

Measuring Agenda Setting in Interactive Political Communication

Erin L. Rossiter University of Notre Dame

Abstract: *Although strategies exist to measure actors' efforts to set policy, media, and lawmaking agendas, political scientists lack a method for identifying and accurately measuring another form of agenda setting that lies under the surface anytime two people talk. Within interactions, such as debates, deliberations, and discussions, actors can set the agenda by shifting others' attention to their preferred topics. In this article, I use a topic model that locates where topic shifts occur within an interaction in order to measure the relative agenda-setting power of actors. Validation exercises show that the model accurately identifies topic shifts and infers coherent topics. Three empirical applications also validate the agenda-setting measure within different political settings: U.S. presidential debates, in-person deliberations, and online discussions. These applications show that successfully setting the agenda can shape an interaction's outcomes, demonstrating the importance of continued research on this form of agenda setting.*

Verification Materials: The materials required to verify the computational reproducibility of the results, procedures, and analyses in this article are available on the *American Journal of Political Science* Dataverse within the Harvard Dataverse Network, at: <http://dx.doi.org/10.7910/DVN/MYM5OB>.

Who sets the agenda during political debates, deliberations, and discussions is fundamental to the study of political discourse but remains an elusive concept to quantify. Although strategies exist to quantify how issues make their way to broader policy and media agendas, we lack a systematic way to measure the agenda setting that occurs when actors *interact* with each other. Yet, measuring agenda setting in political interactions is important because their purpose is often to confine the set of issues relevant for downstream stages of politics. Thus, interactive communication is an ideal venue for actors to advance issues on their personal agendas. In this article, I seek to fill this gap in the study of agenda setting by drawing on research from computer science to measure actors' agenda setting during political interactions. Throughout my applications, I demonstrate the importance of studying agenda setting in debates, de-

liberations, and discussions by showing that it can shape the outcomes of these political interactions.

Agenda setting in interactive political communication is relevant to both formal and informal political settings. Presidential debates, Congressional committee hearings, and Supreme Court oral arguments are all examples of formalized interactions that are embedded in the framework of American government to facilitate decision making. Indeed, research shows these interactions have consequences on political outcomes, such as the ability of a lawyer's oral argument to sway Supreme Court votes (Johnson, Wahlbeck, and Spriggs 2006). Moreover, informal interactions are ubiquitous in politics, for example, as lawmakers and lobbyists interact during the process of crafting and voting on legislation. Citizens also experience politics via informal interactions with others, and research shows that engaging in political

Erin L. Rossiter, Assistant Professor, Department of Political Science, University of Notre Dame, 2077 Jenkins Nanovik Halls, Notre Dame, IN 46556 (erossite@nd.edu).

Valuable feedback for this project was provided by Political Data Science Lab members at Washington University in St. Louis. I thank participants at VIM 2019, New Faces in Political Methodology XI, Text as Data 2018, APSA 2018, PolMeth 2018, and SPSA 2018 for helpful comments. I am very grateful to Christopher Karpowitz, Hans Hassell, Taylor Carlson, Jaime Settle, Amber Boydston, Rebecca Glazier, and Matthew Pietryka for generously sharing data. Funding for elements of this project was provided by the National Science Foundation (SES-1558907, SES-1423788, SES-1938811) and the Weidenbaum Center on the Economy, Government, and Public Policy at Washington University in St. Louis. The human subjects research involving online discussions was approved by the Washington University in St. Louis Institutional Review Board (ID #201904055).

American Journal of Political Science, Vol. 00, No. 0, XXXX 2021, Pp. 1–15

©2021, Midwest Political Science Association

DOI: 10.1111/ajps.12653

discussion can influence subsequent behaviors like vote choice (Beck et al. 2002).

Because political interactions are an opportunity to set the stage for subsequent decisions, outcomes, and behaviors, I expect actors will seek to exercise power during them. It has long been understood that agenda setting is fundamental to understanding the political process (e.g., Cobb and Elder 1972; Kingdon 1984; Schattschneider 1960) and an important source of political power (e.g., Bachrach and Baratz 1962). And in this article, I consider how agenda setting is one form of power actors can exercise within political interactions, as well.

Specifically, actors can set the agenda during an interaction through attempts to shift the discussion to preferred topics. Like familiar forms of agenda setting in the literature, power is derived from shaping what issues do (and do not) receive attention by others. But unlike previously studied venues for agenda setting, interactions are uniquely a social experience. When engaging in a debate, deliberation, or discussion, a complex social exercise is underway as actors negotiate who is speaking and what is being discussed. Therefore, actors have the opportunity to influence the agenda as it develops in real time.

Take, for example, an interactive setting in which the fight over the agenda is particularly evident—U.S. presidential debates. Candidates seek to set the agenda by shifting the debate toward topics they “own” (Boydston, Glazier, and Pietryka 2013; Petrocik 1996). The following lines from the first 2016 general election presidential debate demonstrate Hillary Clinton shifting the agenda to a preferred topic.

Holt: We are at—we are at the final question.

Clinton: Well, one thing. One thing, Lester.

Holt: Very quickly, because we’re at the final question now.

Clinton: You know, he tried to switch from looks to stamina. But this is a man who has called women pigs, slobes and dogs, and someone who has said pregnancy is an inconvenience to employers, who has said...

Lester Holt, as the moderator, tried to introduce his final debate question. Yet, Clinton overpowered his efforts and successfully steered the final debate minutes toward an issue on *her* agenda—Trump’s history of degrading women. Clinton’s skill at setting the agenda not only affected what was discussed during the debate, but it also influenced subsequent media converge. News outlets re-

ported that this was a memorable moment from the first debate (Mason 2016; Ross 2016).

Although it is easy to see why measuring agenda setting in political interactions is important, standard approaches in political science offer no way to quantify Clinton’s role in setting the debate’s agenda. And more broadly, political scientists lack a systematic way to measure the relative agenda-setting power of actors in interactive settings. To be sure, political scientists have quantified other features of interactions. A common approach is to count easily observable quantities, such as the number of words spoken by a participant. Scholars have applied these measures to deliberations (e.g., Karpowitz and Mendelberg 2014), legislative committee hearings (e.g., Kathlene 1994), and Supreme Court oral arguments (e.g., Epstein, Landes, and Posner 2010). Although these count-based measures are useful for studying participation patterns, they fall short when the goal is a systematic measure of what issues make it on the agenda and which actors succeeded in getting them there.

In this article, I build upon these previous efforts in political science to quantify how political actors interact with each other, and I focus my efforts on agenda setting. Specifically, I leverage the text of interactions as data, and I use the parametric Speaker Identity for Topic Segmentation (SITS) model from the computer science topic segmentation literature (Nguyen, Boyd-Graber, and Resnik 2012; Nguyen et al. 2014). SITS extends Latent Dirichlet Allocation (LDA) (Blei, Ng, and Jordan 2003), a topic model used widely in political science, to simultaneously estimate three sets of latent quantities of interest: What topics are on the agenda, where shifts in the agenda occur, and each actor’s agenda-setting power (Nguyen 2015).

I proceed by first comparing agenda setting within political interactions to previously studied forms of agenda setting by the media or policy makers. I then outline SITS and present three validation exercises to show that the model can accurately identify where shifts in the agenda occur and can infer coherent topics. Then I use SITS to investigate agenda setting in three contexts. First, I replicate and extend findings in the literature about which candidates include the economy on their agendas during presidential debates (Boydston, Glazier, and Pietryka 2013; Vavreck 2009). Then, with in-person deliberations, I assess how agenda setting relates to several common measures of participation. I find that deliberators who set the agenda are more likely to shape the deliberation’s outcome, but that this relationship does not hold with the participation measures. Lastly, I perform a more direct test of my claim that agenda setting can shape the outcomes of political interactions, and I show that agenda setting correlates with achieving one’s

desired outcome in an online discussion. Taken together, these applications provide evidence for the validity of the agenda-setting measure across debates, in-person deliberations, and online discussions. I conclude with a discussion of the usefulness of studying agenda setting in interactions as an important form political power, and I provide suggestions for further areas of inquiry enabled by this method.

