

Copyright
by
Lindsay Virginia Dun
2021

The Dissertation Committee for Lindsay Virginia Dun
certifies that this is the approved version of the following dissertation:

**“Everything we think we know is wrong” (or is it?):
Modeling voter decision-making in primary elections**

Committee:

Daron Shaw, Supervisor

Bethany Albertson

Stuart Soroka

Natalie Stroud

Christopher Wlezien

**“Everything we think we know is wrong” (or is it?):
Modeling voter decision-making in primary elections**

by

Lindsay Virginia Dun

DISSERTATION

Presented to the Faculty of the Graduate School of

The University of Texas at Austin

in Partial Fulfillment

of the Requirements

for the Degree of

DOCTOR OF PHILOSOPHY

THE UNIVERSITY OF TEXAS AT AUSTIN

December 2021

Dedicated to my husband, Nathan and my parents, Andrew and Carolyn.
All of them have supported me throughout graduate school and have always
encouraged me to achieve my goals.

Acknowledgments

This project owes a great debt to the data and mentorship provided by Daron Shaw throughout the course of my graduate school career. I also am deeply grateful for the mentorship and advice I received from each of my other committee members, both on matters related and unrelated to the dissertation.

Abstract

“Everything we think we know is wrong” (or is it?): Modeling voter decision-making in primary elections

Lindsay Virginia Dun, Ph.D.
The University of Texas at Austin, 2021

Supervisor: Daron Shaw

How do voters make decisions in primary elections? In this project, I argue that voters first narrow down a large field of primary election candidates using information and viability cues, and then weight the remaining candidates more rigorously in an expected utility framework. I first expand upon a theory first proposed by Stone, Rapoport, and Atkeson (1995), both by incorporating new variables I think are important as well as explicitly incorporating over-time dynamics. My first two empirical chapters focus primarily on the “winnowing” or “narrowing the field” process for voters, analyzing what aggregate campaign and contextual variables influence aggregate indicators of opinion formation and viability. I find that media attention is a

particularly important driver of both processes, and that debate performance and ad spending are also related to my aggregate indicator of viability (poll support). My third empirical chapter focuses on individual decision-making in an expected utility framework, and finds a particularly strong influence of electability perceptions on vote choice. Issue emphasis and candidate traits, however, also are significant predictors of vote choice even when controlling for electability perceptions. All three empirical chapters defend my theory against the “projection” criticism common to a good deal of work on campaign effects.

Table of Contents

List of Tables	x
List of Figures	xii
Chapter 1. Introduction	1
1.1 A Model of Primary Election Vote Choice	6
1.1.1 Step 1: Narrow the Field	9
1.1.2 Step 2: Evaluate candidates	15
1.2 Plan for the Dissertation	22
Chapter 2. Time Series Properties of Primary Election Variables	25
2.1 Data and Measurement	26
2.2 Time Series Analysis of Independent and Dependent Variables	33
2.3 The story of the 2020 primary, told as time series	41
Chapter 3. Narrowing the Field	45
3.1 Analysis Strategy	47
3.2 Results	50
3.2.1 Pooled Models	50
3.2.2 Candidate-Specific Models	55
3.3 Discussion	60
Chapter 4. Evaluating Candidates	63
4.1 Data and Measurement	63
4.2 Results	70
4.2.1 Evaluating Hypothesis 1	70
4.2.2 Evaluating Hypothesis 2	78
4.3 Discussion	84

4.3.1	Primary Elections and Intra-party Issue Ownership . . .	85
4.3.2	Future directions	87
Chapter 5.	Conclusion	89
5.1	Why study primaries?	91
5.2	Applicability of this work to other types of US primaries . . .	95
5.3	Directions for future research	96
Appendices		98
Appendix A.	Supplemental Analyses for Chapter 2	99
A.1	Unit-root tests	99
A.1.1	Polling time series	99
A.1.2	Don't know time series	105
A.1.3	Decision set time series	108
A.1.4	Media attention time series	112
A.1.5	Positive debate coverage time series	116
A.1.6	Endorsement points time series	120
Appendix B.	Supplemental Analyses for Chapter 3	123
B.1	Candidate Specific Results, ADL with one lag	123
Appendix C.	Supplemental Analyses for Chapter 4	128
C.1	Sanders Panel Model, Estimated using OLS and Heteroskedasticity- Robust SE	128
References		130

List of Tables

3.1	Candidate Pooled Models	51
3.2	Opinion/Willingness-to-rate Results	56
3.3	Viability/Poll Support Results	58
3.4	Decision Set Results	59
4.1	Descriptive Statistics, Biden and Sanders Expected Utility Metrics	70
4.2	Test of H1a	72
4.3	Biden Expected Utility Models	73
4.4	Sanders Expected Utility Models	76
4.5	Student Panel Data, First-Difference Expected Utility Results	80
4.6	Student Panel Data, Lag Model Results	82
A.1	ADF Test, Midpoint-Aggregated Biden Poll Series	100
A.2	ADF Test, Midpoint-Aggregated Sanders Poll Series	101
A.3	ADF Test, Midpoint-Aggregated Warren Poll Series	102
A.4	ADF Test, Midpoint-Aggregated Buttigieg Poll Series	103
A.5	ADF Test, Biden Don't Know Time Series	105
A.6	ADF Test, Sanders Don't Know Time Series	105
A.7	ADF Test, Warren Don't Know Time Series	106
A.8	ADF Test, Buttigieg Don't Know Time Series	107
A.9	ADF Test, Biden Decision Set Time Series	108
A.10	ADF Test, Sanders Decision Set Time Series	109
A.11	ADF Test, Warren Decision Set Time Series	110
A.12	ADF Test, Buttigieg Decision Set Time Series	110
A.13	ADF Test, Biden Media Attention Time Series	112
A.14	ADF Test, Sanders Media Attention Time Series	112
A.15	ADF Test, Warren Media Attention Time Series	114

A.16 ADF Test, Buttigieg Media Attention Time Series	115
A.17 ADF Test, Biden Debate Coverage Time Series	116
A.18 ADF Test, Sanders Debate Coverage Time Series	117
A.19 ADF Test, Warren Debate Coverage Time Series	117
A.20 ADF Test, Buttigieg Debate Coverage Time Series	118
A.21 ADF Test, Biden Endorsement Point Time Series	120
A.22 ADF Test, Sanders Endorsement Point Time Series	120
A.23 ADF Test, Warren Endorsement Point Time Series	121
A.24 ADF Test, Buttigieg Endorsement Point Time Series	122
 B.1 Biden Models	 124
B.2 Sanders Models	125
B.3 Warren Models	126
B.4 Buttigieg Models	127
 C.1 Sanders Panel Model	 129

List of Figures

1.1	Model proposed in Stone, Rapoport, and Atkeson (1995) . . .	7
2.1	Midpoint-aggregated poll time series	35
2.2	Daily count of news stories including candidate’s name	37
2.3	Proportion of debate coverage including positive words	38
2.4	Endorsement Points	39
2.5	Proportion of respondents selecting “haven’t heard enough” in the favorability question	40
2.6	Proportion of respondents including candidate in decision set .	42
4.1	Biden Predicted Probabilities, Voter Analysis Data	74
4.2	Sanders Predicted Probabilities, Voter Analysis Data	77
A.1	Histogram of Respondent Age	99
A.2	Pooled poll time series	100
A.3	Biden Poll ACF and PACF	101
A.4	Sanders Poll ACF and PACF	102
A.5	Warren Poll ACF and PACF	103
A.6	Buttigieg Poll ACF and PACF	104
A.7	Biden Don’t Know ACF and PACF	105
A.8	Sanders Don’t Know ACF and PACF	106
A.9	Warren Don’t Know ACF and PACF	106
A.10	Buttigieg Don’t Know ACF and PACF	107
A.11	Biden Decision Set PACF and ACF	108
A.12	Sanders Decision Set ACF and PACF	109
A.13	Warren Decision Set ACF and PACF	110
A.14	Buttigieg Decision Set ACF and PACF	111
A.15	Biden Media Attention PACF and ACF	113
A.16	Sanders Media Attention PACF and ACF	113

A.17 Warren Media Attention PACF and ACF	114
A.18 Buttigieg Media Attention PACF and ACF	115
A.19 Biden Debate Coverage PACF and ACF	116
A.20 Sanders Positive Debate Coverage PACF and ACF	117
A.21 Warren Positive Debate Coverage PACF and ACF	118
A.22 Buttigieg Positive Debate Coverage PACF and ACF	119
A.23 Biden Endorsement Point PACF and ACF	120
A.24 Sanders Endorsement Point PACF and ACF	121
A.25 Warren Endorsement Point PACF and ACF	121
A.26 Buttigieg Endorsement Point PACF and ACF	122

Chapter 1: Introduction

Since the 1972 campaign – when the power to choose the party nominees was shifted from national convention delegates to voters in state primaries and caucuses – Democrats have rarely had a front-runner as dominant as Clinton.

Gallup, 10/22/2007

If Trump is nominated, then everything we think we know about presidential nominations is wrong.

Larry Sabato, 8/22/2015

Voting in American primary elections remains a fairly opaque political process. We know that certain variables matter, particularly at the presidential level: electability in the general election, viability in the primary, and candidate quality remain prominent in the literature as predictors of vote choice. Public attention and party endorsements have also been raised as predictors of eventual primary outcomes. Yet, it is difficult to understand how all of these variables fit together, because we lack a unified model of primary election choice.

The study of primary elections is critically important for our understanding of contemporary American politics. First, one could argue that understanding voting behavior in primary elections is as important as understanding behavior in general elections. Primaries, after all, determine the

choice set available to voters in general elections. The literature on presidential elections demonstrates the importance of that choice set; though these elections are not necessarily perfectly predictable, the aggregate distribution of the vote is heavily influenced by environmental variables like the strength of the economy and the distribution of partisanship in the electorate (Erikson and Wlezien 2012c). These variables are often called the “fundamentals” of general elections. Though it is unclear how well we can forecast elections using fundamentals alone¹, subjective and objective economic indicators can “account for at least half the variance in the final vote” (Erikson and Wlezien 2012c, p. 123). An argument could be made that, at least in elections with fundamentals that clearly favor one party over the other, the U.S. president is effectively chosen in the primaries.

Second, it is clear that theory used to understand vote choice in general elections does not always apply to primary election choice. Party ID, a variable almost determinative of individual general election vote for moderate and strong partisans, cannot be used to predict vote choice in a primary. What the “fundamentals” of these elections are is also far less clear. Nate Silver suggested that fundraising, endorsements, and experience in elected office were possible primary election fundamentals in his first-ever primary election forecast in 2020.² Fundraising and endorsements, however, are at least in part endogenous to initial candidate support. While potentially useful in forecast-

¹<https://fivethirtyeight.com/features/models-based-on-fundamentals-have-failed-at-predicting-presidential-elections/>

²<https://fivethirtyeight.com/features/how-fivethirtyeight-2020-primary-model-works/>

ing, they do little to explain on their own why this support exists in the first place. Silver himself admits that some shifts in primary candidate support appear to be stochastic, and that change in support is unpredictable.

This uncertainty is reflected in the relative polling variability throughout primary elections compared to general elections. Erikson and Wlezien (2012c) find that, in years where early polling data is available for both candidates in a presidential election, the leader in the polls one full year before the election indeed won the election 8 out of 11 times (though, it is important to note, the distribution of preferences measured in polls certainly changes meaningfully over the course of many of these elections). In a FiveThirtyEight analysis of primary election polls taken in the first half of the year before the primary election,³ only 8 out of 16 poll frontrunners actually won the nomination.⁴ Using polls from the second half of the year before the primary election (or, starting a little more than six months ahead of the Iowa caucuses) only marginally improves predictions: 9/16 frontrunners in such polls ended up winning the nomination. Lastly, in FiveThirtyEight’s most recent “State of the Polls” analysis, Silver finds that the weighted average error of polls taken in the final 21 days before the primary election range from 7.6 to 10.1 from 2000–2016, compared to 3.2 to 4.8 for general elections over the same period.⁵ Primary elections, therefore, are more variable, harder to understand, yet ar-

³excluding primaries in which an incumbent president was a candidate

⁴<https://fivethirtyeight.com/features/what-more-than-40-years-of-early-primary-polls-tell-us-about-2020-part-1/> ; <https://fivethirtyeight.com/features/what-more-than-40-years-of-early-primary-polls-tell-us-about-2020-part-2/>

⁵<https://fivethirtyeight.com/features/the-state-of-the-polls-2019/>

guably just as important, as general elections. Questions also persist as to who effectively “chooses” nominees; do elites maneuver and manipulate behind the scenes, or can the public override elite support? The relative lack of scholarship in this area poses a problem for our understanding of contemporary election dynamics and of whether the reforms intended to transfer support to the average partisan have worked as intended.

The observation that it can be difficult to predict the outcome of presidential primary elections is not new (Traugott and Wlezien 2009). Yet, the nomination surprises that occurred in both 2016, with the Republican selection of Donald Trump, and 2020, with the resurgence of Joe Biden’s floundering campaign, suggest that this observation remains largely unaddressed. Important work on primary elections has certainly taken place in the intervening years, but studies have tended to focus on one particular variable known to be important (Steger 2008), and usually focus only on macro-level phenomena (Clinton et al. 2019) or micro-level phenomena (Mutz 1997) without bridging the two. Brady (1993) notably does attempt to create a more unified model of primary vote choice, though he relies exclusively on formal modeling, focuses on only a few variables, and does not integrate activity that occurs in primary elections before voting begins. We are thus left with pieces of a larger puzzle, without a clear picture of how everything fits together. Endorsements matter (Steger 2008), issue emphasis and electability in the general election matter (Rickershauser and Aldrich 2007), candidate traits matter (Barker et al. 2006), perceived viability in the primary matters (Abramson et al. 1992), and so on.

This project attempts to bring these pieces of the puzzle together, using new sources of data to shed light on how the important variables of the primary election literature combine to explain shifts in candidate preferences during the presidential primary cycle. I hope to use these analyses to answer the following question: how do voters arrive at a vote choice in primary elections?

This study will contribute to political science literature by providing and testing a new, unified framework within which to situate the various research that has been done on presidential primaries. I accomplish this in my empirical chapters by assessing how voters narrow down a complex set of candidate choices and how they choose between the candidates they do consider. Taken together, I hope to provide a newly comprehensive picture of primary election preference formation and change.

I will be able to use both individual and aggregate data to answer my research question, primarily (though not entirely) using the 2020 Democratic presidential race as a test case. Most work on primary elections focuses on the race for the presidential nomination, though lessons learned from these studies need not apply only to that context. In particular, an assessment of how voters weight candidate traits, issue emphases, and electability in a general election should carry over to other down-ballot races. I will explore what we can (and cannot) generalize in a concluding section.

1.1 A Model of Primary Election Vote Choice

Long before the Iowa caucuses, many likely voters can express a preference for a particular candidate in the presidential primary cycle. In a poll conducted by The Hill from 11/30 to 12/1/2019, for instance, only 13% of those queried were unable to express a preferred candidate in the 2020 Democratic primary.⁶ So, where do these preferences come from, when voters can't use party identification as a decision rule, and no elections have yet been held? A useful model of early primary election preference formation was proposed by Stone et al. (1995), which I expand upon in this project. Their model is reproduced in Figure 1.

I propose that the Stone, Rapoport, and Atkeson (SRA) model is the best starting point for my work because it is both one of the most comprehensive models in the primary election literature and acknowledges that many voters are unlikely to invest more than minimal cognitive resources in forming an initial preference. Drawing upon Tversky's (1972) "elimination by aspects" model of choice, SRA propose a mixed model in which voters simplify their decision space by first eliminating poor alternatives. This is the "narrowing the field" stage of their model; voters are assumed to eliminate candidates if they do not meet a viability threshold,⁷ if they are not in the voter's party, or if the voter does not have enough information about the candidate to cal-

⁶<https://thehill.com/hilltv/rising/472629-bloomberg-overtakes-harris-in-new-poll>

⁷Viability, in the primary elections literature, refers to a candidate's chances of winning the primary contest; electability refers to a candidate's projected chances of winning the general election

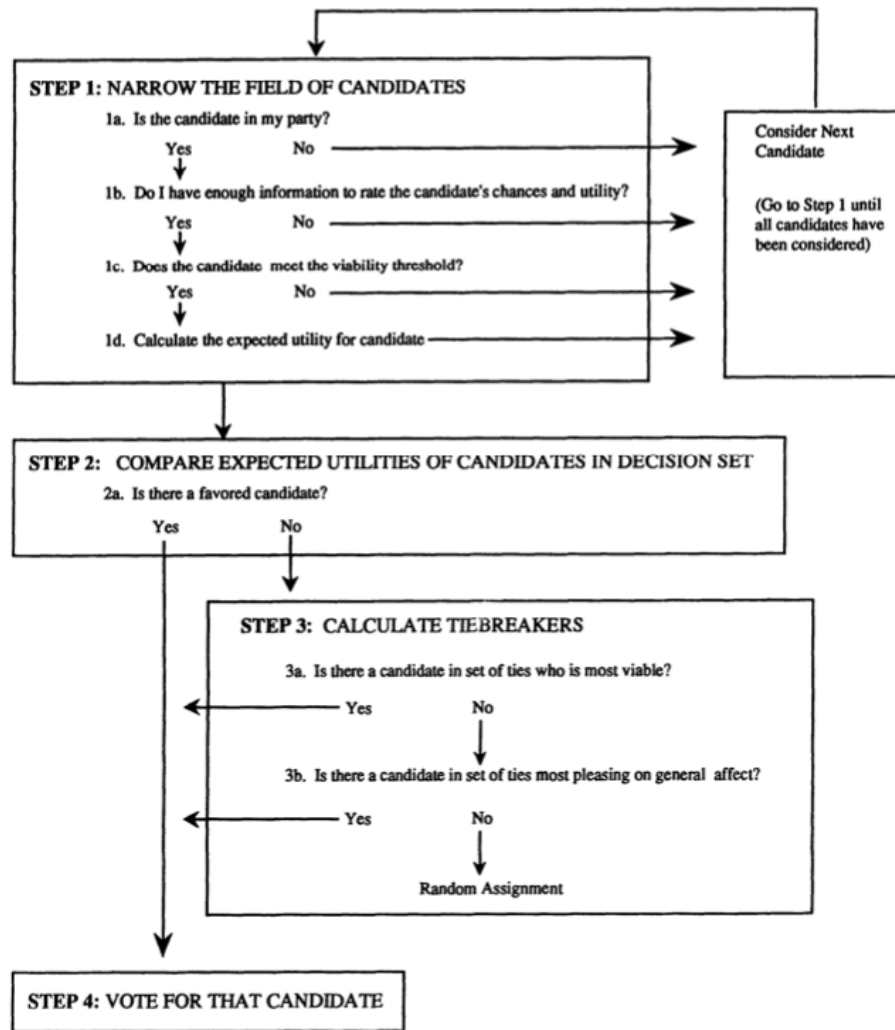


Figure 1.1: Model proposed in Stone, Rapoport, and Atkeson (1995)

culate chances/utility. The remaining candidates in a voter's decision set are evaluated more rigorously, in a traditional expected utility framework. The candidate with the highest expected utility is selected as the vote choice.

Though useful and comprehensive, Stone, Rapoport, and Atkeson note that their model “fails to incorporate an overt dynamic component to a dynamic process” (p. 158). Their model is, in other words, linear; though obviously intended to be a stylized representation of reality, the model does not provide for preference updating or change. One contribution this project seeks to make is to do just this—I will propose how we can expand the SRA model to incorporate the dynamics of the primary election season. First, I will incorporate the campaign and media activity that occur in the months before the Iowa caucuses into empirical analyses of primary election preferences. Second, I do not presuppose either theoretically or empirically that candidate decisions are static. Many voters clearly update candidate preferences in primaries, and I attempt to model how and why preferences change.

I will test whether important variables suggested by prior literature can be incorporated into a unified SRA framework. These tests should, hopefully, help provide new clarity as to which environmental and psychological variables dominate the primary election decision-making process.

The first two empirical chapters of this project will test how voters narrow the field of candidates by first exploring the 2020 electoral context and the time-series properties of primary election variables, and second by exploring multi-variate relationships between variables theorized to be important in SRA’s step 1. The third empirical chapter will introduce new variables to an expected utility model of vote choice to more comprehensively assess how voters arrive at their initial, pre-Iowa vote intention. Future iterations of

this project will further explore how election results from early contests influence vote intention, but for now those analyses are beyond the scope of the dissertation.

The first two stages of primary vote choice – narrowing the field and evaluating candidates – will thus form the heart of the three empirical chapters in my dissertation. In each of the sections to follow, I propose how I specifically will expand upon what we know about these two stages of the process by combining insights from literature in the field. Though the scope of the project is broad, there are clear gaps in the literature that this dissertation proposes to address.

1.1.1 Step 1: Narrow the Field

Voters arrive at some decision set of candidates in a step that Stone, Rapoport, and Atkeson call “narrowing the field”. In other words, before voters can evaluate candidates, voters have to decide which candidates are worth evaluating.

Voters should be adept at assessing point 1a, or the party of the candidate, in this step. However, figuring out how voters determine point 1b (whether they have enough information to evaluate candidate chances and utility) and point 1c (does the candidate meet the viability threshold) is a bit trickier, and I am aware of no study thus far that has attempted to tease this process out empirically. Assessing the relative importance of several different environmental variables at this stage of primary election decision making will

be the contribution of my first empirical chapter. This project will focus mostly on modeling point 1b, or voter information levels, due to data availability at this stage.

The information threshold required for candidate evaluation is probably different for every voter, though we can almost certainly assume that most voters do not take the time to thoroughly research each candidate (Popkin 1991). Prior research can shed some light on which cues voters might use as information shortcuts to help them narrow the field. For instance, activity during the “invisible primary” communicates information that voters can use in step 1 judgments. Cohen et al. (2008) would argue that elites help narrow down the candidates under serious consideration in the “invisible primary” before any voting begins. Endorsements from key groups or politicians signal to their followers who to support and raise public awareness of certain candidates. Public attention to candidates (conceptualized in prior work as google searches), though likely connected to many other early information signals, has also been shown to be predictive of Iowa and New Hampshire outcomes and is argued to play a role in the “invisible primary” (Swearingen et al. 2019).

Media attention also clearly raises the profile of some candidates seeking the nomination. Bartels (1988) has argued that the media’s focus on the political horse race means that the media lavishes coverage on frontrunners in polls, and that this information about the race sticks with the public. In order for the horse race to be interesting, the media needs to determine that certain candidates are potential challengers to the frontrunner. This means those

closest to the first place candidates in the polls also get a fair amount of media attention (Bartels 1988; Paolino and Shaw 2001). And, lastly, campaigns and campaign-related events can raise the profile of candidates (Popkin 1991). Strong debate performances can lead to increased media coverage, for example, which may lead to greater awareness.

Though certainly related to media attention, poll support is likely related to early viability judgments. Majority cues tend to move public opinion in the direction of the cue when people have low levels of commitment to their original views (Mutz 1992). People are also more likely to rely on polling cues specifically when they are making a difficult, or low-information choice (Boudreau and McCubbins 2010). Both of these findings are particularly relevant in primary elections, in which voters have to choose from a wide variety of candidates and when rapid opinion shifts suggest weak commitments to preferences (at least for a large subsection of voters). In this piece, I consider poll averages to be a measurement of viability perceptions, since this information would be the best available cue to primary election voters about a candidate's chances in an election prior to any official contests. Of course, an even better measure would be to ask voters to estimate a probability of primary election victory for each candidate, though at this stage of the project those data do not exist.

I follow Stone, Rapoport, and Atkeson in assuming that viability, or a candidate's chances at winning the primary election, enter into voter considerations in step 1. SRA make no explicit claims about when voters narrow the

field, but I suggest that this could begin long before any primary contests are held. Perceptions of viability will certainly change as election results begin to roll in, and the potential effects of those changing perceptions of primary chances will be explored in a later section. But, initial viability perceptions are likely formed in the months before the Iowa caucuses, and integrating these perceptions in an empirical analysis of early preferences has not, to my knowledge, been attempted.

Bartels (1988) argued that “winnowing” (his term for “narrowing the field”) can only usefully occur after voters have election performance information, as voters need to use electoral success to judge candidate “seriousness” (p. 60). Though he does convincingly argue that voter levels of information, a key component of step 1 in the SRA model, are closely related to electoral performance, one could argue that primary voters in recent elections have higher levels of information about candidates, earlier on in the primary season, than voters did in 1984. The first Democratic primary debate in the 1984 presidential primary, for instance, was held on January 15, 1984.⁸ The first Democratic debate for the 2020 cycle was held on June 26, 2019. Democratic candidates also spent a record \$870 million on ads during the fourth quarter of 2019; even without the spending of Tom Steyer and Michael Bloomberg, two billionaires who personally financed their campaigns, the ad spending of the rest of the field doubled that of the same period in the 2008 cycle.⁹ These are only a few

⁸<https://sites.dartmouth.edu/primaries/history/1984-democratic-debate/>

⁹http://www.cfinst.org/press/releases_tags/20-02-04/Historic_presidential_campaign_spending.It_s.about_to.be

indicators, among the many I could list, to suggest that campaign seasons are now longer, and that campaigns spend increasing amounts of money on efforts to increase voter information and name recognition.

The first two chapters of my dissertation will seek to explore how winnowing occurs, and when it occurs, using environmental and individual-level data to produce a clearer picture of this process than has been available previously. In addition to being one of the first (if not the first) study to model voter information levels in the fall before any elections are held, this study will also attempt to connect activity in the invisible primary to an individual, psychological model of vote choice. The focus of other prominent primary election models tends to be on the vote choice stage, rather than this winnowing stage (Bartels 1988; Brady 1993). There are two main dependent variables suggested to be important in step 1 of the SRA model, and which could plausibly be impacted by activity in the early stage of the primaries: voter information and perceived viability. These two variables should then predict which candidates make it into voter decision sets. This two-stage process will be what I attempt to analyze in the first empirical chapter.

The invisible primary (suggested in the literature to consist of endorsements, exposure, fundraising, and media coverage), and early campaign events like debates, presumably serve to increase voter information about candidate utility (value) and increase perceived viability of candidates. The general election vote choice literature reinforces the importance of including campaign events as possible predictors of opinion shifts at this early stage (Shaw 1999).

The current literature does not illuminate which environmental variables primarily serve to influence information levels or viability assessments—do endorsements, for example, signal viability to voters, or do they provide utility-relevant information instead? To what extent are both perceptions updated? The only exception to this relative ambiguity are expectations regarding media and polling effects. The media seem to focus on horse-race information, which should primarily influence early viability judgments; polls should serve as an indicator of candidate “seriousness” in the primary.

