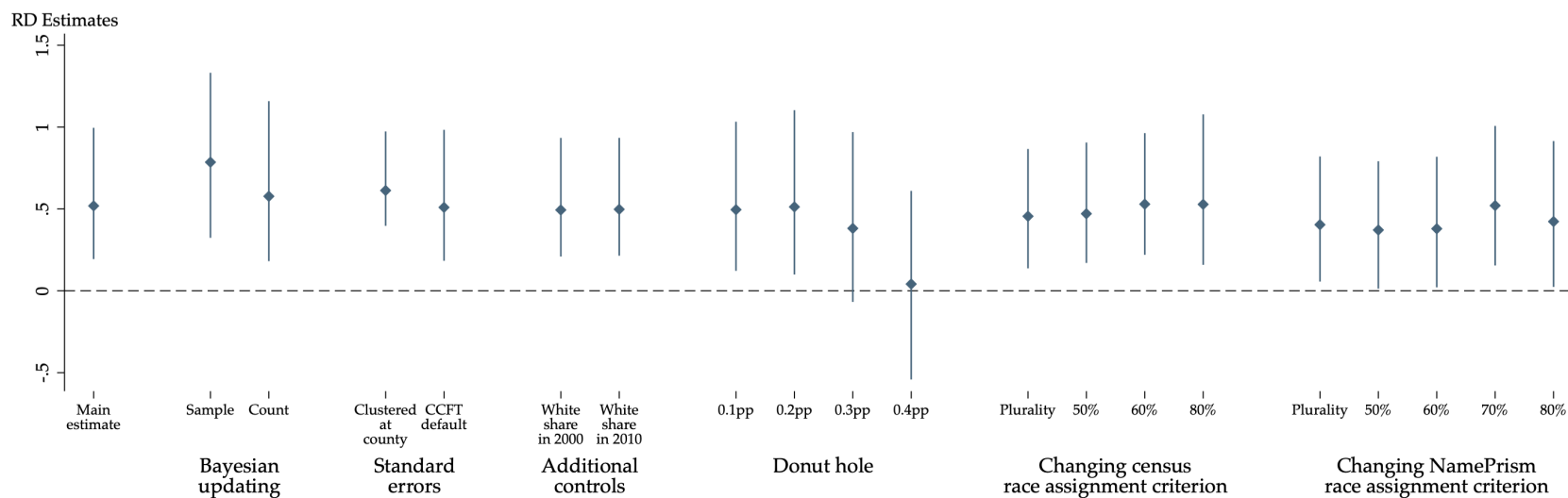
Notes: The figures above display the coefficient estimates from regressions with bandwidths ranging from 2% to 15%. For each chosen bandwidth (h) used for point estimation, I calculate an auxiliary bandwidth (b) using the the ratio of the optimal data-driven bandwidths for the given polynomial order. The auxiliary bandwidth is then used to create the robust bias-corrected (RBC) confidence interval. Note that the RBC confidence intervals are not centred at the MSE-optimal point estimate; they are centred at the bias-corrected point estimate. Each figure includes a vertical line indicating the MSE-optimal data-driven bandwidth for point estimation derived using the procedure developed by Calonico et al. (2017). Figures (a), (b), and (c) use a linear, quadratic, and cubic specification on either side of the cut-off, respectively. The estimation includes year fixed effects and uses a triangular kernel.

Figure A4: Changing Bandwidths and Polynomial Order - White Newcomers Elected per Seat