Conceptualizing Agenda Setting in Interactions

Before introducing the SITS model of agenda setting in the “A Model of Agenda Setting in Interactions” section, I will first discuss the broader study of agenda setting in political science, how agenda setting manifests in interactive communication, and a framework to motivate measurement of this important concept.

As mentioned, the concept of agenda setting is important to several political science literatures. For example, scholars investigate how the media’s agenda setting influences what issues the mass public perceives as important (e.g., McCombs and Shaw 1972). Moreover, there is a vast literature identifying the power of the media, Congress, the President, and other groups to set the policymaking agenda, and how issues rise to and fall from this agenda (e.g., Baumgartner and Jones 1993; Kingdon 1984). Scholars also investigate formal agenda-setting processes of specific institutions, such as how legislative parties seek control over which bills are considered on the floor by gaining agenda-setting powers through Congressional offices (Cox and McCubbins 2005) or how Supreme Court justices vote strategically on which cases are granted review (Black and Owens 2009). Regardless of the specific literature, “agenda setting” can be thought of as an influence over the set of issues that are (and are not) receiving attention, which in turn can influence the set of issues relevant for downstream outcomes and decision making. As such, agenda setting has been viewed as an important source of political power (e.g., Bachrach and Baratz 1962).

Building on these various traditions, I argue agenda setting is one form of power actors seek to exercise during one of the most basic, ubiquitous political activities—talking with others. Agenda setting is a source of power within interactions because of its twofold impact. First, agenda setting is an influence over the set of issues that are receiving attention by others, and therefore leads to an immediate control over the specific topics of discussion. And because of this, agenda setting can shape sub-

sequent outcomes and decision making (e.g., Riker 1986; Schattschneider 1960). For example, successfully setting the agenda in a deliberation might result in keeping certain issues off the table, which then obviates the risk of the issue rising to a vote. Or, conversely, effective agenda setting during a debate might raise the status of an otherwise overlooked issue, which then shapes what the media reports the next day.

However, an analogy to previously studied forms of agenda setting is limited when it comes to how, exactly, actors set the agenda during an interaction. That is because interactions are uniquely a social game. As actors navigate the waters of a social setting, who is speaking and what is being discussed are negotiated in real time. This inherent negotiation of the agenda presents an opportunity for an actor to set the agenda—to gain the floor, introduce their preferred topic, and maintain others’ attention on it.

To be sure, what agendas and agenda setting look like will vary depending on whether the interaction is an adversarial debate or a thoughtful policy deliberation, which have different goals and norms. For example, in U.S. presidential debates, two candidates may seek to shift attention to two distinct sets of issues. Or, in a parliamentary-style debate, actors make seek to shift focus to their argument. However, in a focused policy deliberation, actors may seek to advance competing framings of only a few issues. Either strategy—shifting attention to a new issue or shifting attention to new attributes of an issue—can constituent agenda setting depending on the purpose of the interaction and breadth of discussion.¹

Relatedly, it is important to note that agenda setting is just one form of power available to actors within political interactions. Certain actors may be perceived as more powerful when an interaction begins (e.g., Karpowitz and Mendelberg 2014). Or, actors may seek power through persuasive arguments (e.g., Wang et al. 2017) or other heresthetical strategies (Riker 1986). Agenda setting may not be the form of power most useful to every actor in every situation. Measuring other sources of power within interactions, and how they may or may not relate to agenda setting, is left to future research.

Prior Approaches

Although agenda setting is an important concept for understanding political interactive communication,

¹Considering both issues and frames parallels theory on the media’s first- and second-level agenda setting, where an emphasis on issues and attributes of those issues are both agenda-setting processes (e.g., Weaver 2007).

approaches currently used in political science to analyze interactions are not well-suited for systematically studying this concept.

First, political scientists have used hand-coding methods to analyze interactions, including the measurement of agenda setting (e.g., Boydston, Glazier, and Phillips 2013; Boydston, Glazier, and Pietryka 2013). Although hand coding is often considered a gold standard approach, it can have significant weaknesses. First, recruiting, adequately training, and compensating the work of research assistants can be prohibitively time consuming and costly, especially with a large corpus. Second, research shows that even high-quality coders can provide estimates that are unreliable (Mikhaylov, Laver, and Benoit 2012).

Quantitative analysis of the text of interactions is a more popular approach. Scholars often count directly observable and quantifiable behaviors such as the number of words spoken by participants (e.g., Epstein, Landes, and Posner 2010; Karpowitz, Mendelberg, and Shaker 2012; Kathlene 1994). Although this approach measures participation patterns, count-based measures are limited when the goal is to assess an interaction's agenda and who is influencing it.

To be sure, automated text analysis has been applied to measure the concept of agenda setting (Eggers and Spirling 2018; Quinn et al. 2010).² However, existing methods are not suited to studying agenda setting behavior of actors in *interactions*. Quinn et al. (2010) conceptualize the agenda as what issues are broadly gaining attention in the political arena and which are not. As a macrolevel measure of the agenda, it is not equipped to measure microlevel agenda setting within an interaction. Moreover, Eggers and Spirling (2016) conceptualize the agenda as the relative importance placed on issues over months and years, and measure an actor's ability to influence this long-term agenda. Thus, this measure is not a good fit for the task at hand, as I am interested in an actor's influence over the specific topics of discussion.

Conceptual Framework

Prior approaches are not well-suited for measurement of agenda setting within interactions. To overcome this, I build a conceptual framework of agenda setting in this section. I consider each speaker's agenda-setting ability, where shifts in the agenda occur, and the topical agenda

itself as latent quantities of interest. Then in the next section, I map this conceptual framework to a measurement strategy for these concepts using the text of the interactions as data.

First, I assume an actor seeks to advance a preferred set of topics in the discussion. Under this assumption, learning that an actor successfully shifts others' attention to a new topic provides information about their power over the interaction's agenda. Therefore, I operationalize agenda setting by examining *who* is setting the agenda by successfully shifting the topic of discussion.

There are three important points to clarify. First, I conceptualize agenda setting as a latent ability of an actor. Second, when exercised, agenda setting is a form of power as discussed in the "Conceptualizing Agenda Setting in Interactions" section. And third, a topic shift, although an *indicator* of an actor's agenda-setting ability, is not in and of itself power. Rather, agenda-setting power is also evidenced by others' attention on one's shifted-to topics (forfeiting the opportunity to set the agenda themselves). In this sense, agenda setting is inherently a relative concept. Therefore, an actor's agenda-setting power ought to be considered relative to the power of others who are fighting within the same limited time over the same space-constrained agenda.

To gain an understanding of each actors' agenda-setting ability, I discussed the need to locate where shifts in the agenda occur as an interaction unfolds. An interaction can be thought of as a sequence of different actors taking turns speaking; therefore, it is useful to look at the speaking-turn level for topic shifts. Additionally, I conceptualize shifts in topic as something that we cannot directly observe, but rather, something latent that needs to be inferred from the textual data. Shifts in topic are not directly observable, in part, because we simultaneously need to have a clear idea of the set of topics being discussed in the corpus.

Therefore, I consider an interaction's agenda—the set of issues and/or frames that arise during the interaction—as an additional latent quantity of interest. To infer the set of "topics" that make up the agenda, I adopt a similar strategy as prior research by using an unsupervised topic model of the text of the interactions to explore the issues (e.g., Grimmer 2010) and/or frames (e.g., Aslett et al. Forthcoming) within a text corpus.

In sum, identifying who can shift the topic and maintain attention on their introduced topics is one way to operationalize agenda setting. But, understanding who has changed the topic requires we know the set of topics being discussed and where they shift. Therefore, these three interrelated concepts all need to be inferred simultaneously from the text, which I turn to next.

² Also see Karpowitz and Mendelberg (2014) for a dictionary-based approach to identifying agenda setting when the relevant set of issues are clearly known and defined a priori.