If a voter has enough information about a candidate to judge utility, and a voter perceives a candidate to meet a minimum viability threshold, that candidate should get included in the voter’s decision set of candidates. Stone, Rapoport, and Atkeson test this stage of the model, and find that these two conditions serve as useful screens before moving to step 2. Their results, however, rely on a simulation approach that does not ask voters themselves who the candidates in their decision sets actually are. Using new survey data, I will be able to operationalize voter decision sets as a concrete dependent variable, and predict inclusion in a voter’s decision set using information and viability judgments. This will help us understand how voters arrive at a conscious decision set.

As a last note, it is worth addressing that, until this point, the theory laid out has assumed that endorsements, polls, media attention, and campaign events are exogenous to voters’ information about candidates and perceptions of viability. This is almost certainly not the case—while these variables quite

plausibly do contribute a causal effect on individual information/perceptions, there are equally plausible reciprocal effects. One could see, for example, how elite actors would want to strategically endorse candidates that voters perceive to be viable. Additionally, Utych and Kam (2014) find that increasing candidate viability can increase the amount of information individuals seek about that candidate, leading to increased support and interest. Box-Steffensmeier et al. (2009) argue that the media and campaigns are responsive to voter preferences in general election settings, and find compelling empirical evidence that campaigns and the media adjust activity in response to public support. For this reason, I interpret causality in these models cautiously, and consider significant predictors better interpreted to mean “leading indicators” than strictly causal variables.

1.1.2 Step 2: Evaluate candidates

Once voters have arrived at a decision set of candidates, they have to calculate the expected utility of each candidate in order to arrive at a vote choice. Several studies have performed a basic expected utility analysis of primary decision-making, and have found that this model consistently outperforms other vote choice decision rules. Several models of strategic voting in the primaries, however, have found only mixed evidence that issue positions matter to primary electorates (Bartels 1988; Stone et al. 1992b). This is a curious finding to observers of primary election campaigns, as these campaigns tend to heavily focus on issues. Why would they do this if issues do not seem to

consistently drive primary election vote choice? I propose that we can expand our understanding of this step of the SRA model (who operationalize utility as a combination of ideological proximity and candidate traits) by incorporating a new measure of issue importance. In other words, I seek to propose an important new variable that can add predictive value to a traditional expected utility model. Undertaking these new tests of the expected utility framework will be the primary contribution of my second empirical chapter.

Utility, in the primary election framework, is defined as a voter's perceived value of a particular candidate winning the presidency. In the SRA model, as well as other expected utility models of this type, voters discount the utility they would receive from a particular primary candidate winning the presidency by the probability that the candidate can beat the other party's opponent in a general election. Step 2 of the SRA model thus requires that voters calculate utility and electability for each candidate. The model effectively assumes that there is some degree of strategic thinking for all primary voters, because electability is assumed to enter into evaluations for everyone. Though this is clearly an oversimplification of reality,¹⁰ the expected utility model performs quite well for a large majority of primary voters (Abramowitz 1989; Stone et al. 1992b, 1995). In a two candidate race, this model is expressed as:

$$E(U) = P_A(U_A) - P_B(U_B)$$

¹⁰Sometimes, for instance, voters will support a candidate despite thinking he/she is unelectable. This may be because they want to highlight a particular issue in the primary, or because they want to cast an expressive vote, among other reasons

Where P = the subjective probabilities that candidates A and B will be elected if nominated, and U = candidate utilities (Stone et al. 1992b). Voters are assumed to arrive at a vote choice/candidate preference after this step in the SRA model (assuming there are no exact ties). Typically, the model is estimated with cross sectional data collected right before the primary election season or right after it begins. I follow suit in this work and estimate expected utility models using data collected right before the Iowa caucuses, as well as data collected at the beginning of the primary season, to test hypotheses derived from the theory laid on in this section.

Though the expected utility model assumes that voters primarily use electability to discount utility, research in this area has revealed that the effects of perceived electability are not necessarily so simple. Electability has an independent, direct effect on stated vote choice, which has been taken to mean primary voters are indeed strategic (Abramowitz 1989; Stone et al. 1992b). Electability also has an indirect effect on preference, through increasing positive affect towards the candidate seen as “electable” (Rickershauser and Aldrich 2007). Though I will not attempt to disentangle the precise causal path between electability and vote choice, electability should have an independent effect on voters’ eventual candidate selection if the expected utility framework is correct.

Polls typically operationalize electability using hypothetical head to head candidate matchups, such as the following Suffolk University/USA Today question: “If the November 2020 general election were held today and the

choices were Republican Donald Trump, Democrat Joe Biden, or a third party candidate, for whom would you vote or lean toward?”¹¹ The closest we can come to the evaluation suggested by the expected utility model, which requires that we can compare an individual’s rankings of candidates on some electability scale, is to ask voters to estimate a candidate’s “probability of winning the general election if elected by his party” (Stone et al. 1992b). We can improve upon this measure further if the candidate from the other party is known in advance, by asking respondents how likely it is candidate X can beat his or her presumptive general election opponent. Conceptually clear measures of electability are critical to test the expected utility model.

Political scientists typically operationalize “utility” in a primary either using a spatial voting-type framework, in which voters seek to select the most ideologically proximate candidate, or with perceptions of candidate traits, such as competence or honesty. There is indeed some evidence that candidate traits independently affect voter preferences (Stone et al. 1992b, 1995). The best way to measure candidate utility, beyond trait perceptions, is far from settled—some scholars incorporate proximity using ideological self-placement scales (Stone et al. 1992b), some argue for proximity measured in issue positions (Stone et al. 1992b), and some use feeling thermometer or favorability scores as a summary evaluation measure (Abramowitz 1989; Abramson et al. 1992).

¹¹USA Today (2020). Suffolk University/USA Today Poll, Question 2 [31117341.00001]. Suffolk University Political Research Center. Cornell University, Ithaca, NY: Roper Center for Public Opinion Research.

Ideological proximity, when compared to other utility dimensions like issue proximity and candidate traits, has been found to be a weaker predictor of choice (Stone et al. 1992b). Some have suggested this is because there is not enough variance in intra-party candidate ideology to usefully distinguish between candidates (Aldrich and Alvarez 1994). Closely related to ideological proximity on a left-right dimension is issue position proximity. Though in a head-to-head matchup, this variable performs better than left-right ideological proximity, it is a considerably weaker performer in vote choice models than candidate traits (Stone et al. 1992b). Outside of the expected utility framework, the evidence that issue positions matter to primary voters is not particularly convincing (Williams et al. 1976; Gopoian 1982; Norrander 1986). These studies also tend to find candidate traits to be a superior predictor of primary vote choice.

In many models of primary vote choice, issue positions are assumed to enter into summary evaluations of candidates, but are not measured independently (Abramowitz 1989; Deltas et al. 2016). In studies of more recent elections, scholars have argued issue positions, like left-right ideology, also don't often vary enough between intra-party candidates to be useful (Jackman and Vavreck 2010). This dissertation project hopes to incorporate a new variable, demonstrated to be important to voters in other primary election research, to revive the possibility that issues could matter to primary electorates in a strategic voting model.

A few scholars have pointed out that issue *emphasis* on the part of

campaigns does vary enough to be useful, that picking up on issue emphasis (or issue priorities) is a fairly straightforward task on the part of voters, and that distinguishing candidates in terms of issue emphasis is easier for voters than distinguishing candidates in terms of proximity in an ideological space (Aldrich and Alvarez 1994; Rickershauser and Aldrich 2007). Rickershauser and Aldrich (2007) test an issue-emphasis + electability model that is quite close to an expected utility model, though their dependent variable of interest is candidate favorability rather than vote choice. Rickershauser and Aldrich (2007) find that candidate favorability increases when less sophisticated voters are told that a candidate is particularly concerned with an issue they care about. For all respondents, the authors find that telling voters a candidate emphasizes an issue owned by their party increases favorability. Though the study strongly suggests that issue emphasis could be a core component of a strategic voting model, the study is limited in a few key ways. First, it relies on a relatively small sample of students at a major university for its analysis. Second, though favorability is closely tied to vote choice, it is not conceptually equivalent. It is not uncommon for voters to feel similarly favorably about several primary candidates, which can then lead to errors in prediction (Wattier 1983). The contribution of this dissertation will be to resolve a few of these issues, by including both non-student and student data in the analyses, modeling vote choice directly, and testing (to the extent possible) the predictive value of traits compared to issue emphasis.¹²

¹²Though I do not have the data I would need to test what variables might moderate

Because many expected utility models are tested using cross-sectional survey data, the idea that voters evaluate candidate utility and electability and *then* choose a first choice may be called into question. Voters may instead simply project positive attributes onto their pre-existing favorite candidate. What then appears to be an expected utility calculation may, in fact, simply reflect these projections. This is a legitimate concern and has been raised as a possibility in primary elections literature before (Bartels 1988; Stone et al. 1992b). Similar criticisms have been leveled at the economic voting and retrospective voting literatures in general elections; at an individual level, there is certainly evidence that vote choice and party ID influence economic perceptions and performance assessments (Wlezien et al. 1997; Pickup and Evans 2013; Wlezien 2016).

expected utility calculations, it is important to note that the theory laid out thus far assumes that all voters interpret and apply electability and utility-relevant information to all candidates in the same way. This obviously will not always be the case, as prior political science research has uncovered that certain traits can unconsciously (or consciously) bias how candidates are evaluated by voters. Research on women in politics and racial and ethnic politics provides some examples of how biases might disadvantage female and minority candidates. Black candidates, for example, are perceived as more liberal than their White counterparts, even when they hold similar issue positions (Jacobsmeier 2015), and they are evaluated more harshly than otherwise equivalent white candidates on feeling thermometer scales (Terkildsen 1993). Evidence from 2008 suggests that Barack Obama’s race slowed down the momentum he was able to achieve from early primary victories, as he both lost support among racially resentful voters and did not win over racially resentful Clinton supporters (Jackman and Vavreck 2010). Dolan (2010) provides a useful review of the ways in which gender stereotypes might influence candidate perceptions, arguing that “numerous experiments and surveys indicate that voters believe female politicians are warmer and more compassionate, better able to handle education, family, and women’s issues, and are more liberal, Democratic, and feminist than men”, while “male politicians are seen as strong and intelligent, best able to handle crime, defense, and foreign policy issues, and more conservative” (p. 71). The independent impact of gender stereotypes on vote choice may be small in general election settings, when these attitudes are subsumed by party id (Dolan 2014), but this offers little comfort in primary election settings without this cue.

As is the case for the relationships proposed in the first chapter, there is good reason to believe that the story is not as simple as affect causing utility judgments. Bartels (1988) finds that individual levels of projection appear to decrease with greater levels of information, Stone et al. (1992b) find that expected utility evaluations still show a strong relationship to choice even when controlling for candidate affect, and Hirano et al. (2015) find that voters learn (accurately) about the ideological positions of candidates throughout the course of state primary elections. Additionally, manipulating electability and issue emphasis in an experimental setting is causally linked to candidate evaluations (Rickershauser and Aldrich 2007). Assuming affect is the main driver of utility judgments also does not provide insight as to where positive affect might come from in the first place. However, the question is worth picking up again in this empirical chapter. I will explore the data sources I plan to use to assess this causal identification problem in the methodology section, as well as the limits of what we can infer about the precise causal mechanisms at work given the data I have.

1.2 Plan for the Dissertation

In my empirical chapters, I hope to expand our knowledge of primary elections by explicitly modeling campaign dynamics in a unified theoretical framework. Empirically, I add to the work that has been done by relying almost entirely on time series cross sectional and panel data. Doing so will allow me to address one key shortcoming of a large amount of seminal work

that's been done on primaries: analyses tend to rely on point-in-time, cross sectional data, even though the theory associated with this work hints at over-time dynamics. Bartels (1988) notably relies on TSCS data collected by the ANES, though his work is unique in this regard, and his theory is distinct from the SRA model I test in this work. Abramson et al (1992) also leverage the rolling cross sectional structure of the 1988 ANES Super Tuesday Study to some extent, though they test a much more limited version of the expected utility model I set up in this work and do not have panel data.

To achieve this goal, I first collect a large amount of time series data from the 2020 primary election cycle, and model the time series of interest descriptively. This is important to set up the multivariate analyses I get to later, as well as to gut-check the primary election dynamics themselves. I assess whether the movement in my time series conform to expectations, and argue that that movement is substantively interesting.

Next, I turn to multivariate analyses, and test which of my predictor variables are leading indicators of my information/opinion time series, my viability/poll support time series, and my decision set time series. These findings are meant to be a dynamic, time series test of Step 1 of the SRA model.

Lastly, I use cross-sectional and panel data to test my expected utility model, which builds on Step 2 of the SRA model. Using my panel data, I am able to test whether prior changes in opinions about electability and issue emphasis influence time t vote intention, which to my knowledge is a unique

contribution to literature on the subject.

To preview what I find empirically in all three chapters: campaign dynamics do matter, and are worth considering in the study of primary elections. My findings suggest that legitimate criticisms of election-based work (namely, that voters do not incorporate relevant indicators of campaign success, and merely project opinions onto already preferred candidates) are not the whole story when it comes to primary voting. Additionally, I find that voters behave in a fashion consistent with expected utility models of vote choice, and argue that this finding is relevant in several comparative contexts.

Chapter 2: Time Series Properties of Primary Election Variables

As discussed in my introductory chapter, primary elections remain a challenge to forecast despite decades of important research in the area. Opinion change over the course of the primary election season is frequent and relatively unpredictable. The next two chapters aim to fill this gap in scholarly understanding by more rigorously testing how voters dynamically “narrow the field” of candidates. I argue that voters convert a large field of candidates into a manageable decision set using informational cues related to campaign activity, public support, and elite support. Per my proposed model of decision-making in the early primary season, cues need to both be strong enough for voters to receive relevant information and need to signal that candidates are viable.

This chapter will explore the time series variables I will use in multivariate analyses in Chapter 3. During the 2020 primary, I gathered polling data, media data, endorsement data, and debate performance data for four candidates: Joe Biden, Bernie Sanders, Elizabeth Warren, and Pete Buttigieg. In the sections to follow, I will describe how I transformed all of these raw inputs into time series variables to be used in hypothesis testing.

As explored in the theory section, other variables are likely important in this stage of the primary season as well. For instance, campaign fundraising may be an alternative cue that is important to study at this stage of the

process. This variable is excluded from this study not because it is deemed to be unimportant, but rather because at this stage of the project data limitations exist. Daily campaign fundraising totals are available for donors who contribute greater than \$200, but not for small dollar donors. Given that these types of donations are sharply rising, and played a large role in the 2020 election cycle (particularly for Democrats)¹, it is unclear how valid the time series of larger-value donations might be as a measure of true support—at least, the type of grassroots support that increasingly matters for Democratic voters. Despite this limitation, these data and analyses shed more light on this process than has been available previously, and reflect the importance of incorporating each of these cues into a single model.

2.1 Data and Measurement

All data examined in this chapter are time series data, measured at the daily level from 7/18/2019 until 3/03/2020, allowing for 230 unique observations for each variable.

Polling time series/primary viability: Not many surveys (to date, I have found none in the current primary cycle) measure who respondents think will win a given presidential primary, let alone some type of respondent estimate of the probability of primary victory for each candidate. However, voters likely use horse-race information to inform viability judgments, and polls stand out

¹<https://www.npr.org/2020/10/22/925892007/fundraising-fuels-democratic-money-advantage-over-gop-in-most-races>

as a clear potential indicator of primary support at any given stage of the election cycle. As such, I am limited in this project to measuring an *indicator* of viability, rather than voters' perceptions of viability. Though these two are likely closely related, the SRA model focuses on perceptions of viability, and thus in future iterations of this work I hope to develop a measure that hews more closely to this theory. I am not aware of prior work that has used poll support as an aggregate measure of viability, so future work will be required to test how closely poll support moves with viability perceptions. This will require, at a minimum, rolling cross-sectional data which incorporates a conceptually appropriate viability question. We are limited to considering poll support to be a plausible proxy for viability in this particular project.

I operationalize poll support in the primary (viability) through daily average poll standing, which is measured using data downloaded from FiveThirtyEight.² FiveThirtyEight publishes their polls policy here; broadly, their database aims to include as many publicly available polls as possible, as long as those polls “attempt to survey a representative sample” and publish basic design and methodology notes (such as the field dates for the poll). Out of the 597 polls that include at least one field date within my range and include all four of my candidates as response options, 3 are funded by a partisan source and 115 are labeled as tracking polls. I include all partisan and tracking polls in my data in order to maximize the data available for candidates who might have been under-pollled at various points throughout the time series, such as

²<https://github.com/fivethirtyeight/data/tree/master/polls> ; accessed 6/12/2020

Elizabeth Warren or Pete Buttigieg.

I develop two polling time series using the FiveThirtyEight data. The first poll series counts each poll only once, following Erikson and Wleizen (1999). In this series, I give each poll a single “date” value, assigned as the midpoint of the field dates. For every poll in my time series, I generate a series of values representing the proportion of support for each candidate divided by the total support of the four candidates in my analyses. For example, in each poll, I divide Biden support by the sum of Biden, Sanders, Warren, and Buttigieg support. Because many polls in the sample asked about different numbers of candidates, this allows me to keep daily poll measures for each candidate on the same scale. FiveThirtyEight represents percent support as the proportion of voters selecting a candidate multiplied by 100. So, if in a single poll Biden has 35% support, Sanders has 17% support, Warren has 15% support, and Buttigieg has 5% support, and I want to calculate relative Biden support, I get $35/72 = 0.49$ as my Biden value for that poll. I repeat this process for each of the four candidates.³ I then simply take the mean value of my relative candidate support variable (in the above example, the 0.49 value for Biden) of the total polls in a given day in my time series. I use this variable in all multiple regression analyses.

The second series, used to best assess how candidate support trends

³This analytical decision removes all polls in my total time series of 597 which do not include the full set of four candidates as response options, leaving me with a final set of 490 polls.

over time, pools polling data using the second method outlined in Erikson and Wlezien (1999). A poll gets included in this time series if it has at least one field date that spans my target date range. Polls are included in each day’s mean if they are in the field on that day. For this series, I calculate a weighted mean of each candidate’s poll values for each day, where my weight is equal to $1/\text{total number of poll field days}$. To continue my example, if the above poll was in the field for 7 days, I would multiply 0.49×0.14 , sum that value with all other weighted poll values for that day, and divide by the total sum of the weights for that day to get my final value for Joe Biden.

Media attention: Media attention is measured using news articles downloaded from the Factiva database.⁴ I sum the total news results containing candidate names each day, per candidate, to get my daily media attention time series.

Debate performance: Debate performance is also measured using news articles downloaded from the Factiva database.⁵ To measure the media’s as-

⁴I queried the database four times, or once per candidate. My search terms were “democratic primary and Joe Biden”, etc, which returned all articles in the database containing both the exact phrase “democratic primary” and “Joe Biden”. I then restricted my search to the relevant dates in my time series. I allowed Factiva’s default duplicate checker to remove all duplicate entries from my data before downloading. I also removed any clarification & correction articles from the data. Otherwise, all sources got counted in the final dataset, including any international coverage of the US primary (though Factiva defaults to an English-language search). In the future, I will want to more closely analyze the frequency of various news outlets (domestic vs international, print vs broadcast) to test more fine-grained hypotheses.

⁵I queried the database just once in this search. My search term was “democratic debate”. Again, I then restricted my search to the relevant dates in my time series and I allowed Factiva’s default duplicate checker to remove all duplicate entries from my data before downloading. Because my search term was broader than my candidate attention search, I

assessment of each candidate’s debate performance, I first pulled out sentences relevant to Joe Biden, Bernie Sanders, Elizabeth Warren, or Pete Buttigieg from my total search. For every candidate sentence, I counted positive words using the Lexicoder Sentiment Dictionary (Young and Soroka 2012). I then aggregated positive words and total words in each sentence containing a candidate reference for every day in my time series. I developed a daily indicator of debate coverage for each candidate by dividing total positive words by total words.

Elite endorsements: Endorsements are measured using FiveThirtyEight’s 2020 Democratic primary endorsement tracker. This dataset captures endorsements from several different political actors, including former presidents/vice presidents, national party leaders, governors, U.S. senators and representatives, mayors, and state legislative leaders. FiveThirtyEight also allocates points to each endorsement based on the political visibility of the endorsement. The data was explicitly developed to test hypotheses related to Cohen et al. (2008), as that work is prominently cited in the endorsement data’s methodology statement. In analyses, I retain FiveThirtyEight’s point system and aggregate endorsement points, per candidate, per day.

Voters’ comfort in rating candidates: The SRA model conceptualizes

went through all articles downloaded to delete irrelevant search results. I wanted to err on the side of collecting too many articles rather than too few. I downloaded these data over two days, which created a few issues with duplicate detection. It appears that, due to the method of my search, certain articles counted as duplicates in the first search were not in the second search. However, these issues only affect two days in the time series, and the issue appears to be limited to a relatively small number of articles.

voters' level of information about candidates as: "Do I have enough information to rate the candidate's chances and utility?". The information voters are meant to access, then, is task-related and specific to each person—a voter simply needs to believe what they think about the candidate is sufficient to judge them. The best available proxy for this is what I call my "opinion" time series, which I generate using NationScape rolling cross section data.⁶ To measure voters' comfort level in rating candidates, I rely on the following survey question: "How favorable is your impression of each of the following people, or haven't you heard enough to say?". This question was asked for all four of the Democratic primary candidates in my study. Response options for this question were: "Very favorable", "Somewhat favorable", "Somewhat unfavorable", "Very unfavorable", and "Haven't heard enough to say". Since the question encourages "haven't heard enough" in its frame, it provides the best measure available of the proportion of voters who have evaluated a candidate in some way, and the proportion of those who feel like they do not have sufficient information to rate candidates on a favorability scale. Thus, the proportion of respondents who express a favorability opinion for each candidate can be thought of as a measure of what proportion of voters do not feel comfortable making the most basic utility calculations for each candidate. Because these data are collected using an online opt-in survey, I estimate weights to ensure that movement in this variable is not driven by a particular demographic pattern of response. In all analyses to follow, the opinion time series

⁶Accessible at <https://www.voterstudygroup.org/publication/nationscape-data-set>

represents the proportion of respondents selecting “haven’t heard enough”, divided by the total N of the study, so higher values indicate a larger proportion of respondents unwilling to rate that candidate.⁷

Voter decision sets: I will also measure voter decision sets using NationScape rolling cross section data. In addition to asking about favorability, the survey included a question asking respondents to rank their first, second, and third choice candidates in the 2020 primary election. Though many decision sets likely include more than three candidates, this is the closest we can come to a measure of which candidates are being considered most seriously by voters. To generate my decision set response variable, I created a dummy

⁷NationScape data weighted respondents at the weekly level, rather than the daily level, so I used their weekly weights to impute daily weights for respondents. First, I determined what potential weighting categories were plentiful enough in daily response data to allow for weighting. I was able to weight on age, race, and hispanic ethnicity. The final weighting groups were: above 35, under 35, White (Hispanic and not Hispanic), Black, and Other Race. These categories produce 16 discrete weighting groups and were determined after extensive exploratory analyses designed to ensure I had enough respondents in each group for valid weights. 35 was chosen as my age cutoff because the distribution of ages skewed young in the survey and splitting the groups at this age allowed me to split the variable close to its mean (around 40) while still allowing me enough respondents in each weighting group to determine weights. Appendix A, Figure A.1 reveals the age distribution in the first field week to show an example of age range in the study. Once I determined my weighting groups, I took the sum of the weekly estimated weights for each individual per group, then divided by total respondents in the weekly sample to get weekly benchmarks for each group. For example, I summed under 35 all non-hispanic White respondent weights in a given week in the survey, and divided by the total respondents in that week, to get my weekly benchmark for under 35 non-hispanic White respondents. Then, I determined a daily weight by matching the daily response distribution to that weekly benchmark. To continue my example, if I expected that about 8% of my weekly sample should be under 35 non-hispanic Whites, I would estimate a daily weight to match the proportion of that group in my daily sample to 8%. Some extreme weights were generated through this procedure for under-represented groups, though they were few in number. As such, I left them as-is. In future iterations of the project I will likely set a max value for weights.

variable where 1 = respondent named candidate X as first, second, or third choice, and 0 = respondent did not name candidate.

2.2 Time Series Analysis of Independent and Dependent Variables

Before I explore multivariate relationships between my independent and dependent variables, it is worth examining the time series features of these variables and uncovering what insights we might draw from univariate analyses.

I will start with the polling time series, as polls serve as a benchmark for public opinion trends for each candidate from 7/18/19 to 3/3/20. The midpoint-aggregated polling series for each candidate contains 20 missing days, out of the 230 day series, and missing values are imputed using linear interpolation. The pooled polling time series has no missing days.

Figure 2.1 plots the relative poll support trends for Biden, Sanders, Warren, and Buttigieg, using the midpoint method of aggregation. This figure is replicated, using the pooled time series, in Appendix Figure A.2. Throughout much the primary election season, Biden was a clear leader in the polls. The story of Biden’s campaign, however, was not generally one of persistent dominance, particularly after the Iowa caucuses and New Hampshire primary (Relman 2020; Haltiwanger 2020b). Both polling time series show that Sanders did close the gap in late February 2020, but for most of the election season he polled well under Biden. Elizabeth Warren briefly experienced a polling

surge in late September/early October, and was even considered a frontrunner at the time (Panetta 2020), though this support did not persist and steadily dropped off as the season wore on. As we now know, she would not go on to be a serious competitor in any primary or caucus. Buttigieg experiences what appears to be a real public opinion “bump” towards the end of 2019, though his support never rivals Biden or Sanders.

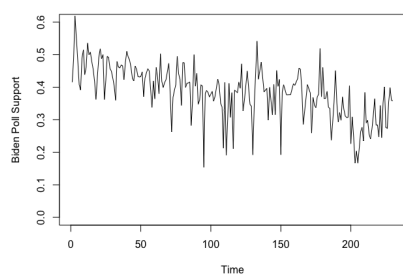
Visual inspection suggests that several of the candidates’ polling time series are trend-stationary. Sanders appears to be the lone candidate who did not experience a period of persistent downward-trending polls during the period, which comports well with the popular narrative that Sanders had a stable, though relatively small, core constituency.⁸

Augmented Dickey-Fuller tests⁹ of the midpoint-aggregated time series reveal that the Biden, Sanders, and Buttigieg series have a significant trend, while the Warren series does not. More specifically, the Biden, Buttigieg, and Sanders series appear to be “trend-stationary”,¹⁰ and the Warren series is likely best represented as an $I(1)$ process (Tables A.1–A.4). All Dickey-Fuller test results, along with ACF and PACF plots for each series, are reported in

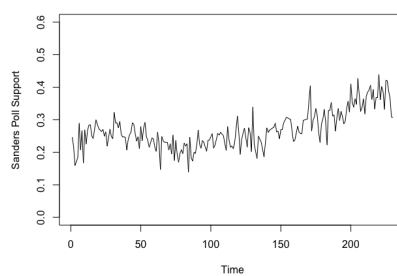
⁸<https://www.washingtonpost.com/politics/2019/12/06/how-sanders-support-compares-his-run/>

⁹Lag length was selected in all ADF tests based on the following criterion: lags were added until the coefficient on the last lag was non-significant. Then, the maximum number of significant lags was included in the test. If no lags were significant, one lag was included. The default lag starting point was 5 lags; if there was no sign that later lags were significant lags were dropped. Some variables were tested with additional lags if the fifth lag was significant in the first model.