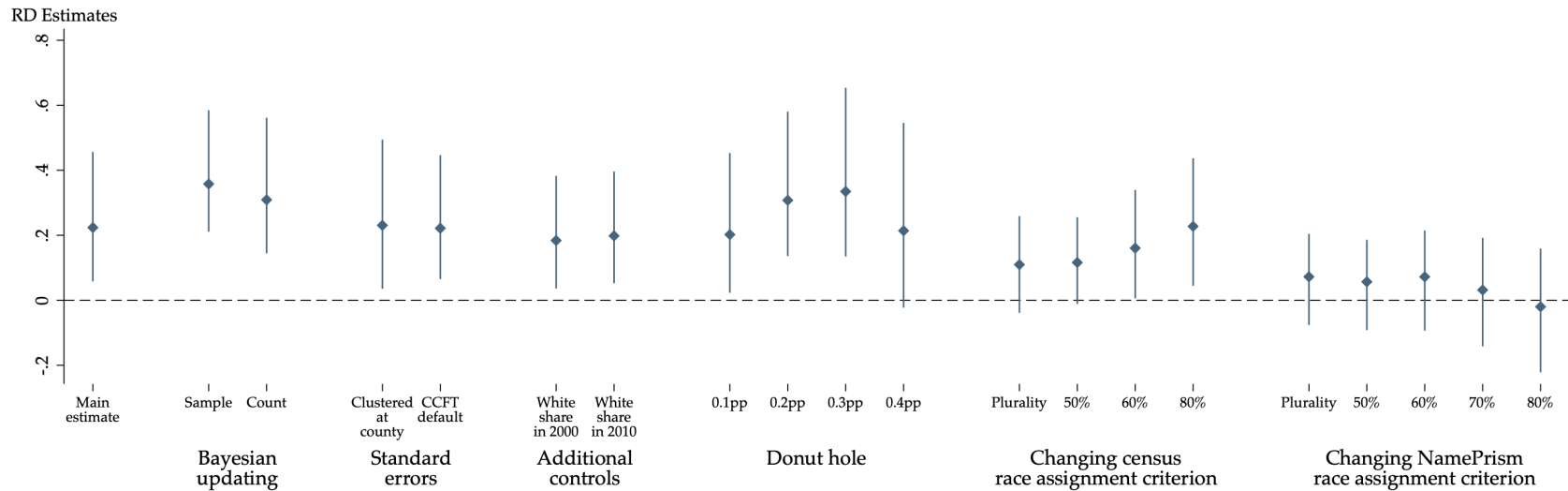
Notes: The figures above display the coefficient estimates from regressions with bandwidths ranging from 2% to 15%. For each chosen bandwidth (h) used for point estimation, I calculate an auxiliary bandwidth (b) using the the ratio of the optimal data-driven bandwidths for the given polynomial order. The auxiliary bandwidth is then used to create the robust bias-corrected (RBC) confidence interval. Note that the RBC confidence intervals are not centred at the MSE-optimal point estimate; they are centred at the bias-corrected point estimate. Each figure includes a vertical line indicating the MSE-optimal data-driven bandwidth for point estimation derived using the procedure developed by Calonico et al. (2017). Figures (a), (b), and (c) use a linear, quadratic, and cubic specification on either side of the cut-off, respectively. The estimation includes year fixed effects and uses a triangular kernel.

Figure A5: Effect on White Candidates per Seat and Changes in Specification or Sample Used



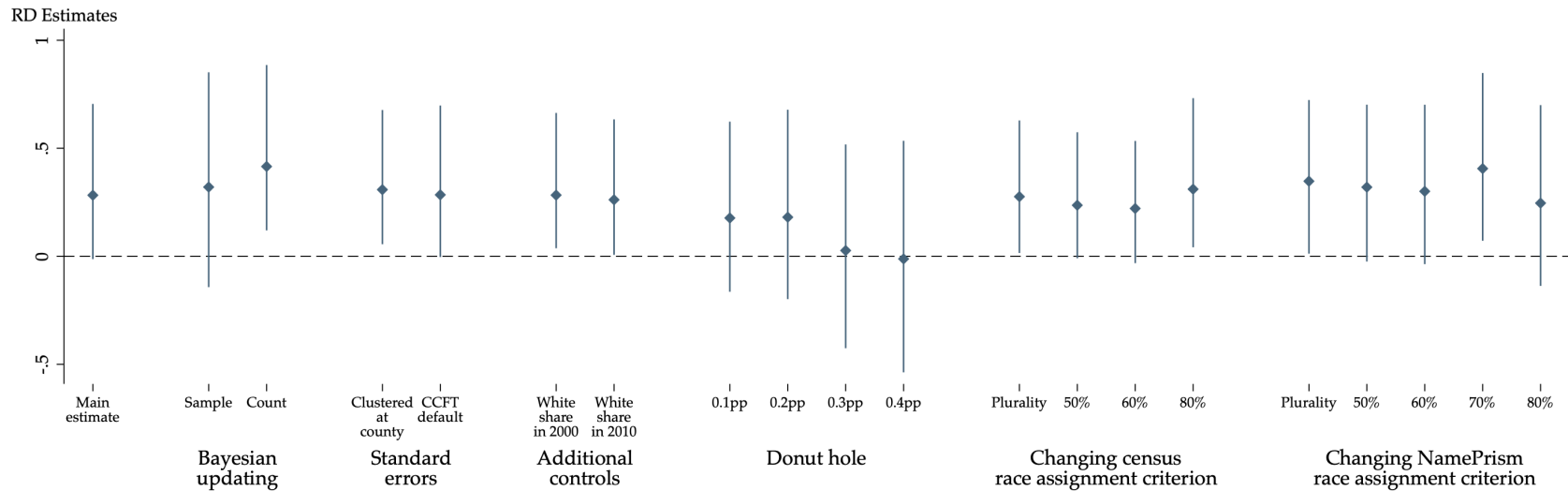
Notes: The figure above presents results for various robustness checks conducted. For each test, the figure includes a 95% robust bias-corrected confidence interval. Each RD specification uses optimal data-driven bandwidths for a cross-sectional specification and includes year fixed effects. The estimation is based on the default settings in the procedure developed by Calonico et al. (2017) - a local linear fit on either side of the cut-off and a triangular kernel. The main bandwidth is used to construct the RD point estimator. The auxiliary bandwidth is then used to create the robust bias-corrected (RBC) confidence interval. Note that the RBC confidence intervals are not centred at the MSE-optimal point estimate; they are centred at the bias-corrected point estimate. Standard errors are clustered at the city level, except for tests where I explicitly cluster standard errors at the county level or use the default standard errors in the procedure developed by Calonico et al. (2017).