FIGURE 1 SITS Data-Generating Process

- **For each speaker $m \in [1, M]$, draw a speaker topic shift probability $\pi_m \sim \text{Beta}(\gamma)$.**
- For each topic $k \in [1, K]$, draw a topic-word distribution $\phi_k \sim \text{Dir}(\beta)$.
- For each turn $t \in [1, T_d]$, in each discussion $d \in [1, D]$ **(with speaker $a_{d,t}$):**
 - **If $t = 1$, set the topic shift $l_{d,t} = 1$, otherwise draw $l_{d,t} \sim \text{Bernoulli}(\pi_{a_{d,t}})$.**
 - **If $l_{d,t} = 0$, set the topic distribution $\theta_{d,t} \equiv \theta_{d,t-1}$, otherwise draw $\theta_{d,t} \sim \text{Dir}(\alpha)$.**
 - For each word index $n \in [1, N_{d,t}]$:
 - Draw a topic $z_{d,t,n} \sim \text{Categorical}(\theta_{d,t})$.
 - Draw a word $w_{d,t,n} \sim \text{Categorical}(\phi_{z_{d,t,n}})$.

Note: Figure adapted from Nguyen et al. (2014). The data-generating process of the parametric Speaker Identity for Topic Segmentation model. Bold text indicates extensions from Latent Dirichlet Allocation.

A Model of Agenda Setting in Interactions

I measure agenda setting within interactions using the parametric SITS model (Nguyen, Boyd-Graber, and Resnik 2012; Nguyen et al. 2014). Specifically, SITS builds upon a familiar topic model in political science, LDA (Blei, Ng, and Jordan 2003), to account for and measure the latent concepts motivated in the “Conceptual Framework” section: the agenda, where shifts in the agenda occur, and the agenda-setting power of actors within a corpus of interactions. For ease of exposition, as I outline the model, I will refer to any single interaction as a “discussion,” each actor participating in a discussion as a “speaker,” and each uninterrupted utterance by a speaker as a “speaking turn.”

SITS Data-Generating Process

The SITS data-generating process is outlined in Figure 1. Because SITS follows in a line of research that extends topic models to estimate additional latent quantities of interest to political scientists (e.g., Grimmer 2010), I use bold text to denote extensions to LDA.³

First, for each speaker $m \in [1, M]$, their agenda-setting ability within the corpus (π_m) is drawn from a symmetric Beta distribution with parameter γ . Then, as with LDA, topics (ϕ_k), or probability distributions over the corpus vocabulary, are drawn for each of $k \in [1, K]$ from a symmetric Dirichlet distribution with param-

eter β . Then, a distribution over topics ($\theta_{d,t}$) needs to be drawn for each speaking turn $t \in [1, T_d]$ for each discussion $d \in [1, D]$. However, this part of the SITS generative process unfolds differently than LDA, because SITS seeks to find “segments,” which are sequences of speaking turns on the same set of topics. Because the first turn of a discussion inherently changes the topic, this is noted by setting a turn-level topic shift binary variable equal to one ($l_{d,t=1} = 1$). For all other turns, whether or not a shift in topic occurs is drawn from a Bernoulli distribution parameterized by the speaker’s agenda-setting measure ($\pi_{a_{d,t}}$, where $a_{d,t}$ is the observed speaker of turn t in discussion d). Therefore, whether or not a speaking turn changes the topic is influenced by its speaker’s latent agenda-setting ability. If a topic change is indicated, a *new* topic distribution is drawn from a symmetric Dirichlet distribution with parameter α . Otherwise the topic distribution from the previous turn *carries over* to the current turn ($\theta_{d,t} \equiv \theta_{d,t-1}$) indicating those speaking turns belong to the same segment. Then, identical to LDA, for each word index $n \in [1, N_{d,t}]$ in the speaking turn, a topic assignment ($z_{d,t,n}$) is drawn given the speaking turn’s distribution over topics and a word ($w_{d,t,n}$) is drawn given its assigned topic.

Note that the agenda-setting measure (π_m) is a speaker-level quantity. It describes the propensity of a speaker to shift topic when speaking. As such, this measure captures the theoretical quantity of interest as it accounts for how often a speaker shifts topic (imagine this as an “uptick” in the numerator) and how often a speaker is willing to maintain attention on others’ topics (an “uptick” in the denominator). Therefore, this measure is most meaningful when comparing the relative agenda-setting abilities of actors who are seeking to influence the same agenda in the same limited amount of time, as

³To make the data-generating process of LDA comparable to a text from an interaction, I consider each “document” in LDA as each speaking turn in an interaction.

motivated in the “Conceptualizing Agenda Setting in Interactions” section.

To estimate SITS in what follows, I use a Gibbs sampler written in Java by Viet An Nguyen that is available to the public (Nguyen 2014). In the Supporting Information (SI), Appendix A (SI p. 1) provides additional details regarding the sampler, Appendix B (SI pp. 1–3) details my preprocessing decisions for each corpus guided by metrics and tools in the text analysis literature (Denny and Spirling 2018), and Appendix C (SI pp. 3–5) details my approach to choosing hyperparameter values by relying on advice in the unsupervised topic modeling literature.

Validation Exercises

Before exploring applications of agenda setting across different political interactions, I present results from three validation exercises. The first exercise provides evidence that SITS can accurately identify where *latent shifts* in the agenda occur. Second, using crowdsourced human judgments, I validate that SITS can infer semantically meaningful *latent topics*. Finally, I assess the interrelated nature of where shifts in topic occur and the topics themselves by examining the resulting *segments* of an interaction. Using crowdsourced human judgments, I find that SITS segments are viewed as more coherent than segments derived from a hand-coding approach.

Latent Topic Shifts

Texts were generously shared by Jaime Settle and Taylor Carlson from a study conducted in the fall of 2015 examining disagreeable political discussion. Participants had an in-person discussion on several topics with a partner for approximately 10 min.⁴ Participants read a prompt on a screen and discussed the prompted topic for a prespecified length of time. At that point, the screen prompted the participants to stop discussing and wait for the next topic, producing sharp shifts between topics with known locations.⁵

This study contains 70 discussions among 140 participants. The conversations were an average of 38 turns long. I preprocessed the text by removing numbers, stemming, and removing infrequent terms. I also transformed all features to lower case and removed all punctuation. I

estimated three SITS chains from the data with randomly drawn starting values. I averaged the posterior mean for each turn-level variable across the three chains.⁶

To assess if SITS can accurately identify where latent shifts in topic occur, I classify a speaking turn as shifting topic if the posterior probability of a shift is greater than or equal to 0.50. I then compare where shifts were inferred by SITS to where topic changes were prompted by the researchers. SITS identifies 81.40% of the locations that begin a new prompted topic segment by the researcher.⁷ I check for an SITS-inferred topic shift within two speaking turns after a researcher-prompted shift because often after reading the prompt, the first few speaking turns would simply answer the prompt’s question by saying “yes,” “no,” or “do you want to go first?” Instead of classifying this as a topic shift, SITS would classify a subsequent speaking turn that actually began discussing the topic at hand as a shift.

This exercise demonstrates that SITS can accurately identify the speaking turns that should be attributed as shifting the agenda.⁸ Moreover, SITS provides a more nuanced view of how topics ebbed and flowed in these discussions than if we considered the locations of researcher-prompted topic changes as ground truth.

Latent Topics

Next I assess whether SITS can infer semantically coherent topics. Influential work in computer science proposes that crowdsourced tasks are more useful than traditional metrics to assess if a topic model returns semantically meaningful and distinct topics (Chang et al. 2009). Therefore, I use the “topic intrusion” task proposed by Chang et al. (2009) to validate the topics from a SITS model estimated on 20 U.S. general election presidential debates held between 1992–2016.⁹

⁶I estimate SITS with $K = 13$, $\alpha = 1/K$, $\beta = 0.1$, and $\gamma = 1$, as outlined in Appendix C (SI pp. 3–4).