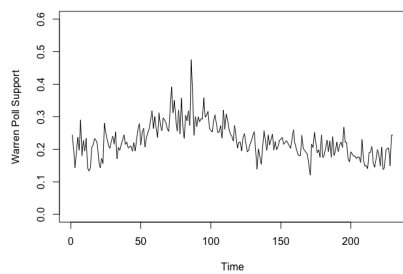
¹⁰Note that the Sanders series just misses the cutoff for the 95% critical value, and as such is treated as stationary in multivariate analyses



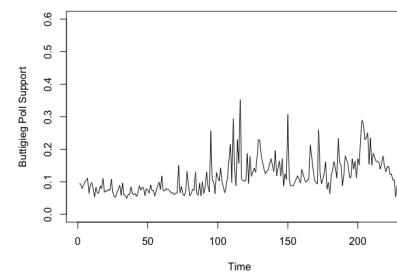
(a) Biden



(b) Sanders



(c) Warren



(d) Buttigieg

Figure 2.1: Midpoint-aggregated poll time series

Appendix A.

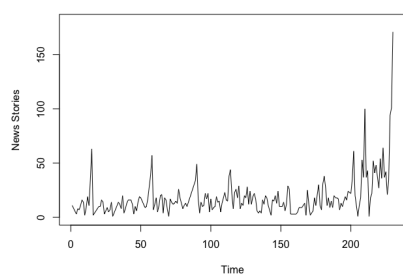
Figure 2.2 plots the daily media attention time series (total number of news stories including a candidate’s name). There are common spikes in attention for each candidate, which correspond to primary debates. The Biden series has 21 missing days, the Sanders series has 22 missing days, the Warren series has 21 missing days, and the Buttigieg series has 30 missing days. Again, all missing values are imputed using linear interpolation.

Graphically, it appears that Sanders slightly edged out Biden in terms of media coverage. The average count of news stories including each candidate’s name confirms this—Sanders and Biden averaged 21.47 and 19.11 news stories per day, respectively. For all candidates, coverage picks up as soon as the first primary contests are held. Likely due to this, all media attention time series fail ADF tests (Tables A.13–A.16) and are thus differenced in the multivariate analyses in the next chapter.

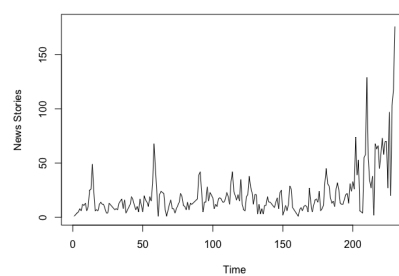
Figure 2.3 plots the time series of debate coverage sentiment for all four candidates. All candidates have a gap in debate coverage during the holidays in my data, which is to be expected.¹¹

The average proportion of debate coverage including positive words was similar for each candidate under study. The Biden series had an average proportion of 0.026 positive words per day, Sanders’ average was 0.027, Warren’s

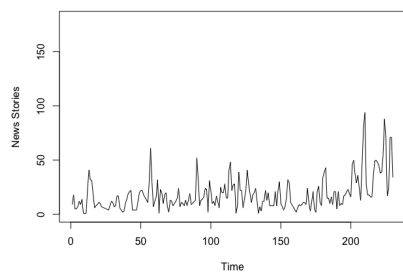
¹¹In total, this time series includes 52 missing days for Biden, Sanders, and Warren, and 59 missing days for Buttigieg, which is by far the largest proportion missing of any time series



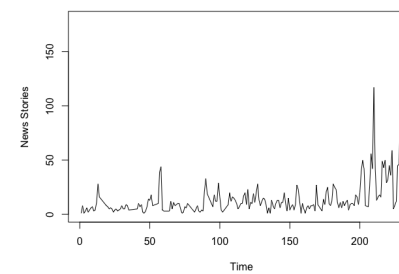
(a) Biden



(b) Sanders

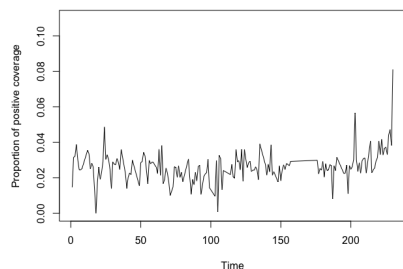


(c) Warren

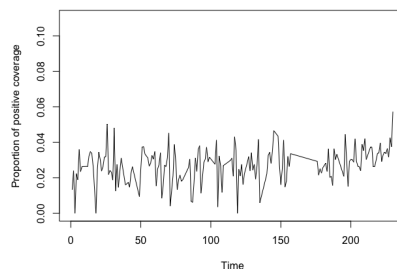


(d) Buttigieg

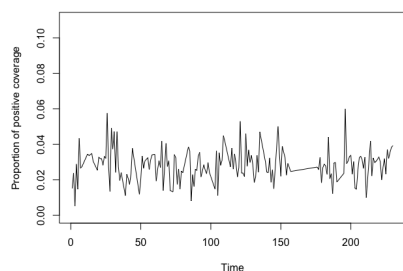
Figure 2.2: Daily count of news stories including candidate's name



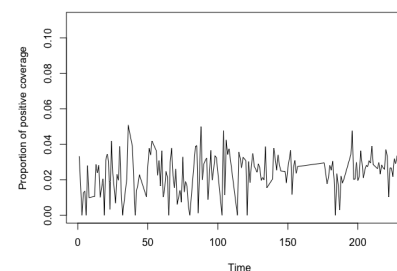
(a) Biden



(b) Sanders



(c) Warren



(d) Buttigieg

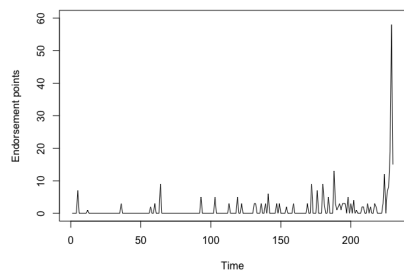
Figure 2.3: Proportion of debate coverage including positive words

average was 0.029, and Buttigieg’s average was 0.024. This comports well with the popular narrative that Warren was a particularly strong debater.¹² These averages also suggest that debate performance might be important, but it is certainly not enough to pull away as a candidate. ADF tests of these series (Tables A.17–A.20) suggest all series should be treated as stationary.¹³

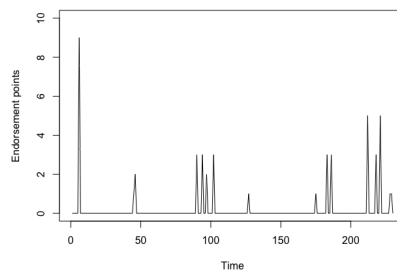
The endorsement series are plotted in figure 2.5. These series, obvi-

¹²<https://www.nytimes.com/2020/02/20/us/politics/elizabeth-warren-debate.html>

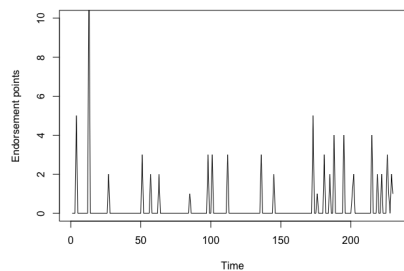
¹³I treat the Biden debate coverage series as stationary, because even with a significant number of lags the ADF critical value is very close to significant



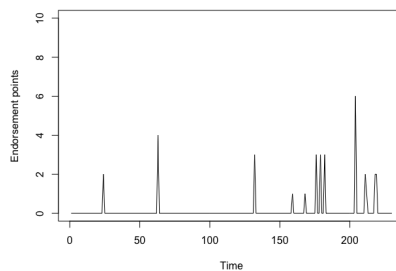
(a) Biden



(b) Sanders



(c) Warren



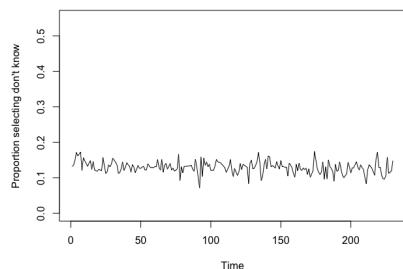
(d) Buttigieg

Figure 2.4: Endorsement Points

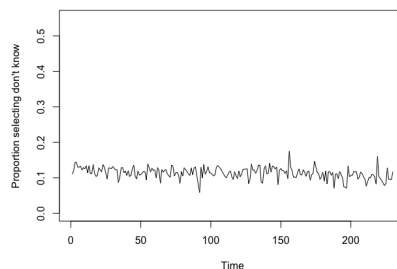
ously, had a large amount of missing days. I treat all missing days as zero endorsements for time series analysis. Biden had accumulated 292 endorsement points by 3/03/2020, Sanders had 46, Warren had 78, and Buttigieg had 33.¹⁴

Next we will turn to my dependent variables. The proportion of respondents who select that they haven't heard enough about a candidate to rate them on a favorability scale follow expected patterns. Biden and Sanders

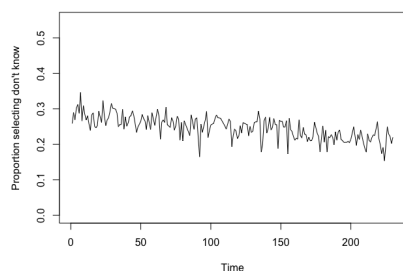
¹⁴ADF Tests can be found in Tables A.21–A.24. Biden is the only candidate who appears to have a nonstationary endorsement series.



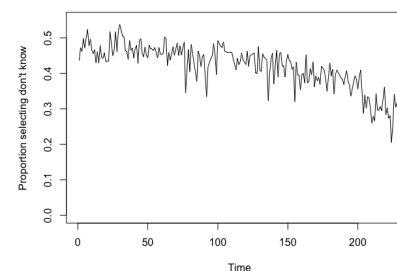
(a) Biden



(b) Sanders



(c) Warren



(d) Buttigieg

Figure 2.5: Proportion of respondents selecting “haven’t heard enough” in the favorability question

had relatively low proportions of “haven’t heard enough” responses with little trend, Warren’s proportion was somewhat higher with a mild downward trend, and Buttigieg’s was the highest with the sharpest downward trend. Each series appears to be stationary in ADF tests (Tables A.5–A.8). In descriptions to follow, I call this time series my “opinion” time series.

The decision set time series show that Biden had a relatively high, stable proportion of respondents selecting him as their first, second, or third choice candidate in the NationScape rolling cross-section. The Sanders series

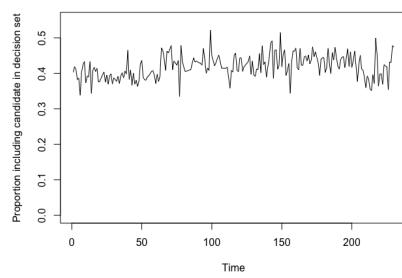
suggest that there was a mild upward trend in the proportion of respondents including him in their decision set, with a possible decay toward the end of the series. Warren had a more substantial bump, though she never quite reached the level of consideration enjoyed by Biden and Sanders, and her support decayed by the end. Buttigieg also had a mild bump in the proportion of respondents including him in their decision set, though that was never quite enough to catch up to the other candidates.

ADF tests (Tables A.9–A.12) suggest it is safe to consider the Biden, Sanders, and Buttigieg decision set series as stationary. The Warren series fails to pass the ADF test, and is thus differenced in the candidate-specific Warren multi-variate analyses.

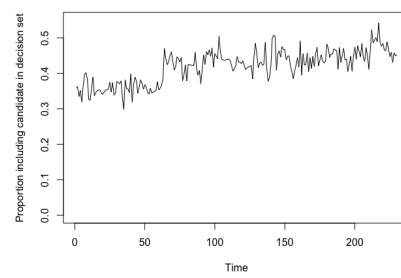
2.3 The story of the 2020 primary, told as time series

Taken together, we can begin to see the story of the 2020 Democratic Primary play out in these time series variables. As aforementioned, Biden enjoyed a persistent polling lead—even when it trended down, his poll support remained above all candidates except for Sanders. The Democratic electorate tended to believe Biden was the most “electable” general election candidate throughout the cycle, even if the 2020 primary electorate did not seem to agree on what “electable” really meant (Brownstein 2019). This likely reinforced his polling lead, which increased perceptions of viability in the primary election as well.

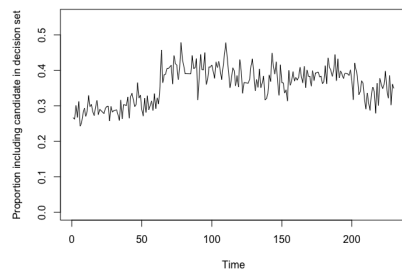
The media attention time series also clearly shows that coverage of all



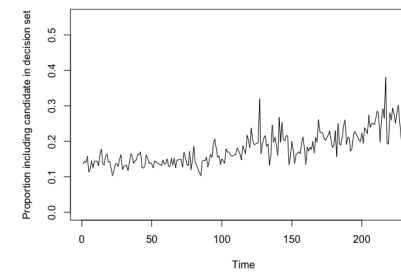
(a) Biden



(b) Sanders



(c) Warren



(d) Buttigieg

Figure 2.6: Proportion of respondents including candidate in decision set

candidates increased over the course of the campaign. Biden and Sanders, the two most credible contenders at the end of the period under study, enjoyed the largest bump in media coverage. Positive debate coverage of Biden also rose towards the end of the time series, which is interesting given that he was not considered a particularly strong debater. It is possible that, if a candidate’s reputation as a debater is already established, simply exceeding those expectations produces enough positive coverage to give a candidate’s campaign a boost.

We already can see a few reasons to doubt “the party decides” hypothesis, at least if measured as endorsements alone. The bulk of endorsement points over the cycle went to Biden, and came in towards the end of the period under study. Multivariate analyses in the next chapter will help assess empirically whether the early endorsements that went to Warren mattered for either her opinion or viability time series. The one notable anecdote from the cycle that suggests strategic endorsements could have mattered was Clyburn’s endorsement of Biden in South Carolina, which appeared to have a real impact on voters in that state (Drezner 2020).

An important contextual feature of the 2020 Democratic primary worth noting is that individual preference ordering among the primary election candidates wasn’t always, or perhaps even primarily, driven by ideological positioning. For instance, polls at the time suggested that Warren voters split fairly evenly among Biden and Sanders when she left the race (Garrison 2020). Candidate traits seemed to motivate many Democratic primary voters, even

though there wasn't much evidence that personality traits mattered to Trump voters in 2016 (Brownstein 2019).

Lastly, a key contextual feature of the 2020 Democratic primary is that many voters did indeed wait until late in the cycle to decide who they'd vote for. This was not a race in which opinions clearly crystalized early on. As late as March 5, 2020, a Business Insider poll suggested that 30% of Democratic voters did not know who they wanted to be the nominee (Haltiwanger 2020a). This suggests that, while some people decide throughout the early primary who they want to support, a solid plurality might wait until they have more information—namely, how early election results shake out.

The next chapter will use the time series variables explored in this chapter to model what predicts voters' level of information about a candidate and the proportion of voters including candidates in their decision set. This should help us uncover more clearly which of these variables matter the most in the dynamics of the primary election season.

Chapter 3: Narrowing the Field

While it's useful to explore the time series properties of primary election variables, uncovering dynamic relationships requires a multivariate model. This chapter takes that focus, and reports a series of pooled and candidate-specific regression models including each variable described in the previous chapter.

As explored in the theory section, step 1 of the SRA model presumes that primary voters narrow down a (typically large) field of candidates by first screening candidates based on whether 1) they have enough information about the candidates to judge their viability and utility, and 2) the candidate meets a viability threshold. In a simulation-based framework, and relying solely on cross-sectional data, Stone, Rapoport, and Atkeson find support for this stage of their model. This chapter builds upon and expands their work by modeling both information and viability indicators early on in the primary season, and by testing how well information and viability can predict voters' decision sets in a dynamic, time series framework.

To recap, literature in this area suggests several cues are useful to study at this stage. Average poll standing, media attention, elite endorsements, and debate performance may provide voters with information they can use to evaluate candidates. As opinions crystallize, we may see that the proportion of respondents willing to rate candidates, media attention, endorsements, and

debate performance influence aggregate perceptions of viability, or overall poll support. And lastly, we should see that willingness to rate candidates and poll support (my operationalizations of information and viability) influences decision sets.

I refrain from making specific predictions as to which variables influence willingness to provide an opinion about candidates, and which influence poll support, because theory suggests that many of these variables likely influence both judgements. In general, we'd expect that endorsements raise the profile of candidates (increasing the proportion of respondents willing to rate candidates) and lead to increased poll support. Positive debate coverage should have the same effects. Media attention should certainly increase willingness to rate candidates, though it's unclear directionally how attention might influence poll support (and the direction of the effect likely depends on the candidate in question). As aforementioned, willingness to rate candidates and poll support should both be positively related to inclusion in respondents' decision sets. I run all analyses using data collected during the 2020 Democratic primary season.

Through the models in the next two chapters, I hope to bring together various theories of how primary election decision-making occurs and unify those theories within a single framework. The analyses in this chapter seeks to integrate Cohen et al.'s work on the "invisible primary", and Bartels's findings on media influence, with the "narrowing down" process proposed by SRA. Using the steps outlined by SRA as my guide, I first assess what drives opinion

formation (willingness to rate candidates) and aggregate viability (poll support) indicators. These models will help me uncover whether endorsements, for example, serve to primarily raise the profile of candidates or whether they increase voters’ perceptions that the endorsed candidates are viable. These findings will be novel contributions to the primary elections literature. Next, I test whether opinion or viability judgments are significant drivers of candidates making it into individual “decision sets”, and whether or not those variables operate differently for different candidates. Operationalization and measurement for all variables in this chapter (except ad spending, detailed below) are described in Chapter 2.

3.1 Analysis Strategy

Models in this chapter estimate time t levels in the opinion time series, the poll support time series, and the decision set time series, using two lags of all relevant predictor variables and two lags of the dependent variable. Many of the time series variables in my data are not clearly exogenous to one another—poll support and media attention, for instance, likely drive each other. For this reason, I take a conservative approach and only test whether certain time series are leading indicators of my time series of interest. Results can thus be interpreted as “causal” in a time-series sense, though future work will be required to more specifically pin down causality in a theoretical sense.

Regression models in this chapter thus take the following form, and are estimated using Ordinary Least Squares:

$$Y_t = \alpha_0 + \alpha_1 Y_{t-1} + \alpha_2 Y_{t-2} + \beta_1 X_{t-1} + \beta_2 X_{t-2} + \epsilon_t$$

There are certainly other possible modeling approaches one could take with these data, such as a more classic auto-regressive distributed lag model or the mathematically-equivalent general error correction model. I report a one-lag ADL approach in Appendix B, Section 3, though I refrain from including those models in the main text of this chapter due to the difficulty inherent in interpreting contemporaneous effects among the variables I test. In this iteration of the project, I do not report a GECM setup, as I do not have an apriori theoretical reason to model change instead of levels (Soroka et al. 2015). These analyses should thus be considered a first step—as the project progresses, and I am able to gather more data, I may find that an alternative specification is more useful than the two-lag approach.

In the models to follow, I also include a variable called “Ad Spending”, which is calculated using weekly ad buy data for Biden, Sanders, Warren, and Buttigieg. The ad buy data covers cable, broadcast, radio, and digital ads. The time series properties of ad buy data are not detailed in Chapter 2 because these data are only observable at the weekly level. Because of this, I divide each candidate’s total ad buy for a given week by 7, and then impute a value for each day of the week using a linear, additive approach. So, for example, if a candidate purchased \$7 worth of ads in a week, they’d be coded as Day 1 = \$1, Day 2 = \$2, Day 3 = \$3, and so on. Because the ad purchases are often

fairly large sums of money, I divide the daily imputed values by 100,000 (all coefficients are thus the effect of spending an additional \$100,000 on ads).

This approach assumes that the effect of ads is cumulative throughout a given week and “resets” in the next week. Given that ad effects are fairly short-lived (Gerber et al. 2011), this approach seemed more reasonable than continuing to sum ad spending over the course of the campaign. It also seemed possible that a concentrated ad campaign could have cumulative effects over the course of several days, which is why I sum spending over the course of the week rather than simply divide total spending by 7. However, I acknowledge that this transformation still requires fairly strong assumptions, and that further work will be required to more rigorously test if it is appropriate in this campaign context. As such, ad spend coefficients should be interpreted as suggestive.

As a final analytical note, I run all analyses using my time series with interpolated missing values. Many time series variables had relatively few missing days, and as such I have reasonable confidence that my results won’t be driven by this interpolation. However, in future iterations of this work, I plan to run all models without interpolated days as a robustness check. I also err on the side of treating a few variables as stationary that are on the border between stationary and non-stationary. An additional robustness check will be to test alternative specifications of those variables in future work.

3.2 Results

3.2.1 Pooled Models

First, I explore potential relationships between all of my time series variables using a stacked dataset including all candidates. Setting up such a model is a bit of a challenge, as the endorsement time series is nonstationary for Biden only, and the poll support and decision set variables are nonstationary for Warren only. Given that media attention appears to be nonstationary for all candidates, I include two lags of differenced media attention in the pooled model. All other variables enter into the model as lagged levels. I also estimate all models with a trend term, candidate fixed effects, and cluster-robust standard errors.

I tested each of the pooled models with three lags, to test whether results became stronger when additional days were added. In general, results looked similar, with somewhat greater levels of significance for the significant predictors in the two-lag model. Because the interpretation of the results remains very similar in the two-lag and three-lag setup, I opt for the simpler two-lag model here.

Overall, the pooled results support the theoretical importance of the opinion time series, the media attention time series, the poll support time series, the debate coverage time series, and the spending series in early primary election dynamics.

First, in the opinion model (Table 3.1, Model 1), we see that the second lag on media attention significantly predicts opinion at time t . Recall that

Table 3.1: Candidate Pooled Models

	<i>Dependent variable:</i>		
	Opinion Series _t	Poll Support/Viability Series _t	Decision Set Series _t
	(1)	(2)	(3)
Decision Set Lag $t - 1$			0.359*** (0.039)
Decision Set Lag $t - 2$			0.231*** (0.059)
Opinion Lag $t - 1$	0.409*** (0.041)	-0.061 (0.052)	0.002 (0.029)
Opinion Lag $t - 2$	0.225*** (0.050)	-0.102 (0.077)	0.007 (0.088)
Δ Media Attention Lag $t - 1$	-0.0001 (0.00005)	-0.0001 (0.0001)	
Δ Media Attention Lag $t - 2$	-0.0002* (0.0001)	0.0001 (0.00003)	
Poll Lag $t - 1$	-0.009 (0.013)	0.386*** (0.064)	0.059*** (0.033)
Poll Lag $t - 2$	-0.029 (0.015)	0.283*** (0.046)	0.030 (0.028)
Debate Performance Lag $t - 1$	-0.102 (0.069)	0.245*** (0.039)	
Debate Performance Lag $t - 2$	0.082 (0.057)	-0.040 (0.142)	
Endorsements Lag $t - 1$	0.001*** (0.0001)	0.0001 (0.001)	
Endorsements Lag $t - 2$	0.001 (0.001)	-0.003* (0.001)	
Ad Spending	0.0004 (0.0002)	0.001*** (0.0004)	
Buttigieg Fixed Effect	0.095** (0.030)	-0.045* (0.021)	-0.077*** (0.011)
Sanders Fixed Effect	-0.011** (0.003)	-0.048*** (0.011)	0.011* (0.005)
Warren Fixed Effect	0.037** (0.013)	-0.035*** (0.007)	-0.011 (0.006)
Trend	-0.0001* (0.0001)	-0.0001 (0.0001)	0.0002*** (0.00002)
Constant	0.076*** (0.018)	0.155** (0.052)	0.118** (0.044)

Note:

*p<0.05; **p<0.01; ***p<0.001

the media attention variable is simply the volume of news articles each day that included a candidate's name, and that this variable is differenced. The interpretation of the coefficient is thus that a one story increase from the day prior significantly predicts a decrease in the proportion of NationScape respondents unwilling to rate a candidate two days later. This finding aligns with expectations—as media attention towards a given candidate increases, it is sensible that more respondents form opinions about that candidate.

The endorsement effect is a bit more puzzling, as its effect runs in the opposite direction we'd expect. This finding is likely driven by Biden's endorsement data (as we will see in the candidate-specific results). Biden's endorsement series is highly heteroskedastic, and time series diagnostics suggest that the variable should be differenced for Biden. This first-lag effect disappears if Biden's endorsement series is differenced, and all other candidate endorsement series are left as levels, suggesting the effect is indeed driven by the particularities of the Biden series. As such, I hesitate to over-interpret what the sign and significance of this coefficient could mean.

Candidate fixed effects are significant and reinforce what we'd expect; that Biden and Sanders started out with a higher proportion of respondents willing to rate them on a favorability scale than Warren and Buttigieg. The trend term also is significant and in the expected direction.

Model 2 reveals that media attention (the second lag is significant at $p < 0.1$), debate performance, and spending are significant predictors of poll support/viability at time t . When all candidate data are pooled into a single

model, increased media attention (relative to the prior day) influences poll support two days later. The finding that the second lag of differenced media attention is either a significant or marginally significant predictor in both the opinion and poll support models is interesting, and suggests media effects might take a bit longer to affect primary campaign dynamics than the other variables tested here. For instance, only the first lag on debate coverage is significant, and spending enters into the model at time t , suggesting positive shocks in those two series are reflected in poll support sooner than a positive shock in media coverage is.

The debate coverage coefficient is substantively quite large, though it is sensible when considering the scale of the debate coverage variable. Because the debate coverage series represents the proportion of positive words, relative to all words, in sentences about candidates' debate performances, that coefficient represents the largest possible effect of debate coverage (moving from 0% to 100% positive). The actual range of positive debate coverage in the data is 0 to 0.08, suggesting that the largest observable effect in these data could be about a 0.02 (or 2%) increase in poll support, relative to the other candidates. While it is interesting to find a significant effect of campaign events at all, in practice we'd expect the effect to be rather modest.

Model 2 also suggests that spending has a significant and positive effect on poll support, even when controlling for a time trend. However, the substantive significance of this effect is again rather modest—the model estimates that spending an additional \$100,000 at time t is related to about a 0.001 (or

0.1%) increase in poll support at the same time. The relationship is notable, but requires further data and analysis before firm conclusions can be drawn. Do candidates start to spend more, for instance, as their poll numbers rise (perhaps due to increased donations)? Or is it the reverse—does their spending correlate with tangible campaign gains?