Figure A6: Effect on Whites Elected per Seat and Changes in Specification or Sample Used



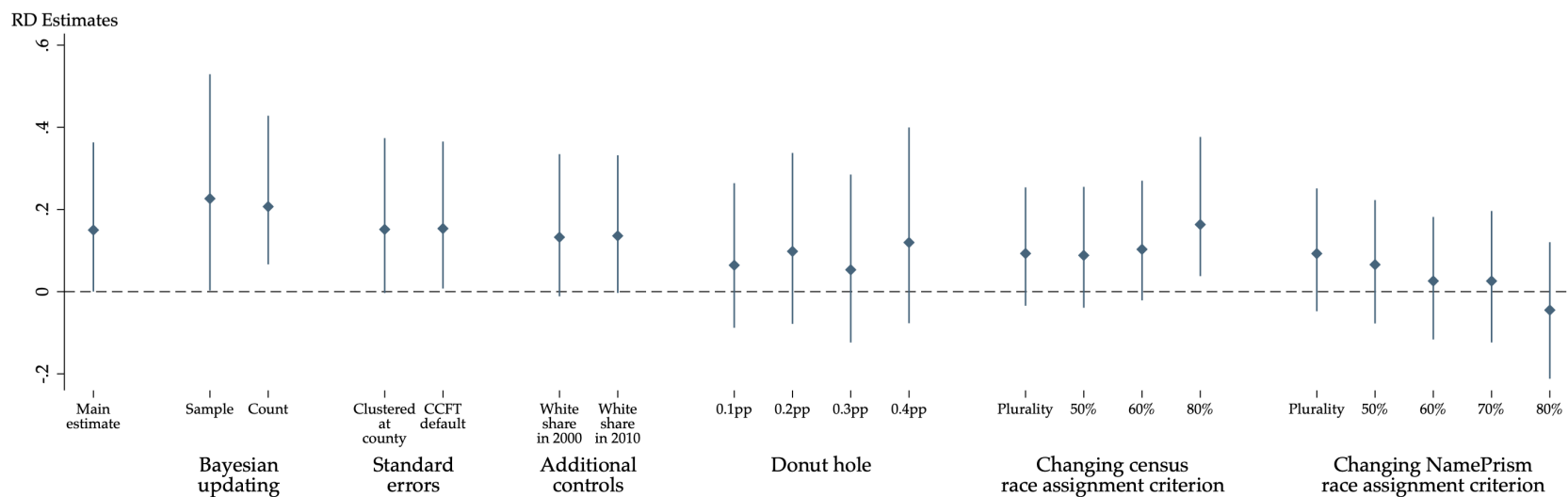
Notes: The figure above presents results for various robustness checks conducted. For each test, the figure includes a 95% robust bias-corrected confidence interval. Each RD specification uses optimal data-driven bandwidths for a cross-sectional specification and includes year fixed effects. The estimation is based on the default settings in the procedure developed by Calonico et al. (2017) - a local linear fit on either side of the cut-off and a triangular kernel. The main bandwidth is used to construct the RD point estimator. The auxiliary bandwidth is then used to create the robust bias-corrected (RBC) confidence interval. Note that the RBC confidence intervals are not centred at the MSE-optimal point estimate; they are centred at the bias-corrected point estimate. Standard errors are clustered at the city level, except for tests where I explicitly cluster standard errors at the county level or use the default standard errors in the procedure developed by Calonico et al. (2017).