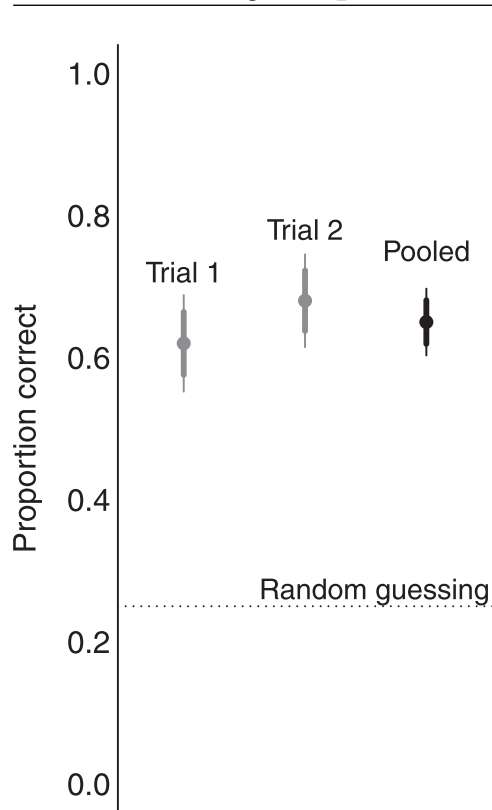
⁷I have no way to know whether participants shifted topic *within* a researcher-prompted topic segment. Participants could, and likely did, change topic when not prompted to. Therefore, for this exercise, I only validate SITS against locations where topic changes were known to occur.

⁸An additional validation study in Appendix D (SI pp. 6–7) compares SITS to automated text analysis methods in political science, and shows that these commonly used methods do not perform well when adapted to the task of identifying where shifts in topic occur within an interaction.

⁹For this exercise, I estimated SITS with $K = 44$, $\alpha = 1/K$, $\beta = 0.1$, and $\gamma = 1$ before available metrics indicated a preferable K would be $K = 52$ for the analysis in section “Electoral Debates.” Substantive results in section “Electoral Debates” replicate in all elections except 2008 using either hyperparameter choice.

⁴Appendix D (SI p. 6) outlines researcher provided topics.

⁵I watched video recordings of each discussion, and the participants complied by discussing the prompted topics.

FIGURE 2 Test of Semantically Meaningful Topics

Note: Proportion of correct answers to the topic intrusion task. Thick line shows 80% confidence interval, and thin line shows 95% confidence interval. Gray bars show the identical repeated trials of 200 tasks each. Black bar represents the result when pooling both trials. A difference of proportions test indicates that Trial 1 and Trial 2 are not significantly different ($p = 0.25$).

The topic intrusion task presents the human judge with a document, in this case a segment inferred by SITS. The judge is also presented with four word sets. Three of these word sets represent the three highest probability topics for the segment. The fourth word set is the intruder, drawn randomly from the segment's low probability topics. Each word set contains the top eight frequent and exclusive (FREX) topwords for the topic (Roberts et al. 2014). I set up the topic intrusion task for 200 randomly drawn segments from the debates. Then, Amazon Mechanical Turk (MTurk) workers were asked to choose which word set was most unrelated to the passage. In line with recent work on validation procedures for topic models by Ying, Montgomery, and Stewart (2019), I ran two trials of the same 200 tasks. Figure 2 plots the results for each trial separately as well as the pooled result.

Workers completed 62% and 68% of the tasks correctly in Trial 1 and Trial 2, respectively. A difference of proportions test indicates that Trial 1 and Trial 2 are not significantly different ($p = 0.25$). This result is comparable to one, and better than three, of four models assessed by Ying, Montgomery, and Stewart (2019) using the topic intrusion task. In all, human coders and SITS largely agree about which topics are and are not associated with the inferred segments of the debates.

Latent Segments

Next, I validate that SITS can identify coherent segments of an interaction. To do so, I follow a procedure proposed by Grimmer and King (2011) to evaluate “cluster quality,” which is the similarity of the documents (here, segments of the debates) estimated to belong to the same cluster (here, having similar topic distributions). Importantly, I evaluate SITS segments against segments derived from a hand-coding approach. Hand-coded data are from Boydston, Glazier, and Pietryka (2013) for the 1992, 2004, and 2008 U.S. general election presidential debates. Boydston, Glazier, and Pietryka hand code several variables from the debate transcripts, including the topic of each question posed to the candidates and the topic of each phrase in the candidates' responses. Then, they deem a candidate as going “off-topic” and thus, engaging in agenda-setting behavior, if the phrase's topic does not correspond to the question's topic.

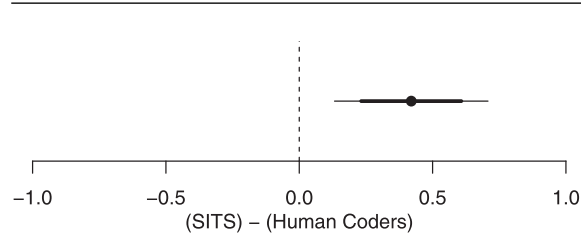
Comparing the segments inferred by SITS to those derived from hand coding required five steps. First, I determined where topic changes occurred (and thus, formed segments of the debates) according to each method. Second, I determined the similarity of these segments according to each method's topic assignments.¹⁰ Third, I set up the exercise outlined by Grimmer and King (2011). Separately with the segments from the SITS and the hand-coding approaches, I drew 25 random pairs of segments with the same most-assigned topic and 25 random pairs of segments with a different most-assigned topic.¹¹ Fourth, four unique MTurk workers rated the similarity of the segments within each pair on a 3-point scale: (1) unrelated, (2) loosely related, or (3) closely related.¹² Of interest is each method's “cluster quality,”

¹⁰Appendix D (SI p. 8) details the specific steps taken to estimate shifts in topic and cluster the resulting segments by topic.

¹¹Appendix D (SI p. 10) shows that the results are not likely to be due to differences in the length of the segments derived by the two approaches.

¹²Appendix D (SI pp. 8–10) provides details of the task's instructions and an example of a pair of segments.

FIGURE 3 Coherence of Topics and Segments Inferred by SITS versus Hand-Coding Approach



Note: Difference in the Grimmer and King (2011) cluster quality measure between the SITS approach and a hand-coded approach to segmentation and topic assignment. Dot shows point estimate, thick line shows 80% confidence interval, and thin line shows 95% confidence interval.

which is “the average similarity of pairs of documents from the same cluster minus the average similarity of pairs of documents from different clusters, as judged by human coders one pair at a time” (Grimmer and King 2011, p. 5).

The fourth and final step is to use difference in means to compare the two methods. Figure 3 plots this point estimate along with the 80% (thick line) and 95% (thin line) confidence interval. The estimated difference between approaches is positive and significant. Therefore, the Grimmer and King (2011) evaluation suggests SITS can infer segments that are even more coherent than those derived from a hand-coding approach.

Applications

I next explore agenda setting within three different interactive political communication. First, with U.S. presidential debates, I show that SITS can be used to test theories about the topics actors should favor in electoral debates. Second, using in-person deliberations, I shift emphasis from the topical agenda to who can set the agenda. I validate that participants who SITS identifies as setting the agenda indeed shift attention to their ideas, and do not do so by simply interrupting or out-talking others. And third, using a novel online discussion study, I show that agenda setting can shape the outcomes of political interactions.

Electoral Debates

I first assess the contribution of SITS to the measurement of agendas and agenda setting in electoral debates.

Boydston, Glazier, and Pietryka (2013) examined how theories of agenda setting throughout presidential campaigns may translate to agenda setting within a debate (see also Boydston, Glazier, and Phillips 2013). This research relied on a rigorous hand-coded content analysis, and therefore reasonably analyzed only three elections. I use SITS to replicate and extend their findings regarding which candidates’ agendas feature the economy by analyzing all general election presidential debates from 1992 to 2016.

Vavreck (2009) crafts a typology to describe when presidential candidates should focus their campaigns on the economy. Clarifying candidates are those that benefit from doing so (outparty candidates in a bad economy and in-party candidates in a good economy), whereas insurgent candidates do not (out-party candidates in a good economy and in-party candidates in a bad economy).