Lastly, we again see a rather puzzling endorsement effect, which completely disappears if the Biden data are removed from the model. The first lag on endorsements becomes significant and positive if the Biden endorsement series is included and differenced, which changes the way we'd interpret this result. Reassuringly, the debate coefficient is significant and positive in both of these alternative specifications. I take the endorsement coefficient sensitivity to suggest that more work is needed to uncover the effect of endorsements in primary elections. Additional data and analyses, in addition to thinking through the conceptualization of the FiveThirtyEight endorsement point time series, will be necessary to determine the role endorsements play in poll support.

The pooled decision set model (Model 3) suggests that poll support is most closely related to including a candidate in one's decision set. And, again, it is worth considering the variable's range when assessing these results, as poll support and the decision set time series are both measured as proportions. So, the largest possible effect that poll support could have is about a 6% increase in the proportion of respondents including a candidate in their decision set. Because the range of the poll support variable in the pooled data is 0.05 to

0.62, the more practical estimate is about a 3% maximum jump. Again, given that this is measured at the daily level, it is reasonable to expect that the effect of a single day poll increase wouldn't be substantively huge, and the finding is still statistically meaningful.

3.2.2 Candidate-Specific Models

Candidate-specific models present a more mixed pattern of results. In all models to follow, the Biden endorsement series, the Warren poll series, and the Warren decision set series are all differenced, per the time series diagnostics explored in Chapter 2.

To begin, we will explore the opinion formation models, detailed in Table 3.5. Not much is related to information levels about Joe Biden (Model 1), which is unsurprising given that the vast majority of voters had opinions about him going into the primary election season already. Just as in the pooled model, the Biden endorsement coefficient is significant and has an unexpected sign. Again, I hesitate to over-interpret this finding for a few reasons—first, the endorsement time series for Biden, even when differenced, is fairly heteroskedastic over the course of the campaign. This occurs because endorsements mainly come in towards the end of the time series. Second, the estimated effect is extremely small—one endorsement “point” in the FiftyEight time series (which, for Biden, ranges from 0 to 58 in levels) is estimated to increase the proportion of respondents saying they “haven’t heard enough” about Biden by 0.0009. Additional modeling to account for the quirks in

Table 3.2: Opinion/Willingness-to-rate Results

	<i>Dependent variable:</i>			
	Biden (1)	Sanders (2)	Warren (3)	Buttigieg (4)
Opinion Lag $t - 1$	0.170* (0.068)	0.147* (0.067)	0.168* (0.069)	0.352*** (0.066)
Opinion Lag $t - 2$	0.004 (0.069)	-0.056 (0.068)	-0.049 (0.069)	0.180** (0.067)
Δ Media Attention Lag $t - 1$	-0.0001 (0.00009)	-0.0001 (0.00007)	-0.00001 (0.0001)	-0.00002 (0.0002)
Δ Media Attention Lag $t - 2$	-0.0001 (0.00009)	-0.0001 (0.00007)	-0.0001 (0.0001)	-0.0006*** (0.0002)
Poll Support Lag $t - 1$	0.010 (0.019)	0.015 (0.030)	-0.057 (0.042)	0.043 (0.052)
Poll Support Lag $t - 2$	0.016 (0.019)	-0.055 (0.029)	-0.001 (0.042)	-0.012 (0.051)
Debate Performance Lag $t - 1$	-0.067 (0.170)	0.081 (0.124)	-0.132 (0.191)	-0.100 (0.221)
Debate Performance Lag $t - 2$	0.028 (0.167)	0.031 (0.124)	0.031 (0.191)	0.244 (0.219)
Endorsements Lag $t - 1$	0.0002 (0.0003)	0.001 (0.001)	0.0004 (0.001)	-0.003 (0.003)
Endorsements Lag $t - 2$	0.0009* (0.0004)	-0.002 (0.001)	0.0003 (0.001)	-0.005 (0.003)
Ad Spending	-0.00001 (0.0004)	-0.000003 (0.0002)	-0.00004 (0.0006)	0.0004 (0.0007)
Trend	-0.0002 (0.00002)	-0.00005* (0.00003)	-0.0003*** (0.00004)	-0.0004*** (0.00008)
Constant	0.101*** (0.017)	0.117*** (0.013)	0.254*** (0.026)	0.231*** (0.036)
Observations	227	227	227	227
R ²	0.107	0.160	0.483	0.726
Adjusted R ²	0.057	0.113	0.453	0.711

Note:

*p<0.05; **p<0.01; ***p<0.001

Biden’s endorsement series, plus additional data, are necessary to dive a bit deeper into why this result shows up as significant in the Biden opinion model. However, there are reasonable data-driven reasons to expect that the true story isn’t that endorsements actually increases the proportion of respondents without a favorability opinion of Biden. No other predictor variables appear to be significant drivers of the Biden opinion series.

The Sanders and Warren opinion models (Model 2 & 3) are also mostly null results. The trend term suggests the series did slightly trend down over time for both candidates, but nothing else shows up as significant. However, we do see a significant effect of the second lag of differenced media attention in the Buttigieg model (Model 4), suggesting that the Buttigieg data was a likely driver of the pooled results.

Viability models (Table 3.3) suggest that candidate data need to be pooled to uncover a debate coverage effect—when the data are separated out, debate coverage is not significant in any candidate model. Ad spending, however, is significant and positive for Sanders (Model 2) and Warren (Model 3). It is notably strong in the Warren model, suggesting that day-over-day changes in her poll series were closely related to her campaign spending.

Decision set models (Table 3.4) reveal one reason why the opinion series was not significant in the pooled models—the second lag of the series has significant, opposite-signed effects in the Warren model (Model 3) and Buttigieg model (Model 4). This makes sense when we consider what occurred in both of these time series over the course of the campaign. Warren’s decision set

Table 3.3: Viability/Poll Support Results

	<i>Dependent variable:</i>			
	Biden	Sanders	Warren, D1	Buttigieg
	(1)	(2)	(3)	(4)
Opinion Lag $t - 1$	-0.124 (0.246)	0.168 (0.152)	-0.108 (0.110)	-0.025 (0.087)
Opinion Lag $t - 2$	-0.455 (0.248)	-0.229 (0.152)	-0.145 (0.111)	0.100 (0.082)
Δ Media Attention Lag $t - 1$	0.0004 (0.0003)	-0.0001 (0.0002)	-0.00003 (0.0002)	0.0001 (0.0003)
Δ Media Attention Lag $t - 2$	-0.00007 (0.0003)	0.00006 (0.0002)	-0.00005 (0.0002)	0.00006 (0.0003)
Poll Support Lag $t - 1$	-0.763*** (0.068)	-0.626*** (0.067)	-0.521*** (0.067)	-0.781*** (0.069)
Poll Support Lag $t - 2$	0.085 (0.068)	0.247*** (0.066)	-0.212** (0.067)	0.263*** (0.068)
Debate Performance Lag $t - 1$	0.550 (0.614)	0.350 (0.279)	0.281 (0.307)	0.170 (0.291)
Debate Performance Lag $t - 2$	0.089 (0.601)	0.176 (0.279)	-0.039 (0.306)	-0.321 (0.287)
Endorsements Lag $t - 1$	0.001 (0.001)	-0.0004 (0.003)	0.001 (0.002)	-0.001 (0.004)
Endorsements Lag $t - 2$	-0.0004 (0.001)	0.003 (0.003)	0.0009 (0.002)	-0.003 (0.004)
Ad Spending	-0.0002 (0.002)	0.0007* (0.0003)	0.002* (0.0001)	0.001 (0.001)
Trend	-0.0005*** (0.0001)	0.0001 (0.00006)	-0.0002* (0.00007)	0.0002* (0.0001)
Constant	0.382*** (0.050)	0.078** (0.029)	0.068 (0.042)	-0.002 (0.048)
Observations	227	227	227	227
R ²	0.389	0.310	0.257	0.397
Adjusted R ²	0.355	0.272	0.216	0.363

Note:

*p<0.05; **p<0.01; ***p<0.001

Table 3.4: Decision Set Results

	<i>Dependent variable:</i>			
	Biden (1)	Sanders (2)	Warren, D1 (3)	Buttigieg (4)
Decision Set Lag $t - 1$	0.243*** (0.068)	0.398*** (0.067)	-0.628*** (0.066)	0.247*** (0.069)
Decision Set Lag $t - 2$	0.125 (0.069)	0.129 (0.067)	-0.279*** (0.064)	0.061 (0.070)
Opinion Lag $t - 1$	0.043 (0.129)	-0.006 (0.115)	-0.029 (0.084)	0.018 (0.059)
Opinion Lag $t - 2$	0.167 (0.127)	0.044 (0.114)	0.244* (0.084)	-0.139* (0.058)
Poll Support Lag $t - 1$	0.043 (0.034)	-0.012 (0.049)	0.003 (0.057)	0.046 (0.043)
Poll Support Lag $t - 2$	-0.011 (0.034)	-0.083 (0.049)	-0.040 (0.056)	0.094* (0.044)
Trend	0.0001** (0.000005)	0.0003*** (0.000006)	0.00007 (0.056)	0.0002*** (0.000006)
Constant	0.209*** (0.044)	0.184*** (0.030)	-0.075* (0.019)	0.037*** (0.036)
Observations	228	228	227	228
R ²	0.186	0.686	0.314	0.652
Adjusted R ²	0.160	0.676	0.295	0.641

Note:

*p<0.05; **p<0.01; ***p<0.001

time series went down over the course of the campaign, even as information about her would have increased. Buttigieg’s decision set series never reaches the same height as Warren’s, but it does not noticeably decline, either. Taken together, Models 3 and 4 suggest that a greater proportion of potential voters with opinions about a candidate does not necessarily translate into inclusion in a greater proportion of decision sets. They also suggest that an alternate specification may be more appropriate as the project progresses forward. For instance, the proportion of respondents with an opinion about a candidate might need to be interacted with poll support for a more accurate picture of decision set dynamics.

3.3 Discussion

Overall, the findings reported in this chapter present some interesting initial results. Media attention appears to be a particularly important variable to focus on in future iterations of the project, as do spending and poll support. Findings also suggest that debate performance might be an underrated campaign variable that helps raise the profile of candidates and increases their poll support. If, theoretically, we feel comfortable equating aggregate polls with an aggregate indicator of viability, the substantive takeaway would be that performing well in a debate increases voter perceptions that a candidate is a “serious” contender in a primary.

The stronger pattern of results in the pooled models suggests that the candidate-specific models may be under-powered, which is another important

consideration as the project progresses. Gathering data from additional candidates, in additional elections, will be critical to check the robustness of the patterns presented here. Doing so will also allow me to test whether certain effects only show up for lesser-known candidates. The results in this chapter provide evidence that there was a stronger pattern of significant results for Buttigieg, who was the only such candidate who stayed in the race long enough for me to compare against the others. However, a clear comparison point to Buttigieg in this race would be Kamala Harris, who shot up in name recognition throughout the early primary. It is worth exploring the possibility that Step 1 of the SRA model applies most strongly to candidates who are relative unknowns at the start of the race.

I report the results of an ADL specification with contemporaneous and lagged effects in Appendix B. Those models suggest that there are contemporaneous relationships between my time series that are not captured in my lagged models. Due to the difficulty in interpreting what a contemporaneous effect means in this case, when many variables tested here plausibly cause changes in each other, I do not include those results in the main text. However, to preview those results, there is a contemporaneous effect of differenced media attention on the Biden and Warren poll support time series, and a contemporaneous effect of positive debate coverage on the Buttigieg opinion time series. The decision set ADL models also reveal likely contemporaneous relationships between the opinion time series and decision sets for Biden and Warren.

Another clear future direction for this research is to dig into possi-

ble effects (or the lack thereof) of endorsements. The heteroskedasticity of these series will clearly need to be accounted for, and the theory underlying the FiveThirtyEight measure will need to be re-examined. Assigning more endorsement “points” for higher office holders is sensible, but the exact multiplier FiveThirtyEight uses may not be quite right. These models don’t necessarily present strong evidence that endorsements are critical variables for voters as they narrow the field, but they also do not entirely rule out the possibility that endorsements matter. In any case, the visualizations of the endorsement series in Chapter 2 suggest endorsements start piling in when candidates are already doing well (such was the case for Biden). Endorsements might influence voters, but the majority of officeholders also appear to wait and see where voters are leaning as well. Before we throw out “the party decides” hypothesis, however, it is worth noting the relatively strong effects of media attention. To the extent that we’d consider the media an “elite”, and to the extent we believe a party can help shape media coverage, there may be evidence that elite actors do indeed shape the course of primaries.

Finally, it is worth noting that I leave out potential effects of other candidates’ variables in the candidate-specific models. Might it be the case, for instance, that Biden or Sanders’ decision set proportions benefitted from a decrease in Warren’s? This is another potential set of questions I hope to address in future iterations of the project.

Chapter 4: Evaluating Candidates

In my final empirical chapter, I turn to Step 2 of the SRA model. This chapter assesses how voters decide between candidates, using an updated expected utility model. In the analyses to follow, I seek to revive policy issues as an important predictor of primary vote choice and underscore the role issues can play in primary elections. Findings in this chapter contribute to the broader literature on the topic by first 1) replicating the predictive validity of an expected utility model in primary vote choice, 2) refining an issue-based measure of utility and using that measure to model choice directly, 3) demonstrating that changes in perceptions of issue emphasis are directly associated with changing vote intentions, using panel survey data. I have two main survey datasets I will use to test my application of the expected utility model: Fox News’s 2020 Primary Voter Analysis Survey, and a panel survey of university students. Again, these data restrict analyses to the 2020 Democratic primary.

4.1 Data and Measurement

The 2020 Primary Voter Analysis Survey was fielded in Iowa, New Hampshire, South Carolina, Alabama, Colorado, Minnesota, Missouri, Arizona, California, Massachusetts, North Carolina, Texas, Virginia, Michigan, Ohio, Florida, and Illinois. Interviews in each state began six days before the

primary election or caucus and ran until polls closed.¹ Importantly, the survey includes questions asking about the importance of candidate traits, who could best handle certain policy issues, and a probability-type electability estimate.

The student panel survey was fielded on 1/29/20 (before the Iowa caucus), 2/12/20 (after the New Hampshire Primary), and 3/4/20 (after Super Tuesday) in an online University of Texas introductory political science course. Because of the study's panel design, the data will be used to help strengthen the causal claims made in the chapter regarding the expected utility model. The survey dataset is also novel in several other ways, as I collect relatively fine-grained probability-type electability perceptions of the kind theorized to be important in an expected utility framework. Taken together, this survey is the first attempt (that I'm aware of) to test the expected utility model of vote choice using a panel design, electability probability estimates, and issue emphasis. The study was explicitly designed to address limitations in prior work and to test and validate new measures of public opinion that we can use to understand primary election voting behavior.

Following the theory outlined in section 1.1.2, my primary hypotheses in this chapter are as follows:

H1a: Issue emphasis X electability (the expected utility model, with an issue focus) will better predict candidate choice than electability alone.

¹Samples in all states included probability samples of 1,750-2,000 voters and 500-900 non-voters. Certain states were supplemented with a non-probability online sample of 600-2,000 self-identified registered voters.

H1b: Issue emphasis will be a significant predictor of vote choice, even when controlling for the perceived importance of candidate traits.

H2a: Changes in perceptions of issue emphasis cause changes in vote intention.

H2b: Changes in perceptions of candidate electability in a general election cause changes in vote intention.

Below are the variables I propose to use to test H1 in the Fox News Voter Analysis Data.

- Measurement of Independent Variables
 - **Candidate electability:** Respondents in all states were asked: “Thinking about the general election in November, do you think each of the following candidates definitely could, probably could, probably could NOT or definitely could NOT beat Donald Trump?” Respondents in Missouri, Michigan, and Mississippi were shown a list including Biden, Sanders, Warren, and Bloomberg. Respondents in Arizona, Florida, Illinois, and Ohio were only asked about Biden’s and Sanders’s electability.
 - **Candidate utility (issues):** Respondents in a subset of states were asked: “Regardless of who you support in the primary, which of the following candidates do you think would be best able to handle: (issue)”. The issues asked about varied by state, as did the

list of candidates offered to respondents. The total set of issues a respondent could be asked about included: the economy, foreign policy, issues related to race, health care, climate change, immigration, gun policy, international trade, and corruption in government. Though this question does not quite tap issue emphasis on the part of candidates, they are modeled after the issue ownership literature, which is conceptually quite similar. Presumably, candidates who have made the clearest connection between their campaign and a particular issue would be considered best able to handle that issue (and, candidate effectiveness on issues as been referenced in studies of issue emphasis in primaries before—see Aldrich and Alvarez (1994)). Though issue ownership developed in the general elections literature, I believe that there is a great deal of crossover between issue ownership in the general election and issue emphasis in the primaries, and will explore this idea further in a concluding section.

- **Candidate utility (traits):** Respondents in all states were asked: “How important is each of the following qualities in the Democratic nominee for president? (will work across party lines, has the best policy ideas, cares about people like you, is a strong leader, has the right experience)”. Unfortunately, we do not get a candidate-specific measure for this item. However, it can still be included in a general expected utility model, as will be explored below.

- Measurement of Dependent Variable
 - **Vote Intention:** Respondents in all states received the question “Who do you plan to vote for in the Democratic caucus/primary election for President?”. I will argue that this variable should incorporate both sincere (utility) and strategic (electability) considerations, since the question specifically asks about vote choice and *not* simply first choice preference. In all analyses to follow using Voter Analysis data, I only model Biden and Sanders vote choice. I do this in order to include as many states and respondents in my analyses as possible, and some respondents were surveyed in states late enough on the primary election schedule that these two candidates were the only two serious contenders left in the race.

Though the expected utility model implies that voters multiply electability and utility estimates, authors modeling vote choice within this framework do not always include an interaction term (Stone et al. 1992b). They find that a multiplicative expected utility term is indeed significant in models of choice, but believe that it imposes a precise metric employed by voters that is not necessary to test the model. I follow suit and estimate a model without an interaction term. However, I also closely follow the analyses presented in Table 2 of Stone et al. (1992b), and test whether a multiplicative electability X issue summary metric better predicts stated vote choice than an electability metric alone. This required coming up with a single utility “score” for

Biden and for Sanders using the Voter Analysis data. To do this, I created a summary issue-utility score for Biden and Sanders (since those are my only candidate-specific utility measures), and multiply those scores by the respective electability likert-type question for each candidate.

My student data will be used to test H2. Analyses incorporate the following variables from my student survey:

- Measurement of Independent Variables
 - **Candidate electability:** Respondents in the student panel survey were asked: “Below are the names of some Democrats who are running for president in 2020. Regardless of who you intend to support in the 2020 Election, how strong do you think each of the following candidates would be in a race against Donald Trump? For each of the following potential candidates, please rate them on a scale from zero to 10, where zero means they would definitely lose to Trump and 10 means they would definitely beat Trump. 5 means you think they would have about a 50/50 chance of beating Trump. If you have not heard of the person, you do not need to rate them.” The list of candidates included Bernie Sanders, Elizabeth Warren, Michael Bloomberg, Pete Buttigieg, Andrew Yang, Amy Klobuchar, Joe Biden, Tulsi Gabbard and Tom Steyer. This question was intentionally designed to try and measure the type

of electability calculation presumed to be in voters' minds in the expected utility model.

- **Candidate utility (issues):** Respondents were asked: “Which of the following 2020 Democratic presidential candidates is paying the most attention to the issue of... (climate change, gun policy, health care, wealth and income inequality, and immigration)”. This question was designed to match the issue emphasis measure employed in Rickershauser and Aldrich (2007).
- Conceptualization and Measurement of Dependent Variable, Student Panel Data
 - **Vote Intention:** Respondents were asked: “In the 2020 Democratic primary/caucus for president, who will you vote for? Your best guess is fine”. This, again, is meant to tap vote choice, not simply preference.²

²In wave 3, I edited the list of candidates respondents saw to reflect candidates who had dropped out of the race, and I edited the vote choice question to read “who *did* you vote for”. These changes were made a few hours after survey launch, which means about 30/599 students saw a slightly different candidate list and vote choice question than the rest of the respondents. However, I do not believe this change meaningfully influences the relationships described in this chapter.

4.2 Results

4.2.1 Evaluating Hypothesis 1

In order to test H1a, I use the Voter Analysis data to assess whether an expected utility model, including an issue-based measure of utility, better predicts vote choice than electability measures alone. I restrict my analyses to Florida, Illinois, Missouri, Mississippi, Michigan, and Arizona respondents, because those were the only states in which issue questions were asked. I do not pool respondents into one general model, because the issue questions varied from state to state. Since data availability limits me to calculating issue utility scores and electability scores solely for Joe Biden and Bernie Sanders, the data used in the analyses behind Tables 4.1 and 4.2 include only those respondents who indicated either a Biden or Sanders vote choice.

Table 4.1: Descriptive Statistics, Biden and Sanders Expected Utility Metrics

State	N Biden Votes	N Sanders Votes	JB Mean Electability	BS Mean Electability	JB Mean Issue Score	BS Mean Issue Score
Florida	2245	703	3.41	2.65	0.64	0.26
Illinois	1696	825	3.29	2.75	0.54	0.33
Missouri	1161	599	3.25	2.67	0.45	0.26
Mississippi	793	183	3.53	2.76	0.62	0.18
Michigan	1378	823	3.26	2.83	0.42	0.28
Arizona	990	573	3.33	2.75	0.52	0.39

To create a state-specific respondent issue utility score, I summed the total number of issues each respondent thought Biden could handle best, and the total number of issues each respondent thought Sanders could handle best, and then divided the Biden and Sanders sums by the total number of policy issues asked about in that state. Doing so provides me with a candidate-specific score that I can multiply by each respondent's electability score for each candidate.

Similar to Stone et al. (1992b), I first test the proportion of respondent vote choices correctly predicted using *solely* an electability measure, and compare that to the proportion of respondent vote choices correctly predicted using an electability X issue utility score (incorrect predictions are counted if the score suggests a vote choice for Biden, and R voted for Sanders, and vice versa OR if the score predicted a tie). The results of these analyses are presented in Table 4.2.

Results of these analyses strongly suggest that an expected utility prediction, based on perceived policy issue competence, outperforms an electability-only prediction. In this study, I use a chi-squared test to assess whether the proportion of correct predictions is meaningfully different in each model (note that this varies from Stone, Rapoport, and Abramowitz, who report a proportional reduction in error statistic). These analyses support H1a and suggest that policy issues can be an important utility consideration when measured as policy emphasis or competence, rather than policy position or ideological closeness. Note that we can successfully predict almost 83% of primary votes among Florida, Illinois, and Arizona respondents using just these two measures.

Tables 4.3 and 4.4 present another test of H1a, as well as a test of H1b, and attempt a more complete specification of an expected utility model. In Table 4.3, the dependent variable is coded as 1 = Expressed intention to vote for Biden, and 0 = Expressed intention to vote for someone else, and in Table 4.4, the dependent variable is coded as 1 = Expressed intention to vote for

Table 4.2: Test of H1a

State	Electability Only	Expected Utility	Chi-squared	p
Florida	0.669	0.825	189.05	<0.001
Illinois	0.624	0.830	268.32	<0.001
Missouri	0.672	0.736	16.504	<0.001
Mississippi	0.593	0.722	35.56	<0.001
Michigan	0.628	0.703	27.442	<0.001
Arizona	0.644	0.827	133.65	<0.001

Sanders, and 0 = Expressed intention to vote for someone else. These models *do not* restrict the sample to Biden and Sanders voters only and are estimated using logistic regression. All candidate trait questions were non-specific, so coefficients on these variables should be interpreted as the general effect of thinking it’s important to “work across party lines” (etc) on preference for Biden or Sanders. Though this is not quite the best expected utility measure (it would be better if I had data assessing whether each respondent thought Biden could work across party lines, for example), it is the best measure I have given the available data.

The Biden expected utility models (Table 4.3) reveal clear patterns that persist across effectively all states in the survey. Biden voters tended to not think it was important for a candidate to have the best policy ideas, but in Arizona, Florida, Michigan, and Missouri they were more likely to believe it was important to work across party lines. In those same states, Biden voters were more likely to say they thought it was important that the nominee “has the right experience”. Every single coefficient on the issue-based questions is significant in the Biden models, reinforcing that perceived issue competence

Table 4.3: Biden Expected Utility Models

	<i>Dependent variable:</i>					
	Vote Choice = Biden					
	(1) Arizona	(2) Florida	(3) Illinois	(4) Michigan	(5) Missouri	(6) Mississippi
Can beat Donald Trump	0.104 (0.169)	-0.086 (0.088)	0.035 (0.110)	0.081 (0.124)	0.242 (0.155)	0.132 (0.180)
Will work across party lines	0.446*** (0.095)	0.278*** (0.079)	0.157 (0.094)	0.368*** (0.085)	0.438*** (0.106)	0.020 (0.151)
Has the best policy ideas	-0.628*** (0.115)	-0.309** (0.101)	-0.623*** (0.120)	-0.595*** (0.110)	-0.950*** (0.137)	-0.130 (0.205)
Cares about people like me	-0.204 (0.113)	-0.288** (0.111)	-0.160 (0.115)	-0.363*** (0.102)	-0.462*** (0.129)	-0.103 (0.241)
Strong leader	0.069 (0.164)	0.211 (0.138)	0.193 (0.165)	0.102 (0.136)	0.112 (0.182)	0.508 (0.319)
Has the right experience	0.467*** (0.122)	0.294** (0.111)	0.130 (0.126)	0.224* (0.111)	0.395** (0.140)	0.226 (0.211)
Likelihood Biden could beat Trump	0.480*** (0.088)	0.459*** (0.065)	0.265** (0.090)	0.841*** (0.084)	0.895*** (0.107)	0.666*** (0.135)
Biden is best able to handle gun policy		0.866*** (0.111)	0.296* (0.145)			
Biden is best able to handle health care (AZ, FL, IL)	1.300*** (0.135)	1.631*** (0.111)	1.471*** (0.145)			
Biden is best able to handle immigration (AZ)	1.286*** (0.145)					
Biden is best able to handle climate change (AZ)	0.312* (0.138)					
Biden is best able to handle economy (IL)			1.755*** (0.141)			
Biden is best able to handle issues related to race (FL, IL)		0.940*** (0.112)	1.074*** (0.139)			
Biden is best able to handle corruption in government			0.895*** (0.170)			
Biden is best able to handle economy (MI)				1.264*** (0.126)		
Biden is best able to handle issues related to race (MO, MS)					1.812*** (0.154)	1.977*** (0.222)
Biden is best able to handle health care (MI, MO, MS)				1.864*** (0.130)	2.386*** (0.218)	2.225*** (0.259)
Constant	-4.262*** (0.922)	-3.065*** (0.588)	-1.967** (0.705)	-2.873*** (0.683)	-2.884*** (0.844)	-5.345*** (1.298)
Observations	1,857	3,124	2,486	2,268	1,775	955
Log Likelihood	-876.526	-1,299.308	-817.082	-991.701	-631.134	-312.159
Akaike Inf. Crit.	1,775.053	2,620.615	1,660.164	2,003.402	1,282.267	644.318
<i>Note:</i>						*p<0.05; **p<0.01; ***p<0.001

has some independent effect on vote choice that is not captured by perceived general election electability. Because electability and trait questions are both measured using four-point likert scales, we can directly compare effect sizes between traits and electability. These results largely confirm prior findings that traits can be a core component of utility, and in some cases the coefficients on certain traits rival the coefficient on electability.