Figure A7: Effect on White Newcomers Running per Seat and Changes in Specification or Sample Used



Notes: The figure above presents results for various robustness checks conducted. For each test, the figure includes a 95% robust bias-corrected confidence interval. Each RD specification uses optimal data-driven bandwidths for a cross-sectional specification and includes year fixed effects. The estimation is based on the default settings in the procedure developed by Calonico et al. (2017) - a local linear fit on either side of the cut-off and a triangular kernel. The main bandwidth is used to construct the RD point estimator. The auxiliary bandwidth is then used to create the robust bias-corrected (RBC) confidence interval. Note that the RBC confidence intervals are not centred at the MSE-optimal point estimate; they are centred at the bias-corrected point estimate. Standard errors are clustered at the city level, except for tests where I explicitly cluster standard errors at the county level or use the default standard errors in the procedure developed by Calonico et al. (2017).

Figure A8: Effect on White Newcomers Elected per Seat and Changes in Specification or Sample Used



Notes: The figure above presents results for various robustness checks conducted. For each test, the figure includes a 95% robust bias-corrected confidence interval. Each RD specification uses optimal data-driven bandwidths for a cross-sectional specification and includes year fixed effects. The estimation is based on the default settings in the procedure developed by Calonico et al. (2017) - a local linear fit on either side of the cut-off and a triangular kernel. The main bandwidth is used to construct the RD point estimator. The auxiliary bandwidth is then used to create the robust bias-corrected (RBC) confidence interval. Note that the RBC confidence intervals are not centred at the MSE-optimal point estimate; they are centred at the bias-corrected point estimate. Standard errors are clustered at the city level, except for tests where I explicitly cluster standard errors at the county level or use the default standard errors in the procedure developed by Calonico et al. (2017).

Appendix B: Race Assignment and Measurement Error

B1: Measurement Error Derivation

For simplicity, consider the local randomization RD framework, where the true model is given by:

$$Y^* = \delta + X^*\beta + \epsilon$$

I have omitted city and time subscripts for simplicity. Here, Y^* is the outcome variable, namely the number of white candidates running for election in a city, X^* is a dummy variable for victory of a nonwhite candidate against a white candidate in the previous election in the same city, and ϵ is an error term that is uncorrelated with X^* .

I do not observe Y^* and X^* . Instead, I observe $Y = Y^* + \eta$, where η represents the measurement error in counting the number of white candidates. Moreover, I observe $X = X^* + \mu$, where μ represents the measurement error from classifying elections with the nonwhite candidate's victory as elections with victory for the white candidate, or vice versa.

Now consider μ . Note that the probability of mis-measuring X^* within a given sample is very low at higher race assignment criteria. Since I measure the margin of victory without any error, this can only occur due to matching the white candidate to a nonwhite race group and identifying the nonwhite candidate as a white candidate. However, I include this measurement error for completeness. Clearly, μ is not a classical measurement error since X^* (and X) is a binary variable:

$$\mu = X - X^* = \begin{cases} -1 & \text{if false negative} \\ 0 & \text{if correct} \\ 1 & \text{if false positive} \end{cases}$$

Aigner (1973) showed that a binary independent variable with measurement error leads to downward biased slope estimate. Using his framework, Let V be the number of observations (N) that are categorized as a nonwhite victory, $L = N - V$ be the number of observations that are categorized as a minority loss. Let α be the false discovery rate, i.e. the proportion of V that are, in fact, nonwhite losses, and let ρ be the false omission rate, i.e. the proportion of L that are, in fact, nonwhite victories. Note that since X is a binary variable, $Var(X) = VL$. Aigner (1973) shows that:

$$\begin{aligned} Cov(X, \mu) &= (\alpha + \rho)VL \\ Cov(X, \mu) &= (\alpha + \rho)Var(X) \end{aligned} \tag{B1}$$

Now consider η . Note that η can be rewritten as $\eta = \eta^{FP} - \eta^{FN}$, where η^{FP} represents the number of false positives, i.e. nonwhite candidates that are counted as white candidates and η^{FN} represents the number of false negatives, i.e. white candidates that are not counted as white candidates. There are two reasons why a white candidate may not be counted. First, it may be because I have misidentified the candidate as nonwhite. Second, because I have not assigned the candidate to any particular race category. The latter could happen because I could not match the last name to the census surname file, or because the name did not reach the threshold I have selected to assign a race to the last name.