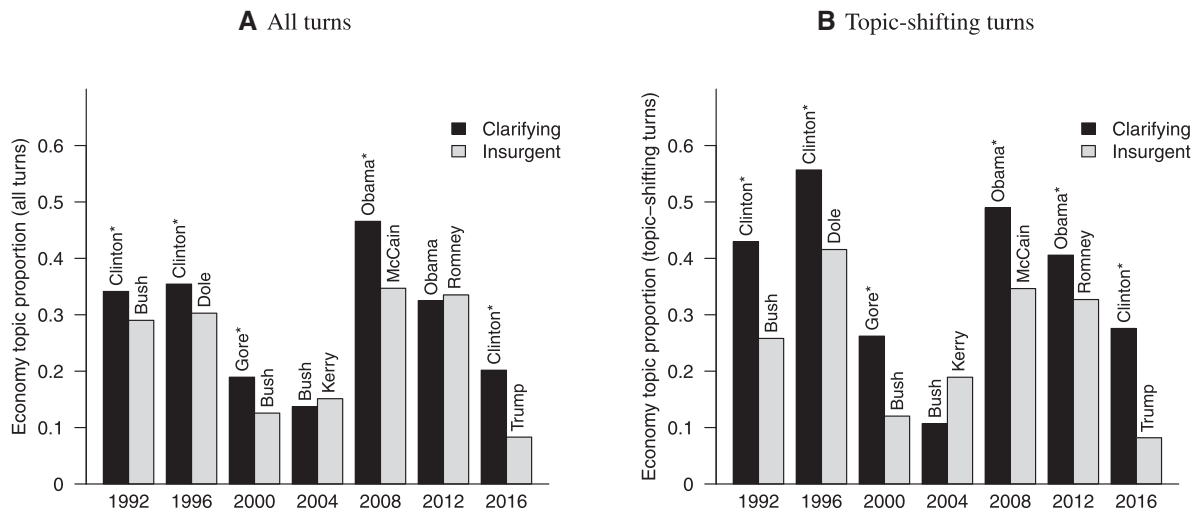
Boydston, Glazier, and Pietryka assess how these patterns manifest in the context of presidential debates. As discussed in the “Latent Segments” section, Boydston, Glazier, and Pietryka hand code the 1992, 2004, and 2008 debates to assess two general hypotheses. In terms of Vavreck’s typology, the two hypotheses are that the clarifying candidate should (1) talk more about the economy than the insurgent candidate and (2) focus more of their agenda-setting efforts on the economy than insurgent candidate.

I preprocessed the text by removing capitalization, punctuation, and numbers. I also remove a set of stopwords that included common English stopwords, annotations to the transcripts (such as “applause”), and the name and titles of all speakers. I also perform stemming and remove infrequent terms.¹³ The resulting corpus contained 944 unique terms used across 3,818 speaking turns in the 20 debates. I estimated three SITS chains from the data with randomly drawn starting values.¹⁴ I use iterations from all three chains to estimate posterior means of turn-level topic shifts, and I use FREX stopwords to describe topics from the best performing model.

Figure 4 describes patterns of how much clarifying and insurgent candidates discussed the economy during the debates. The barplots show topic proportions for

¹³Appendix B (SI pp. 1–3) discusses the preprocessing in detail. Because preprocessing decisions may have consequences for downstream, substantive results, I use the `preText` R package to assess this (Denny and Spirling 2018). I find that results may be sensitive to removing a common set of stopwords, so I replicate all results when retaining stopwords in Appendix E (SI pp. 10–11).

¹⁴I estimate SITS with $K = 52$, $\alpha = 1/K$, $\beta = 0.1$, and $\gamma = 1$, as outlined in Appendix C (SI pp. 3–5). Appendix E (SI pp. 10–11) provides details on convergence and model selection.

FIGURE 4 Clarifying Candidates Talk More about the Economy

Note: Topic proportions for economy topics in all of a candidate's speaking turns (a) and all of the turns in which they shifted topic (b). One-sided difference of proportions tests indicate that the clarifying candidate spoke about and set the agenda to the economy more (* $p < 0.05$), except in the 2004 and 2012 elections.

economy topics for all of a candidate's speaking turns in (a) and all topic-shifting turns in (b).¹⁵ I use Vavreck's classification of candidates as clarifying or insurgent for the 1992–2008 elections (Vavreck 2009, p. 38). Across those elections, we see support for the hypotheses that clarifying candidates talk more about the economy in (a), and that they engage in agenda setting in order to shift the course of the debate to the economy in (b), except in 2004 and 2012.

In particular, the results for 1992, 2004, and 2008 are in line with results from the hand-coded content analysis. Boydston, Glazier, and Pietryka (2013) note that in 1992 and 2008, the economy was a salient issue, but in 2004, defense was uniquely more important to the public than the economy. These patterns hold in Figure 4, as we see the 2004 election discussed the economy less than any other election.

Finally, the 2012 and 2016 elections were both outside the purview of the Boydston, Glazier, and Pietryka (2013) and Vavreck (2009) analyses. In 2012, both candidates were stressing the economy and running what looked like a clarifying campaign (Sides and Vavreck 2014), even if Obama was technically the clarifying candidate as the incumbent in a slowly growing economy. We see that holds in Figure 4 as the candidates spoke about and shifted attention to the economy in nearly

equal amounts. Finally, in 2016, Clinton would be considered the clarifying candidate (the in-party candidate in a growing economy), even if the economic concerns were not the main campaign issue or motivation for voters (Sides, Tesler, and Vavreck 2019). We see she also spoke about the economy more than Trump—both in total and when she was setting the agenda.

In-Person Deliberations

To date, political scientists have lacked a systematic way to measure the agenda-setting features of interactions. Instead, participation is often measured in these contexts using count-based measures, such as the number of words spoken by a participant. However, participating in an interaction is a necessary but not sufficient component of setting its agenda. Setting the agenda requires participation that also shifts the set of topics currently receiving attention. Therefore, I assess the relationship between status quo, count-based measures of participation, and the proposed measure of agenda setting. Then, I examine how participation and agenda setting correlate with successfully shaping the outcome of the deliberation. I find that agenda setting positively correlates with shaping the deliberation's outcome, although I fail to find any correlation between the outcome and participation measures.

Deliberation texts were generously shared by Christopher Karpowitz and Hans Hassell from a pilot study conducted in June of 2016 examining the

¹⁵From the FREX topwords, I judged which topic(s) were on the economy, including unemployment, taxes, and more. Appendix E (SI pp. 12–13) presents FREX topwords and aggregate topic-shifting patterns for all 52 topics.

TABLE 1 Agenda Setting Does Not Correlate with Quantity of Participation

	Agenda setting	Proportion of comments	Proportion of talk time	Proportion of interruptions	Composite measure
Agenda setting	1.00	−0.07	−0.03	−0.33*	−0.19
Proportion of comments		1.00	0.88*	0.21	0.93*
Proportion of talk time			1.00	−0.04	0.81*
Proportion of interruptions				1.00	0.52*
Composite measure					1.00

Note: *Significant correlations, $p < 0.05$.

effect of stress on deliberation participation.¹⁶ For this study, Brigham Young University (BYU) students were recruited to discuss the BYU Dress and Grooming Standards (DGS), a specific set of rules governing the appearance of all students and staff at the university.

The study included 10 discussion groups, each composed of four members. Participants first completed a prediscussion survey regarding the DGS. Participants then engaged in a 25-minute discussion and were tasked with agreeing on up to two recommended changes to the DGS. Participants voted on their group's recommendations after the discussion. Recommendations gaining a majority of the postdiscussion votes would be sent to the Honor Code office, with no guarantee that the changes would be implemented.

Although these deliberations did not involve a conventional policy issue, this application is a useful case study as it involves an important issue for participants with relatively strong and conflicting attitudes. BYU students have petitioned and recently protested the Honor Code and how it is enforced, including the DGS portion, with their efforts even making national news (Turkewitz 2014; Levin 2019). Moreover, not only do pretreatment survey responses suggest participants held conflicting views, but participants often expressed these conflicting perspectives during the deliberation. Although the implications of this application are unclear for deliberative settings, which may involve participants with weaker attitudes and a lower stakes outcome, I will turn to this kind of setting in the "Online Discussions" section.