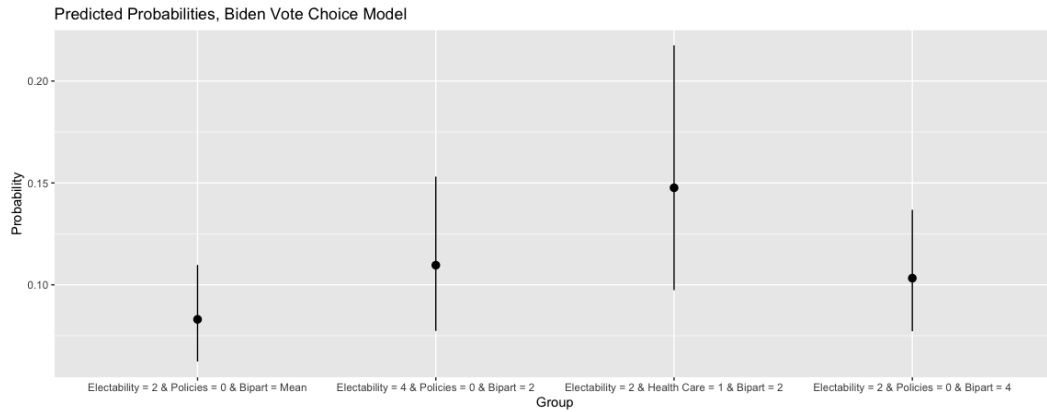


Figure 4.1: Biden Predicted Probabilities, Voter Analysis Data

A critical test in this work is to assess if issue emphasis can compete with candidate traits as a utility measure. We can use regression results from Table 4.3 to assess whether varying the importance of candidate traits matters more or less than varying perceived issue competence. Figure 4.1 transforms the regression results from the Arizona model into predicted probabilities of vote choice at varying levels of the expected utility variables. The first point represents the predicted probability of voting for Biden when a hypothetical respondent rates him as 2/4 on the electability scale, 0 on best able to handle all policies, and ranks the importance of bipartisanship in general at the sample

mean (about 3.5/4). The second point alters this hypothetical respondent to rank Biden at 4/4 on the electability scale and decreases their perceived importance of bipartisanship to 2/4. Next, electability and bipartisanship are both set at 2 and Biden is coded as best able to handle the issue of health care, and last the respondent rates Biden as 2/4 on electability, not best able to handle any issues, but rates the importance of bipartisanship as 4/4.

The takeaway from the predicted probabilities plotted in Figure 4.1 suggest that the biggest jump in support comes from varying Biden’s perceived ability to handle health care. Strongly believing in the importance of bipartisanship and rating Biden as highly electable predict similar probabilities of voting for Biden, all else equal. Error bars in these cases suggest we should interpret predicted probability differences in these hypothetical respondents with caution, but are directionally interesting nonetheless.

The Sanders expected utility models (Table 4.4) also reveal persistent patterns. Interestingly, traits mattered in the Sanders models as well, but in a very different way. Sanders voters were significantly *less* likely to think it was highly important that the nominee could beat Donald Trump, and were also almost uniformly less likely to say it was highly important that the nominee “has the right experience”. In Illinois, Michigan, and Missouri, however, they were more likely than others to say it was important that the nominee “has the best policy ideas”. Though Sanders voters generally perceived it to be less important that the nominee could beat Trump, it is worth noting the electability coefficient—thinking Sanders could beat Trump, in other words,

Table 4.4: Sanders Expected Utility Models

	<i>Dependent variable:</i>					
	Vote Choice = Sanders					
	(1) Arizona	(2) Florida	(3) Illinois	(4) Michigan	(5) Missouri	(6) Mississippi
Can beat Donald Trump	-0.806*** (0.174)	-0.341** (0.122)	-0.787*** (0.115)	-0.697*** (0.119)	-0.972*** (0.149)	-0.484* (0.209)
Will work across party lines	-0.136 (0.106)	-0.236* (0.104)	0.146 (0.108)	-0.412*** (0.098)	-0.561*** (0.120)	-0.234 (0.180)
Has the best policy ideas	0.255 (0.158)	0.125 (0.150)	0.297* (0.139)	0.402** (0.132)	0.828*** (0.163)	0.334 (0.243)
Cares about people like me	0.311 (0.171)	0.492** (0.161)	0.023 (0.136)	-0.207 (0.123)	0.013 (0.147)	0.127 (0.289)
Strong leader	-0.237 (0.202)	-0.037 (0.193)	-0.059 (0.189)	-0.055 (0.164)	-0.347 (0.202)	-0.825* (0.383)
Has the right experience	-0.268 (0.139)	-0.666*** (0.155)	-0.360* (0.140)	-0.529*** (0.132)	-0.586*** (0.155)	-0.819*** (0.237)
Likelihood Sanders could beat Trump	0.672*** (0.103)	0.664*** (0.090)	0.517*** (0.097)	0.910*** (0.092)	1.242*** (0.113)	0.670*** (0.151)
Sanders is best able to handle gun policy		1.399*** (0.154)	0.535** (0.174)			
Sanders is best able to handle health care (AZ, FL, IL)	1.993*** (0.194)	2.640*** (0.167)	1.663*** (0.168)			
Sanders is best able to handle immigration (AZ)	1.560*** (0.174)					
Sanders is best able to handle climate change (AZ)	1.370*** (0.194)					
Sanders is best able to handle economy (IL)			1.966*** (0.172)			
Sanders is best able to handle issues related to race (FL, IL)		1.737*** (0.154)	1.182*** (0.166)			
Sanders is best able to handle corruption in government			1.261*** (0.164)			
Sanders is best able to handle economy (MI)				2.262*** (0.213)		
Sanders is best able to handle issues related to race (MO, MS)					2.485*** (0.214)	2.203*** (0.288)
Sanders is best able to handle health care (MI, MO, MS)				2.262*** (0.141)	2.046*** (0.171)	2.649*** (0.247)
Constant	-2.132* (1.058)	-3.384*** (0.814)	-2.245** (0.783)	0.847 (0.726)	0.510 (0.828)	2.137 (1.601)
Observations	1,833	3,092	2,469	2,258	1,769	943
Log Likelihood	-528.913	-679.168	-626.043	-724.653	-511.162	-233.895
Akaike Inf. Crit.	1,079.826	1,380.335	1,278.086	1,469.305	1,042.324	487.791
<i>Note:</i>					*p<0.05; **p<0.01; ***p<0.001	

was still linked to voting for him. Similar to the Biden models, we again find that every single policy coefficient is significant.

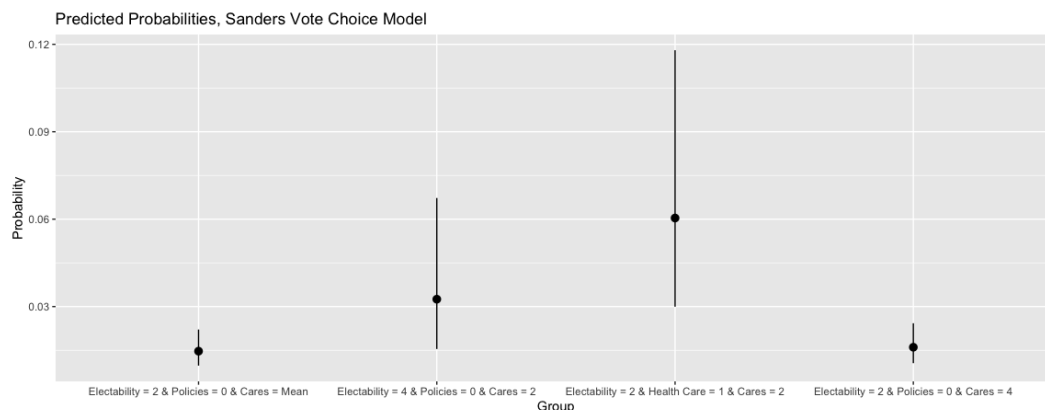


Figure 4.2: Sanders Predicted Probabilities, Voter Analysis Data

Figure 4.2 is set up the exact same way as the Biden chart, but in this case I vary a hypothetical respondent’s perceived importance of “caring about people like me”, since that was the strongest trait coefficient in the Sanders Arizona model. Figure 4.2 is also suggestive of the finding that strongly valuing candidate traits mattered less to Sanders voters than to Biden voters. There is also more statistical confidence that respondent profiles 2 & 3 (which vary electability and health care competence, respectively) would be more likely to vote for Sanders than profile 1 (intended to represent the “baseline” respondent). Again, the largest jump in the predicted probability of a Sanders vote comes from believing Sanders is best able to handle health care.

The results in this section have one clear takeaway: issue emphasis/competence seems to matter, and an expected utility model does a reason-

able job at characterizing how voters make decisions in primaries. Candidate traits, the dominant utility predictor in prior work, appears to be less important and more context-specific than issue emphasis. This is a novel finding in the expected utility literature and suggests a new, improved way to measure candidate utility going forward. However, these models leave open the “projection” criticism, since they rely entirely on cross-sectional data. So, while these are encouraging results, using panel data could help give me more causal leverage.

4.2.2 Evaluating Hypothesis 2

In this section, I turn to my causal hypotheses. Using my student panel data, I set out to test whether changes in perception of issue emphasis and changes in perceptions of electability are related to changes in vote intention. Then, I test a model that is more explicitly causal, and find evidence that prior changes in attitudes about electability and issue emphasis are related to vote intentions in later time periods.

Table 4.5 sets up a general first-differences style regression model, in which I only model whether a respondent’s stated vote intention changed wave-over-wave ($1 =$ respondent changed vote intention from time t to $t + 1$, $0 =$ respondent did not change vote intention from time t to $t + 1$). I also estimate these change models using logistic regression. In each model, I predict change in vote intention as a function of change from time t to $t + 1$ regarding who was paying the most attention to climate change, gun control, health care,

income inequality, and immigration. These are also all dummy coded where 1 = candidate believed to be paying the most attention to issue j changed from time t to $t + 1$, 0 = no change). Additionally, I include changes in electability perceptions as a predictor of changing vote intention, which is coded as the perceived electability of the intended vote choice candidate at time t - perceived electability of *that same candidate* at time $t + 1$. A positive value, in this case, represents that the respondent thought their initial intended vote choice at time t was *more electable* than they believed that candidate to be at time $t + 1$ (in other words, their perception of their initial vote choice candidate's electability decreased from time t to $t + 1$). I purposefully set up the model to not be candidate-specific, in order to try leverage all of the variation that occurred in the data.

In total, 114/280 students who took both waves 1 & 2 of my panel reported switching their intended vote over this period, 57/175 of wave 2 & 3 students reported switching their vote, and 98/191 of wave 1 & 3 students reported switching their vote.

Results of this model suggest that revising opinions about the electability of one's original intended vote choice candidate is most consistently and strongly related to changing one's vote choice to someone else. The effect is consistent in size and is highly significant across all waves. Changing opinions regarding who was best able to handle issues were also related to changes in vote intention, though the pattern is less consistent. Hypothesis 2b is supported by these results, but they provide only partial support for Hypothesis

Table 4.5: Student Panel Data, First-Difference Expected Utility Results

	<i>Dependent variable:</i>		
	$\Delta\text{Vote Wave 1} - 2$	$\Delta\text{Vote Wave 2} - 3$	$\Delta\text{Vote Wave 1} - 3$
$\Delta\text{Attention to Climate Change}$	0.041 (0.281)	0.119 (0.424)	0.306 (0.306)
$\Delta\text{Attention to Gun Control}$	0.308 (0.285)	-0.871* (0.438)	-0.592 (0.357)
$\Delta\text{Attention to Health Care}$	0.725* (0.288)	0.943* (0.456)	0.670 (0.370)
$\Delta\text{Attention to Income Inequality}$	0.655* (0.290)	0.551 (0.440)	0.245 (0.388)
$\Delta\text{Attention to Immigration}$	0.055 (0.291)	0.353 (0.410)	0.848* (0.368)
$\Delta\text{Initial Candidate Electability}$	0.318*** (0.078)	0.372*** (0.109)	0.351*** (0.089)
Constant	-1.562*** (0.284)	-1.756*** (0.372)	-1.742*** (0.369)
Observations	280	175	191
Log Likelihood	-158.746	-81.284	-101.613
Akaike Inf. Crit.	331.491	176.568	217.227
<i>Note:</i>		*p<0.05; **p<0.01; ***p<0.001	

2a.

Though the results presented in Table 4.5 are interesting, they are still open to potential challenges regarding causality. Showing that all of these variables change at the same time leaves open the possibility of the hypothesized causal relationship, but it does not rule out a possible projection relationship (students change their vote intention, and then fit candidate attitudes to align with that intention). In order to set up a more precise causal model, I leverage the three-wave panel structure of my data. I do this by modeling vote intention for Bernie Sanders (by far the most popular vote choice in all waves in my student sample) in wave 3, using attitudes from the prior two waves as predictors.

The model is set up as follows: the dependent variable in my causal model is reported Sanders vote intention in wave 3 (1 = expressed intention to vote for Sanders, 0 = did not), and I control for reported Sanders vote intention in waves 1 & 2 (which are coded the same way). All other variables in the model are coded such that positive values indicate changing attitudes that benefit Sanders. Issue emphasis variables are dummy-coded, where 1 = did not believe Sanders was paying the most attention to the issue in time t , but *did* believe Sanders was paying the most attention to the issue in time $t + 1$, and 0 = no change in the attitude wave over wave, or change away from Sanders. Change in Sanders electability is measured as perceived Sanders electability at time $t + 1$ – perceived Sanders electability at time t .

The results of this analysis are reported in Table 4.6. As expected, vote

Table 4.6: Student Panel Data, Lag Model Results

	<i>Dependent variable:</i>
	Wave 3 Reported Sanders Vote
Wave 2 Sanders Vote	3.132** (0.641)
Wave 1 Sanders Vote	1.473* (0.707)
Δ Belief Sanders is emphasizing Climate Change $w2, w3$	0.962 (0.692)
Δ Belief Sanders is emphasizing Gun Control $w2, w3$	2.165* (0.938)
Δ Belief Sanders is emphasizing Health Care $w2, w3$	-0.340 (0.646)
Δ Belief Sanders is emphasizing Income Inequality $w2, w3$	0.524 (0.652)
Δ Belief Sanders is emphasizing Immigration $w2, w3$	0.129 (0.688)
Δ Sanders Electability $w2, w3$	0.258 (0.158)
Δ Belief Sanders is emphasizing Climate Change $w1, w2$	0.664 (0.878)
Δ Belief Sanders is emphasizing Gun Control $w1, w2$	-0.302 (0.811)
Δ Belief Sanders is emphasizing Health Care $w1, w2$	0.937 (0.791)
Δ Belief Sanders is emphasizing Income Inequality $w1, w2$	0.415 (0.771)
Δ Belief Sanders is emphasizing Immigration $w1, w2$	1.733* (0.834)
Δ Sanders Electability $w1, w2$	0.339* (0.155)
Constant	-2.108** (0.489)
Observations	193
Log Likelihood	-58.854
Akaike Inf. Crit.	147.708
<i>Note:</i>	*p<0.05; **p<0.01

intentions for Sanders in waves 1 and 2 are significantly related to intentions to vote for Sanders in wave 3. The most critical variables to my causal story are those measuring change across waves 1–2, because changes in those attitudes very clearly occur before wave 3 responses are measured. Indeed, we do find that lagged attitude changes that benefit Sanders are significantly related to Sanders vote intention in wave 3, and relationships are in the expected direction. Believing that Sanders was more electable in wave 2 (Feb 12, 2020) than in wave 1 (Jan 29, 2020) was positively related to Sanders vote intention in wave 3 (March 4, 2020), even when controlling for change in that attitude from wave 2 to 3 as well as vote choice in waves 1 and 2. The same relationship is positive and significant for immigration emphasis attitudes. Given the relatively long time gap in between waves, I find these results encouraging and believe that they clearly demonstrate that reverse causality/projection isn't the definitive story regarding attitude updating in the primaries.

No other lagged issue emphasis change variables are significant, so H2a again receives only partial support. However, given the rather conservative nature of this test, this model doesn't rule out that we'd see more significant results if the panel data had either more waves or more fine-grained time measurements. H2b is again supported. To check the robustness of these results, I run the same model, estimated with OLS and heteroskedasticity-robust standard errors and report those results in Appendix C. This specification does not change the sign or significance or any of my predictors. Ideally, in the future, I can collect enough panel waves to also control for respondent-level

fixed effects. It will also be necessary to test if the findings observed from a relatively small student sample are generalizable to a broader population. The student panel survey was primarily set up to test whether or not my expected utility model is empirically supported, and I hope to use these results to request funding for a larger, more representative study.

4.3 Discussion

The analyses presented in this chapter build upon prior work and extend theory about primary election voting in a few key ways. Like the analyses presented in the first and second chapters, they present a *dynamic* picture of attitude change that is often missing in prior work on the subject. Second, they help address the causal identification issues present in many foundational pieces on expected utility voting in the primaries. And lastly, they introduce issue emphasis into an expected utility framework and demonstrate that primary voters do balance issue emphasis against more strategic considerations like general election electability.

Notably, when issue emphasis and trait variables are tested together in the same model, issue emphasis seems to be more strongly and consistently related to vote choice. Issue emphasis and electability perceptions are both significantly and consistently related to vote choice, highlighting the usefulness of my expected utility framework as a lens through which we can understand primary elections. However, while analysis of my panel data suggests that electability perceptions are causally linked to vote intentions, I cannot

entirely rule out potential projection relationships for issue emphasis. More data will be necessary to test the degree to which projection might occur on issue emphasis/issue competence in primaries.

4.3.1 Primary Elections and Intra-party Issue Ownership

In many ways, the version of the expected utility model I’ve tested here is an intra-party issue ownership model. Aldrich and Alvarez’s conceptualization of “issue emphasis” in a primary is very similar to Petrocik’s theory of issue ownership in general elections (Petrocik 1996). In both theories, issues function to “frame the vote choice as a decision to be made in terms of problems facing the country that (a candidate) is better able to ‘handle’ than (their) opponent(s)” (Petrocik 1996, p. 826). In other words, issues in campaigns function (at least in part) as signals to voters regarding who “cares” the most about an issue, or who will resolve a problem.

Theoretically, this suggests that primary election campaigns might function as framing contests in a similar manner to general election campaigns. Future work may uncover that candidates in primaries focus on issues that are owned by their party, but are also meant to appeal to a “lane” or subgroup of voters within their party. Note, for instance, that the relationship between believing Sanders is best able to handle health care is almost uniformly more strongly related to a Sanders vote than the equivalent relationship is for Biden in Tables 4.3 & 4.4. In an intra-party issue ownership framework, this makes sense, as Sanders was more likely to be seen as innovative on this issue among

the Democratic party base, and emphasizing that issue allowed him to solidify support among a large group of Democrats who didn't like the Affordable Care Act status quo. This is not to say that Biden voters didn't care about health care, or even that they didn't think Biden was best able to handle health care (because they generally did). Rather, the takeaway would be that emphasizing health care might not have been as advantageous to Biden as emphasizing other issues he could "own", and would thus draw additional support his way. The data do not allow me to test Petrocik's framework specifically, but they are suggestive that issue ownership is a theory with applications beyond a general election context.

I would argue, therefore, that these results do not simply inform our knowledge about how voters make decisions in primaries. They inform us about how parties themselves work by revealing a clear link between issue emphasis on the part of campaigns and voters' decisions. Because primaries are intra-party contests, we can start to draw a line between effective political campaigns and policy emphases on the part of the party at large. My results suggest that successful campaigns in primaries either accurately assess the policy priorities of their party's electorate *or* they successfully convince the majority of their party that issues they emphasize are the most important. This chapter is thus a story of how voters think through their decisions, but is also a story about the ability of both the electoral context and an effective campaign to shape the policy priorities of a party. That is a powerful effect, particularly if that candidate wins the general election. A similar can-

didate/campaign effect in primaries has been raised by Herrera (1999) with respect to party ideology.

How voters judge electability is much less clear at this stage, and requires further research. While voters appear to be strategic, it is unclear which indicators inform strategic assessments. I will return to the importance of this work in the concluding chapter.

4.3.2 Future directions

Though these models use attitude change over time to better assess causality, there is clearly more work to be done. An interesting future direction for this project will be to gather time series data from a larger sample, at more time periods, to test other hypotheses related to the dynamics of expected utility considerations in primary elections. For instance, one could imagine that different variables start to matter more in expected utility calculations as the field of candidates narrows. Additional time series data, designed to test expected utility hypotheses, could assess whether different considerations matter more/less at different stages of the campaign cycle.

These results also inform how we can interpret results from the prior chapter. Tests of Hypothesis 1 again show that different variables matter more and less to different groups of voters. It would appear, for instance, that Sanders voters prioritized different candidate traits than Biden voters, and even valued electability in different ways. Taken together, the results from the prior two chapters suggest that theories of primary voting are incom-

plete if they don't account for campaign-specific and time-specific variation in decision-making.

Chapter 5: Conclusion

This dissertation has explored how we can unify recent findings in the primary elections literature into a single theoretical framework which more accurately captures the dynamics of the primary election season. Throughout the project, I've sought to answer the following question: *How do voters arrive at a vote choice in presidential primaries?*

I tackled this question empirically throughout Chapters 2–4. In Chapter 2, I collected time series data on poll support/viability, media attention, debate performance, endorsements, opinions, and decision sets for four different candidates. Substantively, the findings revealed that potential voters did indeed seem to begin to form opinions about lesser-known candidates as the campaign wore on. The campaign environment, in other words, increased the willingness of voters to rate candidates on an electability scale over time. Media attention to each of the four candidates studied also increased as the campaign season progressed, suggesting it is one likely driver of name recognition and willingness to evaluate candidates. The debate performance time series, interestingly, didn't suggest that positive coverage was completely tied to being the best debater. Warren, who was particularly strong in this regard, had only a slightly higher average of positive coverage than other candidates studied. And Biden, who was not necessarily considered to be a strong debater, was covered fairly positively towards the end of the campaign. Lastly,

it was clear that endorsements coalesced around the eventual winner of the primaries, and that most endorsements came in at the very end of the period studied—after that candidate had won several contests.

In Chapter 3, I found additional empirical support that media attention is a key driver of both willingness to rate candidates on a favorability scale and poll support. Lagged values of differenced media coverage both increased willingness to rate and foreshadowed increased poll support. Chapter 3 analyses added the additional nuance that media effects take a bit more time to be reflected in willingness to rate and poll support than other variables. I also found evidence that spending was associated with poll support in pooled models, suggesting that diving deeper into the effects of ad spending in primaries is a useful avenue for future research. The candidate-specific models presented a more mixed pattern of results, though supported the general finding that media attention mattered in the opinion time series and spending mattered in the poll time series.

Chapter 4 turned away from the more aggregate focus of Chapters 2 and 3 and assessed how individuals weighed different considerations in an expected utility framework. Rickershauser and Aldrich’s finding that campaign issue emphasis is a key driver of favorability in the primaries was supported in models of vote choice. Issue emphasis appeared to add additional explanatory power to both static and dynamic models of primary vote intention. However, general election electability was also a very clear driver of vote intention, and was the strongest predictor of vote intention change in panel regression models.

The expected utility model appears in general to be a very good descriptor of individual-level decision making, once voters have narrowed down the overall field of candidates. My results suggest ways to make that expected utility model even stronger in future work.

5.1 Why study primaries?

Upon first glance, this project may appear to be a relatively narrow investigation into a peculiarity of the American political system. However, I would argue that the study of primaries can tell us something about voting behavior more generally, and that these findings are not limited to voting behavior in the United States.

First, Bartels makes the point in his 1988 work that studying the dynamics of primary elections results in a rather profound critique of neoliberal political theory and rational choice literature, which takes voter preferences to be fixed and exogenous. Primary elections, in other words, present a particularly interesting case study in which to stress-test many classic economic models of voting. It is evident in both Bartels's work and this project that voter preferences in these types of elections are, at the very least, not fixed. Candidates rise and fall in popularity, and voters respond in meaningful ways to the campaign environment. The degree to which voter preferences are exogenous is a tougher research question to answer. My work points to at least some level of preference exogeneity, as voters do not appear to simply project electability and utility perceptions onto a preferred candidate. The opinions

we'd expect to change before preferences are updated do indeed change. But it's hard to believe that a given voter's expected utility perceptions are independent of both social factors and the electoral context. While it may be less radical to question rational choice assumptions today than it was in 1988 (and others had proposed prominent alternatives even earlier—like Fiorina in *Retrospective Voting*), it is useful to consider how models developed in a rational choice framework do and do not apply to voting behavior, and primary elections present a unique opportunity to do so.

For instance, though the expected utility model is very clearly linked to the rational choice tradition, I believe it is still a helpful tool we can use to study primary elections. It would, of course, be a mistake to suggest that a point-in-time expected utility framework accurately describes the way voters make decisions in these elections, and many scholars who use this model in their work acknowledge this (either implicitly or explicitly) (Abramowitz 1989; Abramson et al. 1992). Yet, I find that expected utility variables are strongly related to vote choice. A balance of strategic considerations and sincere preferences is likely how many voters consciously consider their choices to be made, especially given the media's heavy focus on electability in the 2020 primaries. However, there are pieces of the puzzle that the expected utility framework leaves out. I acknowledge that the theory underpinning the model assumes voter preferences are formed exogenously, and I've noted that I doubt that's entirely the case. Expected utility results are also subject to change as voters' opinions change. The takeaway is thus not that voters are “rational”, per se.

Future work will be required to better understand how well voters pick up on issue emphasis on the part of campaigns, and what exactly drives electability perceptions. We may indeed find that voters who appear to be rational, or are acting as though they balance strategic and sincere considerations, are driven by inputs that don't align with objective indicators of electability or issue emphasis. The takeaway here is that the study of primary elections presents us with a unique opportunity to test how rational voters are, in an election without dominant information cues like party ID.

The second reason it is worthwhile to study primaries is that many parties around the world are increasingly adopting primaries for candidate selection, particularly in Latin America (Carey and Polga-Hecimovich 2006). Research suggests that one reason for this is that candidates chosen via primary elections are stronger than candidates selected via other means (ibid). Australia¹ and the UK² (among others) have also experimented (or are experimenting) with moves towards a US-style primary system. If the trend towards the democratization of candidate selection continues, then it will continue to be useful to study how voters make decisions in intra-party contests, and how generalizable results from the US might be. The US will likely remain the only country that holds sequential primaries, but the theory outlined in this dissertation has focused almost entirely on the dynamics of public opinion *be-*

¹https://www.aph.gov.au/About_Parliament/Parliamentary_Departments/Parliamentary_Library/FlagPost/2019/07/primary-primers-the-us-primary-style-contest-for-the-next-uk-prime-minister-is-the-worst-of-both-worlds/

²<https://blogs.lse.ac.uk/usappblog/2019/07/05/primary-primers-the-us-primary-style-contest-for-the-next-uk-prime-minister-is-the-worst-of-both-worlds/>

fore the first election is held. A fascinating next step in this research would be to test some of the theory outlined here in non-US contexts. Doing so would continue to strengthen our ability to shed light on how voters reason in elections without a powerful information shortcut, which again may provide us with better analytical leverage to test classic theories of voting.