Let C be the total number of candidates running in the election, which implies that there are $C - Y^*$ nonwhite candidates running for election. Let γ^{NW} be the probability that a nonwhite candidate is misidentified as white, leading to false positives. Let γ^W be the probability that a white candidate is either misidentified as nonwhite or remains unassigned, leading to false negatives. I can then rewrite η in the following way:

$$\begin{aligned}\eta &= \eta^{FP} - \eta^{FN} \\ &= \gamma^{NW}(C - Y^*) - \gamma^W Y^*\end{aligned}\tag{B2}$$

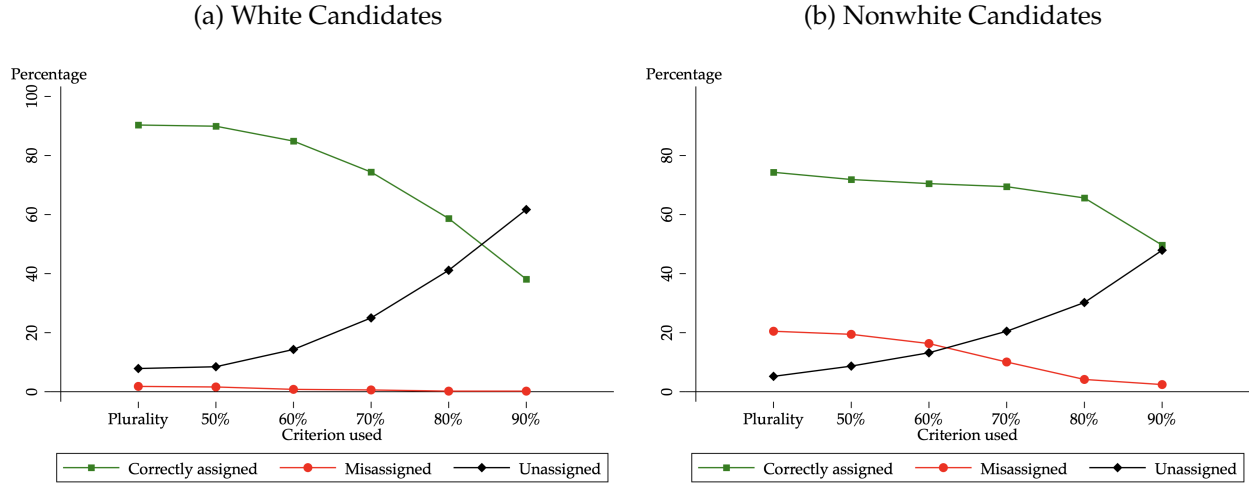
Substituting (A1) and (A2) in the OLS estimate:

$$\begin{aligned}\hat{\beta} &= [X'X]_{-1}[X'Y] \\ &= [X'X]_{-1}[X'(Y^* + \eta)] \\ &= [X'X]_{-1}[X'(\delta + X^*\beta + \epsilon + \eta)] \\ &= [X'X]_{-1}[X'(\delta + X\beta - \mu\beta + \epsilon + \eta)] \\ &= [X'X]_{-1}[X'(\delta + X\beta - \mu\beta + \epsilon + \eta^{FP} - \eta^{FN})] \\ &= [X'X]_{-1}[X'(\delta + X\beta - \mu\beta + \epsilon + \gamma^{NW}(C - Y^*) - \gamma^W Y^*)] \\ &= [X'X]_{-1}[X'(\delta + X\beta - \mu\beta + \epsilon + \gamma^{NW}(C - \delta - X^*\beta - \epsilon) - \gamma^W(\delta + X^*\beta + \epsilon))] \\ &= [X'X]_{-1}[X'(\delta + X\beta - \mu\beta + \epsilon - \gamma^{NW}(\delta + X^*\beta + \epsilon) - \gamma^W(\delta + X^*\beta + \epsilon) + \gamma^{FNR-NW}C)] \\ &= [X'X]_{-1}[X'(\delta + X\beta - \mu\beta + \epsilon - \gamma^{NW}(\delta + X\beta - \mu\beta + \epsilon) - \gamma^W(\delta + X\beta - \mu\beta + \epsilon) + \gamma^{FNR-NW}C)] \\ &= [X'X]_{-1}[X'((1 - \gamma^{NW} - \gamma^W)(\delta + X\beta - \mu\beta + \epsilon) + \gamma^{NW}C)]\end{aligned}$$