I preprocessed the text removing capitalization, punctuation, and numbers. I also perform stemming and remove infrequent terms.¹⁷ The resulting corpus con-

tained 667 unique terms used across 899 speaking turns in 10 deliberations. I estimated three SITS chains from the data with randomly drawn starting values.¹⁸ I use iterations from all three chains to estimate posterior means of the agenda-setting measures.

I first investigate how agenda setting may correlate with commonly used participation measures. Table 1 visualizes a correlation matrix between agenda setting and three count-based participation measures: proportion of the group's (1) comments, (2) speaking time, and (3) interruptions made by a participant. I also include a composite measure created by averaging a standardized version of each count-based measure.

First, there is no evidence of a correlation between the first two count-based measures. Second, we see a negative correlation between how often a participant interrupts others and their agenda setting. Taken together, these results suggest agenda setting is not achieved by out-talking or interrupting others. Finally, we see that aggregating the count-based measures does not provide additional leverage for measuring agenda setting.

To further explore agenda setting and its impact in these deliberations, I assess how agenda setting relates to the deliberation's outcome—the group's DGS policy recommendations. For each deliberation, I note what was eventually recorded as a group recommendation, and I then trace the recommendation back to which participant introduced it during the deliberation.¹⁹ I then examine if agenda setting and participation are correlated with shaping the deliberation's outcomes.

Table 2 presents coefficients from logistic regression models with clustered standard errors at the group level in parentheses. Additionally, I include a variable indicating group membership to estimate an intercept shift

¹⁶I find no evidence of a treatment effect; therefore, I do not include the treatment as a variable in subsequent analyses.

¹⁷See Footnote 13. I find that results may be sensitive to removing stopwords, so I replicate all results when retaining stopwords in Appendix F (SI p. 17).

¹⁸I estimate SITS with $K = 18$, $\alpha = 1/K$, $\beta = 0.1$, and $\gamma = 1$, as outlined in Appendix C (SI pp. 3–5). Appendix F (SI p. 15) provides details on convergence.

¹⁹I do this exercise blind to the agenda-setting measures.

TABLE 2 Agenda Setters More Likely to Shape Deliberation Outcome

	Introduced group proposal				
	(1)	(2)	(3)	(4)	(5)
Proportion of comments	2.59 (2.18)				
Proportion of talk time		3.76 (2.24)			
Proportion of interruptions			−3.14 (2.67)		
Composite measure				1.97 (1.18)	
Agenda setting					6.00* (2.35)
Strong DGS attitudes	1.50 (0.84)	1.23 (0.99)	0.93 (0.92)	1.61 (0.96)	0.87 (1.03)
Constant	−2.77 (1.51)	−2.94 (1.35)	0.05 (1.37)	−2.60 (1.82)	−2.90* (1.41)
Observations	40	40	40	40	40
AIC	68.59	64.94	69.02	69.59	63.70

Note: * $p < 0.05$. Coefficients from logistic regressions with clustered standard errors at the discussion-group level in parentheses. Dependent variable is introducing an idea included as a group policy proposal ($y = 1$) or not ($y = 0$). Explanatory variables are those from Table 1, each scaled to range from 0 to 1. Each regression includes group indicators.

for each group. The outcome of each model is introducing one of the ideas included as a group policy proposal ($y = 1$) or not ($y = 0$). Each model assesses the correlation between the outcome and a count-based measure of participation or agenda setting, all scaled to range between 0 and 1. It may be the case that those with strongly held DGS attitudes prioritized the outcome of the deliberation more than others and for that reason were more interested in seeking their preferred outcome. Therefore, I also control for pretreatment DGS attitude strength in the models by creating an indicator for strong DGS attitudes.²⁰

First we see the coefficients on each count-based measure are not distinct from zero in Models 1–4. Therefore, I find no evidence of a correlation between the participation measures and success in including one's ideas as policy recommendations. However, the positive coefficient on the agenda-setting measure in Model 5 suggests that participants who succeeded at setting the agenda

were also more likely to have their ideas included as a group recommendation.²¹

The agenda-setting measure has potential payoffs for empirical scholars of deliberation. For example, it can provide new insight into an ideal of deliberative democracy—equal consideration—which pertains to all points of view receiving attention (e.g., Karpowitz and Mendelberg 2014). As a measure of attention at its core, the agenda-setting measure proposed here may aid in identifying the conditions under which this ideal is more or less achieved.

Online Discussions

Lastly, I validate that agenda setting is a source of power available to actors when they interact with each other. To do so, I fielded a novel online discussion study. Importantly, I constructed the discussions to feature disagreement, and I incentivize participants with additional compensation to achieve their preferred outcome.

²⁰Before the discussion, participants were asked to rate their agreement with 24 questions regarding the purpose and fairness of the DGS. I used these questions to create an additive index of pretreatment DGS attitude strength. I then consider a participant to have “strong” attitudes if they fell at or below (above) the first (third) quantile.

²¹The large agenda-setting coefficient makes sense given how I measured the dependent variable—introducing a proposal eventually adopted by the group (the dependent variable) is itself a measure of agenda setting.

Specifically, my goal is to build upon the previous application and more directly test whether setting a discussion's agenda correlates with shaping its outcome.

The study involved four stages. First, participants took an online, prediscussion survey. During the survey, participants learned about five prominent charities and indicated which charity they would prefer to receive a \$1 donation from the researchers.²² At the conclusion of this survey, participants were asked if they would be willing to return for an optional follow-up task at a specific time within the hour. Second, participants who indicated they were willing to return were randomly assigned a partner who disagreed about which charity should receive the donation.²³ Third, participants returned to the online platform at the prespecified time and engaged in a 10-minute online, written discussion with their assigned partner. Fourth, participants answered a short postdiscussion survey.

I fielded the study between April 2016 and June 2020 on MTurk. Participants were paid \$1 for the prediscussion survey and \$3 for completing the follow-up discussion task. Participants also received a \$1 bonus payment if their charity was chosen to receive the researcher's donation in the postdiscussion survey by *both* participants. The bonus payment was intended to incentivize participants to pursue their preferred outcome. In addition to the bonus payment, the participants were further incentivized to agree because the donation would not be made unless participants indicated agreement in the postdiscussion survey.

MTurk is an online labor market increasingly used in social science research to quickly and inexpensively recruit samples (see Berinsky, Huber, and Lenz 2012). It is important to consider the use of the MTurk subject pool for this application. First and foremost, ethical concerns have been raised regarding the compensation of MTurk participants. This work's compensation followed Williamson's (2016) guidance, with payment above federal minimum wage at approximately \$14 per hour. Additionally, research suggests that MTurk participants are not influenced by monitoring cues in their donations to charity (Saunders, Taylor, and Atkinson 2016), are not receptive to experimenter demand effects (Mumolo and Peterson 2019), and are more attentive and cooperative than participants from other online samples (Boas, Christenson, and Glick 2020), making MTurk a

useful subject pool for this specific discussion task requiring a high level of engagement. A limitation of the MTurk subject pool is their high digital literacy, which may moderate what I observe about how the participants interact with each other online (Munger et al. 2021).

This procedure yielded 91 discussions; however, 10 did not reach an agreement. I preprocessed the text by transforming all words to lowercase, stemming, and removing infrequent terms. I also selectively removed punctuation, keeping exclamation marks, question marks, and punctuation meant to mimic emojis. I kept this punctuation because it is an important part of the communication in an online environment.²⁴ I estimated three SITS chains from the data with randomly drawn starting values.²⁵ I use iterations from all three chains to estimate posterior means of the speaker-level agenda-setting parameter.

The agenda-setting measure in this application had a mean of 0.38 and standard deviation of 0.16. Recall the agenda-setting measure is interpreted as the probability a participant will shift topic when speaking. That these participants, on average, shift topic every third speaking turn makes sense in this context. The 10-minute discussions were short, so participants shifted topic fairly rapidly to make sure they came to a decision.