Lastly, it is possible that the findings presented here apply to elections beyond primaries. Scholars have applied the same expected utility model tested here to two-stage plurality elections, and have found that it describes voter decision-making in those contexts well (Blais et al. 2011). Results presented in Chapter 4, then, may inform our ability to understand expected utility inputs in countries with two-stage plurality elections like France, Italy, Argentina, Austria, Finland, Ghana, and many others. Voters in these types of elections need to be strategic in ways quite like voters in primaries—they need to balance electability considerations (who could win in the second-stage contest) against preference considerations (which candidate, if elected, provides the most utility). It is clear from my work that candidate traits and issue emphasis seem to influence utility judgments for US voters. Future work can help identify how true this is in other electoral contexts. And, the potential rational choice critique very much carries over to voters in two-stage plurality elections as well. For instance, even if voters appear to be strategic, are they forming electability and utility judgments in the way the rational choice framework would expect? Tackling these types of questions in a comparative context would be another useful avenue for the work I’ve begun in this project.

5.2 Applicability of this work to other types of US primaries

Though my data are limited to that of presidential primary campaigns, I do not have a reason to believe that the basic decision framework proposed here is fundamentally different in other types of primary elections. In a gubernatorial primary, for example, voters probably narrow down the field of candidates based on who they think has a chance to win and their level of information about the candidate, and then make a basic expected utility calculation to determine their choice. Again, though other types of US primary elections are one-shot rather than sequential, all of the empirics in this work focus on activity that occurs before any elections are held.

One other main difference between presidential and non-presidential primaries is the level of information voters have about candidates. This could play out in an expected utility framework in one of two ways: either 1) fewer candidates make it into voters' decision sets or 2) the information threshold voters decide they require to consider candidates is lower in non-presidential primaries. Both scenarios raise the possibility that there is an increased chance that voters choose a sub optimal candidate, either because they eliminate too many candidates at the beginning of the decision or because they reason with a great deal of uncertainty.

Both of these extensions to other types of primary elections can, of course, be tested. I raise them here to argue that the theory I propose is not limited to presidential primaries and can be used as a blueprint for future

work.

5.3 Directions for future research

Throughout the dissertation, I have raised potential avenues for future work. The clearest next step in this research will be to gather data from additional primary elections to test the robustness of some of the patterns found here. It will be critical, for example, to gather data from a competitive Republican primary. It is possible that Republican voters react differently to campaign variables, and consider candidates using different criteria, than Democratic voters. This was also an election context in which electability loomed particularly large, which may have (at least in part) driven the electability results presented in Chapter 4. I consider this work to be an important first step towards a full test of my theory, though I acknowledge the limitations inherent in analyzing data from a single election.

In future iterations of this work, I also would like to collect data for more candidates. There was enough evidence that candidate-specific constituencies behaved uniquely to suggest that it is worthwhile to compare expected utility models at the candidate level. And, of course, it will be helpful to have panel data with additional waves, so that I can appropriately control for potential individual-level idiosyncrasies.

I do not believe that these limitations take away from what we can learn just from the 2020 Democratic primary, using data that was already available. This project, despite its limitations, is able to present a more dynamic picture

of primary election decision-making than has been available previously. The primary contribution of that new, dynamic picture is to help reinforce the causal importance of campaign-related variables in a primary election context. The theory laid out at the beginning of the work also tended to be supported more often than it was challenged, suggesting it will be a useful jumping-off point for future work on the subject.

Appendices

Appendix A: Supplemental Analyses for Chapter 2

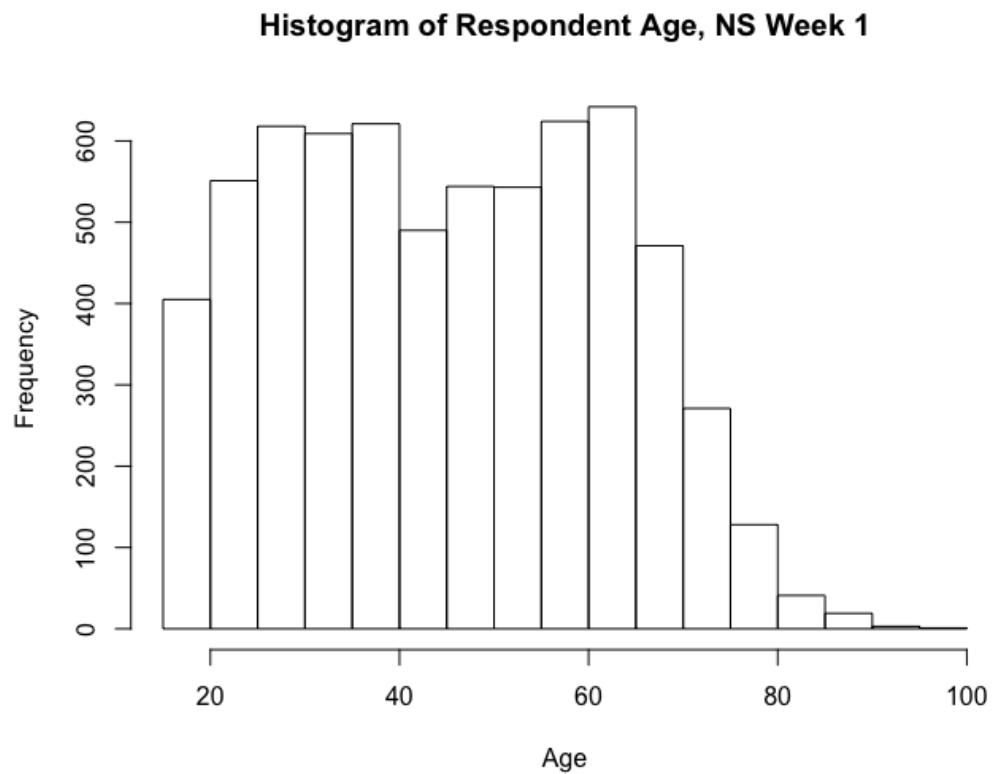
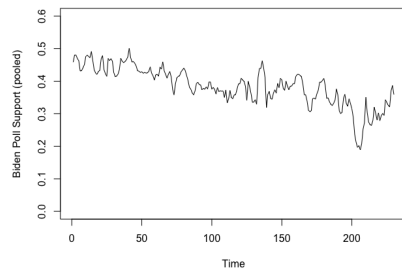


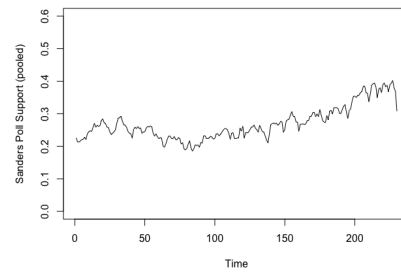
Figure A.1: Histogram of Respondent Age

A.1 Unit-root tests

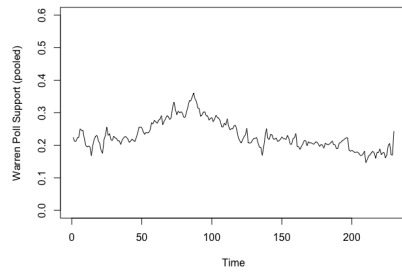
A.1.1 Polling time series



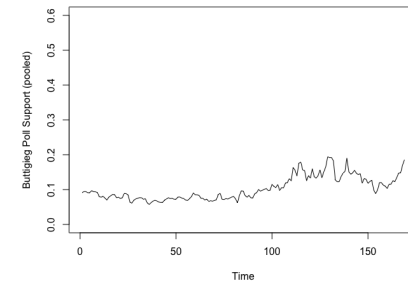
(a) Biden



(b) Sanders



(c) Warren



(d) Buttigieg

Figure A.2: Pooled poll time series

Table A.1: ADF Test, Midpoint-Aggregated Biden Poll Series

	Coefficient	Std. Error	Test Statistic
Constant	0.308	0.039	7.981
Biden Poll Lag ₁	−0.653	0.080	−8.209
Trend	−0.0005	0.00009	−5.675
Biden Poll First Difference Lag ₁	−0.076	0.067	−1.139

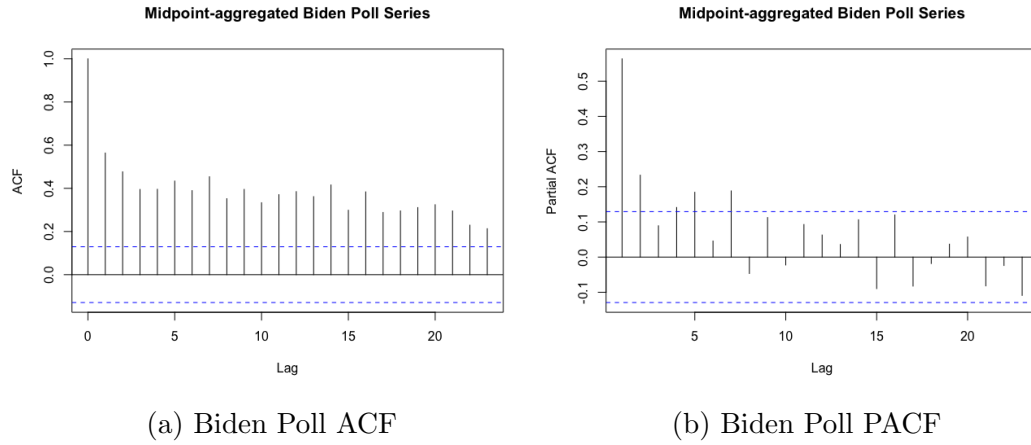
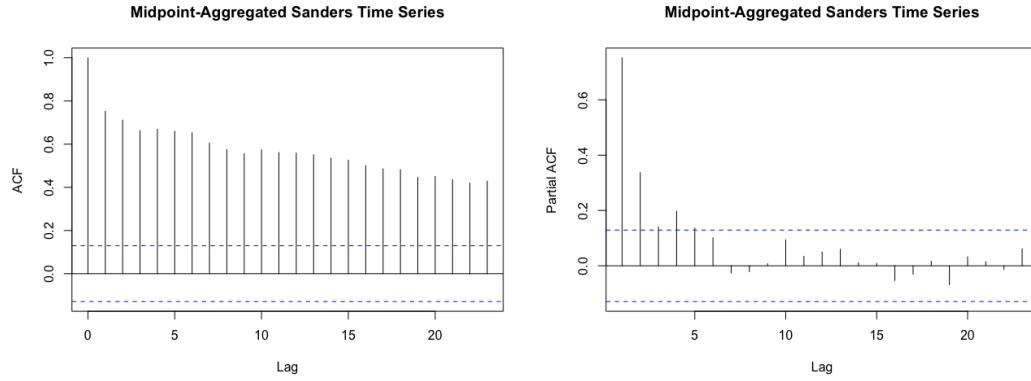


Figure A.3: Biden Poll ACF and PACF

Table A.2: ADF Test, Midpoint-Aggregated Sanders Poll Series

	Coefficient	Std. Error	Test Statistic
Constant	0.038	0.014	2.750
Sanders Poll Lag ₁	-0.177	0.065	-2.735
Trend	0.0001	0.00005	1.873
Sanders Poll First Difference Lag ₁	-0.505	0.082	-6.133
Sanders Poll First Difference Lag ₂	-0.336	0.084	-4.020
Sanders Poll First Difference Lag ₃	-0.319	0.079	-4.019
Sanders Poll First Difference Lag ₄	-0.173	0.067	-2.569



(a) Sanders Poll ACF

(b) Sanders Poll PACF

Figure A.4: Sanders Poll ACF and PACF

Table A.3: ADF Test, Midpoint-Aggregated Warren Poll Series

	Coefficient	Std. Error	Test Statistic
Constant	0.033	0.017	1.934
Warren Poll Lag ₁	−0.121	0.062	−1.967
Trend	−0.00004	0.00004	−1.079
Warren Poll First Difference Lag ₁	−0.566	0.082	−6.895
Warren Poll First Difference Lag ₂	−0.421	0.087	−4.842
Warren Poll First Difference Lag ₃	−0.354	0.084	−4.214
Warren Poll First Difference Lag ₄	−0.377	0.081	−4.637
Warren Poll First Difference Lag ₅	−0.285	0.078	−3.669
Warren Poll First Difference Lag ₆	−0.241	0.066	−3.655

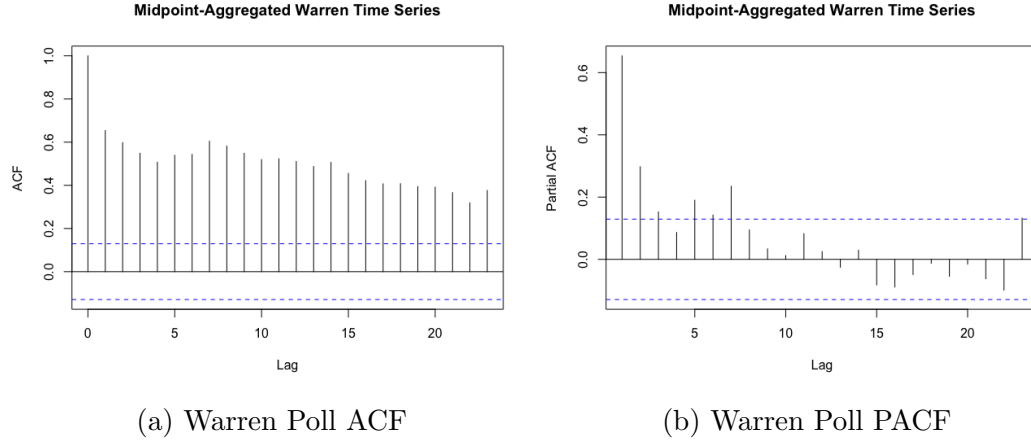
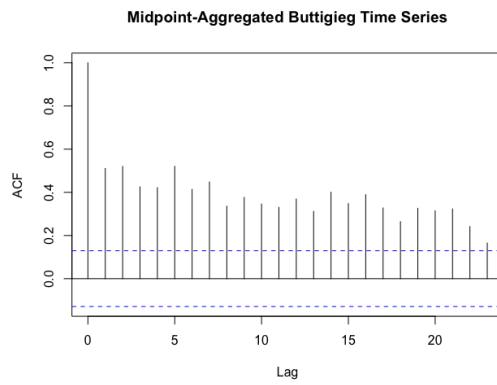


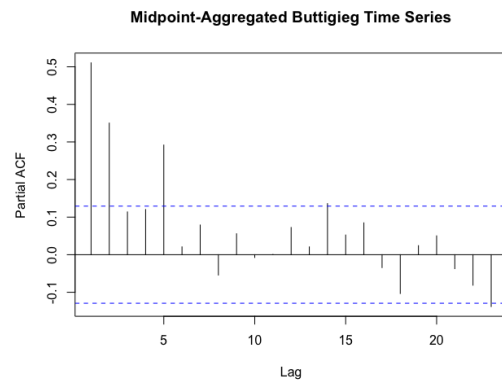
Figure A.5: Warren Poll ACF and PACF

Table A.4: ADF Test, Midpoint-Aggregated Buttigieg Poll Series

	Coefficient	Std. Error	Test Statistic
Constant	0.022	0.008	2.681
Buttigieg Poll Lag ₁	-0.325	0.098	-3.325
Trend	-0.0001	0.00006	2.123
Buttigieg Poll First Difference Lag ₁	-0.467	0.099	-4.723
Buttigieg Poll First Difference Lag ₂	-0.243	0.094	-2.585
Buttigieg Poll First Difference Lag ₃	-0.267	0.086	-3.103
Buttigieg Poll First Difference Lag ₄	-0.257	0.067	-3.840

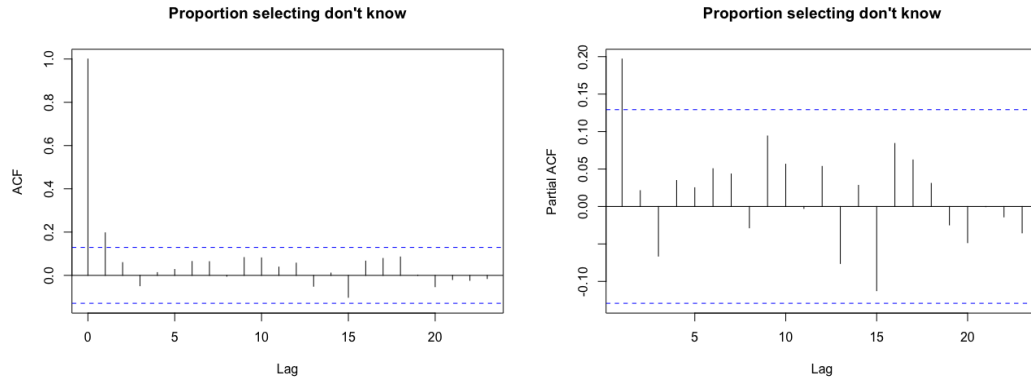


(a) Buttigieg Poll ACF



(b) Buttigieg Poll PACF

Figure A.6: Buttigieg Poll ACF and PACF



(a) Biden Don't Know ACF

(b) Biden Don't Know PACF

Figure A.7: Biden Don't Know ACF and PACF

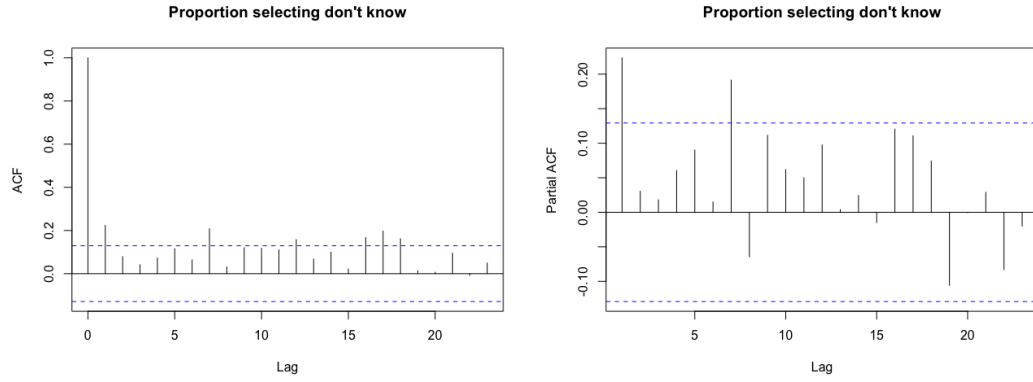
A.1.2 Don't know time series

Table A.5: ADF Test, Biden Don't Know Time Series

	Coefficient	Std. Error	Test Statistic
Constant	0.119	0.012	9.724
Biden DK Lag ₁	-0.864	0.088	-9.913
Trend	-0.00005	0.00002	-2.983
Biden DK First Difference Lag ₁	0.020	0.067	0.302

Table A.6: ADF Test, Sanders Don't Know Time Series

	Coefficient	Std. Error	Test Statistic
Constant	0.112	0.011	10.161
Sanders DK Lag ₁	-0.911	0.088	-10.351
Trend	-0.00007	0.00002	-4.321
Sanders DK First Difference Lag ₁	0.047	0.067	0.708



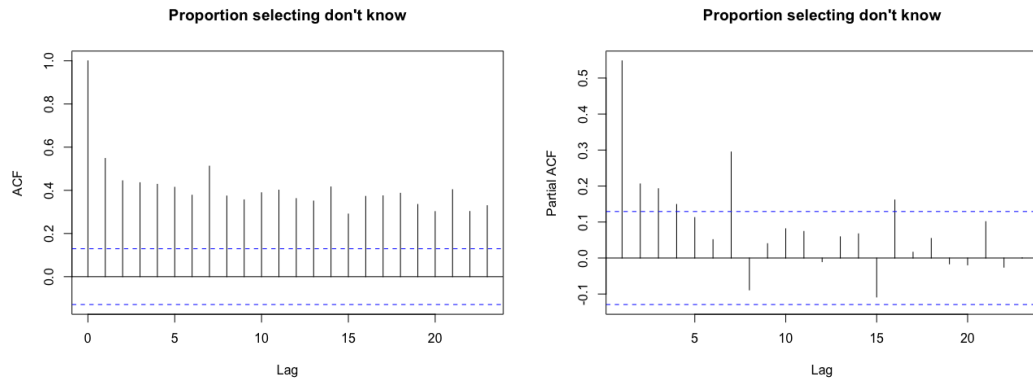
(a) Sanders Don't Know ACF

(b) Sanders Don't Know PACF

Figure A.8: Sanders Don't Know ACF and PACF

Table A.7: ADF Test, Warren Don't Know Time Series

	Coefficient	Std. Error	Test Statistic
Constant	0.253	0.025	10.118
Warren DK Lag ₁	-0.885	0.087	-10.217
Trend	-0.0003	0.00004	-7.899
Warren DK First Difference Lag ₁	0.049	0.067	0.743



(a) Warren Don't Know ACF

(b) Warren Don't Know PACF

Figure A.9: Warren Don't Know ACF and PACF

Table A.8: ADF Test, Buttigieg Don't Know Time Series

	Coefficient	Std. Error	Test Statistic
Constant	0.147	0.042	3.514
Buttigieg DK Lag ₁	−0.290	0.082	−3.527
Trend	−0.0002	0.00007	−3.365
Buttigieg DK First Difference Lag ₁	−0.420	0.091	−4.602
Buttigieg DK First Difference Lag ₂	−0.366	0.090	−4.077
Buttigieg DK First Difference Lag ₃	−0.182	0.082	−2.222
Buttigieg DK First Difference Lag ₄	−0.122	0.067	−1.820

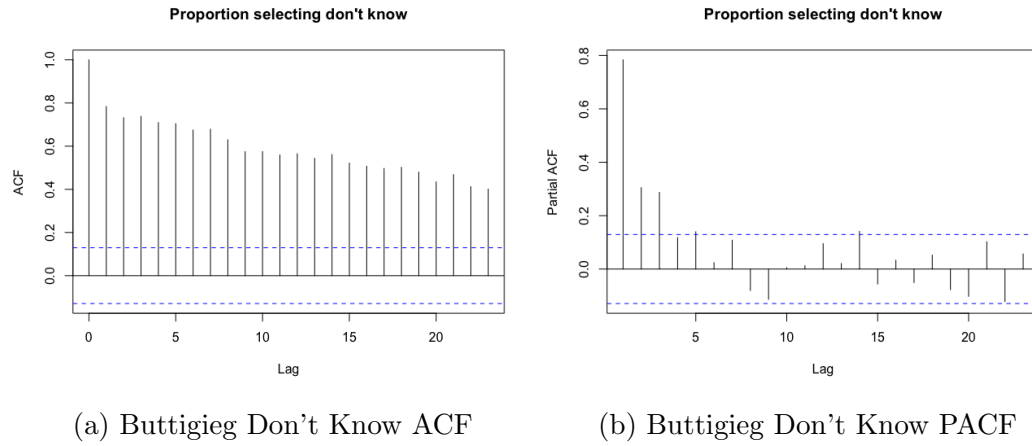


Figure A.10: Buttigieg Don't Know ACF and PACF

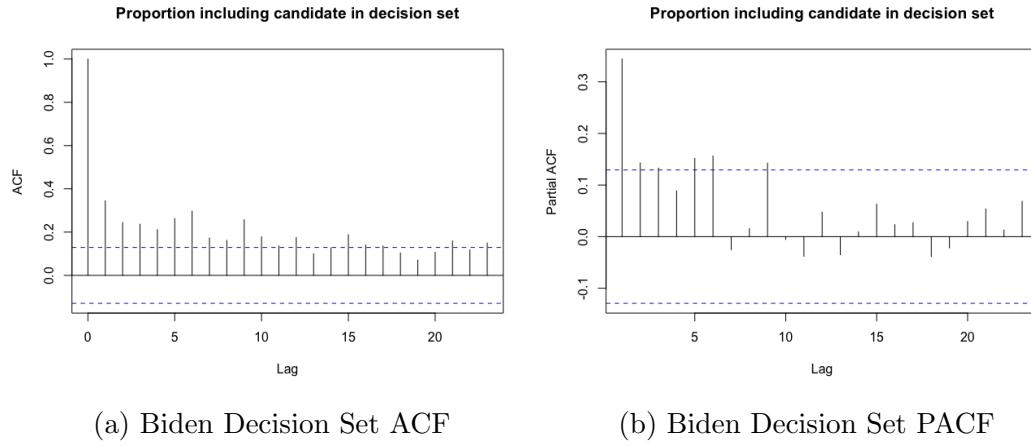


Figure A.11: Biden Decision Set PACF and ACF

A.1.3 Decision set time series

Table A.9: ADF Test, Biden Decision Set Time Series

	Coefficient	Std. Error	Test Statistic
Constant	0.160	0.043	3.691
Biden DS Lag ₁	-0.394	0.107	-3.671
Trend	0.00005	0.00003	1.496
Biden DS First Difference Lag ₁	-0.404	0.110	-3.699
Biden DS First Difference Lag ₂	-0.352	0.105	-3.357
Biden DS First Difference Lag ₃	-0.284	0.098	-2.904
Biden DS First Difference Lag ₄	-0.251	0.086	-2.911
Biden DS First Difference Lag ₅	-0.153	0.068	-2.243

Table A.10: ADF Test, Sanders Decision Set Time Series

	Coefficient	Std. Error	Test Statistic
Constant	0.111	0.029	3.830
Sanders DS Lag ₁	−0.303	0.081	−3.755
Trend	0.0001	0.00005	2.905
Sanders DS First Difference Lag ₁	−0.313	0.089	−3.552
Sanders DS First Difference Lag ₂	−0.238	0.085	−2.794
Sanders DS First Difference Lag ₃	−0.140	0.078	−1.787
Sanders DS First Difference Lag ₄	−0.154	0.067	−2.294