$$\begin{aligned}\Rightarrow \text{plim } \hat{\beta} &= (1 - \gamma^{NW} - \gamma^W)\left[\frac{\text{Var}(X) - \text{Cov}(X, \mu)}{\text{Var}(X)}\right]\beta + \gamma^{NW}\frac{\text{Cov}(X, C)}{\text{Var}(X)} \\ &= (1 - \gamma^{NW} - \gamma^W)(1 - \alpha - \rho)\beta + \gamma^{NW}\beta_{CX}\end{aligned}\tag{B3}$$

I cannot sign the bias. The first term shows that any measurement error will bias my estimate downward. However, the second term could lead to an upward bias in my estimate because of counting some nonwhite candidates as white. Note that β_{CX} in this case refers to the RD estimate

Figure B1: Quality of Race Assignment



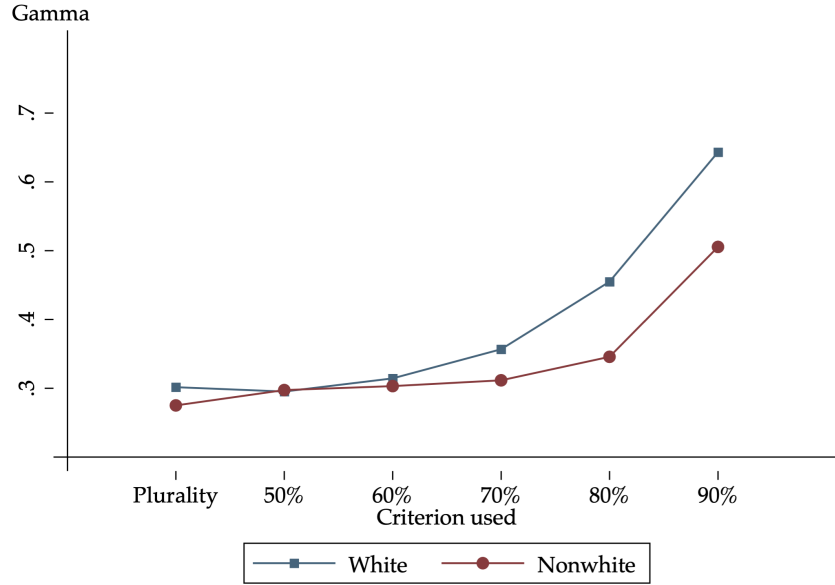
Notes: These figures provide an assessment of the quality of race assignment using the census surname files for white and nonwhite candidates by comparing against the candidate races collected by Beach and Jones (2017). The sample consists of 496 white candidates and 288 nonwhite candidates (as identified by Beach and Jones (2017)).

with the total number of candidates running in the election as the outcome variable within the given sample (i.e. X ; not X^*). Since the CEDA dataset includes the total number of candidates, I can get a consistent estimate of β_{CX} . As part of my robustness checks, I verify that the total number of candidates does not change at the discontinuity, implying that any measurement error within my sample is leading to a downward bias in my estimates. Furthermore, note that in the second term, β_{CX} is multiplied by γ^{NW} , which is the probability that a nonwhite candidate is mistakenly classified as white. As I move to a stricter criterion, the γ^{NW} , goes down, lowering the positive component of the bias.

B2: Assessing the Quality of Race Assignment

To assess the quality of race assignment using the census surname files, I use data published along with a recent study by Beach and Jones (2017). The authors also used the CEDA dataset for election years between 2005 and 2011 to study the effects of increasing ethnic diversity in the city council on public spending. They collected data on city council candidates primarily by posting the candidates' photographs on Amazon's Mechanical Turk platform and asking workers to report the candidate's ethnicity based on their name and the photograph. For a subsample, they were also able to obtain data directly from the city council's office, and for another subsample of Asian and Hispanic candidates, they contacted National Association of Latino Elected and Appointed Officials and the Asian Pacific American Institute for Congressional Studies to complement their

Figure B2: Quality of Race Assignment



Notes: This figure shows the measurement error for a given sample based on the quality of race assignment. Higher values of γ for white (nonwhite) candidates will lead to a more downward biased treatment effect for the number of white (nonwhite) candidates in the next election.

data collection.