Unlike Application 2, here I am able to more directly assess the discussion's outcome—who's preferred charity was chosen. To test if agenda setting during the discussion correlates with achieving one's preferred outcome, I conduct a two-sided, paired Wilcoxon signed-rank test with the 81 discussions that reached an agreement. Results suggest that the partner who was more successful at setting the discussion's agenda was significantly more likely to achieve their preferred outcome ($p = 0.02$). As a robustness check, I find results of a two-sided, paired t -test are consistent.²⁶

To assess the robustness of this finding, I test if other discussion tactics—speaking first or speaking more—explain who achieved their preferred outcome. I fail to find that speaking first is associated with achieving one's preferred outcome (two-sided t -test, $p = 0.27$). Moreover, I fail to find evidence that suggests the number of words used in the conversation was meaningfully

²²Appendix G (SI pp. 18–19) shows charity information given to participants.

²³Participants answered an open-ended question asking why they chose their preferred charity. Participants were not considered for the discussion stage of the study if this answer was of poor quality.

²⁴See Footnote 13. Appendix B (SI pp. 2–3) shows that results are not likely to be sensitive to the choice to use these common preprocessing steps.

²⁵I estimate SITS with $K = 17$, $\alpha = 1/K$, $\beta = 0.1$, and $\gamma = 1$ as outlined in Appendix C (SI pp. 3–5). Appendix G (SI p. 19) provides details on convergence.

²⁶The mean difference between the partner who achieved their preferred outcome and the partner who did not was 0.06 ($p = 0.02$), which is 0.41 standard deviations.

different for partners that achieved their preferred outcome and those that did not (two-sided, paired Wilcoxon signed-rank test $p = 0.19$).

This application provides evidence that agenda setting is a form of power—participants who set the agenda get what they want out of the discussion. Although these discussions were largely civil and deliberative, with participants engaging with each others' views, political discussions and comments online can often be uncivil (Coe, Kenski, and Rains 2014) or even spread misinformation (Anspach and Carlson 2020). SITS could help extend this research by identifying when comments are derailed and what kinds of people are better at shifting others' attention in a negative way.

Conclusion

In this article, I introduced a measure of agenda setting applicable to the countless interactions that occur across the political sphere. Importantly, I validated the measurement of agenda setting across a diverse set of discursive settings. Debates are oppositional, strategic, and have the goal of identifying a “winner” and “loser.” Conversely, ideal deliberations are characterized by collaboration and thoughtful consideration of all perspectives. Online discussions are often informal, even anonymous, and lack body language or other interpersonal cues to guide the communication. I validated SITS in each of these environments.

There are of course limitations to the measure of agenda setting proposed here. First, this measure cannot explain how an agenda setter maintains attention on newly shifted-to topic. It could be their personality, others' agreement, social norms, or other factors that successfully shifts attention to one's preferred topic. Relatedly, the purpose of this measure is to capture relative agenda-setting power of those engaging in an interaction. Therefore, it may fail to account for broader power dynamics that shape how the interaction unfolds, such as deference to the preferred agenda of a powerful person who is not even in the room.

Moreover, this article intended to validate the agenda-setting measure and stress its importance to the study of political interactions but leaves many questions unanswered regarding the theoretical role of agenda setting in different interactions. For example, under what conditions does agenda-setting power equate to perceived power? Depending on the social norms of the setting, setting the agenda could be perceived as rude, controlling, or irritating to others, so exercising agenda-

setting power may harm an individual's perceived power. Understanding the relationships between agenda setting as a form of power, perceived power, influence, persuasion, and more is left for future work.

SITS measures agenda setting by extending LDA, and future work could extend SITS to measure additional latent quantities of interest. For example, SITS could be altered to account for a speaker's tendency to bring up similar topics over time. If Trump brings up ISIS in the first presidential debate, he is likely to bring it up in future debates. If SITS is extended to account for this, one could further imagine detecting a latent coalition of speakers that shift to similar topic distributions.

Finally, SITS is just one of many automated topic segmentation methods. SITS is particularly useful in an interactive setting because it uses speaker identity to inform the segmentation task. However, other topic segmentation methods have promise for researchers needing to find coherent topic segments in their corpora (see Purver 2011). For example, topic segmentation might be useful for researchers trying to segment a large document into more manageable sizes or theoretically useful quantities to then be coded by crowdsourced workers. Topic segmentation methods provide a principled way to approach tasks like this where the text is not already structured in the most theoretically or practically useful way.

In sum, this article set out to quantify the agenda-setting dynamics that lie under the surface anytime two people talk. Agenda setting in debates, discussions, and deliberations is often overlooked by empirical researchers because we lack methodological tools suited for these settings. This article helps shift the quantitative study of agenda setting to interactions by offering a technique for measuring this important concept across a variety of settings. The hope is that the measure provides opportunity for additional theoretical development and principled analysis regarding the agenda-setting dynamics in political interactions.

References

- Anspach, Nicolas M., and Taylor N. Carlson. 2020. “What to Believe? Social Media Commentary and Belief in Misinformation.” *Political Behavior* 42(3):697–718.
- Aslett, Kevin, Nora Webb Williams, Andreu Casas, Wesley Zuidema, and John Wilkerson. “What Was the Problem in Parkland? Using Social Media to Measure the Effectiveness of Issue Frames.” *Policy Studies Journal*. Forthcoming.
- Bachrach, Peter, and Morton S. Baratz. 1962. “Two Faces of Power.” *American Political Science Review* 56:947–952.