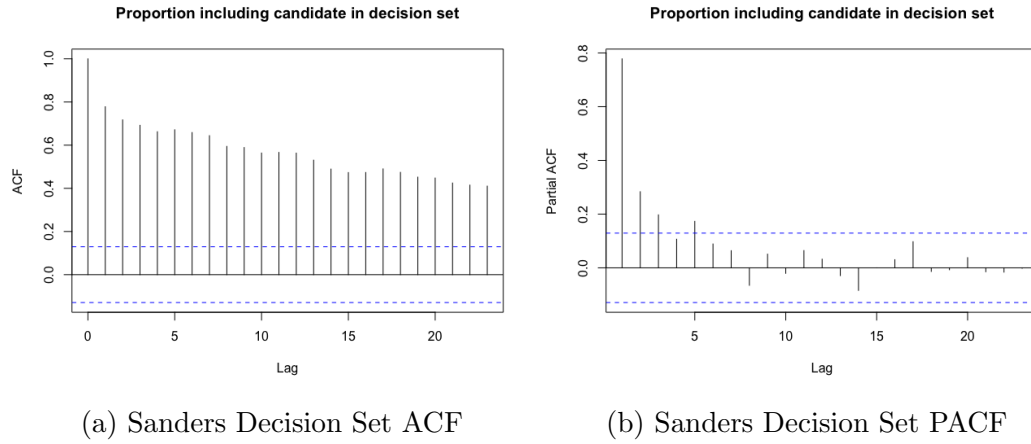


Figure A.12: Sanders Decision Set ACF and PACF

Table A.11: ADF Test, Warren Decision Set Time Series

	Coefficient	Std. Error	Test Statistic
Constant	0.040	0.018	2.209
Warren DS Lag ₁	−0.111	0.054	−2.045
Trend	0.000009	0.00004	0.249
Warren DS First Difference Lag ₁	−0.626	0.078	−8.047
Warren DS First Difference Lag ₂	−0.451	0.084	−5.354
Warren DS First Difference Lag ₃	−0.296	0.081	−3.671
Warren DS First Difference Lag ₄	−0.211	0.057	−3.154

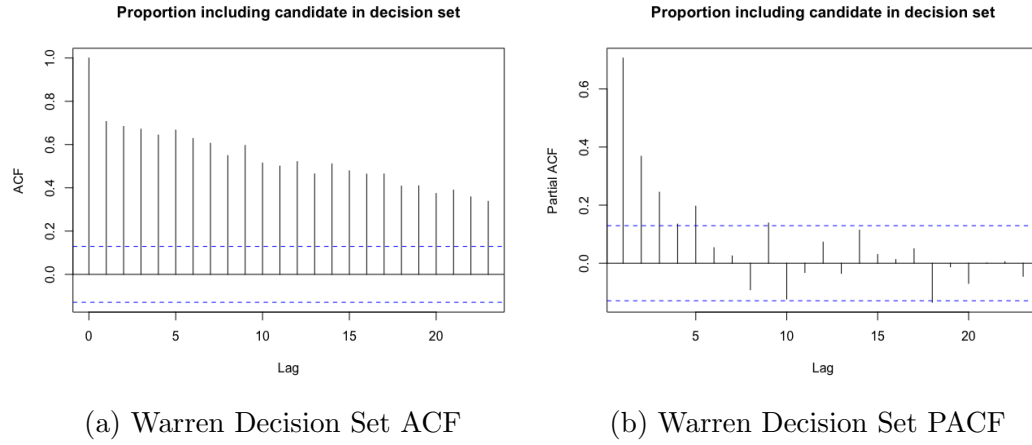
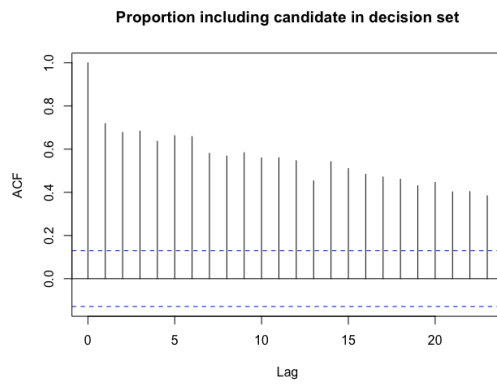


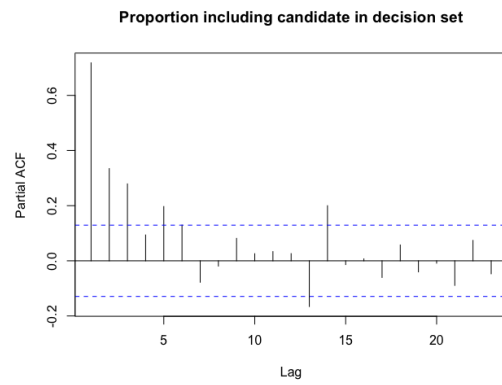
Figure A.13: Warren Decision Set ACF and PACF

Table A.12: ADF Test, Buttigieg Decision Set Time Series

	Coefficient	Std. Error	Test Statistic
Constant	0.059	0.011	5.359
Buttigieg DS Lag ₁	−0.501	0.090	−5.584
Trend	0.0003	0.00006	4.683
Buttigieg DS First Difference Lag ₁	−0.244	0.084	−2.916
Buttigieg DS First Difference Lag ₂	−0.155	0.067	−2.299



(a) Buttigieg Decision Set ACF



(b) Buttigieg Decision Set PACF

Figure A.14: Buttigieg Decision Set ACF and PACF

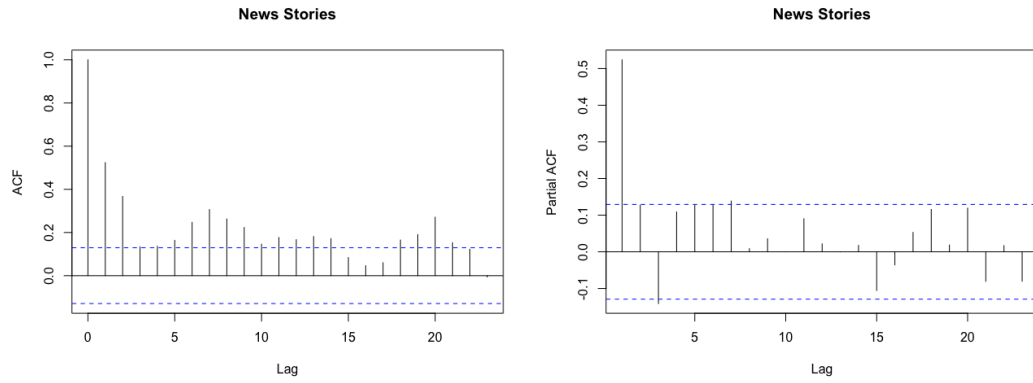
A.1.4 Media attention time series

Table A.13: ADF Test, Biden Media Attention Time Series

	Coefficient	Std. Error	Test Statistic
Constant	−2.132	2.142	−0.996
Biden MA Lag ₁	0.047	0.118	0.393
Trend	0.021	0.016	1.304
Biden MA First Difference Lag ₁	−0.510	0.129	−3.944
Biden MA First Difference Lag ₂	−0.237	0.120	−1.974
Biden MA First Difference Lag ₃	−0.455	0.111	−4.115
Biden MA First Difference Lag ₄	−0.411	0.100	−4.094
Biden MA First Difference Lag ₅	−0.430	0.095	−4.525
Biden MA First Difference Lag ₆	−0.252	0.083	−3.051

Table A.14: ADF Test, Sanders Media Attention Time Series

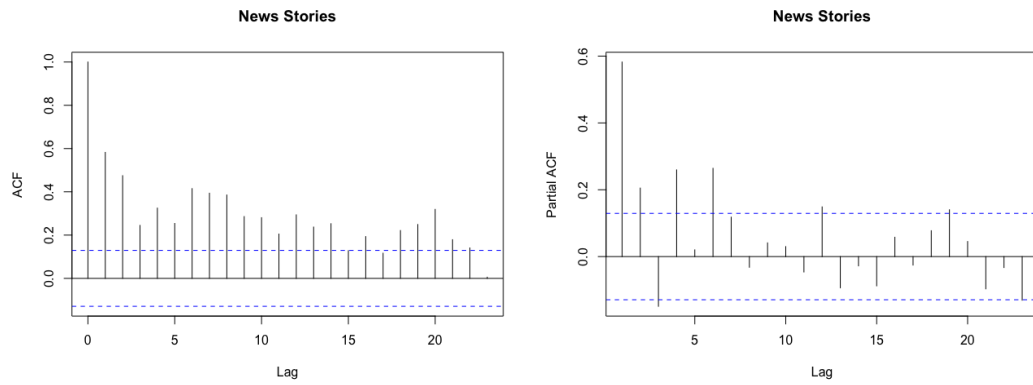
	Coefficient	Std. Error	Test Statistic
Constant	−1.649	2.044	−0.807
Sanders MA Lag ₁	0.006	0.088	0.071
Trend	0.024	0.018	1.347
Sanders MA First Difference Lag ₁	−0.495	0.100	−4.921
Sanders MA First Difference Lag ₂	−0.212	0.095	−2.229
Sanders MA First Difference Lag ₃	−0.456	0.085	−5.364
Sanders MA First Difference Lag ₄	−0.346	0.085	−4.069
Sanders MA First Difference Lag ₅	−0.446	0.079	−5.666



(a) Biden Media Attention ACF

(b) Biden Media Attention PACF

Figure A.15: Biden Media Attention PACF and ACF



(a) Sanders Media Attention ACF

(b) Sanders Media Attention PACF

Figure A.16: Sanders Media Attention PACF and ACF

Table A.15: ADF Test, Warren Media Attention Time Series

	Coefficient	Std. Error	Test Statistic
Constant	1.548	1.750	0.885
Warren MA Lag ₁	−0.246	0.092	−2.686
Trend	0.029	0.014	2.051
Warren MA First Difference Lag ₁	−0.093	0.102	−0.910
Warren MA First Difference Lag ₂	−0.435	0.096	−4.512
Warren MA First Difference Lag ₃	−0.300	0.093	−3.232
Warren MA First Difference Lag ₄	−0.221	0.086	−2.561
Warren MA First Difference Lag ₅	−0.219	0.077	−2.844
Warren MA First Difference Lag ₆	−0.161	0.074	−2.178

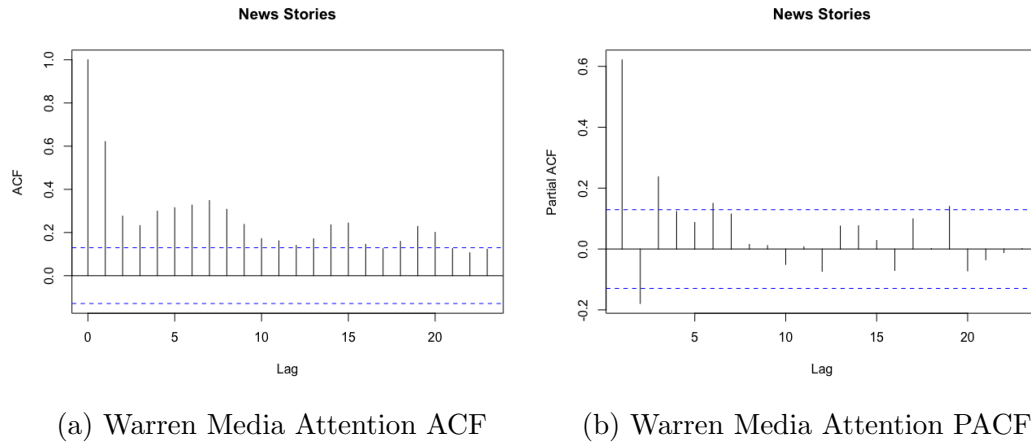


Figure A.17: Warren Media Attention PACF and ACF

Table A.16: ADF Test, Buttigieg Media Attention Time Series

	Coefficient	Std. Error	Test Statistic
Constant	−0.608	1.616	−0.376
Buttigieg MA Lag ₁	−0.172	0.111	−1.543
Trend	0.031	0.015	2.008
Buttigieg MA First Difference Lag ₁	−0.271	0.118	−2.296
Buttigieg MA First Difference Lag ₂	−0.177	0.110	−1.610
Buttigieg MA First Difference Lag ₃	−0.267	0.099	−2.705
Buttigieg MA First Difference Lag ₄	−0.364	0.091	−3.987
Buttigieg MA First Difference Lag ₅	−0.392	0.087	−4.513
Buttigieg MA First Difference Lag ₆	−0.251	0.078	−3.213

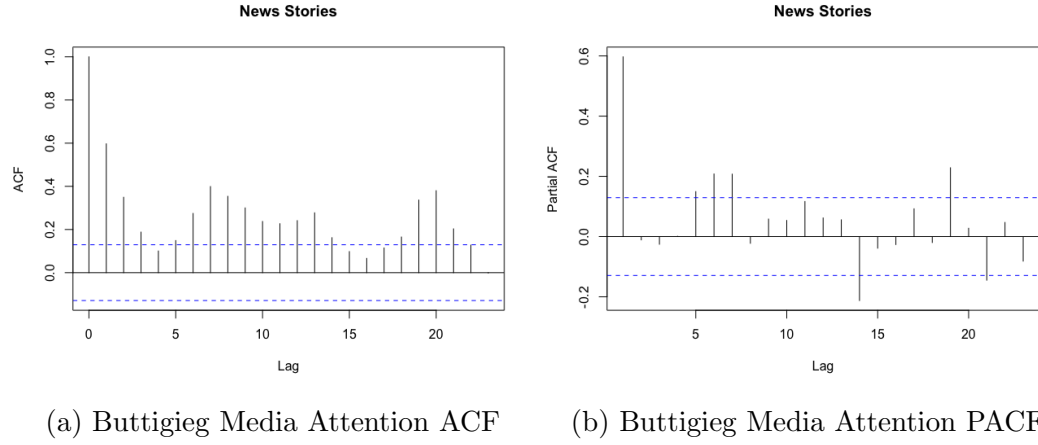


Figure A.18: Buttigieg Media Attention PACF and ACF

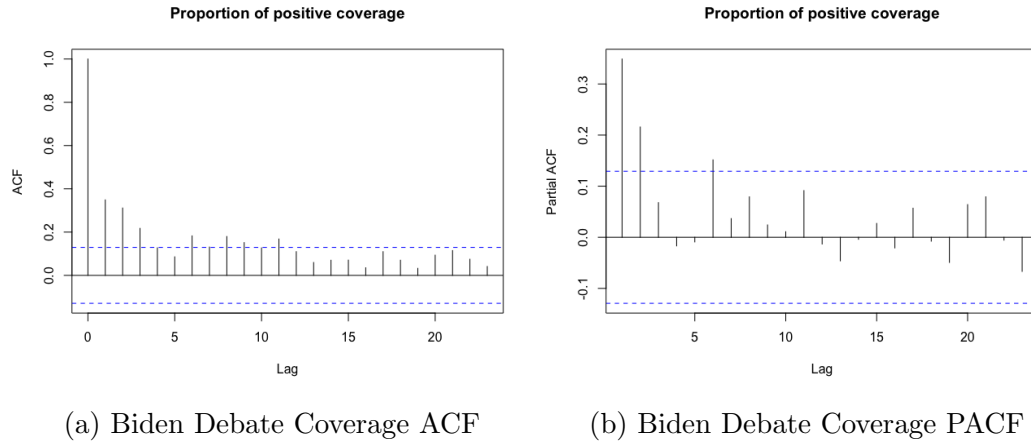


Figure A.19: Biden Debate Coverage PACF and ACF

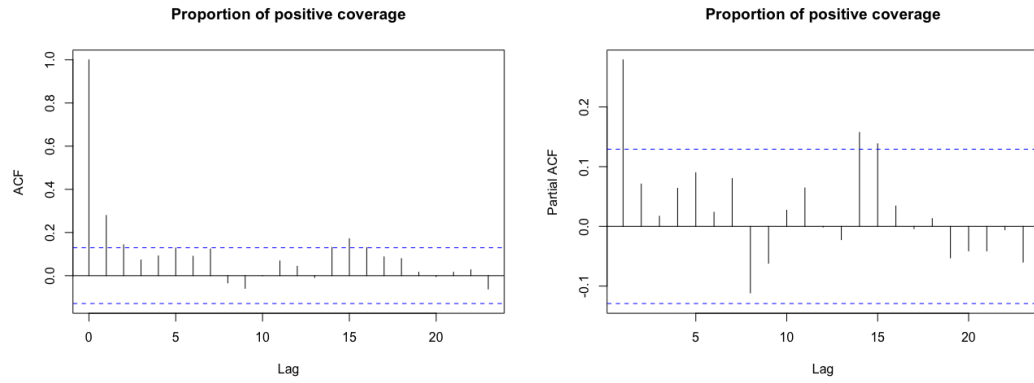
A.1.5 Positive debate coverage time series

Table A.17: ADF Test, Biden Debate Coverage Time Series

	Coefficient	Std. Error	Test Statistic
Constant	0.007	0.003	2.282
Biden Debate Lag ₁	−0.321	0.116	−2.770
Trend	0.00002	0.000008	2.224
Biden Debate First Difference Lag ₁	−0.386	0.136	−3.397
Biden Debate First Difference Lag ₂	−0.166	0.108	−1.532
Biden Debate First Difference Lag ₃	−0.072	0.101	−0.713
Biden Debate First Difference Lag ₄	−0.114	0.091	−1.244
Biden Debate First Difference Lag ₅	−0.196	0.073	−2.672

Table A.18: ADF Test, Sanders Debate Coverage Time Series

	Coefficient	Std. Error	Test Statistic
Constant	0.017	0.002	7.536
Sanders Debate Lag ₁	−0.759	0.085	−8.911
Trend	0.00003	0.000009	3.506
Sanders Debate First Difference Lag ₁	−0.022	0.068	−0.325



(a) Sanders Debate Coverage ACF (b) Sanders Debate Coverage PACF

Figure A.20: Sanders Positive Debate Coverage PACF and ACF

Table A.19: ADF Test, Warren Debate Coverage Time Series

	Coefficient	Std. Error	Test Statistic
Constant	0.020	0.003	7.585
Warren Debate Lag ₁	−0.713	0.085	−8.367
Trend	0.000001	0.000008	0.130
Warren Debate First Difference Lag ₁	−0.132	0.063	−1.992

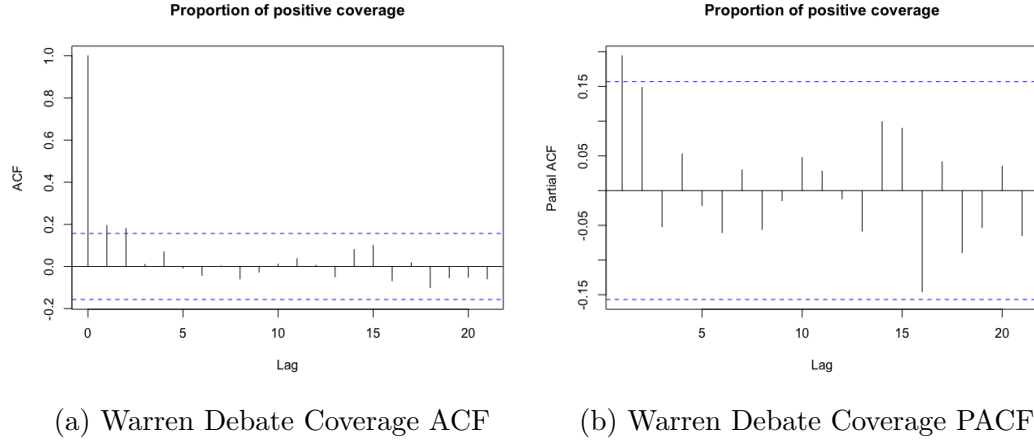
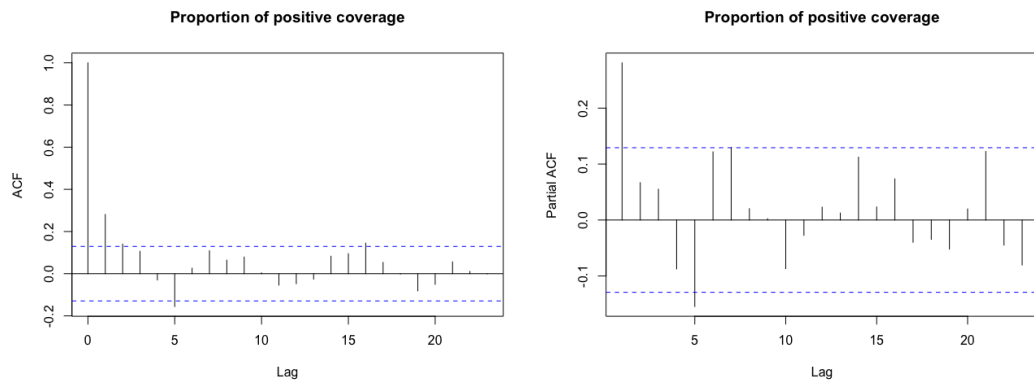


Figure A.21: Warren Positive Debate Coverage PACF and ACF

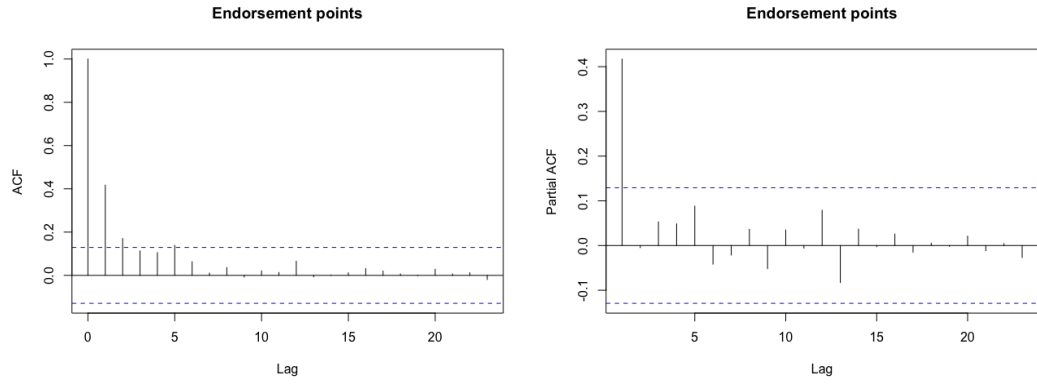
Table A.20: ADF Test, Buttigieg Debate Coverage Time Series

	Coefficient	Std. Error	Test Statistic
Constant	0.016	0.002	6.327
Buttigieg Debate Lag ₁	−0.787	0.107	−7.457
Trend	0.00003	0.00001	2.712
Buttigieg Debate First Difference Lag ₁	0.020	0.098	0.210
Buttigieg Debate First Difference Lag ₂	0.062	0.084	0.740
Buttigieg Debate First Difference Lag ₃	0.117	0.066	1.766



(a) Buttigieg Debate Coverage ACF (b) Buttigieg Debate Coverage PACF

Figure A.22: Buttigieg Positive Debate Coverage PACF and ACF



(a) Biden Endorsement Point ACF (b) Biden Endorsement Point PACF

Figure A.23: Biden Endorsement Point PACF and ACF

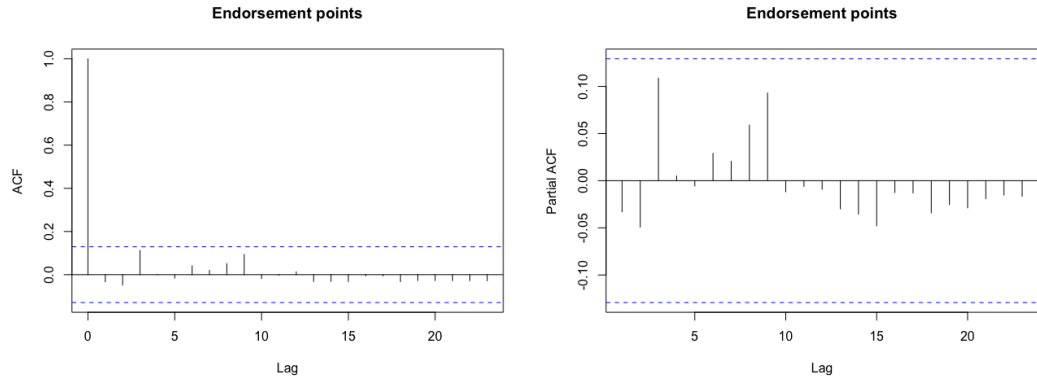
A.1.6 Endorsement points time series

Table A.21: ADF Test, Biden Endorsement Point Time Series

	Coefficient	Std. Error	Test Statistic
Constant	−0.601	0.541	−1.110
Biden Endorsement Lag ₁	0.707	0.268	2.639
Trend	0.003	0.005	0.535
Biden Endorsement First Difference Lag ₁	−1.472	0.292	−5.037
Biden Endorsement First Difference Lag ₂	−1.209	0.254	−4.756
Biden Endorsement First Difference Lag ₃	−0.903	0.202	−4.459
Biden Endorsement First Difference Lag ₄	−0.556	0.137	−4.072

Table A.22: ADF Test, Sanders Endorsement Point Time Series

	Coefficient	Std. Error	Test Statistic
Constant	0.146	0.124	1.180
Sanders Endorsement Lag ₁	−1.088	0.096	−11.331
Trend	0.0006	0.0009	0.680
Sanders Endorsement First Difference Lag ₁	0.051	0.067	0.767

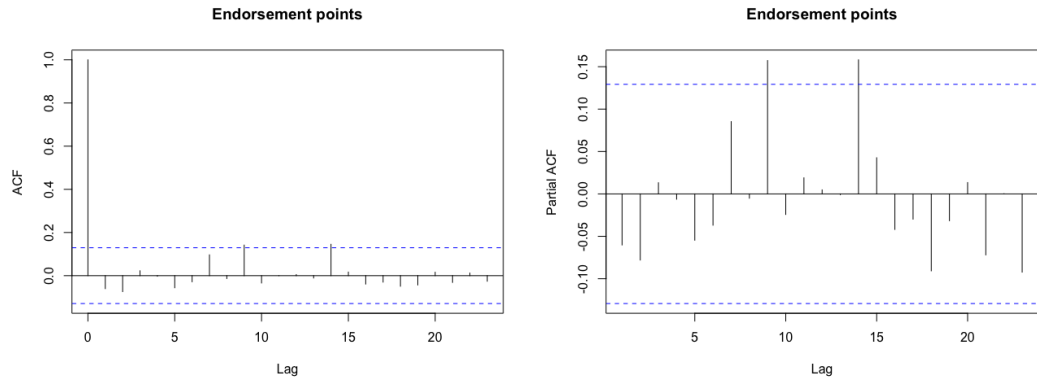


(a) Sanders Endorsement Point ACF (b) Sanders Endorsement Point PACF

Figure A.24: Sanders Endorsement Point PACF and ACF

Table A.23: ADF Test, Warren Endorsement Point Time Series

	Coefficient	Std. Error	Test Statistic
Constant	0.302	0.162	1.862
Warren Endorsement Lag ₁	-1.148	0.097	-11.776
Trend	0.0008	0.001	0.644
Warren Endorsement First Difference Lag ₁	0.080	0.067	1.200