The authors published ethnicity data for a subset of the candidates they obtained the ethnicities of. These were candidates who won a close election against a candidate from the modal ethnicity for that particular city, and who ran for re-election later. In this sample, I have 496 candidates that were identified as white by Beach and Jones (2017), and 288 candidates that were identified as nonwhite (black, Hispanic, or Asian).

Figures B1a and B1b show the quality of race assignment using the census surname files for white and nonwhite candidates, respectively. Even using a rather weak plurality criterion, the accuracy of race assignment is very high for white candidates. However, a high percentage of nonwhite candidates are misclassified as white. Using a 70% criterion, the percentage of nonwhite candidates misclassified as white falls to about 10% as more of such candidates remain unassigned. More detailed tables by separate race categories are provided in Tables B1a and B1b for the plurality and the 70% criterion, respectively.

Another way to assess the quality is by looking at the amount of measurement error for a given sample. Consider the term γ^{NW} and γ^W in equation (B3). Let $\gamma = \gamma^{NW} + \gamma^W$. For higher values of γ , my estimate of the white candidates running in the next election in cities with nonwhite

Table B1a: Comparing Census and Beach & Jones (2017) Using the Plurality Criterion

Beach & Jones (2017)	Census Surname Files					Total
	White	Black	Hispanic	Asian	Unassigned	
White	448	3	1	5	39	496
Black	39	3	0	1	1	44
Hispanic	14	1	165	1	9	190
Asian	6	0	1	42	5	54
Total	507	7	167	49	54	784

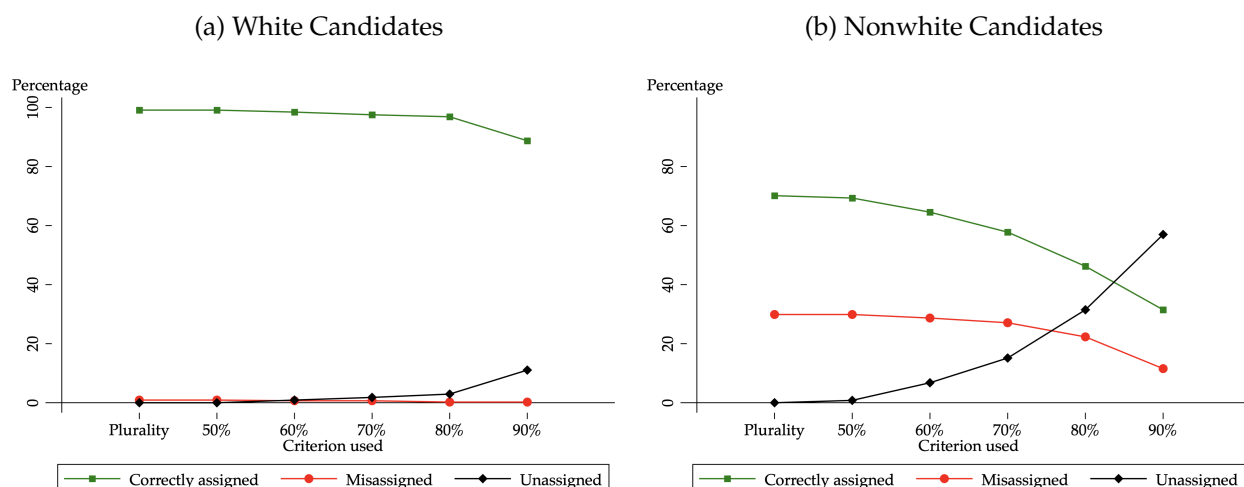
Table B1b: Comparing Census and Beach & Jones (2017) Using the 70% Criterion

Beach & Jones (2017)	Census Surname Files					Total
	White	Black	Hispanic	Asian	Unassigned	
White	369	0	1	2	124	496
Black	17	1	0	1	25	44
Hispanic	9	0	160	1	20	190
Asian	3	0	1	36	14	54
Total	398	1	162	40	183	784

victories will be more downward biased. Figure B2 plots γ for white and nonwhite candidates. We see that γ remains close to 0.3 for both white and nonwhite candidates as the criterion is tightened from plurality to 70%. Within this range, a tighter criterion is removing bad matches, increasing the number of unassigned candidates. Past 70%, we see big jumps in γ as a lot of good matches are dropped into the unassigned category. Based on the figures 1 and 2, I will use the 70% criterion for as my preferred criterion for race assignment.