- Baumgartner, Frank R., and Bryan D. Jones. 1993. *Agendas and Instability in American Politics*. Chicago, IL: University of Chicago Press.
- Beck, Paul Allen, Russell J. Dalton, Steven Greene, and Robert Huckfeldt. 2002. "The Social Calculus of Voting: Interpersonal, Media, and Organizational Influences on Presidential Choices." *American Political Science Review* 96(1):57–73.
- Berinsky, Adam J., Gregory A. Huber, and Gabriel S. Lenz. 2012. "Evaluating Online Labor Markets for Experimental Research: Amazon.com's Mechanical Turk." *Political Analysis* 20(3):351–68.
- Black, Ryan C., and Ryan J. Owens. 2009. "Agenda Setting in the Supreme Court: The Collision of Policy and Jurisprudence." *The Journal of Politics* 71(3):1062–75.
- Blei, David M., Andrew Y. Ng, and Michael I. Jordan. 2003. "Latent Dirichlet Allocation." *Journal of Machine Learning Research* 3(Jan):993–1022.
- Boas, Taylor C., Dino P. Christenson, and David M. Glick. 2020. "Recruiting Large Online Samples in the United States and India: Facebook, Mechanical Turk, and Qualtrics." *Political Science Research and Methods* 8(2):232–50.
- Boydston, Amber E., Rebecca A. Glazier, and Claire Phillips. 2013. "Agenda Control in the 2008 Presidential Debates." *American Politics Research* 41(5):863–99.
- Boydston, Amber E., Rebecca A. Glazier, and Matthew T. Pietryka. 2013. "Playing to the Crowd: Agenda Control in Presidential Debates." *Political Communication* 30(2):254–77.
- Chang, Jonathan, Sean Gerrish, Chong Wang, Jordan L. Boyd-Graber, and David M. Blei. 2009. Reading Tea Leaves: How Humans Interpret Topic Models. Y. Bengio, D. Schuurmans, J. Lafferty, C. K. I. Williams, and A. Culotta, editors. In *Advances in Neural Information Processing Systems*. 288–96. Curran Associates, Inc.
- Cobb, Charles D., and Roger W. Elder. 1972. *Participation in American Politics: The Dynamics of Agenda Building*. Boston, MA: Allyn and Bacon.
- Coe, Kevin, Kate Kenski, and Stephen A. Rains. 2014. "Online and uncivil? Patterns and Determinants of Incivility in Newspaper Website Comments." *Journal of Communication* 64(4):658–79.
- Cox, Gary W., and Mathew D. McCubbins. 2005. *Setting the Agenda: Responsible Party Government in the US House of Representatives*. Cambridge: Cambridge University Press.
- Denny, Matthew J., and Arthur Spirling. 2018. "Text Preprocessing for Unsupervised Learning: Why It Matters, When It Misleads, and What to Do about It." *Political Analysis* 26(2):168–89.
- Eggers, Andrew C., and Arthur Spirling. 2018. "The Shadow Cabinet in Westminster Systems: Modeling Opposition Agenda Setting in the House of Commons, 1832–1915." *British Journal of Political Science* 48(2):343–367.
- Epstein, Lee, William M. Landes, and Richard A. Posner. 2010. "Inferring the Winning Party in the Supreme Court from the Pattern of Questioning at Oral Argument." *The Journal of Legal Studies* 39(2):433–67.
- Grimmer, Justin. 2010. "A Bayesian Hierarchical Topic Model for Political Texts: Measuring Expressed Agendas in Senate Press Releases." *Political Analysis* 18(1):1–35.
- Grimmer, Justin, and Gary King. 2011. "General Purpose Computer-Assisted Clustering and Conceptualization." *Proceedings of the National Academy of Sciences* 108(7):2643–50.
- Johnson, Timothy R., Paul J. Wahlbeck, and James F. Spriggs. 2006. "The Influence of Oral Arguments on the US Supreme Court." *American Political Science Review* 100(1):99–113.
- F Karpowitz, Christopher, and Tali Mendelberg. 2014. *The Silent Sex: Gender, Deliberation, and Institutions*. Princeton, NJ: Princeton University Press.
- Karpowitz, Christopher F., Tali Mendelberg, and Lee Shaker. 2012. "Gender Inequality in Deliberative Participation." *American Political Science Review* 106(3):533–47.
- Kathlene, Lyn. 1994. "Power and Influence in State Legislative Policymaking: The Interaction of Gender and Position in Committee Hearing Debates." *American Political Science Review* 88(3):560–76.
- Kingdon, John W. 1984. *Agendas, Alternatives, and Public Policies*. 2nd ed. New York: HarperCollins.
- Levin, Dan. 2019. "Brigham Young Students Value Their Strict Honor Code. But Not the Harsh Punishments." *New York Times*. April 12. <https://www.nytimes.com/2019/04/12/us/byu-honor-code.html>. Last accessed on May 10, 2020.
- Mason, Melanie. 2016. "The Most Memorable Moments from the First Presidential Debate." *Los Angeles Times*. September 26. <https://www.latimes.com/politics/la-na-pol-debate-moments-20160926-snap-htmlstory.html>. Last accessed on September 25, 2019.
- McCombs, Maxwell E., and Donald L. Shaw. 1972. "The Agenda-Setting Function of Mass Media." *Public Opinion Quarterly* 36(2):176–87.
- Mikhaylov, Slava, Michael Laver, and Kenneth R. Benoit. 2012. "Coder Reliability and Misclassification in the Human Coding of Party Manifestos." *Political Analysis* 20(1):78–91.
- Mummolo, Jonathan, and Erik Peterson. 2019. "Demand Effects in Survey Experiments: An Empirical Assessment." *American Political Science Review* 113(2):517–29.
- Munger, Kevin, Ishita Gopal, Jonathan Nagler, and Joshua Tucker. 2021. "Accessibility and Generalizability: Are Social Media Effects Moderated by Age or Digital Literacy?" *Research & Politics*. 8(2):1–16.
- Nguyen, Viet-An. 2014. "Speaker Identity for Topic Segmentation (SITS)." GitHub repository. <https://github.com/vietansean/sits>. Last accessed September 25, 2019.
- Nguyen, Viet-An. 2015. "Guided Probabilistic Topic Models for Agenda-Setting and Framing." PhD dissertation. University of Maryland, College Park. <https://drum.lib.umd.edu/handle/1903/16600>.
- Nguyen, Viet-An, Jordan Boyd-Graber, and Philip Resnik. 2012. SITS: A Hierarchical Nonparametric Model Using Speaker Identity for Topic Segmentation in Multiparty Conversations. In *Proceedings of the 50th Annual Meeting of the Association for Computational Linguistics: Long Papers—Volume 1*. Association for Computational Linguistics pp. 78–87.
- Nguyen, Viet-An, Jordan Boyd-Graber, Philip Resnik, Deborah A. Cai, Jennifer E. Midberry, and Yuanxin Wang. 2014. "Modeling Topic Control to Detect Influence in

- Conversations Using Nonparametric Topic Models.” *Machine Learning* 95(3):381–421.
- Petrocik, John R. 1996. “Issue Ownership in Presidential Elections, with a 1980 Case Study.” *American Journal of Political Science*. 40(3):825–50.
- Purver, Matthew. 2011. Topic Segmentation. In Gokhan Tur and Renato De Mori, editors, *Spoken Language Understanding: Systems for Extracting Semantic Information From Speech*. Wiley, Hoboken NJ, 291–317.
- Quinn, Kevin M., Burt L. Monroe, Michael Colaresi, Michael H. Crespin, and Dragomir R. Radev. 2010. “How to Analyze Political Attention with Minimal Assumptions and Costs.” *American Journal of Political Science* 54(1):209–28.
- Riker, William H. 1986. *The Art of Political Manipulation*. New Haven, CT: Yale University Press.
- Roberts, Margaret E., Brandon M. Stewart, Dustin Tingley, Christopher Lucas, Jetson Leder-Luis, Shana Kushner Gadarian, Bethany Albertson, and David G. Rand. 2014. “Structural Topic Models for Open-Ended Survey Responses.” *American Journal of Political Science* 58(4):1064–1082.
- Ross, Janell. 2016. “Trump on ‘Fat Slobs,’ Housekeepers and Women who Don’t Have that ‘Presidential Look’.” *The Washington Post*. September 27. <https://www.washingtonpost.com/news/the-fix/wp/2016/09/27/trump-on-fat-slobs-housekeepers-and-women-who-dont-have-that-presidential-look/>. Last accessed on September 25, 2019.
- Saunders, Timothy J., Alex H. Taylor, and Quentin D. Atkinson. 2016. “No Evidence that a Range of Artificial Monitoring Cues Influence Online Donations to Charity in an MTurk Sample.” *Royal Society Open Science* 3(10):150710.
- Schattschneider, Elmer E. 1960. *The Semi-Sovereign People: A Realist’s View of Democracy in America*. Belmont, CA: Wadsworth Publishing.
- Sides, John, and Lynn Vavreck. 2014. *The Gamble: Choice and Chance in the 2012 Presidential Election*. Princeton, NJ: Princeton University Press.
- Sides, John, Michael Tesler, and Lynn Vavreck. 2019. *Identity Crisis: The 2016 Presidential Campaign and the Battle for the Meaning of America*. Princeton, NJ: Princeton University Press.
- Turkewitz, Julie. 2014. “At Brigham Young, Students Push to Lift Ban on Bears.” *New York Times*. November 17. <https://www.nytimes.com/2014/11/18/us/campaigning-to-change-the-cleanshaven-look-at-brigham-young-university.html>. Last accessed on May 10, 2020.
- Vavreck, Lynn. 2009. *The Message Matters: The Economy and Presidential Campaigns*. Princeton, NJ: Princeton University Press.
- Wang, Lu, Nick Beauchamp, Sarah Shugars, and Kechen Qin. 2017. “Winning on the Merits: The Joint Effects of Content and Style on Debate Outcomes.” *Transactions of the Association for Computational Linguistics* 5:219–32.
- Weaver, David H. 2007. “Thoughts on Agenda Setting, Framing, and Priming.” *Journal of Communication* 57(1):142–47.
- Williamson, Vanessa. 2016. “On the Ethics of Crowdsourced Research.” *PS: Political Science & Politics* 49(1):77–81.
- Ying, Luwei, Jacob M. Montgomery, and Brandon M. Stewart. 2019. “Inferring Concepts from Topics: Towards Procedures for Validating Topics as Measures.” *PolMeth XXXVI*, Cambridge, MA. Society for Political Methodology. Last accessed on June 5, 2020. <https://polmeth.mit.edu/sites/default/files/documents/YingMontgomeryStewart%20-%20Main.pdf>.

Supporting Information

Additional supporting information may be found online in the Supporting Information section at the end of the article.

Appendix A: Sampler

Appendix B: Text preprocessing

Appendix C: Hyperparameters and model selection

Appendix D: Validation exercises

Appendix E: Presidential debates

Appendix F: In-person deliberations

Appendix G: Online discussions