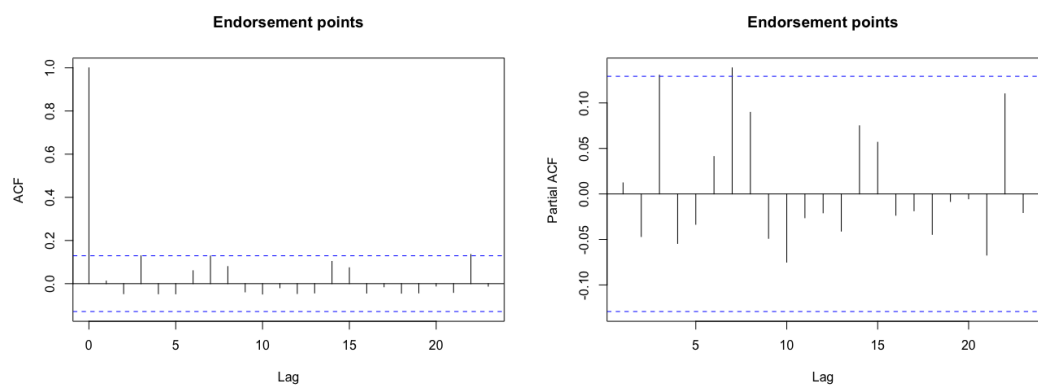


(a) Warren Endorsement Point ACF (b) Warren Endorsement Point PACF

Figure A.25: Warren Endorsement Point PACF and ACF

Table A.24: ADF Test, Buttigieg Endorsement Point Time Series

	Coefficient	Std. Error	Test Statistic
Constant	-0.038	0.090	-0.427
Buttigieg Endorsement Lag ₁	-1.085	0.095	-11.434
Trend	0.002	0.0007	2.442
Buttigieg Endorsement First Difference Lag ₁	0.073	0.067	1.089



(a) Buttigieg Endorsement Point ACF (b) Buttigieg Endorsement Point PACF

Figure A.26: Buttigieg Endorsement Point PACF and ACF

Appendix B: Supplemental Analyses for Chapter 3

B.1 Candidate Specific Results, ADL with one lag

Table B.1: Biden Models

	<i>Dependent variable:</i>		
	Don't Know Time Series	Poll Support Time Series	Decision Set Time Series
	(1)	(2)	(3)
Decision Set Lag			0.373*** (0.064)
Don't Know Lag	0.169* (0.067)	0.221 (0.260)	0.100 (0.126)
Δ Media Attention	0.0001 (0.0001)	0.001* (0.001)	
Δ Media Attention Lag	0.00001 (0.0001)	-0.0004 (0.0003)	
Poll Support	-0.007 (0.018)		0.019 (0.032)
Poll Support Lag	0.039* (0.017)	0.490*** (0.059)	-0.031 (0.031)
Positive Debate Coverage	0.059 (0.162)	0.067 (0.623)	
Positive Debate Coverage Lag	-0.042 (0.169)	0.290 (0.650)	
Δ Endorsements	-0.001 (0.001)	0.005* (0.002)	
Δ Endorsements Lag	0.001 (0.0004)	-0.002 (0.002)	
Don't Know		-0.111 (0.261)	-0.401** (0.125)
Ad Spending	-0.000 (0.000)	-0.00006*** (0.00000)	
Constant	0.095*** (0.011)	0.185*** (0.049)	0.306*** (0.037)
Observations	228	228	228
R ²	0.086	0.380	0.167
Adjusted R ²	0.044	0.351	0.149
Residual Std. Error	0.017 (df = 217)	0.064 (df = 217)	0.031 (df = 222)
F Statistic	2.040* (df = 10; 217)	13.281*** (df = 10; 217)	8.923*** (df = 5; 222)
<i>Note:</i>			*p<0.05; **p<0.01; ***p<0.001

Table B.2: Sanders Models

	<i>Dependent variable:</i>		
	Don't Know Time Series	Poll Support Time Series	Decision Set Time Series
	(1)	(2)	(3)
Decision Set Lag			0.724*** (0.046)
Don't Know Lag	0.185** (0.067)	0.015 (0.157)	-0.061 (0.123)
Δ Media Attention	0.00001 (0.0001)	-0.0002 (0.0003)	
Δ Media Attention Lag	0.0001 (0.0001)	0.0001 (0.0002)	
Poll Support	-0.009 (0.029)		0.003 (0.051)
Poll Support Lag	-0.033 (0.030)	0.555*** (0.059)	0.072 (0.051)
Positive Debate Coverage	-0.070 (0.125)	0.434 (0.289)	
Positive Debate Coverage Lag	0.033 (0.127)	0.446 (0.293)	
Endorsements	-0.002 (0.001)	0.002 (0.003)	
Endorsements Lag	0.001 (0.001)	-0.002 (0.003)	
Don't Know		-0.049 (0.158)	-0.216 (0.122)
Ad Spending	-0.000 (0.000)	0.000014*** (0.000)	
Constant	0.106*** (0.011)	0.090** (0.030)	0.127*** (0.030)
Observations	228	228	228
R ²	0.120	0.616	0.626
Adjusted R ²	0.080	0.598	0.617
Residual Std. Error	0.016 (df = 217)	0.037 (df = 217)	0.029 (df = 222)
F Statistic	2.964** (df = 10; 217)	34.802*** (df = 10; 217)	74.213*** (df = 5; 222)
<i>Note:</i>		*p<0.05; **p<0.01; ***p<0.001	

Table B.3: Warren Models

	<i>Dependent variable:</i>		
	Don't Know Time Series	Poll Support Time Series	Decision Set Time Series
	(1)	(2)	(3)
Δ Decision Set Lag			0.490*** (0.021)
Don't Know Lag	0.429*** (0.062)	-0.031 (0.061)	0.077 (0.051)
Δ Media Attention	0.00002 (0.0002)	-0.0004* (0.0002)	
Δ Media Attention Lag	0.0001 (0.0001)	0.0002 (0.0001)	
Δ Poll Support	0.149 (0.076)		0.083 (0.057)
Δ Poll Support Lag	-0.031 (0.045)	0.493*** (0.022)	-0.041 (0.034)
Positive Debate Coverage	-0.041 (0.207)	-0.064 (0.184)	
Positive Debate Coverage Lag	-0.088 (0.208)	0.047 (0.185)	
Endorsements	-0.002 (0.001)	-0.0002 (0.001)	
Endorsements Lag	-0.0003 (0.001)	0.002 (0.001)	
Don't Know		0.118 (0.060)	-0.118* (0.051)
Ad Spending	-0.00002*** (0.000)	0.000 (0.000)	
Constant	0.154*** (0.018)	-0.025 (0.018)	0.010 (0.012)
Observations	228	228	228
R ²	0.377	0.726	0.745
Adjusted R ²	0.348	0.713	0.739
Residual Std. Error	0.026 (df = 217)	0.023 (df = 217)	0.020 (df = 222)
F Statistic	13.106*** (df = 10; 217)	57.386*** (df = 10; 217)	129.655*** (df = 5; 222)
<i>Note:</i>			*p<0.05; **p<0.01; ***p<0.001

Table B.4: Buttigieg Models

	<i>Dependent variable:</i>		
	Don't Know Time Series	Poll Support Time Series	Decision Set Time Series
	(1)	(2)	(3)
Decision Set Lag			0.314*** (0.065)
Don't Know Lag	0.706*** (0.049)	-0.111 (0.084)	-0.005 (0.055)
Δ Media Attention	-0.0002 (0.0003)	-0.0003 (0.0004)	
Δ Media Attention Lag	0.00005 (0.0002)	-0.0001 (0.0003)	
Poll Support	-0.031 (0.055)		0.128** (0.041)
Poll Support Lag	-0.057 (0.057)	0.356*** (0.066)	0.089* (0.042)
Positive Debate Coverage	-0.700** (0.240)	0.276 (0.301)	
Positive Debate Coverage Lag	0.097 (0.244)	0.081 (0.301)	
Endorsements	-0.002 (0.004)	0.005 (0.005)	
Endorsements Lag	-0.005 (0.004)	0.0003 (0.005)	
Don't Know		-0.047 (0.084)	-0.310*** (0.052)
Ad Spending	-0.000 (0.000)	0.00002* (0.000)	
Constant	0.154*** (0.026)	0.122*** (0.034)	0.230*** (0.030)
Observations	228	228	228
R ²	0.654	0.342	0.649
Adjusted R ²	0.638	0.311	0.641
Residual Std. Error	0.037 (df = 217)	0.045 (df = 217)	0.028 (df = 222)
F Statistic	41.077*** (df = 10; 217)	11.259*** (df = 10; 217)	81.942*** (df = 5; 222)
<i>Note:</i>		*p<0.05; **p<0.01; ***p<0.001	

Appendix C: Supplemental Analyses for Chapter 4

C.1 Sanders Panel Model, Estimated using OLS and Heteroskedasticity-Robust SE

Table C.1: Sanders Panel Model

	<i>Dependent variable:</i>
	Wave 3 Sanders Vote
Wave 2 Sanders Vote	0.480** (0.074)
Wave 1 Sanders Vote	0.158* (0.060)
Δ Belief Sanders is emphasizing Climate Change $w2, w3$	0.109 (0.074)
Δ Belief Sanders is emphasizing Gun Control $w2, w3$	0.176* (0.065)
Δ Belief Sanders is emphasizing Health Care $w2, w3$	-0.030 (0.084)
Δ Belief Sanders is emphasizing Income Inequality $w2, w3$	0.052 (0.068)
Δ Belief Sanders is emphasizing Immigration $w2, w3$	0.028 (0.063)
Δ Sanders Electability	0.025 (0.013)
Δ Belief Sanders is emphasizing Climate Change $w1, w2$	0.061 (0.050)
Δ Belief Sanders is emphasizing Gun Control $w1, w2$	-0.019 (0.069)
Δ Belief Sanders is emphasizing Health Care $w1, w2$	0.076 (0.069)
Δ Belief Sanders is emphasizing Income Inequality $w1, w2$	0.049 (0.077)
Δ Belief Sanders is emphasizing Immigration $w1, w2$	0.177* (0.065)
Δ Sanders Electability $w1, w2$	0.035* (0.018)
Constant	0.211** (0.055)
Observations	193
<i>Note:</i> p<0.1; *p<0.05; **p<0.01	

References

- Abramowitz, A. I. (1989). Viability, Electability, and Candidate Choice in a Presidential Primary Election: A Test of Competing Models. *The Journal of Politics*, 51(4):977–992.
- Abramson, P. R., Aldrich, J. H., Paolino, P., and Rohde, D. W. (1992). Sophisticated voting in the 1988 presidential primaries. *American Political Science Review*, 86(1):55–69.
- Aldrich, J. H. and Alvarez, R. M. (1994). Issues and the presidential primary voter. *Political Behavior*, 16(3):289–317.
- Alvarez-Melis, D. and Saveski, M. (2016). Topic modeling in twitter: Aggregating tweets by conversations. In *Tenth international AAAI conference on web and social media*.
- Ansolabehere, S. and Iyengar, S. (1994). Of Horseshoes and Horse Races: Experimental Studies of the Impact of Poll Results on Electoral Behavior. *Political Communication*, 11:413–430.
- Atkeson, L. R. (1999). Sure, I voted for the winner! Overreport of the Primary Vote for the Party Nominee in the National Election Studies. *Political Behavior*, 21(3):197 – 215.

- Barberà, P., Casas, A., Nagler, J., Egan, P. J., Bonneau, R., Jost, J. T., and Tucker, J. A. (2019). Who Leads? Who Follows? Measuring Issue Attention and Agenda Setting by Legislators and the Mass Public Using Social Media Data. *American Political Science Review*, 113(4):883–901.
- Barker, D. C. (2005). Values, frames, and Persuasion in Presidential Nomination Campaigns. *Political Behavior*, 27(4):375–394.
- Barker, D. C. and Lawrence, A. B. (2006). Media Favoritism and Presidential Nominations: Reviving the Direct Effects Model. *Political Communication*, 23(1):41–59.
- Barker, D. C., Lawrence, A. B., and Tavits, M. (2006). Partisanship and the dynamics of candidate centered politics in american presidential nominations. *Electoral Studies*, 25(3):599–610.
- Bartels, L. M. (1988). *Presidential Primaries and the Dynamics of Public Choice*. Princeton University Press, Princeton, New Jersey.
- Berelson, B. R., Lazarsfeld, P. F., and McPhee, W. N. (1954). *Voting: A Study of Opinion Formation In a Presidential Campaign*. University of Chicago Press, Chicago, Illinois.
- Blais, A., Labbé-St-Vincent, S., Laslier, J.-F., Sauger, N., and Van Der Straeten, K. (2011). Strategic vote choice in one-round and two-round elections: An experimental study. *Political Research Quarterly*, 64(3):637–645.

- Boudreau, C. and McCubbins, M. D. (2010). The blind leading the blind: Who gets polling information and does it improve decisions? *Journal of Politics*, 72(2):513–527.
- Box-Steffensmeier, J., Darmofal, D., and Farrell, C. A. (2009). The aggregate dynamics of campaigns. *The Journal of Politics*, 71(1):309–323.
- Brady, H. E. (1993). Knowledge, strategy, and momentum in presidential primaries. *Political Analysis*, 5:1–38.
- Brownstein, R. (2019). How pundits may be getting electability all wrong. *The Atlantic*.
- Campbell, A., Converse, P. E., Miller, W. E., and Stokes, D. E. (1960). *The American Voter*. University of Chicago Press, Chicago, Illinois.
- Carey, J. M. and Polga-Hecimovich, J. (2006). Primary elections and candidate strength in latin america. *Journal of Politics*, 68(3):530–543.
- Carman, C. J. and Barker, D. C. (2005). State political culture, primary frontloading, and democratic voice in presidential nominations: 1972-2000. *Electoral Studies*, 24(4):665–687.
- Ceci, S. J. and Kain, E. L. (1982). Jumping on the bandwagon with the underdog: The impact of attitude polls on polling behavior. *Public opinion quarterly*, 46(2):228–242.

- Clinton, J. D., Engelhardt, A. M., and Trussler, M. J. (2019). Knockout Blows or the Status Quo? Momentum in the 2016 Primaries. *The Journal of Politics*, 81(3):997–1013.
- Cohen, M., Karol, D., Noel, H., and Zaller, J. (2008). *The Party Decides: Presidential Nominations Before and After Reform*. University of Chicago Press, Chicago, Illinois.
- Collingwood, L., Barreto, M. A., and Donovan, T. (2012). Early primaries, viability and changing preferences for presidential candidates. *Presidential Studies Quarterly*, 42(2):231–255.
- Conway, B. A., Kenski, K., and Wang, D. (2013). Twitter use by presidential primary candidates during the 2012 campaign. *American Behavioral Scientist*, 57(11):1596–1610.
- Conway, B. A., Kenski, K., and Wang, D. (2015). The Rise of Twitter in the Political Campaign: Searching for Intermedia Agenda-Setting Effects in the Presidential Primary. *Journal of Computer-Mediated Communication*, 20(4):363–380.
- Deltas, G., Herrera, H., and Polborn, M. K. (2016). Learning and Coordination in the Presidential Primary System. *The Review of Economic Studies*, 83(4):1544–1578.
- Dolan, K. (2010). The impact of gender stereotyped evaluations on support for women candidates. *Political Behavior*, 32:69–88.

- Dolan, K. (2014). Gender stereotypes, candidate evaluations, and voting for women candidates: What really matters? *Political Research Quarterly*, 67(1):96–107.
- Dowdle, A. J., Adkins, R. E., and Steger, W. P. (2009). The Viability Primary: Modeling Candidate Support before the Primaries. *Political Research Quarterly*, 62(1):77–91.
- Drezner, D. (2020). Can the democratic party still decide? *The Washington Post*.
- Edgerly, S., Thorson, K., and Wells, C. (2018). Young citizens, social media, and the dynamics of political learning in the US Presidential Primary Election. *American Behavioral Scientist*, 62(8):1042–1060.
- Erikson, R. S. and Wlezien, C. (2008). Are political markets really superior to polls as election predictors? *Public Opinion Quarterly*, 72(2):190–215.
- Erikson, R. S. and Wlezien, C. (2012a). Markets vs. polls as election predictors: An historical assessment. *Electoral Studies*, 31(3):532–539.
- Erikson, R. S. and Wlezien, C. (2012b). Markets vs. polls as election predictors: An historical assessment. *Electoral Studies*, 31(3):532–539.
- Erikson, R. S. and Wlezien, C. (2012c). *The Timeline of Presidential Elections: How Campaigns Do (and Do Not) Matter*. University of Chicago Press, Chicago, Illinois.

Garrison, J. (2020). Elizabeth warren ends her presidential campaign, holds off on endorsement. *USA Today*.

Geer, J. G. (1988). Assessing the representativeness of electorates in presidential primaries. *American Journal of Political Science*, 32(4):929–945.

Gerber, A. S., Gimpel, J. G., Green, D. P., and Shaw, D. R. (2011). How large and long-lasting are the persuasive effects of televised campaign ads? results from a randomized field experiment. *American Political Science Review*, 105(1):135–150.

Goldman, S. K. (2018). Fear of Gender Favoritism and Vote Choice during the 2008 Presidential Primaries. *The Journal of Politics*, 80(3):786–799.

Gopoian, J. D. (1982). Issue preferences and candidate choice in presidential primaries. *American Journal of Political Science*, 26(3):523–546.

Grafstein, R. (2003). Strategic Voting in presidential primaries: Problems of explanation and interpretation. *Political Research Quarterly*, 56(4):513–519.

Haltiwanger, J. (2020a). 30nominee. *Business Insider*.

Haltiwanger, J. (2020b). Bernie sanders has defied powerful critics and risen from long-shot 2016 candidate to 2020 democratic frontrunner. *Business Insider*.

- Haynes, A. A., Gurian, P.-H., Crespín, M. H., and Zorn, C. (2004). The Calculus of Concession: Media Coverage and the Dynamics of Winnowing in Presidential Nominations. *American Politics Research*, 32(3):310–337.
- Haynes, A. A. and Rhine, S. L. (1998). Attack politics in presidential nomination campaigns: An examination of the frequency and determinants of intermediated negative messages against opponents. *Political Research Quarterly*, 51(3):691–721.
- Herrera, R. (1999). The origins of opinion of american party activists. *Party Politics*, 5(2):237–252.
- Hirano, S., Lenz, G. S., Pinkovskiy, M., and Snyder Jr., J. M. (2015). Voter learning in state primary elections. *American Journal of Political Science*, 59(1):91–108.
- Huckfeldt, R. and Sprague, J. (1991). *Citizens, Politics, and Social Communication: Information and Influence in an Election Campaign*. Cambridge University Press, New York, New York.
- Jackman, S. and Vavreck, L. (2010). Primary Politics: Race, Gender, and Age in the 2008 Democratic Primary. *Journal of Elections, Public Opinion and Parties*, 20(2):153–186.
- Jacobsmeier, M. L. (2015). From black and white to left and right: Race, perceptions of candidates’ ideologies, and voting behavior in u.s. house elections. *Political Behavior*, 37:595–621.

- Khan, U. and Lieli, R. P. (2018). Information flow between prediction markets, polls and media: Evidence from the 2008 presidential primaries. *International Journal of Forecasting*, 34(4):696–710.
- Knight, B. and Schiff, N. (2010). Momentum and Social Learning in Presidential Primaries. *Journal of Political Economy*, 118(6):1110–1150.
- Lau, R. R. (2013). Correct Voting in the 2008 U.S. Presidential Nominating Elections. *Political Behavior*, 35(2):331–355.
- Maestas, C. D., Buttice, M. K., and Stone, W. J. (2014). Extracting Wisdom from Experts and Small Crowds: Strategies for Improving Informant-based Measures of Political Concepts. *Political Analysis*, 22(3):354–373.
- McCann, J. A., Partin, R. W., Rapoport, R. B., and Stone, W. J. (1996). Presidential Nomination Campaigns and Party Mobilization: An Assessment of Spillover Effects. *American Journal of Political Science*, 40(3):756.
- McCann, J. A., Rapoport, R. B., and Stone, W. J. (1999). Heeding the Call: An Assessment of Mobilization into H. Ross Perot’s 1992 Presidential Campaign. *American Journal of Political Science*, 43(1):1.
- McGregor, S. C., Mourão, R. R., and Molyneux, L. (2017). Twitter as a tool for and object of political and electoral activity: Considering electoral context and variance among actors. *Journal of Information Technology & Politics*, 14(2):154–167.

- Mirhosseini, M. R. (2015). Primaries with strategic voters: trading off electability and ideology. *Social Choice and Welfare*, 44(3):457–471.
- Muddiman, A., McGregor, S. C., and Stroud, N. J. (2019). (Re)Claiming Our Expertise: Parsing Large Text Corpora With Manually Validated and Organic Dictionaries. *Political Communication*, 36(2):214–226.
- Mutz, D. (1992). Impersonal influence: Effects of representations of public opinion on political attitudes. *Political Behavior*, 14(2):89–122.
- Mutz, D. C. (1997). Mechanisms of Momentum: Does Thinking Make It So? *The Journal of Politics*, 59(1):104–125.
- Norpoth, H. and Perkins, D. F. (2011). War and momentum: The 2008 presidential nominations. *PS: Political Science & Politics*, 44(3):536–543.
- Norrander, B. (1986). Correlates of vote choice in the 1980 presidential primaries. *Journal of Politics*, 48(1):156–166.
- Norrander, B. (1996). Presidential Nomination Politics in the Post-Reform Era. *Political Research Quarterly*, 49(4):875–915.
- Panetta, G. (2020). The 2020 democratic primary may look like a three-way race, but many presidential nominees started at the back of the pack. *Business Insider*.
- Paolino, P. and Shaw, D. R. (2001). Lifting the hood on the straight-talk express. *American Politics Research*, 29(5):483–506.

Pasek, J. and Dailey, J. (2018). Why don't tweets consistently track elections? lessons from linking twitter and survey data streams. In Stroud, N. J. and McGregor, S., editors, *Digital Discussions: How Big Data Informs Political Communication*.

Petrocik, J. R. (1996). Issue ownership in presidential elections, with a 1980 case study. *American Journal of Political Science*, 40(3):825–850.

Pickup, M. and Evans, G. (2013). Addressing the endogeneity of economic evaluations in models of political choice. *Public Opinion Quarterly*, 77(3):735–754.

Popkin, S. L. (1991). *The Reasoning Voter*. University of Chicago Press, Chicago, Illinois.

Rapoport, R. B., Stone, W. J., and Abramowitz, A. I. (1991). Do Endorsements Matter? Group Influence in the 1984 Democratic Caucuses. *American Political Science Review*, 85(1):193–203.

Relman, E. (2020). Southern black voters carry joe biden to huge victories on super tuesday. *Business Insider*.

Rickershauser, J. and Aldrich, J. H. (2007). It's the electability, stupid, or maybe not? electability, substance, and strategic voting in presidential primaries. *Electoral Studies*, 26(2):371–380.

- Rogstad, I. (2016). Is Twitter just rehashing? Intermedia agenda setting between Twitter and mainstream media. *Journal of Information Technology & Politics*, 13(2):142–158.
- Russell, F. M., Hendricks, M. A., Choi, H., and Stephens, E. C. (2015). Who Sets the News Agenda on Twitter?: Journalists’ posts during the 2013 US government shutdown. *Digital Journalism*, 3(6):925–943.
- Ryan, J. M. (2018). Partisan Dynamics in Presidential Primaries and Campaign Divisiveness. *American Politics Research*, 46(5):834–868.
- Shaw, D. R. (1999). A study of presidential campaign event effects from 1952 to 1992. *The Journal of Politics*, 61(2):387–422.
- Soroka, S. N., Stecula, D. A., and Wlezien, C. (2015). It’s (change in) the (future) economy, stupid: Economic indicators, the media, and public opinion. *American Journal of Political Science*, 59(2):457–474.
- Steger, W. P. (2007). Who Wins Nominations and Why?: An Updated Forecast of the Presidential Primary Vote. *Political Research Quarterly*, 60(1):91–99.
- Steger, W. P. (2008). Interparty Differences in Elite Support for Presidential Nomination Candidates. *American Politics Research*, 36(5):724–749.
- Steinskog, A. O., Therkelsen, J. F., and Gambäck, B. (2017). Twitter Topic Modeling by Tweet Aggregation. In *Proceedings of the 21st Nordic Con-*

ference of Computational Linguistics, pages 77–86, Gothenburg, Sweden. Linköping University Electronic Press.

Stone, W. J., Atkeson, L. R., and Rapoport, R. B. (1992a). Turning On or Turning Off? Mobilization and Demobilization Effects of Participation in Presidential Nomination Campaigns. *American Journal of Political Science*, 36(3):665.

Stone, W. J. and Rapoport, R. B. (1994). Candidate Perception among Nomination Activists: A New Look at the Moderation Hypothesis. *The Journal of Politics*, 56(4):1034–1052.

Stone, W. J., Rapoport, R. B., and Abramowitz, A. I. (1992b). Candidate Support in Presidential Nomination Campaigns: The Case of Iowa in 1984. *The Journal of Politics*, 54(4):1074–1097.

Stone, W. J., Rapoport, R. B., and Atkeson, L. R. (1995). A Simulation Model of Presidential Nomination Choice. *American Journal of Political Science*, 39(1):135.

Swearingen, C. D., Stiles, E., and Finneran, K. (2019). Here’s looking at you: Public versus elite driven models of presidential primary elections. *Social Science Quarterly*.

Terkildsen, N. (1993). When white voters evaluate black candidates: The processing implications of candidate skin color, prejudice, and self-monitoring. *American Journal of Political Science*, 37(4):1032–1053.

- Traugott, M. W. and Wlezien, C. (2009). The Dynamics of Poll Performance during the 2008 Presidential Nomination Contest. *The Public Opinion Quarterly*, 73(5):866–894.
- Tversky, A. (1972). Elimination by aspects: A theory of choice. *Psychological Review*, 79(4):281–299.
- Utych, S. M. and Kam, C. D. (2014). Viability, Information Seeking, and Vote Choice. *The Journal of Politics*, 76(1):152–166.
- Wattier, M. J. (1983). The simple act of voting in 1980 democratic presidential primaries. *American Politics Quarterly*, 11(3):267–291.
- Wells, C., Cramer, K. J., Wagner, M. W., Alvarez, G., Friedland, L. A., Shah, D. V., Bode, L., Edgerly, S., Gabay, I., and Franklin, C. (2017). When We Stop Talking Politics: The Maintenance and Closing of Conversation in Contentious Times: Political Talk in Contentious Times. *Journal of Communication*, 67(1):131–157.
- Williams, D. C., Weber, S. J., Haaland, G. A., Mueller, R. H., and Craig, R. E. (1976). Voter decisionmaking in a primary election: An evaluation of three models of choice. *American Journal of Political Science*, 20(1):37–49.
- Wlezien, C. (2016). On causality in the study of valence and voting behavior: An introduction to the symposium. *Political Science Research and Methods*, 4(1):195–197.

Wlezien, C., Franklin, M., and Twiggs, D. (1997). Economic perceptions and vote choice: Disentangling the endogeneity. *Political Behavior*, 19(1):7–17.

Wolfers, J. and Zitzewitz, E. (2004). Prediction markets. *The Journal of Economic Perspectives*, 18(2).

Young, L. and Soroka, S. (2012). Affective news: The automated coding of sentiment in political texts. *Political Communication*, 29(2):205–231.