Appendix C: NamePrism

Figure C1: Quality of Race Assignment Using NamePrism



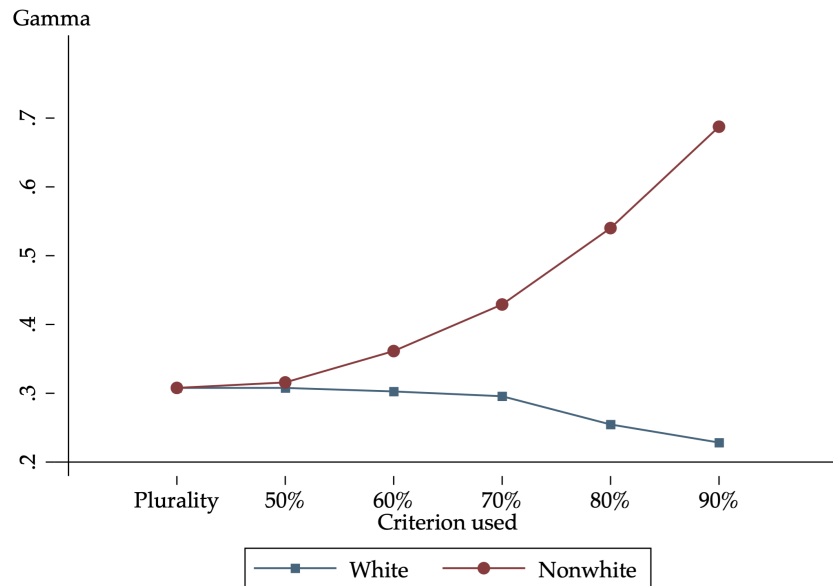
Notes: These figures provide an assessment of the quality of race assignment using NamePrism for white and nonwhite candidates by comparing against the candidate races collected by Beach and Jones (2017). The sample includes 443 white and 251 nonwhite candidates (as identified by Beach and Jones (2017)).

NamePrism uses a training data set of over 57 million contact lists from a major email company combined the phenomenon of homophily in communication to predict nationality. Homophily is the tendency of people to associate with people of similar age, language, and location. NamePrism then uses U.S. census surname file as the ground truth to derive a posterior ethnicity distribution for the first names. The final product can be used to get a probability distribution based on first and last names across the same six race categories as the census surname files.

Figures C1a and C1b compare the quality of race assignment using NamePrism against the data from Beach and Jones (2017). A caveat: These graphs are based on a slightly smaller subset of candidates than Figures B1a and B1b for two reasons. First, I used the NamePrism API after removing all candidates that ran for mayor. Second, due to an error in coding, I did not obtain the NamePrism probabilities for any candidates running in 2015. Note that the latter could also affect the replication of the main results of the paper using NamePrism shown below. The final sample includes 443 candidates identified as white by Beach and Jones (2017), and 251 candidates that were identified as nonwhite (black, Hispanic, or Asian).

Similar to the Census surname files, even using a rather weak plurality criteria delivers very accurate race assignment for white candidates. On the other hand, NamePrism misclassifies a very high percentage of nonwhite candidates as white. Even with an 80% criterion, the percentage of

Figure C2: Gamma



Notes: This figure shows the measurement error for a given sample based on the quality of race assignment using NamePrism. Higher values of γ for white (nonwhite) candidates will lead to a more downward biased treatment effect for the number of white (nonwhite) candidates in the next election.

nonwhite candidates misclassified as white remains above 20%.

Figure C2 plots the value of γ , as well as the analogous term for nonwhite candidates. Recall that a higher the value of γ leads to a more downward biased estimate of the main treatment effect (number of white candidates). For white candidates, the value of γ is lower using NamePrism than the census surname files, and falls as the assignment threshold is increased past 70%. On the other hand, the value of γ increases dramatically for nonwhite candidates as I move to a stricter criterion.